



nr 5. maj 2025

www.mlodytechnik.pl



Tu przejrzyj
i kupisz ten numer

NEWS 24/7
przeładowanie
na swoim smartfonie

mlody **m.technik**

Ciekawi świata są zawroty głowy



CHIPTELIGENCJA

Krzemowo-komputerowy zawrót głowy



ISSN 0462-9760 Indeks 365408

9 477046219762501

cena: **14,90 zł** (w tym 8% VAT)

Sieci neuronowe typu transformer, część 3

Normalizacja

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

Czy to już ostatnie soki, jakie da się wycisnąć z krzemu? Jakie są alternatywy? Co oznacza sztuczna inteligencja dla branży półprzewodników i komputerów?

Półprzewodnikowy przedmiot pożądania

Wojny krzemowe wkraczają na nowy etap. Takie zdanie można napisać bezpiecznie w dowolnym momencie i będzie prawdziwe. Raz po raz bowiem pojawiają się nowe obszary i fronty rywalizacji w dziedzinie techniki półprzewodnikowej.

Toczone są zresztą nie tylko pomiędzy mocarstwami: Stanami Zjednoczonymi, Chinami, Rosją, także Europą, która stara się nadrobić to, co przez dekady zaniedbała. Równie bezwzględna jest konkurencja pomiędzy producentami, która przecina granice zarówno wrogich sobie, jak i sojusznicznych państw. Czasem idzie w parze z po-

lityką, najczęściej wtedy, gdy politycy to wymuszają, ale czasem nie.

Fronty wojen procesorowych przecinają granice państw

O tym, że w wojnach tych linie frontu przebiegają w sposób skomplikowany, nie zawsze taki, jak mogłoby się nam wydawać,

dobrze świadczy kontrowersja związana z wykluczeniem przez USA Polski z pakietu krajów, do których dozwolony jest eksport najbardziej zaawansowanych chipów AI. Nas to bulwersowało, jednak wyjaśnienie tej decyzji jest logiczne i oparte na danych, co nie znaczy, oczywiście, że miłe dla Polski.

À propos tzw. chipów AI. Właśnie ogólnie tak określony obszar technologiczny budzi najwięcej emocji. Panuje przekonanie, że przodownictwo i dominacja w dziedzinie sztucznej inteligencji zapewni dominację w ogóle, czyli także polityczną, militarną i ekonomiczną. Do tego, jak się uważa, trzeba mieć dostęp do najbardziej zaawansowanych technik procesorowych. Jednak przekonanie to podkopało wejście chińskiego DeepSeeka, dającego porównywalne do czołowych amerykańskich modeli wyniki, przy domniemanym wykorzystaniu „gorszego” sprzętu do szkolenia (do „lepszego” Chiny oficjalnie dostępu nie mają z powodu sankcji USA).

Rzecz cała wymaga jeszcze pełnego wyjaśnienia, ale w środowiskach zajmujących się rozwojem AI od dłuższego czasu nie jest tajemnicą, że wcale nie chodzi o to, by kupić jak najwięcej drogiego sprzętu, kart NVIDIA, do szkolenia modeli, lecz by osiągnąć lepsze wyniki przy zużywaniu mniejszych zasobów, pieniędzy, sprzętu, energii i przede wszystkim danych. I ten wymiar wojen krzemowych wydaje się w najbliższych latach najważniejszy.

Mirosław Usidus

Spis treści

Temat numeru: Chipteligencja.

Krzemowo-komputerowy zawrót głowy

- 26 • AI w świecie komputerów. Czy nadchodzi era super „Alcetów”?
- 32 • Czy to już ostatnie soki, jakie da się wycisnąć z krzemu? Jakże są alternatywy? Od prawa Moore’a do pojedynczych atomów
- 38 • Komputery kwantowe – są czy ich nie ma i czy w ogóle są potrzebne? W kubitach płata się więcej pytań niż odpowiedzi
- 46 • Wojny krzemowe – kolejny etap. Bój do ostatniego nanometra

Technika

- 8 Info Zoom
- 16 Dodaj do obserwowanych Horyzonty mgłą spowite
- 17 • Czy okulary AR mogą już zastąpić smartfony? Inteligencja na nosie
- 21 • Prawdopodobieństwo uderzenia asteroidy w 2032 roku. Ferie z armagedonem
- 23 • Teoria martwego internetu żywa jak nigdy. Maszyny przejęły naszą sieć?

m.technik

- 52 Mobilne aplikacje. Test aplikacji: Fitness i zdrowie

Fantastyka naukowa w „Młodym Techniku”

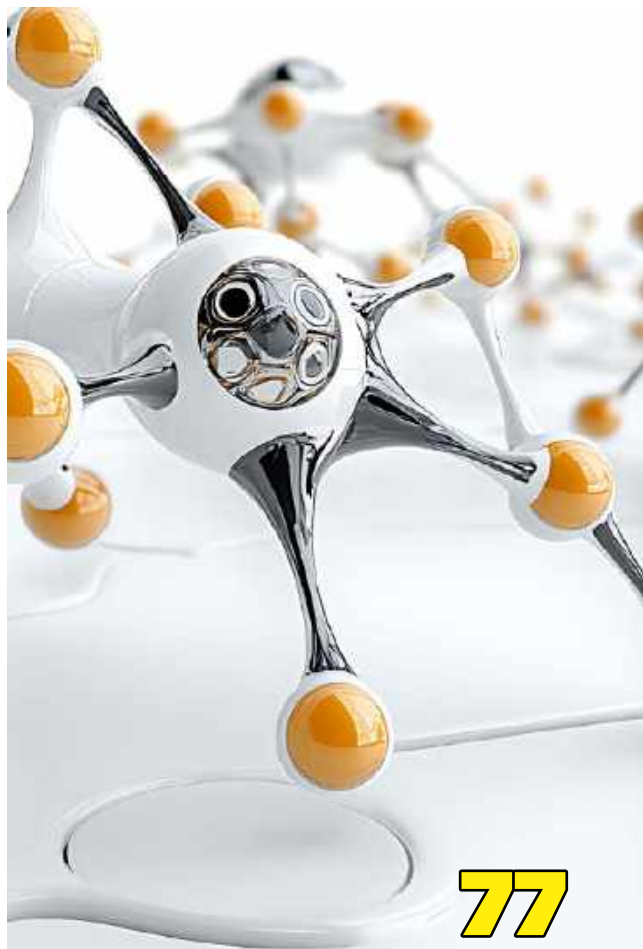
- 54 Raport z obserwacji DM-S-3

Szkoła

- 58 MT studiuje: Inżynieria materiałowa
- 60 Chemia inna niż w szkole: Drugie życie plastiku, część 3
- 64 Fizyka bez granic: Czy kwarki naprawdę są kolorowe?
- 68 Koniec i co dalej: Filmy na płytach. Chmury zamiast regatów
- 71 Pomysły genialne, zwirowane i takie sobie
- 72 Matematyka z ludzką twarzą: Do Lizbony przez Stambuł i Alaskę
- 77 Sieci neuronowe typu transformer, część 3. Normalizacja i połączenia rezydualne, interpretacja działania bloków feed-forward, blok dekodera i mechanizm uwagi krzyżowej
- Klub i Szkoła Wynalazców
- 84 • Szkoła Wynalazców – dozwolone do lat 15
- 84 • Klub Wynalazców – bez ograniczeń wieku
- 85 • Vademecum Młodego Wynalazcy
- 88 Na warsztacie: Silnik zaburtowy
- Odkryj historię wynalazków
- 92 • Cyfrowe bliźniaki
- 96 • Klasyfikacja i najbardziej znane zastosowania cyfrowych bliźniaków

Hobby

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



77

Sieci neuronowe typu transformer

W tym wydaniu MT m.in.:

- **Horyzonty mgłą spowite:**
Inteligencja na nosie.
Czy okulary AR mogą już zastąpić smartfony?
- **Koniec i co dalej:**
Filmy na płytach. Chmury zamiast regatów
- **Test aplikacji:**
Fitness i zdrowie

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:
Mirosław 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.mlodYTECHNIK.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

Chipteligencja

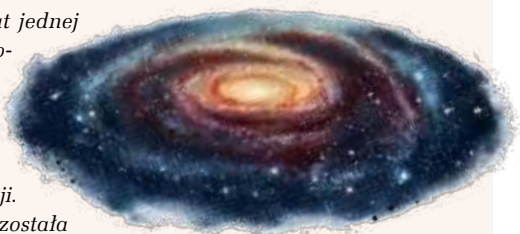
Krzemowo-komputerowy zawrót głowy

Zapowiedzi „końca fizycznej możliwości miniaturyzacji” układów półprzewodnikowych opartych na krzemie i rozważania, czy da się jeszcze coś z krzemu wycisnąć, opisujemy w MT od lat. Za każdym razem okazuje się, że da się coś jeszcze zrobić, zmniejszyć i upakować, podnosząc wydajność na nowe poziomy, mimo ogólnej zgody, że prawo Moore’a już nie obowiązuje.

List miesiąca

Uniwersum jako symulacja

Poruszyliście Państwo na swoich łamach niedawno temat jednej z najbardziej intrygujących koncepcji współczesnej nauki i filozofii – hipotezy Wszechświata jako symulacji. Teoria ta, choć kontrowersyjna, zyskuje coraz większą popularność zarówno w kręgach naukowych, jak i w kulturze popularnej. Chciałbym przedstawić to, jak ja widzę argumenty za i przeciw tej hipotezie, aby umożliwić czytelnikom głębsze zrozumienie jej implikacji.



Hipoteza symulacji, znana również jako „teoria symulacji”, została po raz pierwszy sformułowana przez filozofa Nicka Bostroma w 2003 roku.

Bostrom przedstawił argument, że jedna z trzech możliwości musi być prawdziwa: żadna cywilizacja nie osiągnie poziomu technologicznego, który pozwoliłby na stworzenie realistycznych symulacji rzeczywistości; cywilizacje, które osiągną taki poziom technologiczny, nie będą zainteresowane tworzeniem symulacji rzeczywistości; żyjemy w symulacji¹. Argument Bostroma opiera się na założeniu, że jeśli zaawansowane cywilizacje są w stanie tworzyć symulacje rzeczywistości, to prawdopodobieństwo, że żyjemy w jednej z takich symulacji, jest niezwykle wysokie.

Jednym z kluczowych elementów hipotezy symulacji jest założenie, że zaawansowane cywilizacje będą w stanie stworzyć symulacje, które są nie do odróżnienia od rzeczywistości. Współczesne technologie komputerowe, takie jak wirtualna rzeczywistość (VR) i rozszerzona rzeczywistość (AR), pokazują, że jesteśmy na dobrej drodze do tworzenia coraz bardziej realistycznych symulacji. Jednakże, aby stworzyć symulację na poziomie całego wszechświata, potrzebne byłyby zasoby i technologie, które obecnie są poza naszym zasięgiem.

Projekty takie jak ****SIBELIUS-DARK**** dowodzą, że współczesna nauka potrafi modelować ewolucję Wszechświata od Wielkiego Wybuchu z wykorzystaniem superkomputerów. Symulacja objęła 130 miliardów cząstek, odtwarzając struktury na skalę 600 milionów lat świetlnych. Choć to kropla w morzu wobec pełnego kosmosu, pokazuje technologiczny potencjał.

Hipoteza symulacji niesie ze sobą szereg implikacji filozoficznych, naukowych i etycznych, które warto rozważyć. Filozoficzne implikacje stawiają pod znakiem zapytania nasze rozumienie rzeczywistości. Jeśli żyjemy w symulacji, to co to oznacza dla naszej tożsamości, wolnej woli i sensu życia? Czy nasze doświadczenia są mniej wartościowe, jeśli są wynikiem symulacji? Te pytania mają głębokie implikacje dla filozofii umysłu i ontologii.

Z drugiej strony, istnieją również argumenty przeciwko hipotezie symulacji. Jednym z głównych argumentów jest to, że modelowanie fizyki naszego wszechświata na nawet największym komputerze jest niemożliwe. Naukowcy odkryli, że nie da się odwzorować wszystkich aspektów rzeczywistości w symulacji komputerowej, co sugeruje, że prawdopodobnie nie żyjemy w symulacji.

Innym argumentem przeciwko hipotezie symulacji jest to, że nasze doświadczenia i percepcje są zbyt złożone, aby mogły być wynikiem symulacji. Argumenty sceptyczne, takie jak solipsyzm czy „mózg w naczyniu”, również podważają hipotezę symulacji, sugerując, że nasze doświadczenia mogą być wynikiem manipulacji przez zewnętrzne siły.

W końcu – jeśli symulacja ma być spójna, dlaczego jej twórcy pozwoliliby nam rozważać jej istnienie? To pytanie pozostaje bez odpowiedzi, podważając logiczną integralność koncepcji.

Hipoteza Wszechświata jako symulacji jest fascynującą koncepcją, która stawia przed nami wiele pytań dotyczących naszej rzeczywistości. Choć istnieją przekonujące argumenty za tym, że żyjemy w symulacji, istnieją również silne argumenty przeciwko tej teorii. Z jednej strony, postęp w symulacjach kosmologicznych i teorii informacyjne Vopsona nadają jej pozory naukowej legitymacji. Z drugiej – brak empirycznego potwierdzenia i fundamentalne ograniczenia technologiczne skłaniają do ostrożności. Warto jednak kontynuować badania nad związkiem fizyki z teorią informacji, które mogą przynieść przełomy niezależnie od statusu naszej rzeczywistości.

Z wyrazami szacunku,

Marcin Surdut, Sękocin

Zmierzch manualnej skrzyni biegów w aucie

Dziękuję, że poruszyliście na łamach „Młodego Technika” temat, który może nie jest specjalnie nieznany, ale warto odnotować to zjawisko i upamiętnić odejście czegoś, na czym się jako kierowca „wychowałem”.

Dzisiejsze skrzynie automatyczne oferują więcej biegów (8–10 przełożeń), szybsze zmiany przełożeń oraz porównywalną, a często nawet lepszą efektywność paliwową niż ich manualne odpowiedniki. Współczesne automatyczne skrzynie biegów, takie jak dual-clutch (DCT) czy bezstopniowe (CVT), dorównują manualnym pod względem sprawności i dynamiki, a nawet je przewyższają. Według analiz NHTSA i EPA, 8-biegowe automaty redukują zużycie paliwa o 6,5–7,8% w porównaniu do 4-biegowych, podczas gdy DCT oferują oszczędności sięgające 12,6%.

Rosnąca popularność napędów elektrycznych i hybrydowych stanowi istotny czynnik zmniejszający zapotrzebowanie na manualne skrzynie biegów. Silniki elektryczne charakteryzują się spłaszczoną krzywą momentu obrotowego i nie wymagają złożonych wielobiegowych przekładni. Większość pojazdów elektrycznych wykorzystuje jednobiegowe przekładnie redukcyjne, eliminując potrzebę manualnej zmiany biegów. W miarę postępującej elektryfikacji floty udział manualnych skrzyń będzie naturalnie spadał.

Rozwój systemów wspomagających kierowcę (ADAS) oraz dążenie do autonomicznej jazdy wymuszają stosowanie przekładni automatycznych. Manualna zmiana biegów jest niekompatybilna z systemami autonomicznymi i stanowi przeszkodę w implementacji zaawansowanych funkcji bezpieczeństwa, jak adaptacyjny tempomat czy asystenci jazdy w korku.

Obserwujemy też globalną zmianę preferencji konsumenckich. Nowe pokolenia kierowców, dorastające w erze smartfonów i automatyzacji, częściej postrzegają prowadzenie jako konieczność, a nie przyjemność. Dla nich wygoda automatycznej skrzyni jest istotniejsza niż satysfakcja z manualnej kontroli. Trend ten widoczny jest nawet na tradycyjnie „manualnych” rynkach europejskich.

Producenci samochodów, dążąc do optymalizacji kosztów produkcji, redukują różnorodność oferowanych podzespołów. Malejący popyt na skrzynie manualne sprawia, że ich rozwój i produkcja stają się coraz mniej

opłacalne. W rezultacie powstaje efekt spirali – mniejsza dostępność zwiększa koszty jednostkowe, co dodatkowo ogranicza popyt.

Patrząc w przyszłość, możemy przewidzieć kilka prawdopodobnych scenariuszy.

W najbliższej dekadzie należy spodziewać się dalszego udoskonalania przekładni automatycznych. Prawdopodobnie zobaczymy powszechne zastosowanie skrzyń 10-biegowych oraz dalszy rozwój technologii dwusprzęgłowych. Przekładnie będą ściślej zintegrowane z systemami zarządzania energią pojazdu, wykorzystując dane z nawigacji satelitarnej, radarów i czujników do predykcyjnego doboru optymalnych przełożeń.

Interesującym kierunkiem rozwoju mogą być systemy hybrydowe, łączące elementy przekładni manualnych i automatycznych. Przykładem są już istniejące rozwiązania z automatycznym sprzęgłem, pozwalające kierowcy na wybór przełożeń bez obsługi pedału sprzęgła. Możliwe jest pojawienie się nowych koncepcji interfejsów użytkownika, umożliwiających większą kontrolę nad automatyczną skrzynią biegów dla kierowców ceniących sportowy styl jazdy.

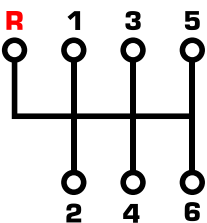
Mimo ogólnego trendu spadkowego, manualne skrzynie biegów prawdopodobnie przetrwają w niszowych zastosowaniach. Można przewidywać, że przez długi czas będą dostępne w samochodach sportowych i entuzjastycznych, gdzie bezpośrednia kontrola nad pojazdem stanowi kluczowy element doświadczenia z jazdy. Paradoksalnie, manualne skrzynie mogą ewoluować w kierunku dóbr luksusowych, dostępnych za dopłatą w modelach premium.

Możliwe jest pojawienie się systemów symulujących wrażenia z manualnej zmiany biegów przy zachowaniu faktycznie automatycznego działania przekładni. Łopatkowe manetki przy kierownicy to zaledwie początek tej tendencji. Przyszłe systemy mogą oferować zaawansowane sprzężenie zwrotne, symulujące opór dźwigni zmiany biegów czy wibracje charakterystyczne dla mechanicznych przekładni.

Jak każda transformacja technologiczna, proces ten niesie zarówno straty, jak i korzyści. Tracimy pewien aspekt bezpośredniej interakcji z maszyną i specyficzną umiejętności motorycznej, zyskujemy natomiast większy komfort, bezpieczeństwo i efektywność energetyczną. Rozwój technologiczny rzadko podąża sentymentalną ścieżką – kieruje się raczej optymalizacją parametrów uznanych w danym czasie za priorytetowe.

Z poważaniem,

Karol Churkowski z Nowego Tomyśla





SYNTEZA TERMOJĄDROWA

22 minuty – nowy rekord stabilności termojądrowej we francuskim tokamaku

Francuski reaktor tokamak CEA (fr. Skrót od „Commissariat à l'énergie atomique et aux énergies alternatives”) WEST pobił rekord w utrzymywaniu plazmy, w której zachodzą reakcje termojądrowe, osiągając czas 22 minut (1337 sekund). Poprzedni rekord należał do chińskiego reaktora EAST, który utrzymał plazmę przez 1066 sekund, czyli bez mała 18 minut.

Jeśli chodzi o syntezę termojądrową, to głównym wyzwaniem jest nie sama fuzja atomów, lecz stworzenie odpowiednich warunków, w których reakcja syntezy jądrowej będzie stabilnie sama się podtrzymywać, produkując energię netto. Oznacza to osiągnięcie temperatury od 100 do 150 milionów stopni Celsjusza, ciśnienia od 5 do 10 atmosfer w punkcie reakcji i utrzymanie wysokoenergetycznej plazmy na stabilnym poziomie przez co najmniej 10 sekund. Wprawdzie w CEA trwało to dłużej, ale nie znaczy to jeszcze, że całość procesu jest opanowana.

„WEST osiągnął nowy kamień milowy w technologii, utrzymując plazmę wodorową przez ponad dwadzieścia minut przez wstrzyknięcie 2 MW mocy grzewczej”, powiedziała w oświadczeniu Anne-Isabelle Etienne, dyrektor ds. badań podstawowych CEA. Choć WEST nigdy nie stanie się reaktorem komercyjnym, zebrane tam informacje zostaną wykorzystane do ulepszenia bardziej ambitnych maszyn, takich jak budowany od szeregu lat na południu Francji wielki Międzynarodowy Eksperymentalny Reaktor Termojądrowy (ITER). ■



Wycieczka po tokamaku
WEST: <https://youtu.be/B1f4QwqlvTU>



Sonda Blue Ghost 1 amerykańskiej firmy Firefly Aerospace wylądowała z powodzeniem na Księżycu. Firma podała, że lądownik, który dotknął księżycowego gruntu w rejonie Mare Crisium, w północno-wschodniej części widocznej strony Księżycza znajduje się w „pionowej, stabilnej” pozycji. To pierwsza misja księżycowa Firefly Aerospace w ramach inicjatywy NASA Commercial Lunar Payload Services (CLPS) i programu powrotu na Księżyc, Artemis. Niestety, o sukcesie lądowania nie może mówić inna prywatna firma, Intuitive Machines, której sonda Athena przewróciła się (po raz drugie w ciągu dwóch lat) podczas lądowania.

Blue Ghost 1 wystartowała w styczniu 2025 roku wspólnie z lądownikiem księżycowym Resilience japońskiej firmy ispace. Na pokładzie lądownika znalazło się dziesięć instrumentów naukowych i technologicznych NASA, które według planu mają działać na powierzchni Księżycza przez jeden dzień księżycowy, co odpowiada około czternastu dniom ziemskim. Firefly Aerospace opracowuje dwie kolejne misje lądownika księżycowego. Firma zdobyła w 2023 r. finansowanie CLPS za Blue Ghost 2, misję lądownika księżycowego na drugą stronę Księżycza, która dostarczy również misję Lunar Pathfinder europejskiej agencji ESA na orbitę wokół Księżycza. NASA przyznała Firefly także fundusze na w ramach CLPS na budowę



KOSMOS

Księżycowe lądowania ze zmiennym szczęściem, eksplozja Starshipa i powrót astronautów z ISS

Blue Ghost 3, lądownika, który ma polecieć do regionu Gruithuisen Domes na widocznej stronie Księżyca.

Niepowodzeniem (częściowym, gdyż dolny człon Super Heavy wylądował pewnie w „szczypcach chwytających”) zakończyła się ósma misja testowa rakiety Starship firmy SpaceX. Wydarzyła się rzecz podobna do tego, co stało się podczas siódmego testu. Po kilku minutach kontrola misji straciła kontakt z rakietą lecącą na orbitę i ponownie doszło do eksplozji, której widowiskowe skutki widać było na niebie. Udany był za to test nowej europejskiej rakiety Ariane 6 w pierwszej jej komercyjnej misji, w której umieściła na orbicie francuskiego satelitę wojskowego. Ta rakietka to krok do europejskiej niezależności w dziedzinie

wynoszenia na orbitę. Musi jednak udowodnić, że korzystanie z niej jest kosztowo opłacalne w porównaniu do transportu, jaki oferuje SpaceX. Ta firma po ciosach jakimi były dwie eksplozywne misje Starshipa, zanotowała w końcu sukces, jakim było ściągnięcie na Ziemię w kapsule Dragon astronautów Suni Williams i Butcha Wilmore’a, którzy utknęli na Międzynarodowej Stacji Kosmicznej po tym, jak NASA nie zgodziła się na ich powrót w przeciekającej kapsule Starliner Boeinga, w czerwcu 2024 r. ■



Montaż filmów ukazujących spadające przez atmosferę szczątki Starshipa po eksplozji podczas ósmego testu: <https://www.youtube.com/shorts/TIZK-AOON4A>



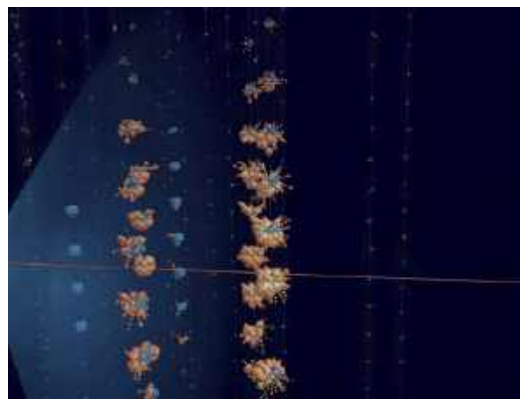
OBRAZOWANIE

Kwantowa technika pozwala uchwycić moment powstania życia – na żywo

Dzięki kamerom zaprojektowanym do pomiarów efektów kwantowych badacze z Uniwersytetu w Adelaide w Australii zdołali uchwycić obrazy rozwijających się embrionów, minimalizując uszkodzenia, które zwykle powodują konwencjonalne techniki obrazowania wykorzystujące oświetlenie. Osiągnięcie to, opisane w „APL Photonics”, stanowi przełom na styku fizyki kwantowej i biologii, umożliwiając obserwację najwcześniejszych etapów życia po raz pierwszym w tak dokładny sposób.

Jeden z autorów pracy badawczej, Kishan Dholakia, podkreśla w publikacji, że uszkodzenia powodowane oświetleniem są poważnym problemem w technikach obrazowania zjawisk w dziedzinie mikrobiologii. Dlatego on i jego koledzy skupili się na optymalizacji ustawień kamery i zastosowaniu technik detekcji inspirowanej dziedzina mechaniki kwantowej. Testując trzy różne typy kamer, ORCA-Flash V3, ORCA-Quest i iXON Life 888 w obserwacjach embrionu myszy, zespół dostosowywał różne ustawienia, znacznie poprawiając jakość obrazu bez zwiększania intensywności światła.

Zastosowane w tych badaniach ulepszenia opierają się do fundamentalnej naturze światła. Przy ekstremalnie niskich natężeniach światło zachowuje się bardziej jak pojedyncze cząstki (fotony) niż fale, co bezpośrednio wynika z mechaniki kwantowej. Kamery wykorzystujące efekty kwantowe mogą wykrywać te pojedyncze pakiety energii świetlnej w każdym pikselu, umożliwiając uzyskanie niespotykanej dotąd czułości. Planowane dalsze prace zespołu z australijskiej uczelni mają wkroczyć dalej w technikę obrazowania kwantowego, w którym stany cząstek światła dostarczą jeszcze dokładniejszych danych i pozwolą na obserwacje o nieznanej dotychczas precyzji. ■



WSZECHŚWIAT

Tak potężnego neutrino jeszcze nie złapaliśmy

Dwudziestokrotnie większą energię niż jakiekolwiek zarejestrowane dotychczas neutrino kosmiczne, czyli około 220 milionów miliardów elektronowoltów, miało neutrino, które uderzyło i przebiło się przez Morze Śródziemne. Zdarzenie to miało miejsce jeszcze w lutym 2023 i zarejestrowane zostało przez częściowo zbudowany Teleskop Neutrinowy Sześciennego Kilometra (KM3NeT). Jednak dopiero po ponad dwóch latach, po analizach i weryfikacji, naukowcy zdecydowali się opublikować o nim informację. Opracowanie ukazało się w „Nature”.

To, co zarejestrowano w detektorze, to mion poruszający się niemal równoległe do horyzontu o wielkiej energii. Na podstawie jego energii i trajektorii naukowcy KM3NeT ustalili, że musiał on zostać wyemitowany przez zderzenie z udziałem neutrino pochodzącego z przestrzeni kosmicznej, a nie przez cząstkę z atmosfery. Symulacje sugerują, że energia neutrino wynosiła około 220 petaelektronowoltów. Poprzedni rekordzista tego typu miał około dziesięć petaelektronowoltów.

Według fizyków, tak potężne energetycznie neutrino mogą pochodzić z supermasywnych czarnych dziur w aktywnych jądrach galaktyk. Detekcja tych cząstek, które bardzo słabo wchodzi w interakcje z materią, wymaga gigantycznych instalacji takich jak IceCube, zbudowanych z czujników zamkniętych w lodzie Antarktydy, lub zanurzonych w wodzie, jak KM3NeT. Ten ostatni to właściwie dwa detektory neutrin, jeden u wybrzeży Sycylii, drugi w pobliżu południowej Francji. Są nadal w budowie, ale już zbierają dane. ■



ASTROFIZYKA

Czy ciemna energia słabnie? Jeśli tak, to mamy naukową rewolucję

Ciemna energia, uznawana za główny, stale oddziałujący czynnik wpływający na ekspansję Wszechświata, może nie jest wcale taka ważna i wcale nie ma charakteru stałego. Odczyty z Instrumentu Spektroskopowego Ciemnej Energii (DESI), opublikowane w marcu 2025 roku, wskazują, że oddziaływanie ciemnej energii nie jest tak silne, jak w przeszłości. Naukowcy uważali, że niekończąca się ekspansja Wszechświata ostatecznie doprowadzi do energetycznej śmierci Wszechświata. Jeśli jednak, jak wskazują nowe wyniki, ciemna energia słabnie z czasem i wygaśnie za miliardy lub biliony lat, Wszechświat może spaść na siebie i zakończyć się odwrotnością Wielkiego Wybuchu, czyli tzw. „Wielkim Skurczem”.

DESI wykorzystuje 4-metrowy teleskop Mayalla w Kitt Peak National Observatory w Arizonie do mapowania Wszechświata w 3D, mając wgląd w jego przeszłość do 11 miliardów lat wstecz. Międzynarodowy zespół badaczy pod wodzą laboratorium w Berkeley, należącym do Departamentu Energii Stanów Zjednoczonych, w ciągu pierwszych trzech lat swojej pięcioletniej misji DESI zmapował prawie piętnaście milionów galaktyk i kwazarów. I właśnie odczyty, wzbudziły wielkie poruszenie w środowisku astrofizyków. Instrument DESI służy do pomiarów efektów znanych jako

barionowe oscylacje akustyczne, które służą do ujawniania rozkładu materii w różnych epokach w historii Wszechświata. Odczyty służą jako „standardowy probiez” do pomiaru siły ciemnej energii.

Hipotetyczna forma energii, zwana „ciemną energią”, ma wypełniać według od lat dominujących poglądów całą przestrzeń, stanowiąc ponad 68 proc. Wszechświata, i wywierać na nią ujemne ciśnienie, przyspieszając tempo rozszerzania się Wszechświata. Badacze przed laty spostrzegli, że efekty ciemnej energii można najłatwiej wyjaśnić, powołując się na stałą kosmologiczną, którą w 1917 r. zaproponował Albert Einstein, choć później porzucił tę koncepcję i podobno nazwał ją swoim „największym błędem”. Odnowiona przez teoretyków ciemnej energii zakładała, że ma ona stałą wartość. Wyniki z DESI sugerują coś innego – że ewoluuje i zmienia się z wiekiem Wszechświata, co rewolucjonizowałoby astrofizykę i modele Wszechświata. Jednak badacze wskazują, że wyniki dotychczas uzyskane nie dają jeszcze wystarczającej pewności i trzeba poczekać, aż DESI zakończy program badań w 2026 r., mapując 50 milionów galaktyk i kwazarów. ■



Reportaż o badaniach
DESI: <https://youtu.be/fQkF55yot5I>



MÓZG

Zdalny odczyt myśli dokładny w czterech piątych

Rekonstrukcja zdań na podstawie myśli z dokładnością do ok. 80 proc. jest możliwa dzięki najnowszej technice firmy Meta, opracowanej we współpracy z Baskijskim Centrum Poznawczym z Hiszpanii. Jest to możliwe dzięki nowemu modelowi AI, który potrafi rekonstruować zdania na podstawie aktywności mózgu rejestrowanej za pomocą nieinwazyjnych metod, magnetoencefalografii (MEG) i elektroencefalografii (EEG).

Model AI został przeszkolony na podstawie zapisów mózgu trzydziestu pięciu uczestników eksperymentu, podczas pisania przez nich zdań. Okazało się, że model pozwala na odtworzenie ludzkich myśli o nowych zdaniach z wiernością około osiemdziesięciu proc. znaków pisarskich. Metoda MEG była z tym mniej więcej dwa razy skuteczniejsza niż EEG. Model AI analizuje zapisy rejestrowanej aktywności, odkrywając, w jaki sposób mózg zmienia abstrakcyjne myśli w słowa, sylaby, a nawet specyficzne ruchy palców podczas pisania. Ważnym odkryciem w tych badaniach jest fakt korzystania przez mózg z „dynamicznego kodu neuronowego”, systemu, który łączy różne etapy rozwoju języka, zachowując dostęp do wcześniejszych informacji.

Metoda ta ma jednak swoje ograniczenia. MEG wymaga magnetycznie ekranowanego środowiska, a uczestnicy muszą być nieruchomi, aby możliwe były precyzyjne odczyty. Technika ta została przetestowana jedynie na zdrowych osobach, zatem jej skuteczność w przypadku osób z urazami mózgu jest niepewna. ■



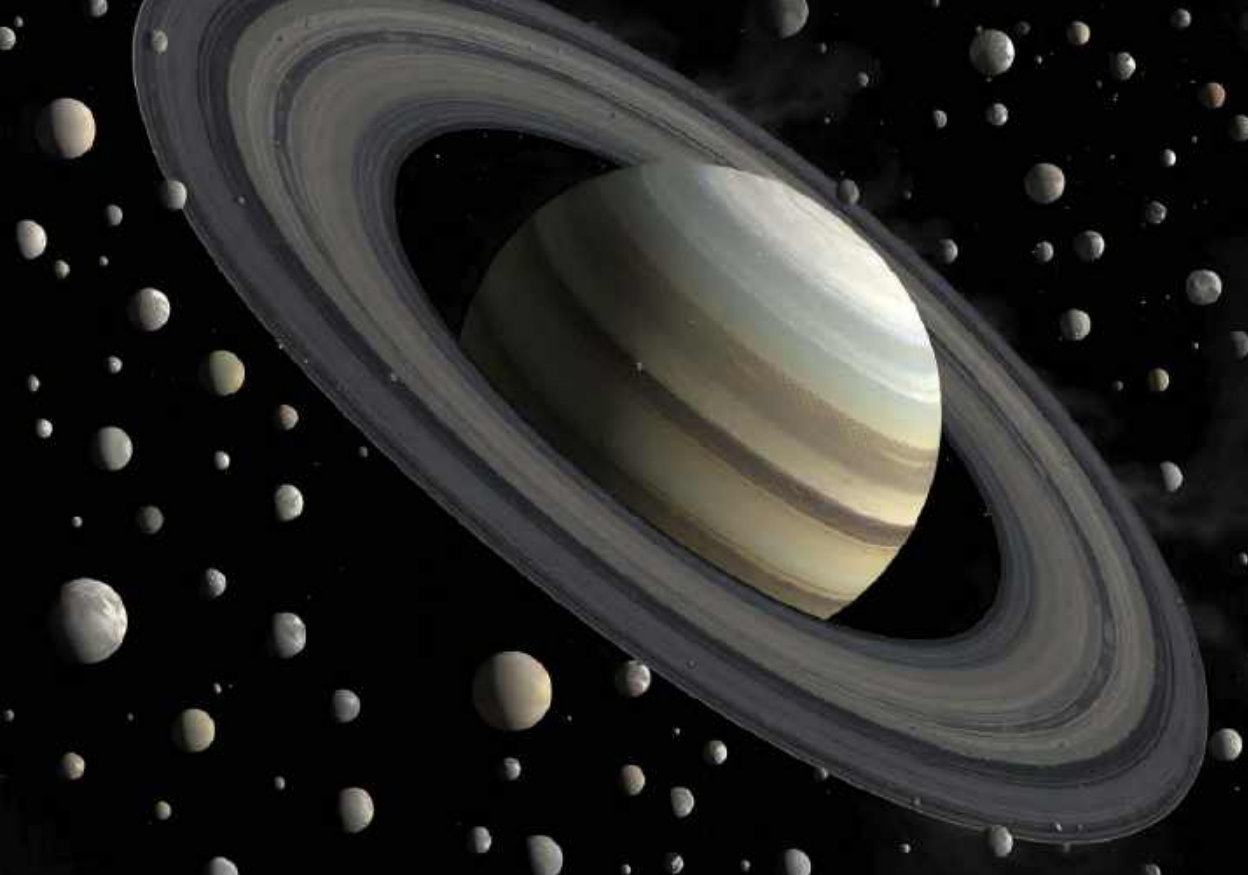
ENERGIA

Organiczne ogniwa słoneczne prawie tak wydajne jak krzemowe

Przekroczenie progu 20 proc. efektywności to najnowsze osiągnięcie w dziedzinie fotowoltaiki organicznej (OPV), czyli takiej, która opiera się nie na krzemie, lecz na materiałach będących związkami węgla. Zdaniem twórców rozwiązania, którzy opisali je w periodyku „Nature”, stanowi to znaczący krok naprzód do powstania realnej alternatywy dla tradycyjnych ogniw krzemowych.

Materiały wykorzystywane w fotowoltaice organicznej do absorpcji światła słonecznego i generowania energii elektrycznej to głównie specjalnie preparowane polimery o małych cząsteczkach, które mogą być przetwarzane w postaci cienkich warstw. W przeciwieństwie do konwencjonalnych krzemowych paneli słonecznych OPV mogą być zintegrowane z półprzezroczystymi oknami, fasadami budynków, a nawet zakrzywionymi powierzchniami.

Przez dziesięciolecia OPV pozostawały w tyle za fotowoltaiką krzemową, która może pochwalić się sprawnością w najlepszych rozwiązaniach przekraczającą 25 proc. Pierwsza generacja ogniw OPV miała sprawność ledwo przekraczającą 10 proc. Ale sytuacja się zmienia. Do przełomu przyczyniły się techniki tzw. akceptorów niefulerenowych (NFA) oraz zaawansowanej inżynierii materiałowej. W modułach słonecznych o powierzchni 15 cm² OPV osiągają potwierdzoną sprawność na poziomie ponad 18 proc., zaś w jeszcze mniejszych ogniwach OPV osiągnęły rekordową efektywność konwersji mocy (PCE) na poziomie 20,43 proc. w małych urządzeniach. ■



UKŁAD SŁONECZNY

Ponad ćwierć tysiąca księżyców Saturna – a może więcej

Astronomowie w publikacji na łamach „Research Notes of the American Astronomical Society” podają, że odkryli 128 nowych księżyców krążących wokół Saturna, powiększając ich sumaryczną liczbę do 274, co daje tej planecie zdecydowaną dominację pod względem liczby naturalnych satelitów w Układzie Słonecznym, przed Jowiszem, który ma 95 znanych księżyców, Uranem z 28 księżycami i Neptunem z szesnastoma. Nowo odkryte obiekty w układzie Saturna są prawdopodobnie wynikiem kosmicznych zderzeń, które pozostawiły szczątki na orbicie planety ok. stu milionów lat temu.

Wiele z nowo odkrytych księżyców Saturna to obiekty o rozmiarze zaledwie kilku kilometrów. Jednak dopóki obserwacje wykazują, że krążą po orbitach wokół większego ciała planetarnego, naukowcy,

k którzy katalogują obiekty w Układzie Słonecznym, definiują je jako księżyce. Prawo do nadania im nazw przysługuje Edwardowi Ashtonowi z Instytutu Astronomii i Astrofizyki Akademii Sinica na Tajwanie, który był głównym autorem publikacji na temat odkrycia. Obecnie tradycja nakazuje nazywać je, czerpiąc z nordyckiej mitologii.

Księżyce zostały odkryte w 2023 roku za pomocą teleskopu Canada France Hawaii Telescope położonego na górze Mauna Kea na Hawajach. Istnienie tak wielu księżyców wokół Saturna wskazuje na liczne kolizje, które wydarzyły się w przeszłości w tej części Układu Słonecznego. Badacze nie wykluczają, że obiektów tego typu krążą w układzie Saturna znacznie więcej, nawet tysiące. ■



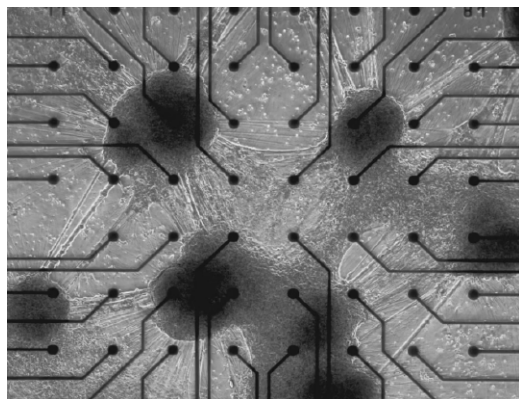
FIZYKA

Światło w ciało nadstałe przeobrażone – po raz pierwszy

Po raz pierwszy udało się przekształcić światło w „ciało nadstałe”, czyli egzotyczny stan materii, w której, według ogólnej definicji, atomy są ułożone w uporządkowany sposób, a przy tym mogą poruszać się bez żadnego tarcia, tak jak w nadciekłej cieczy. Publikacja na ten temat ukazała się w „Nature”.

Ciało nadstałe to pojęcie ze świata kwantowego. Do tej pory rozumiano je jako tworzone przy użyciu atomów, charakteryzujących się zerową lepkością w ekstremalnie zimnych środowiskach i uformowanych w struktury krystaliczne, podobne do układu atomów w kryształach soli. W celu przekształcenia światła w materiał o takich właściwościach naukowcy skierowali laser na próbkę arsenku galu, która została ukształtowana tak, by miała grzbiety na strukturze powierzchni. Gdy światło uderzało w grzbiety, interakcje między nim a materiałem prowadziły do powstania hybrydowych cząstek, które były uwięzione przez owe grzbiety w kontrolowany sposób, co prowadziło do powstania struktury analogicznej z ciałem nadstałym.

Niewielkiemu międzynarodowemu zespołowi nanotechnologów, inżynierów i fizyków udało im się wykazać, że powstała tak struktura jest zarówno ciałem stałym, jak i cieczą oraz że nie ma lepkości, co dowodzi, iż rzeczywiście mamy do czynienia z taką właśnie fazą materii. ■



SZTUCZNA INTELIGENCJA

Biokomputer na ludzkich komórkach mózgowych już do kupienia

Australijska firma Cortical Labs wprowadziła na rynek pierwszy na świecie „biologiczny komputer”, który łączy ludzkie komórki mózgowe z krzemowym sprzętem, tworząc działające sieci neuronowe. Nowy rodzaj inteligencji obliczeniowej nazwany CL1, będący typem „syntetycznej inteligencji biologicznej” (ang. „synthetic biological intelligence”, SBI), jest, jak twierdzą twórcy, bardziej dynamiczny, zrównoważony i energooszczędny niż poprzednie rozwiązania tego rodzaju.

Sieci neuronowe z ludzkich komórek są stale ewoluującym organicznym komputerem, a inżynierowie, którzy go stworzyli, twierdzą, że uczy się on tak szybko i elastycznie, że całkowicie przewyższa oparte na krzemie chipy sztucznej inteligencji używane do szkolenia znanych dużych modeli językowych (LLM), takich jak ChatGPT. Oferując usługę „Wetware-as-a-Service” (z ang. „wilgotny (biologiczny) sprzęt jako usługa”, WaaS), firma proponuje klientom wykorzystanie biokomputera CL1 bezpośrednio lub w sposób zdalny za pośrednictwem chmury.

Według artykułu na temat CL1, który ukazał się w piśmie „Neuron”, biokomputery z komórek mózgowych mogą zrewolucjonizować wiele procesów, od badań nad opracowywaniem leków i testów klinicznych po budowanie „inteligencji” w robotach, umożliwiając szerokie możliwości personalizacji w zależności od potrzeb. CL1 będzie powszechnie dostępny dla wszystkich zainteresowanych klientów w drugiej połowie 2025 roku. ■



ROBOTYKA

Niesamowite ruchy robota z hydraulicznymi mięśniami

Clone Robotics, firma z siedzibą w Polsce, zaprezentowała prototyp humanoidalnego robota o nazwie Protoclone, konstrukcję wyposażoną w ponad tysiąc sztucznych mięśni, wykorzystujących technikę nazwaną Myofiber, która opiera się na koncepcji pneumatycznych mięśni opracowaną przez amerykańskiego fizyka i inżyniera Josepha McKibbina. Na filmie prezentującym to osiągnięcie widać zawieszono go na linkach humanoida, którego kończyny drżą i poruszają się.

Sztuczne mięśnie konstrukcji oparte są na sieci rurek z balonami, które kurczą się po napełnieniu płynem hydraulicznym, naśladując funkcję ludzkich mięśni. Robot Protoclone ma polimerowy szkielet, który odtwarza 206 ludzkich kości, a sercem systemu jest elektryczna pompa o mocy 500 watów, która przepompowuje płyn z prędkością 40 litrów na minutę.

Na jego system sensoryczny składają się cztery kamery głębi umieszczone w czaszce, siedemdziesiąt czujników inercyjnych do śledzenia pozycji stawów oraz trzysta dwadzieścia czujników ciśnienia, które zarządzają „siłą mięśni”. Dzięki temu systemowi robot reaguje na bodźce wizualne i uczy się, obserwując ludzi wykonujących zadania.

Protoclone jest na razie prototypem wymagającym zawieszenia pod sufitem w celu zachowania stabilności. Clone Robotics wcześniej demonstrowało komponenty tej technologii, rękę robotyczną i tors, które również wykorzystują system mięśni Myofiber. Firma ma plany uruchomienia wersji komercyjnej swojego systemu, która nazywana jest Clone Alpha. ■



Prezentacja robota
Protoclone:
<https://youtu.be/H7dhwFcuUn0>



BADANIA KOSMOSU

♦ Badacze z amerykańskiej agencji DARPA pracują nad sposobami hodo-
wania w kosmosie masywnych obiektów o charakterze biologicznym, które będą tworzyć anteny teleskopów, ciągną wind kosmicznych czy ogromne sieci do wyłapywania śmieci orbitalnych.



♦ Kształt spiralnego dysku może mieć wewnętrzna część Obłoku Oorta otaczającego nasz Układ Słoneczny – uważa międzynarodowy zespół naukowców, który opublikował niezrecenzowaną jeszcze pracę na ten temat z repozytorium ArXiv. ♦

EKO-TECHNOLOGIE

♦ Niewymagającą kosztownej pirolizy pleksiglasu (polimeta-
krylanu metylu) metodę przetwarzania, utylizacji i recyklingu tego tworzywa opracowali naukowcy ze szwajcarskiej politechniki ETH, a polega ona na dodaniu rozpuszczalnika dichlorobenzenowego do pokruszonego pleksiglasu w temperaturze między 90 a 150°C a następnie naświetleniu promieniowaniem ultrafioletowym, co roz-
bija łańcuchy polimerowe tworzywa. ♦ Honda, wspólnie z General Motors, opracowała nowy system wodorowych ogniw paliwowych, które są tańsze w produkcji, mają większą moc i większą trwałość w bardziej kompaktowej obudowie niż poprzednie konstrukcje, zastosowane w modelu Honda CR-V e-FCEV, hybrydowym crossoverze typu plug-in na wodór i wykazały wzrost sprawności do ok. 60 proc. – plan zakłada wprowadzenie ich na rynek do 2027 r. ♦

KOMPUTERY KWANTOWE

♦ Na łamach „Physical Review Letters” ukazała się publikacja na temat Zuchongzhi-3, opartego na nadprzewodnikach prototypu komputera kwantowego ze 105 kubitami, działającego, jak twierdzą jego twórcy – zespół badawczy z Chińskiego Uniwersytetu Nauk i Technologii (USTC), z prędkością 10^{15} razy większą niż najszybszy obecnie dostępny superkomputer i milion razy większą niż wyniki układu kwantowego zbudowanego ostatnio przez Google. ♦ Microsoft, po siedemnastu latach inwestycji w badania i prace laboratoryjne, ogłosił opracowanie swojego nowego procesora kwantowego o nazwie Majorana 1, który, jak oficjalnie głosi firma, ma skrócić czas potrzebny na zbudowanie potężnych komputerów kwantowych z dziesięcioleci do zaledwie kilku lat, a opiera się na budowie „topoprzewodnika”, nowego rodzaju materiału opracowanego we współpracy z naukowcami (którzy swoje odkrycia ogłosili wcześniej w „Nature”), składającego się z arsenku indu i aluminium, umożliwiającego kontrolowanie tzw. cząstek Majorany i tworzenia bardziej stabilnych kubitów. ♦

TECHNIKA WOJSKOWA

♦ Chińscy naukowcy twierdzą w publikacji na łamach chińskiego periodyku „Acta Aeronautica et Astronautica Sinica”, że wstrzykiwanie helu do silników rakietowych pocisków wojskowych sprawia, że są one znacznie trudniejsze do wykrycia i przechwycenia, przy czym, co ciekawe, Chińczycy mówią, iż nowe rozwiązanie jest inspirowane technicznymi problemami (wyciekami helu) misji statku Starliner Boeinga, przez które astronauta „utknęli” na ISS. ♦ Według chińskich badaczy z Instytutu Optyki, Mechaniki i Fizyki w Changhun, sterowiec stratosferyczny wyposażony w zaawansowane systemy wykrywania w zakresie podczerwieni mógłby identyfikować samoloty stealth, takie jak amerykański myśliwiec F-35, z odległości prawie dwóch tysięcy km. ■

M.U.



1. Prezentacja Google Glass

Czy okulary AR mogą już zastąpić smartfony?

Inteligencja na nosie

O tym, że w sektorze smartfonów zapanowała stagnacja i że gadżetem, który miałby zająć ich miejsce w sercach (i w kieszeniach) użytkowników, są inteligentne okulary rozszerzonej rzeczywistości, pisaliśmy na łamach MT już lata temu. Po drodze ideę smartfonów ożywić próbowały oferty telefonów składanych i zginanych, chyba nieskutecznie. A co z tymi okularami?

Okulary takie były głośną nowinką w latach 2012/13, gdy Google zademonstrował swój projekt Glass (1), urządzenie łączące się z siecią, wyposażone w kamerę, wyświetlacz, obsługiwane głosem i już wtedy uważane za następcę smartfona. Na początku 2015 sprzedaż Google Glass w wersji Explorer została zakończona. Uznano technikę okularów AR za niedojrzałą, a Google i innym producentom konkurencyjnych produktów, które szybko pojawiły się w obiegu, zarzucono brak pomysłu na zabezpieczenie prywatności użytkowników i ich otoczenia.

Potentaci nie odpuszczają rozszerzonej rzeczywistości

Po tamtej porażce Google Glass Google przejął m.in. firmy VR Tilt Brush i Owlchemy Labs, kontynuując prace nad urządzeniami noszonymi na głowie, w tym zestawami słuchawkowymi wirtualnej rzeczywistości Google Cardboard i Google Daydream VR, które ostatecznie też zostały wycofane. W 2021 roku Google znów podjął wysiłek opracowania XR (ang. Skrót od „extended reality”, innej nazwy techniki rozszerzonej rzeczywistości), pracując nad Project Iris.



2. Prezentacja okularów Google opartych na Android XR

Chodziło o zestaw słuchawkowy AR zasilany przez nowy system operacyjny. Jednak odłożono projekt na półkę po tym, jak Apple zaprezentowało zestaw słuchawkowy Vision Pro VR (niedawno Apple zaprzestało sprzedaży tych gogli).

Pod koniec 2024 r. Google ogłosił wdrożenie systemu operacyjnego Android XR jako „duchowego następcę” Project Iris. Zaprojektowany do obsługi urządzeń XR, w tym zestawu słuchawkowego Samsung Project Moohan oraz pary inteligentnych okularów opracowanych przez Google DeepMind Android XR (2) jest silnie zintegrowany z modelem AI Google o nazwie Gemini. Uważa się, że przewagą systemu Android XR jest fakt, że to pierwsza nowa platforma, która pojawiła się w erze sztucznej inteligencji.

Wciąż jednak, i to odnosi się nie tylko do projektów opartych na rozwiązaniach Google, ale także do tego, co proponuje Meta, Magic Leap i inne firmy, nie zostało rozwiązanych wiele innych, poza prywatnością, problemów i zastrzeżeń, które zgłaszano wcześniej wobec okularów rozszerzonej rzeczywistości. Należą do nich między innymi żywotność baterii, wygoda w noszeniu, estetyczny wygląd i dostępność elementów sterujących.

Wielkie firmy wydały miliardy na zakup i budowę technologii okularów rozszerzonej rzeczywistości w nadziei, że w końcu uda im się wyprodukować produkt o rozmiarze i kształcie zwykłych okularów i wszystkich możliwościach prawdziwie inteligentnej konstrukcji, bez wymienionych wad. Przykładem tego dążenia mogą być Ray-Bany firmy Meta, które wprawdzie potrafiły robić zdjęcia i odtwarzać muzykę do ucha, ale nie były jeszcze prawdziwymi okularami

AR czy XR. Takie prawdziwe Meta planuje udostępnić do sprzedaży dopiero w 2027 r. Zobaczmy.

Jesienią 2024 r. podczas konferencji Meta Connect Mark Zuckerberg zademonstrował okulary rozszerzonej rzeczywistości (AR) Orion, które opisał jako „najbardziej zaawansowane okulary, jakie kiedykolwiek widział świat”. Choć ma grube oprawki, urządzenie wygląda jak okulary (3), a nie jak gogle VR czy Vision Pro firmy Apple, które nazwane zostały goglami AR. Orion wykorzystuje niewielkie projektory wbudowane w zauszuki okularów do wyświetlania obrazu przed oczami użytkownika. Zuckerberg zapewniał podczas prezentacji, że „te okulary istnieją, są niesamowite i stanowią przeblysk przyszłości, która moim zdaniem będzie ekscytująca”. Jednak, jak dodał, zespół Meta musi jeszcze sporo „dopracować”, zanim dojdzie do przekształcenia ich w oficjalny produkt konsumencki. Okulary Orion mają być sterowane za pomocą „interfejsu neuronowego”, który pojawi się dzięki przejęciu przez Meta rozwiązań CTRL-labs, firmy opracowującej projekt opaski na rękę, która będzie kompatybilna z innymi urządzeniami, w tym jak należy rozumieć, okularami. W praktyce rozwiązanie to działać miało tak, że użytkownicy mogą gestykulować z opaską na nadgarstku, by poruszać się po aplikacjach na sparowanych okularach Orion.

Parada innowacji

W ośrodkach badawczych i ekosystemie startupów trwają poszukiwania nowych rozwiązań, które udoskonaliłyby działanie inteligentnych



Film o okularach G1:
<https://youtu.be/tBH7mczkIJY>



3. Mark Zuckerberg w okularach Orion

okularów. Na przykład badacze z jednego z laboratoriów Uniwersytetu Stanforda pracują nad zaawansowaną techniką obrazowania holograficznego, która może poprawić możliwości tego rodzaju urządzeń AR, sprawiając, że sprzęt będzie jednocześnie lżejszy i dawać obraz o wyższej jakości niż znane rozwiązania. Podobnie jak w znanych z rynku okularach AR, uczeni ze Stanforda wykorzystują falowodody, które są elementem prowadzącym światło przez okulary do oczu użytkownika. Jednak ich falowód wykorzystuje „meta-powierzchnie nanotechnologiczne”, co mogłoby „wyeliminować potrzebę stosowania nieporęcznej optyki”. Badacze zastosowali też „wyszkolony fizyczny model falowodu”, który wykorzystuje algorytm sztucznej inteligencji do poprawy jakości obrazu. Według publikacji modele „są automatycznie kalibrowane przy użyciu informacji zwrotnych z kamery”. Jeden ze współautorów publikacji w „Nature” na ten temat, Gun-Yeal Lee twierdzi, że nie istnieje inny system AR, który byłby porównywalny zarówno pod względem możliwości, jak i kompaktowości z ich rozwiązaniem.

Ciekawą innowacją są okulary firmy Even Realities oznaczone G1. Wykorzystują one technologię falowodową w monochromatycznym panelu micro-OLED, czyli wyświetlaczu na soczewkach od wewnętrznej strony przed oczami użytkownika. Ramki okularów mają wbudowany procesor Snapdragon, jak również głośniki i kamery, będąc jednocześnie konstrukcją na tyle podobną do zwykłych okularów, że można ją nosić przez cały dzień. Równowaga między wydajnością a dyskrecją formy jest czymś, co innym firmom często trudno osiągnąć. Gdy użytkownik G1 patrzy prosto przed siebie, ma wrażenie, jakby nosił

standardowe okulary blokujące niebieskie światło. Wystarczy jednak, by lekko przechylił głowę, a zmateriałizuje się przed jego oczami wyświetlacz pokazujący godzinę, temperaturę, powiadomienia, a nawet minimapę bieżącej lokalizacji (4). Powiadomienia pojawiają się w czasie rzeczywistym. Połączenia, wiadomości tekstowe, alerty aplikacji są dostępne, ale zaprojektowane tak, by za bardzo nie przeszkadzać. Jednak nie ma sposobu, by odpowiedzieć na wiadomości za pomocą jedynie



Jak używać Even G1 do nawigacji:
<https://youtu.be/miaSMf7tRBg>



4. Tekst wyświetlany na soczewce okularów G1 firmy Even Realities



5. Okulary T1 firmy Lenovo

okularów. Do tego wciąż potrzebny jest smartfon. Zatem G1 bardziej przypominają pager niż telefon.

Rozwiązanie to mogliby docenić nie tylko mówcy, którzy otrzymują w G1 wygodny i dyskretny prompter, ale również osoby podróżujące. Wbudowany asystent AI, zasilany przez Perplexity lub ChatGPT, może dawać szybkie odpowiedzi na pytania, wyświetlane bezpośrednio na wyświetlaczu. Nieocenione, jeśli rzeczywiście działają, mogą się okazać tłumaczenia w czasie rzeczywistym i konwertowanie języka mówionego na czytelny, tłumaczony tekst, choć takie rozwiązanie nie jest zupełną nowością – już wiele lat temu tłumaczenie na żywo za pomocą okularów zaproponowała japońska firma Docomo. G1 obsługuje również sterowanie głosowe i interakcje oparte na gestach. Synchronizuje się zarówno z urządzeniami z systemem iOS, jak i Android. Według specyfikacji G1 działa około półtora dnia na jednym ładowaniu. Ładowane i ładujące etui wydłuża ten czas do prawie tygodnia.

Nieco inne rozwiązanie, choć wciąż należące do kategorii rozszerzonej rzeczywistości i inteligentnych okularów, to wyświetlacz przed oczami użytkownika, ale nie na wewnętrznej stronie soczewki, lecz np. ponad monitorem (5), już w 2022 r. zaprezentowany przez firmę Lenovo jako Glasses (z ang. po prostu „Okulary”) T1. Mogą łączyć się z telefonami i komputerami, umożliwiając oglądanie wideo, granie w gry lub pracę

na większym wirtualnym ekranie, wyświetlanym jako swoista projekcja przed użytkownikiem noszącym okulary. Okulary T1 wyposażono w parę wyświetlaczy Micro OLED o rozdzielczości Full HD i częstotliwości odświeżania 60 Hz. Mogą one wyświetlać treści z telefonów, laptopów, komputerów PC lub innych urządzeń z systemem Windows, Android lub MacOS.

Technika wyświetlaczy to jedno, ale wygoda i dopasowanie okularów też ma znaczenie. W nowy sposób do problemu wygody i designu okularów podeszli projektanci z firmy Bezier, proponując drukowane w 3D oprawki okularów ze struktur podobnych do plastra miodu, które dopasowują się indywidualnie do twarzy użytkownika. Firma tworzy trójwymiarową mapę jego twarzy za pomocą wideo, wyznaczając ważne punkty, które pomogą mu zaprojektować najlepiej dopasowane oprawki.

Czy wreszcie pojawią się okulary na tyle wygodne i na tyle wolne od wszystkich wad, o których była tu mowa? Niektórzy producenci uważają, że ich produkty już spełniają te wszystkie wymogi. Nawet jeśli tak jest naprawdę, to trzeba jeszcze przekonać nas, że rzeczywiście takiego gadżetu potrzebujemy i rozwiązuje on lepiej niż inne urządzenia nasze problemy. ■

Mirosław Usidus



Używanie
okularów G1
do tłumaczenia:
https://youtu.be/PWqHVERl4_U



1. Obraz z obserwatorium astronomicznego ze wskazaniem obiektu 2024 YR4

Prawdopodobieństwo uderzenia asteroidy w 2032 roku

Ferie z armagedonem

Na początku marca 2025 roku, po dokładnych analizach danych z teleskopu Subaru ogłoszono, że asteroida 2024 YR4 (**1**) raczej nie uderzy w Ziemię 22 grudnia 2032 roku. Ryzyko impaktu spadło do 0,004% z poziomu ponad 3% ogłoszonego jeszcze kilka tygodni wcześniej. To, co działo się przez kilka tygodni, gdy raz po raz NASA i inne czynniki podwyższały ryzyko, jest pouczające.

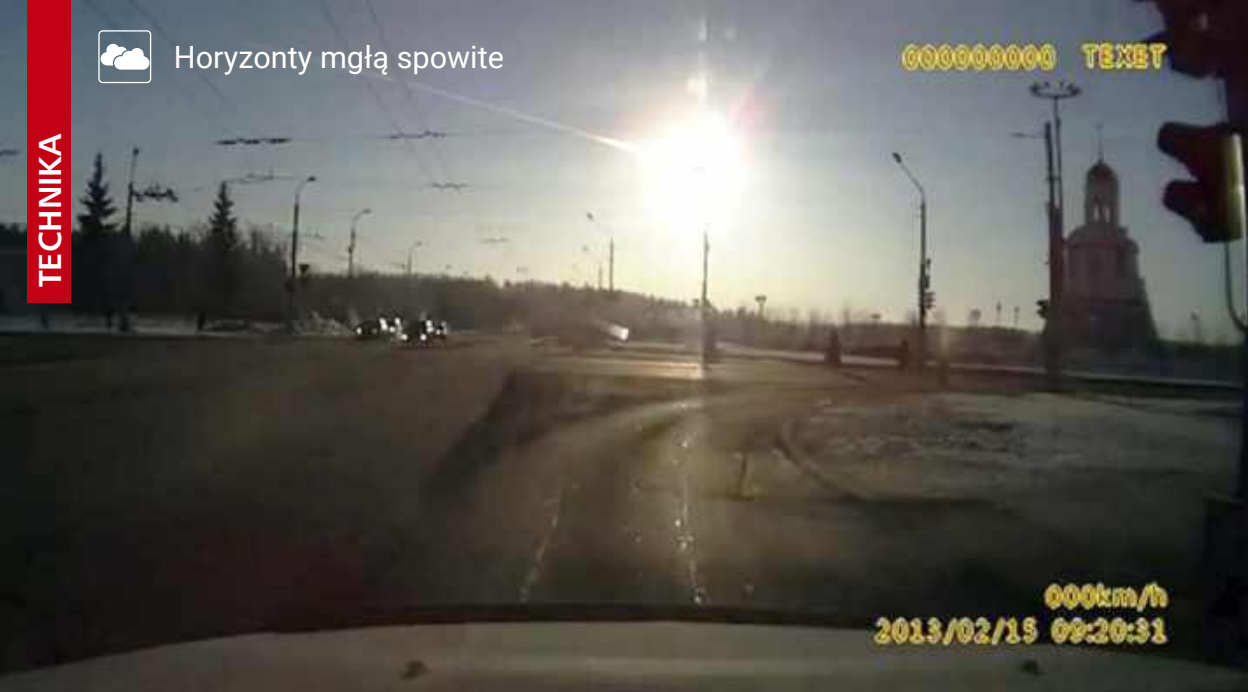
Wnioski i cenna nauka płyną np. z zalewu panikarskiej dezinformacji zwłaszcza w mediach społecznościowych, przedstawiającej podawane przez NASA i ESA nieco wyższe niż przeciętne prawdopodobieństwo uderzenia kawałka kosmicznej skały o rozmiarach szacowanych na najwyżej 90 metrów, jako zapowiedź totalnej zagłady i „armagedon”. W rzeczywistości, nawet gdyby doszło do uderzenia, to zniszczenia miałyby charakter lokalny.

Przy bardzo niesprzyjającym zbiegu wydarzeń meteoryt mógłby uderzyć w gęściej zaludnioną okolicę, np. miasto, niszcząc je. Dotyczy to jednak miasta średniej wielkości

i musiałyby być spełnione pewne warunki, np. 2024 YR4 nie zdeintegrowałby się w ziemskiej atmosferze, jak to miało miejsce z meteorytem (lub precyzyjniej mówiąc – meteoroidem, gdyż do pojedynczego impaktu nie doszło – na Ziemię spadły jedynie odłamki) czelabińskim w 2013 roku, od którego 2024 YR4, według oszacowań, jest tylko nieco większy.

Bolid, który rozpadł się nad Rosją

Tamto wydarzenie sprzed dwunastu lat daje nam pojęcie, jak mogłoby wyglądać spotkanie naszej



2. Jedno ze zdjęć bolidu przelatującego nad Czelabińskiem w lutym 2013 r.

planety z obiektem o takich rozmiarach, choć trzeba wziąć pewną poprawkę na to, że 2024 YR4 jest, według obserwacji, być może dwa razy większy.

Bolid czelabiński miał średnicę około 17–20 metrów i masę do dziesięciu tysięcy ton. Przeleciał przez atmosferę ziemską nad południowym Uralem 15 lutego 2013 roku około godziny 9.20 czasu lokalnego (2). Wszedł do atmosfery pod kątem około 20° do poziomu, z prędkością około 19 km/s (64 tys. km/h). Po wejściu w atmosferę rozpadł się po 32,5 sekund lotu na wysokości 29,7 kilometra nad powierzchnią Ziemi nad obwodem czelabińskim. Powstała w wyniku przełotu i eksplozji bolidu silna fala uderzeniowa spowodowała znaczne straty (uszkodzonych zostało ponad 7,5 tysiąca budynków), a także obrażenia u ponad tysiąca pięciuset osób.

Eksplozja meteoru została zarejestrowana przez siedemnaście z 45 stacji pomiarów infradźwięków działających w ramach programu Comprehensive Nuclear-Test-Ban Treaty Organization. Był to najpotężniejszy wybuch uchwycony przez sensory CTBTO. Na podstawie pomiarów tych stacji szacuje się, że energia wytracona w atmosferze odpowiadała energii wybuchu blisko 0,5 megatony TNT (blisko 40 razy większa od energii wybuchu bomby atomowej zrzuconej na Hiroshimę). Fala uderzeniowa pochodząca od lecącego pod niewielkim kątem do poziomu bolidu rozchodziła się niemal prostopadle do powierzchni ziemi. Na podstawie zniszczeń szacuje się, że jej ciśnienie było 10–20 razy większe od atmosferycznego. Fala uderzeniowa wywołała fale sejsmiczne, które zostały

zarejestrowane przez stacje sejsmiczne i miały moc 2,7 w skali Richtera.

Po upadku meteorytu znaleziono bardzo wiele jego fragmentów i wciąż się je znajduje. Większość z nich ma wielkość kilku milimetrów, ale zdarzają się też większe kawałki. Na powierzchni lodu pokrywającego jezioro Czebarkul znaleziono otwór o średnicy 6 metrów, który mógł zostać wybity przez upadający fragment bolidu, ale nurkowie dotychczas nie odkryli żadnych jego szczątków. Dwa dni po wydarzeniu znaleziono dwa niewielkie fragmenty meteorytu w pobliżu krawędzi przerebła wybitego w lodowej pokrywie jeziora. Zostały one zidentyfikowane jako należące do grupy chondrytów zwyczajnych.

Uważa się, że meteor czelabiński był największym znanym obiektem kosmicznym, jaki zderzył się z Ziemią od czasu katastrofy tunguskiej w 1908 roku. Nie został dostrzeżony przez astronomów przed jego uderzeniem. Był zbyt mały, aby odkryto go w ramach kosmicznych programów obserwacyjnych, śledzących ruchy obiektów bliskich Ziemi i potencjalnie niebezpiecznych. Ponadto meteoroid przyleciał od strony Słońca, przez co trudniej było go dostrzec w trakcie zbliżania się do Ziemi.

Zderzenia meteoroidów takiej wielkości jak meteor czelabiński zdarzają się statystycznie według różnych szacunków co kilkadziesiąt do stu lat, a o rozmiarach większych niż 100 metrów – nie częściej niż co tysiąc lat (przypominamy, że rozmiar 2024 YR4 jest szacowany na maksymalnie 90 m). Astronomowie szacują, że w samej grupie Apolla, z której pochodził meteoroid

czelabiński i do której należy 2024 YR4, znajduje się około 80 milionów obiektów wielkości tego, który rozpadł się nad Rosją w 2013 r.

Fakty nieważne – koniec świata daje społecznościowe zasięgi

Według aktualizacji, o której pisaliśmy, prawdopodobieństwo uderzenia w Ziemię 2024 YR4, którego zaobserwowano po raz pierwszy w obserwatorium w Rio Hurtado w Chile w grudniu 2024 r., jest znikome, więc roztrząsanie, co może się stać na podstawie doświadczeń z 2013 roku, jest raczej niepotrzebne. Ciekawe jest natomiast prześledzenie pierwszych doniesień o asteroidzie i budowaniu atmosfery zagrożenia.

W połowie lutego 2025 r. naukowcy z NASA i Europejskiej Agencji Kosmicznej poinformowali, że szanse na uderzenie w Ziemię tego obiektu pod koniec 2032 roku wynoszą 1,5 proc., co oznaczało wzrost w porównaniu z pierwszymi szacunkami prawdopodobieństwa ze stycznia, wynoszącymi 1,3 proc. Jednocześnie w mediach pojawił się „korytarz ryzyka”

możliwych lokalizacji uderzeń 2024 YR4, który rozciągał się na obszarach od Oceanu Spokojnego przez północną część Ameryki Południowej, Ocean Atlantycki, środkową Afrykę, południową krawędź Półwyspu Arabskiego i północne Indie

Potem poziom ryzyka podawany w niektórych komunikatach wzrósł do 2,1 a następnie – 2,3 proc. W obliczu potencjalnego zagrożenia uderzenia w asteroidę, NASA zapowiedziała wykorzystanie teleskopu kosmicznego Webba do monitorowania 2024 YR4. W okolicach 18 lutego strona Jet Propulsion Laboratory podawała wartość 3,1 proc. jako prawdopodobieństwo impaktu. To było apogeum. Potem poziom ryzyka w prognozach spadał, by na początku marca zbliżyć się do zera.

Przez czas zbliżony długością do zimowych ferii wiele profili społecznościowych zapowiadało koniec świata. Jedne z przerażeniem, inne – z radością, zależnie od bieżącego nastroju. Znamienne, że mało kogo obchodził rzeczywisty stan rzeczy. I to zasługuje na chwilę zadumy. ■

Mirosław Usidus

Życie w sieci nie jest prawdziwym życiem. To życie botów i różnego rodzaju obiektów będących dziełem maszynowym, głównie wygenerowanych przez algorytmy sztucznej inteligencji. Podawana data „śmierci” internetu to zazwyczaj około 2016 lub 2017 r. Od tego czasu sieć już ma nie być „żywa”. Ludzi w niej nie ma (1) lub już nie mają znaczenia.

Teoria martwego internetu żywa jak nigdy

Maszyny przejęły naszą sieć?

Jedna z definicji tej brzmiącej z pozoru dziwnie, ale niepozbowionej podstaw, myśli, podawana przez australijski Uniwersytet Nowej Południowej Walii brzmi tak: „Teoria martwego internetu zasadniczo głosi, że aktywność i treści w internecie, w tym konta w mediach społecznościowych, są w przeważającej mierze tworzone i automatyzowane przez agentów sztucznej inteligencji. Agenci ci mogą szybko tworzyć posty wraz z obrazami generowanymi przez sztuczną inteligencję, zaprojektowanymi w celu zwiększenia zaangażowania (kliknięcia, polubienia, komentarze) na platformach takich jak Facebook, Instagram i TikTok”.

Można na to różnie patrzeć, bo faktycznie duża część aktywności i treści w internecie jest obecnie automatyczna i syntetyczna, o czym „Młody Technik” pisał rok temu. Jest bardziej kontrowersyjna teoria spiskowa powiązana z tym konceptem, która uważa, że to skoordynowana i celowa strategia bliżej nie sprecyzowanych sił w celu kontrolowania populacji i minimalizowania organicznej aktywności człowieka. Według niej, boty społecznościowe zostały stworzone celowo, aby można było manipulować algorytmami choćby w wyszukiwarkach internetowych, modyfikując wyniki wyszukiwania w celu manipulowania konsumentami.



1. Teoria Martwego internetu

Niektórzy jej propagatorzy oskarżają agencje rządowe o wykorzystywanie botów do manipulowania opinią publiczną.

Teoria martwego internetu w tym radykalniejszym wydaniu zakłada, że Google i inne wyszukiwarki cenzurują sieć, filtrując treści, które nie są pożądane, ograniczając to, co jest indeksowane i prezentowane w wynikach wyszukiwania. Choć może się wydawać, że istnieją miliony wyników wyszukiwania dla zapytania, wyniki dostępne dla użytkownika nie odzwierciedlają tego. Problem ten pogłębia zjawisko znane jako rotacja linków, uszkodzanie i znikanie odsyłaczy, a tym samym treści. Konsekwencją jest idea, że Google jest w istocie wioską potiomkinowską, a przeszukiwalna sieć jest znacznie mniejsza, niż nam się wydaje.

Wiele podawanych przez zwolenników teorii martwego internetu faktów i danych da się policzyć, zmierzyć i sprawdzić. Raczej nikt nie wątpi, że od lat następuje wzrost liczby i aktywności botów w internecie. Trudno jednak znaleźć potwierdzenie tej teorii w sensie całościowym a już na pewno – jej radykalnych, spiskowych, elementów.

Dokładne pochodzenie teorii martwego internetu jest trudne do ustalenia. W 2021 roku pojawił się post zatytułowany „Dead Internet Theory: Most Of The Internet Is Fake”, opublikowany na forum ezoterycznym Agora

Road's Macintosh Cafe przez użytkownika o pseudonimie „IlluminatiPirate”, twierdzący, że opiera się na wcześniejszych postach z tego samego forum i z Wizardchan. Teoria weszła do kultury publicznej dzięki szerokiemu zasięgowi i była omawiana na różnych znanych kanałach YouTube. Zyskała większą uwagę dzięki artykulowi w „The Atlantic” zatytułowanemu „Maybe You Missed It, but the Internet ‘Died’ Five Years Ago” opublikowanemu w sierpniu 2021 r.

W 2024 r. Google poinformował, że jego wyniki wyszukiwania są zalewane stronami internetowymi, które „sprawiają wrażenie, jakby zostały stworzone dla wyszukiwarek, a nie dla ludzi”. W korespondencji z Gizmodo rzecznik Google potwierdził rolę generatywnej sztucznej inteligencji w szybkim rozprzestrzenianiu się takich treści i że może ona wyprzeć bardziej wartościowe alternatywy stworzone przez człowieka.

W styczniu 2025 r., po oświadczeniach firmy Meta dotyczących planów wprowadzenia nowych autonomicznych kont opartych na sztucznej inteligencji, zainteresowanie tą teorią ponownie wzrosło. Meta z pomysłu potem się wycofała – tak przynajmniej ogłoszono oficjalnie. Niepokój jednak i wrażenie, że z tą teorią martwego internetu „jest coś na rzeczy”, pozostały. ■

Mirosław Usidus



Chipteligencja

Krzemowo-komputerowy zawrót głowy



1. Stacja DGX Spark AI firmy NVIDIA

Sztuczna inteligencja jest dziś przeważającym refrenem pieśni o rozwoju technik półprzewodnikowych i procesorowych. Z jednej strony jako główna beneficjentka kolejnych „przełomów”, bo to „dla AI” powstają dziś najbardziej zaawansowane chipy. Z drugiej – jako technika aktywnie wykorzystywana w budowie „najwydajniejszego krzemu”.

AI w świecie komputerów

CZY NADCHODZI ERA SUPER „AICETÓW”?

Założyciel i dyrektor generalny firmy NVIDIA, Jensen Huang, ujawnił światu w marcu 2025 r. nowe dzieło swojej firmy, „superkomputer” DGX Spark AI. Ta kompaktowa maszyna (mniej więcej wielkości Apple M4 Mac mini) jest zasilana przez niewielki układ NVIDIA GB10 Grace Blackwell system-on-chip, który ma procesor graficzny wykorzystujący rdzenie Tensor piątej generacji i może pochwalić się obsługą FP4. Dla przypomnienia,

układ ten w połączeniu z 128 GB zunifikowanej pamięci systemowej może wykonać 1000 bilionów operacji obliczeniowych AI na sekundę.

Wraz z superkomputerem DGX Spark, firma ujawniła również DGX Station (1), która jest bardziej rozwiązaniem korporacyjnym dla profesjonalistów. Jest on wyposażony w układ Grace Blackwell Ultra Desktop Superchip GB300, który oferuje 20 petaflopsów wydajności i 784 GB zunifikowanej pamięci systemowej. Rozwiązanie otrzymuje 72-rdzeniowy procesor Grace Neoverse V2 i obsługuje złącze FP4. Maszyna ta będzie wyposażona w kartę interfejsu sieciowego ConnectX-8 SuperNIC obsługującą przepustowość 800 Gb/s. To w porównaniu z tradycyjnymi pecetami biurowymi ogromna ilość mocy, ale zakłada się, że stanie się to normą w nadchodzących latach, biorąc pod uwagę wymagania inteligentnych systemów opartych na sztucznej inteligencji.

Zwróćmy uwagę na używany tu termin „superkomputer”. Być może jesteśmy świadkami narodzin nowej

linii produktów, urządzeń o rozmiarach mniej więcej dobrze znanego stacjonarnego peceta z zastosowaniami podobnym do centrum przetwarzania AI, swobodnego „Alceta”. DGX Spark AI może działać lokalnie lub w akcelеровanym centrum danych w chmurze, umożliwiając twórcom modeli sztucznej inteligencji, badaczom danych, twórcom robotów, analitykom różnego rodzaju i studentom tworzenie aplikacji lub prowadzenie badań nad generatywną sztuczną inteligencją lokalnie, niejako „u siebie” w domu lub w biurze. Jensen powiedział podczas prezentacji m.in.: „Jest oczywiste, że pojawi się nowa klasa komputerów, zaprojektowana dla programistów natywnych dla sztucznej inteligencji i do uruchamiania aplikacji natywnych dla sztucznej inteligencji”. Według niego ten osobisty komputer AI będzie rozwijał sztuczną inteligencję w chmurze, na pulpicie lub w aplikacjach brzegowych.

Kilka dni później pojawiła się w mediach informacja, że Asus wprowadza urządzenie o nazwie Ascent GX10, potężny „superkomputer AI”, który mieści się w dłoni i przypomina białe pudełko (2). Jednak w sercu tego „konkurencyjnego” urządzenia znajduje się nic innego jak Grace Blackwell GB10 firmy NVIDIA. W marketingowych komunikatach Asusa czytamy, że nowy komputer ma przynieść „transformacyjną moc na wyciągnięcie ręki każdego dewelopera”, może dać moc „superkomputera AI w skali petaflopsów”, wyposażając użytkowników w możliwość tworzenia nowych narzędzi AI, w domu za zwykłym biurkiem.

Ascent GX10 z 20-rdzeniowym procesorem NVIDIA o architekturze ARM ma moc obliczeniową 1000 AI TOPS (skrót od ang. „Trillions Operations per Second”, „biliony operacji na sekundę”) i 128 GB pamięci systemowej. Dla porównania, aby można było uzmysłowić sobie, co to znaczy – układy Qualcomm Snapdragon X Elite mogły pochwalić się 45 TOPS, zaś GPU NVIDIA GeForce RTX 4090 ma ponad 1300 TOPS (najnowsza NVIDIA GeForce RTX 5090 ponad 3300 TOPS). Asus umożliwia Ascent GX10 płynną integrację z innymi systemami. Jest to możliwe dzięki zintegrowanym kartom sieciowym (NIC) NVIDIA ConnectX, które umożliwiają połączenie dwóch systemów GX10. Pozwala to, według firmy, „obsłużyć nawet większe modele, takie jak Llama 3.1 o 405 miliardach parametrów”. Asus nie podał jeszcze, kiedy GX10 zostanie

wprowadzony na rynek ani jaka będzie jego struktura cenowa.

„Superkomputery AI” na biurku to stosunkowo nowy nurt. Już wcześniej eksplozja popularności AI wpływała na ofertę producentów krzemu i urządzeń, które wciąż umownie nazywamy komputerami. W lutym 2024 r. NVIDIA zaprezentowała np. chatbota na urządzeniu, co już wcześniej było zaproponowane przez firmę Apple, która chciała, by usługa taka działała na iPhone’ach, zamiast używania zewnętrznych serwerów do przetwarzania zapytań. Zdaniem zwolenników takiego rozwiązania, znacznie przyspiesza ono działanie modeli AI i zwiększa prywatność, NVIDIA przyjęła to podejście, z chatbotem działającym na komputerze z systemem Windows. Chat With RTX to aplikacja demonstracyjna, która pozwala spersonalizować duży model językowy GPT (LLM) przez połączenie go z własnymi treściami, dokumentami, notatkami, filmami lub innymi danymi. Wykorzystując techniki generowania z rozszerzonym odzyskiwaniem (RAG), TensorRT-LLM i akcelerację RTX, można wysłać zapytania do spersonalizowanego, niestandardowego chatbota, by szybko uzyskać kontekstowo trafne odpowiedzi. Chat with RTX obsługuje różne formaty plików, w tym tekst, pdf, doc/docx i xml. Wystarczy wskazać aplikacji folder zawierający pliki, a ona załaduje je do biblioteki w ciągu kilku sekund.

GPU i długość

Przewaga kart graficznych GPU w dziedzinie AI nie polega jedynie na przetwarzaniu równoległym, co jest najczęściej podawanym powodem. Chodzi również o przepustowość pamięci. Procesory typu CPU potrafią

2. Ascent GX10 firmy Asus





3. Wizualizacja CPU i GPU

pobrać (fetching) bardzo szybko, ale niewielkie ilości danych. GPU, mające znacznie większe opóźnienia, działają w takich zadaniach wolniej. Jednak, gdy w grę wchodzi odczytywanie dużych porcji danych, GPU radzi sobie nieporównywalnie lepiej, osiągając nawet 750 GB/s przepustowości, gdy najlepsze procesory CPU mają maks. 50 GB/s szybkości odczytu z pamięci (3).

GPU składają się z tysięcy rdzeni, a każde kolejne wczytanie danych jest w nich szybsze dzięki równoległości wątków. Są w stanie przechowywać informacje w pamięci cache oraz plikach rejestru w celu ponownego użycia, co oznacza wydajność w procesach głębokiej nauki maszynowej. Pierwsza część procesu polega na pobraniu danych z dysku lub pamięci RAM i wgraniu ich na pamięć karty lub cache L1 (pamięć operacyjna) i rejestr. Rejestry są połączone bezpośrednio z jednostką obliczeniową, dla GPU jest to procesor strumieniowy, a dla CPU rdzeń. To tutaj mają miejsce wszystkie obliczenia. Zwykle dąży się do tego, by L1 i rejestr były możliwie blisko jednostki obliczeniowej i jednocześnie były na tyle małe, żeby można było uzyskać do nich szybki dostęp. Każda jednostka obliczeniowa GPU ma pakiet rejestrów nawet trzydzieści razy większych od tych w CPU, a jednocześnie dwukrotnie szybszych. Rezultatem jest nawet 14 MB zarezerwowanych dla pamięci rejestru pracujących z prędkością 80 TB/s. Dla porównania – przeciętna pamięć cache L1 procesora CPU pracuje z prędkością nie większą niż 5 TB/s, a rejestr rzadko ma więcej niż 128 kB.

O wiele mniej rozpowszechniony niż GPU firmy NVIDIA, ale znany w świecie AI, jest Tensor Processing Unit (TPU), układ scalony opracowany przez Google do uczenia maszynowego sieci neuronowych, wykorzystujący autorskie oprogramowanie Google TensorFlow. Google zaczął używać TPU wewnętrznie w 2015 roku, a w 2018 roku udostępnił je do użytku stron trzecich,

zarówno jako część swojej infrastruktury chmury, jak też oferując mniejszą wersję układu na sprzedaż. W porównaniu z procesorem graficznym, TPU są przeznaczone do dużej liczby obliczeń o niskiej precyzji (np. zaledwie 8-bitowej) z większą liczbą operacji wejścia/wyjścia na dżul. Różne typy procesorów nadają się do różnych typów modeli uczenia maszynowego. TPU są dobrze przystosowane do neuronowych sieci konwolucyjnych (CNN). GPU mają zalety dla w pełni połączonych sieci neuronowych, a CPU sprawdzają się w sieciach rekurencyjnych (RNN). W kontekście „sprzętu do AI” najczęściej mówi się o kartach graficznych NVIDIA lub TPU Google’a, ale nie są to jedyne konstrukcje na rynku, które mogą obsługiwać obliczenia wymagane przez modele i generatory. Na przykład firma AMD zaprezentowała w styczniu 2023 r. procesory AMD Ryzen z serii 7040, które są pierwszymi jednostkami o architekturze x86 z wbudowanym akceleratorem sztucznej inteligencji. Układy są produkowane w 4-nanometrowym procesie technologicznym, oferując najmocniejszy w swoim czasie na świecie zintegrowany układ graficzny.

System AI wykorzystuje różne rodzaje obliczeń podczas szkolenia modelu. W skrócie proces ten zaczyna się od załadowania ogromnych ilości danych, które następnie są przetwarzane w iteracjach nauki maszynowej. Tysiące rdzeni GPU przetwarzają ogromne zbiory danych, takich jak obrazy lub wideo. Jeśli wyniki szkolenia mają być osiągnięte szybko, to najprostszym, ale drogim sposobem jest wynajęcie tylu jednostek GPU w chmurze, ile potrzeba (i – na ile stać zarządzających projektem). Procesy wnioskowania AI wymagają mniejszych mocy obliczeniowych na początku, jednak potem ogromna liczba obliczeń i przewidywań potrzebnych do podjęcia decyzji wymaga znacznie więcej energii niż szkolenie w dłuższej perspektywie.

Zbliżyć przetwarzanie do pamięci

Problemu wydajności to stały repertuar dyskusji o CPU i GPU, zaś w kontekście rozwoju AI – rozważań o energochłonności procesu wnioskowania kontra szkolenia modelu. Ma to odniesienia do walki o „przewycięzenie prawa Moore’a” i fizycznych ograniczeń w pakowaniu większej liczby tranzystorów na matrycach o większych rozmiarach. Jest jeszcze kwestia połączeń pomiędzy układami jednostek obliczeniowych. Tradycyjnie taniej było produkować procesory i układy pamięci oddzielnie, a przez wiele lat prędkość zegara procesora była kluczowym czynnikiem decydującym o wydajności. Obecnie to połączenia między układami scalonymi hamują rozwój mocy obliczeniowej. „Kiedy pamięć i przetwarzanie są oddzielone, połączenie komunikacyjne, które łączy te dwie domeny, staje się głównym wąskim gardłem systemu”, wyjaśnia w jednej z wypowiedzi Jeff Shainline z amerykańskiego Narodowego Instytutu Standardów i Technologii (NIST). Podobne opinie formułuje wielu innych specjalistów.

Jednym z kierunków poszukiwania oszczędności w tej dziedzinie są metody próbkowania i tzw. pipeliningu w przyspieszaniu głębokiej nauki przez zmniejszanie ilości przetwarzanych danych. Do grupy takich rozwiązań należy SALIENT (for SAMpling, SLicing, and data moveMeNT), opracowany przez naukowców z Massachusetts Institute of Technology i IBM w celu przewycięzania hardware’owych wąskich gardeł. Podejście to ma radykalnie zmniejszyć wymagania dotyczące uruchamiania sieci neuronowych na dużych zbiorach danych. Niespodzianką nie jest jednak, że metody te ograniczają dokładność i precyzję, czyli nie mogą zostać zastosowane do zadań np. w medycynie czy bezpieczeństwie.

Apple, NVIDIA, Intel i AMD już budują procesory ze specjalnie dla AI tworzonymi rozwiązaniami wbudowanymi w tradycyjne procesory. Amazon dla swoich usług AWS proponuje nowy procesor Inferentia2. Ale te rozwiązania zwykle wciąż wykorzystują tradycyjną architekturę von Neumanna procesorów, zintegrowanej pamięci SRAM i zewnętrznej pamięci DRAM, wszystkie wymagają energii elektrycznej do przenoszenia danych do i z pamięci. Opracowywane w wielu firmach i ośrodkach tzw. obliczenia w pamięci (in-memory computing lub IMC) rozwiązują opisane wyżej problemy, uruchamiając obliczenia AI bezpośrednio w module pamięci, unikając przesyłania danych przez magistralę. IMC sprawdza się w przypadku wnioskowania AI, ponieważ obejmuje ono względnie statyczny (ale duży) zbiór danych o wagach, do którego wielokrotnie uzyskuje się dostęp. Zawsze istnieje potrzeba

przesłania pewnych danych do i na zewnątrz układu obliczeniowego, jednak IMC eliminuje większość kosztów transferu energii i opóźnień związanych z przemieszczaniem danych poprzez utrzymywanie danych w tej samej jednostce fizycznej, gdzie mogą być efektywnie używane i ponownie wykorzystywane do wielu obliczeń. Prace nad przeniesieniem obliczeń bliżej pamięci RAM są zaawansowane, choć nie można powiedzieć, że są już normalną rynkową ofertą. Raczej to wciąż sfera badań i testów.

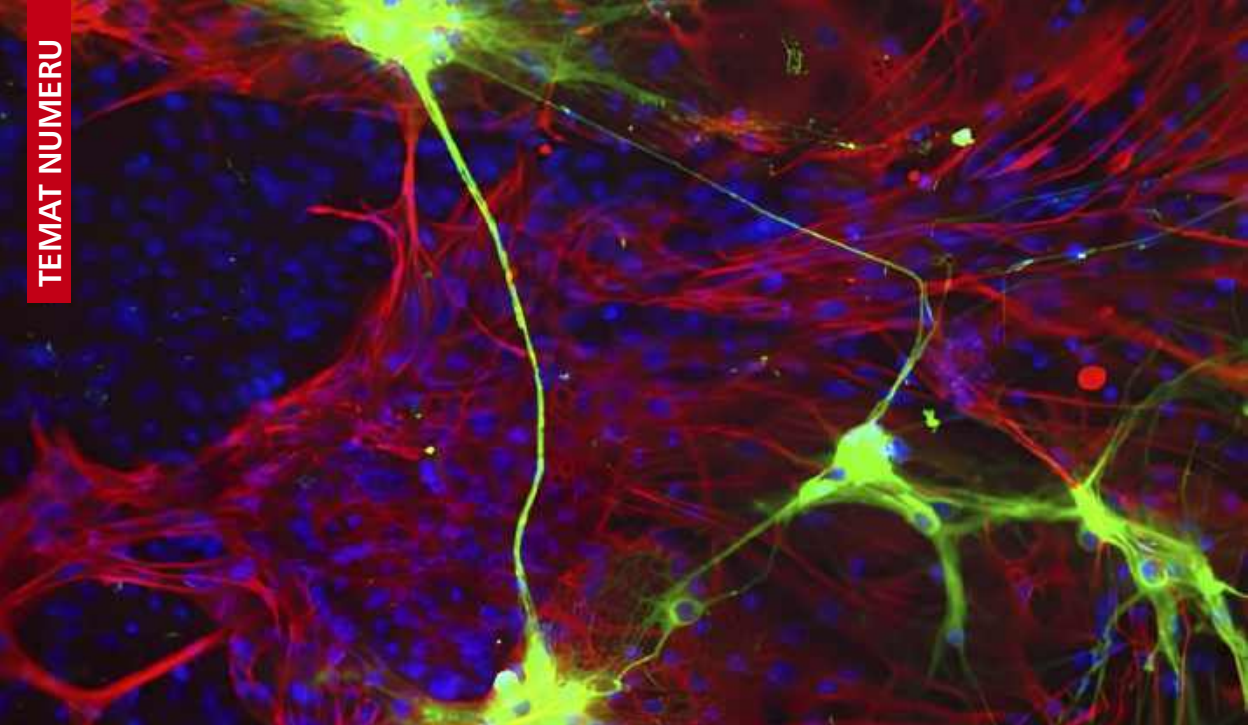
AI projektuje chipy dla AI i naśladuje mózg

Jak powiedział Bill Dally z NVIDIA, jeszcze w kwietniu 2022 roku na konferencji GTC, w jego firmie sztuczna inteligencja uczestniczy w projektowaniu nowych procesorów. Zatem nie jest wykluczone, że ostatecznie to sama AI zaprojektuje i zoptymalizuje sprzęt na swoje potrzeby, taki, który rozwiązywałby wszystkie wyżej opisane problemy i ograniczenia.

W eksperymentach naukowców z grupy badawczej na Uniwersytecie Princeton pod kierownictwem Kaushika Sengupty sztuczna inteligencja zaprojektowała chipy komputerowe, których ludzki umysł czasem nie jest w stanie zrozumieć. Jednak badacze nie wykluczają, że te trudne do ogarnięcia projekty mają przełomowy charakter, jeśli dostatecznie dobrze je przeanalizujemy. W założeniu projektowanie za pomocą AI ma działać tak, że konwolucyjna sieć neuronowa (CNN) analizuje pożądane właściwości chipów, a następnie projektuje je metodą „inżynierii wstecznej”. Jednak algorytmy komputerowe nie wymagają tej samej liniowości lub struktury co ludzki mózg, więc ustalona przez inżynierów-ludzi kolejność lub kształt komponentów chipa nie muszą być wiążące dla sztucznej inteligencji. „Klasyczne projekty starannie łączą te obwody i elementy elektromagnetyczne, kawałek po kawałku, tak aby sygnał płynął w sposób, w jaki chcemy, aby płynął w chipie. Zmieniając te struktury, wprowadzamy nowe właściwości. Teraz [po włączeniu AI – przyp. red.] jest znacznie więcej opcji dla takich działań”, wyjaśnia Sengupta w komunikacie



4. Układ CL1 © Cortical Labs



5. Mikroskopowy obraz neuronowej komórki obliczeniowej BBB z przepływającymi przez nią informacjami © BBB

Uniwersytetu Princeton. Jak dodał jednak, nawet najlepsza sieć neuronowa nadal ma problemy, takie jak halucynacje. Może więc sugerować elementy, które z jakiegoś powodu nie są możliwe w rzeczywistych konstrukcjach. Dlatego badacze są ostrożni w ocenie tego, co tworzy AI.

Inny nurt, pokrewny sztucznej inteligencji a także... ludzkiej inteligencji, to poszukiwania rozwiązań imitujących struktury neuronowe mózgu w konstrukcjach procesorów. Lata temu firma International Business Machines (IBM) podała, że udało jej się skonstruować chip przeprowadzający obliczenia w sposób imitujący funkcje neuronów, synaps i innych „obwodów” w ludzkim mózgu. Jej przedstawiciele twierdzili, że konstrukcja ma charakter przełomowy, a układ w zadaniach takich jak rozpoznawanie kształtów i klasyfikowanie obiektów zużywał znacznie mniej energii niż tradycyjna elektronika. Konstrukcja znana była jako TrueNorth. Została zbudowana dla IBM przez wytwórnię Samsung Electronics z wykorzystaniem tej samej linii montażowej, z której korzysta się do produkcji tradycyjnych mikroprocesorów do urządzeń mobilnych. Chip był odpowiednikiem miliona neuronów i 256 milionów synaps mózgowych.

Później, wraz z wejściem sztucznej inteligencji (AI) do głównego nurtu, jednostka przetwarzania neuronowego (NPU) dołączana do wielu procesorów stała się częścią układów przetwarzających w wielu modelach pecetów lub laptopów nowej generacji. NPU wykorzystują m.in. procesory Qualcomm Snapdragon X

Elite. Procesory Apple z serii A i M mają wbudowaną jednostkę NPU (tzw. Apple Neural Engine). Jednostki NPU można też dodawać do nowszych układów Raspberry Pi. Obecnie pod nazwą NPU kryje się wyspecjalizowany procesor wykorzystywany do przyspieszania operacji sieci neuronowych, w tym zadań obliczeniowych sztucznej inteligencji i uczenia maszynowego (ML). Obejmuje on specjalne optymalizacje sprzętowe, które zwiększają jego wydajność przy jednoczesnym osiągnięciu wysokiej efektywności energetycznej. Jednostki NPU mają możliwości przetwarzania równoległego (mogą wykonywać wiele operacji jednocześnie), a dzięki optymalizacjom architektury sprzętowej mogą wydajnie wykonywać zadania AI i ML, takie jak wnioskowanie i szkolenie. Jednostki NPU mogą być wykorzystywane do wykonywania różnych zadań AI, takich jak rozpoznawanie twarzy, a nawet do szkolenia systemów AI.

NPU nie są jeszcze w sensie ścisłym naśladowaniem struktury mózgowej sieci neuronowej w elektronice. Poszukiwania w tej dziedzinie wciąż trwają, np. badania przeprowadzone i upublicznione w 2022 r. przez zespół kierowany przez naukowców z Los Alamos National Laboratory dotyczyły stosowania algorytmów wstecznej propagacji AI na sprzęcie neuromorficznym. W demonstracji zespół zastosował implementację algorytmu wstecznej propagacji („On-Chip Neuromorphic Backpropagation Algorithm”) w postaci sieci neuronowej do klasyfikacji cyfr i elementów odzieży z dwóch zestawów danych. Pracując

w pełni na chipie, bez komputera w pętli, w układzie zademonstrowano możliwość wykorzystania procesorów neuromorficznych do implementacji nowoczesnych aplikacji głębokiej nauki maszynowej, przy niskim poborze mocy. Algorytm wstecznej propagacji zespołu osiągnął średnią dokładność wnioskowania powyżej 96 proc. Wynik ten i test porównawczy na podobnym zbiorze danych były konkurencyjne z wynikami dokładności wnioskowania tego samego algorytmu zaimplementowanego na standardowym sprzęcie komputerowym. Algorytm przetwarzzał zbiory danych z taką samą szybkością jak tradycyjny układ z komputerem, ale zużywał tylko 2,5 proc. mocy potrzebnej w takim układzie.

Nad wykorzystaniem, nawet nieinspirowanych tylko wzrost biologicznych struktur do przetwarzania AI od lat pracuje Cortical Labs, startup z siedzibą w australijskim Melbourne. Buduje miniaturowe „mózgi” z naturalnych, biologicznych neuronów osadzonych na wyspecjalizowanym chipie komputerowym. Firma uczy hybrydowe minimózgi wykonywania zadań podobnych do tych, które może wykonywać sztuczna inteligencja, przy minimalnym zużyciu energii. Firma uczyła m.in. swoje układy hybrydowe grać w grę zręcznościową Atari Pong. Cortical Labs stosuje dwie metody konstrukcji swojego urządzenia. Pierwsza polega na pobieraniu neuronów z zarodków myszy. Druga polega na technice, w której ludzkie komórki skóry są przekształcane

z powrotem w komórki macierzyste, a następnie pobudzane do wzrostu w formę ludzkich neuronów. Te są następnie osadzone w środowisku ciekłym na specjalnym układzie scalonym zawierającym tlenek metalu, składającym się z 22 tys. minielektrod, które umożliwiają programistom dostarczenie energii elektrycznej do neuronów, a także odbiór ich sygnałów w drugą stronę. Australijska firma wprowadziła ostatnio na rynek pierwszy na świecie „biologiczny komputer”; nowy rodzaj inteligencji obliczeniowej, nazwany CL1, będący typem „syntetycznej inteligencji biologicznej” (ang. „synthetic biological intelligence”, SBI). Inżynierowie, którzy go stworzyli, twierdzą, że uczy się on tak szybko i elastycznie, że całkowicie przewyższa oparte na krzemie chipy sztucznej inteligencji używane do szkolenia znanych dużych modeli językowych (LLM), takich jak ChatGPT. Oferując usługę „Wetware-as-a-Service” (z ang. „wilgotny (biologiczny) sprzęt jako usługa”, WaaS), firma proponuje klientom wykorzystanie biokomputera bezpośrednio lub w sposób zdalny za pośrednictwem chmury.

Na biologicznym podłożu pracują też eksperymentalne układy Bionode amerykańskiej firmy Biological Black Box (BBB), która hoduje w laboratorium neurony w celu zaprzęgnięcia ich do przetwarzania AI (5). Pytanie tylko, czy takie biologiczne systemy to jeszcze sztuczna, czy może już „naturalna” inteligencja? ■

Mirosław Usidus

Tolkien w XXI wieku.

Co dziś znaczy dla nas Śródziemie

Nick Groom

Wydawnictwo Insignis, liczba stron: 600, cena sugerowana: 64,99 zł

Finalistka nagrody Tolkien Society Best Book Award.

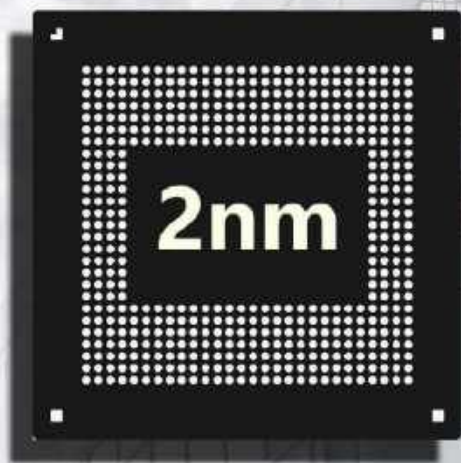
Niezwykle oryginalna i dająca do myślenia literacka podróż po świecie wykreowanym przez J.R.R. Tolkiena.

Co sprawia, że Śródziemie i jego mieszkańcy zawładnęli wyobraźnią milionów na całym świecie? Skąd bierze się to, że wizjonerskie dzieło Tolkiena wciąż fascynuje i inspirowane jest niemalże dziewięćdziesiąt lat po jego premierze?

Tolkien w XXI wieku to wciągające, zaskakujące i świeże spojrzenie na twórczość pisarza. Nick Groom śledzi życiorys Tolkiena, wydobywając z jego kluczowych momentów genezę i inspiracje dla jednych w swoim rodzaju książek profesora, bada ich społeczny odbiór na przestrzeni lat oraz analizuje wpływ, jaki wywarły one na naszą kulturę. Omawia i interpretuje późniejsze adaptacje filmowe (w tym filmy Petera Jacksona i nowy serial *Pierścienie Władzy*), muzyczne i literackie oraz dzieła inspirowane Śródziemiem, które ugruntowały reputację tej fantastycznej krainy jako niekwestionowanego i ponadczasowego fenomenu kulturowego.

Zagłębiając się w tematy takie jak przyjaźń, różnorodność, środowisko czy eukatastrofa, Groom zabiera nas w pełną niespodzianek podróż przez Śródziemie i udowadnia, że twórczość Tolkiena jest dziś bardziej aktualna i żywa, niż sam autor mógłby kiedykolwiek przypuszczać.





1. Symboliczny obraz chipa 2 nm na tle siedziby TSMC

Zapowiedzi „końca fizycznej możliwości miniaturyzacji” układów półprzewodnikowych opartych na krzemie i rozważania na temat, czy da się jeszcze coś z krzemu wycisnąć, opisujemy w MT od lat. Za każdym razem okazuje się, że da się coś jeszcze zrobić, zmniejszyć i upakować, podnosząc wydajność na nowe poziomy, mimo ogólnej zgody, że prawo Moore’a już nie obowiązuje.

**Czy to już ostatnie soki, jakie da się wycisnąć z krzemu?
Jakie są alternatywy?**

OD „PRAWA” MOORE’A DO POJEDYNCZYCH ATOMÓW

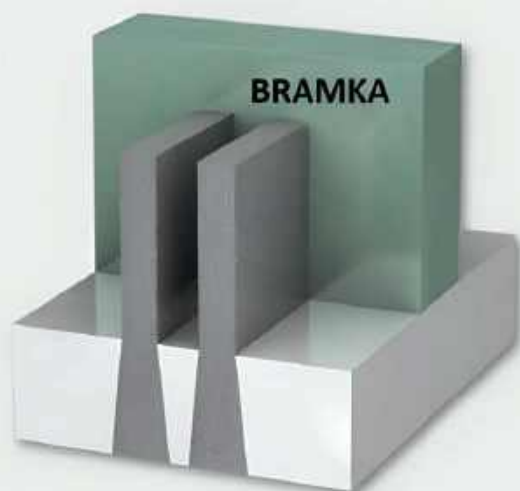
Czasem postęp jest szybszy, niż niedawno jeszcze przewidywano. Prezentacja jednego z partnerów Intela podczas międzynarodowego spotkania przedstawicieli branży urządzeń elektronicznych IEEE (IEDM) w 2019 r. ujawniała plan osiągnięcia procesu litografii układów w chipach o rozmiarach bramki 1,4 nanometra do 2029 roku. Z prezentacji wynikało, że Intel spodziewa się, że miniaturyzacja jego techniki litograficznej będzie postępować w cyklu dwuletnim, począwszy od 10 nm w 2019 r., następnie

7 nm w 2021 r., 5 nm w 2023 r., 3 nm w 2025 r., 2 nm w 2027 r. i 1,4 nm w 2030. Według ekspertów proces 1,4 nm oznacza szerokość równą dwunastu atomom krzemu.

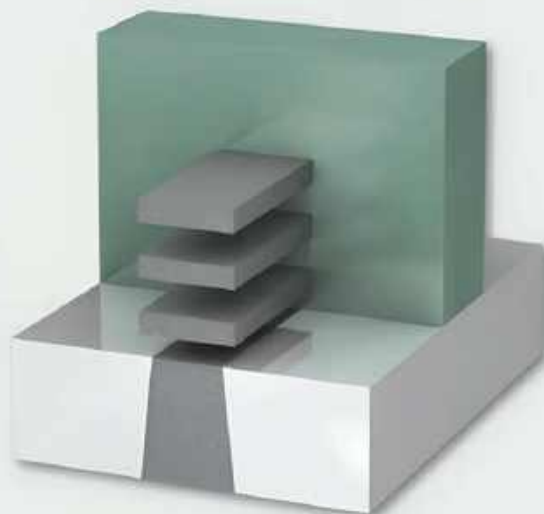
Intel ma mnóstwo kłopotów, nie tylko z miniaturyzacją bramek. Tymczasem tajwański producent krzemowych układów, TSMC przygotowało się już w 2024 r. do próbnej produkcji układów 2 nm (1), z wykorzystaniem procesu wspomaganej sztuczną inteligencją. Tajwański dziennik „Economic Daily” podał w ubiegłym roku, że TSMC rozpoczęło prace przedprodukcyjne nad półprzewodnikami 2 nm. Według doniesień, TSMC postawiło przed inżynierami i naukowcami cel wyprodukowania tysiąca wafli w tym procesie jeszcze w 2024 r. i rozpoczęcia masowej produkcji w 2025 roku.

Nieoficjalne źródła tych doniesień podają, że proces produkcyjny TSMC wykorzystuje metodę wspomaganą sztuczną inteligencją o nazwie AutoDMP, zasilaną przez DGX H100 firmy NVIDIA. Sztuczna inteligencja ma optymalizować projekty chipów trzydzieści razy szybciej niż inne techniki. W ramach procesu

FinFET



Gate-All-Around



2. Poglądowe przedstawienie różnicy pomiędzy architekturą FinFET a GAAFET (gate-all-around)

2 nm TSMC przejdzie na tranzystory GAAFET (gate-all-around), tranzystory z bramką otaczającą (SGT), które mają zwiększyć gęstość jednostek na chipie i poprawić wydajność w stosunku do układów 3 nm o 10–15 procent przy tym samym poziomie mocy lub zużywać o 20–25 procent mniej energii przy tej samej wydajności. GAAFET jest podobny w koncepcji do MOSFET – FinFET (o strukturze 3D zamiast planarnej), ale materiał bramki otacza w nim obszar kanału ze wszystkich stron (2). W zależności od konstrukcji, tranzystory FET typu gate-all-around mogą mieć dwie lub cztery efektywne bramki. Zostały z powodzeniem wytrawione na nanodrutach z arsenku indowo-galowego InGaAs, który ma wyższą ruchliwość elektronów niż krzem.

Tranzystor MOSFET z otaczającą bramką (GAA) został po raz pierwszy zademonstrowany w 1988 roku przez zespół badawczy firmy Toshiba, w tym Fujio Masuokę, Hiroshi Takato i Kazumasa Sunouchi. Masuoka, najbardziej znany jako wynalazca pamięci flash, później opuścił firmę Toshiba i założył Unisant Electronics w 2004 roku, aby wraz z Uniwersytetem Tohoku pracować na technikę otaczającej bramki. W 2006 roku zespół koreańskich naukowców opracował tranzystor 3 nm oparty na technologii GAA FinFET. GAAFET zostały wykorzystane przez IBM do zademonstrowania technologii procesowej 5 nm. Od 2020 r. Samsung i Intel mówią o planach masowej produkcji tranzystorów GAAFET. TSMC wtedy informowało, że będzie nadal używać FinFET w swoim węźle 3 nm, choć w rzeczywistości pracowało już nad tranzystorami GAAFET. Intel chciał startu swojego 2-nanometrowego węzła w 2024 roku,

a Samsung planował rozpocząć masową produkcję 2 nm w 2025 roku. Oba układy miały wykorzystywać GAAFET i tylne zasilanie w celu zwiększenia gęstości logicznej i zmniejszenia wycieków mocy. Były plany wybiegające dalej w przyszłość, Samsung zamierza osiągnąć 1,4 nm do 2027 roku (TSMC planuje w tej perspektywie 1 nm, czyli bramkę z dziesięciu atomów).

Intel nie udostępnił 2-nanometrowego węzła w planowanym czasie a o Samsungu media pisały, że ma ogromne problemy i musiał swoje 2 nm przełożyć na 2026 r. Realnie więc zostaje tylko TSMC, które rozpoczęło produkcję układów 2 nm na początku 2025 r. Na potwierdzenie, że wyprodukowane chipy spełniają wymogi, przyjdzie jeszcze zapewne poczekać. Taiwan Semiconductor Manufacturing stara się zaspokoić popyt ze strony głównego klienta Apple na chipy szczeblek niższe pod względem miniaturyzacji – 3 nm. Według analityków ankietowanych przez EE Times, problemy firmy z narzędziami i wydajnością utrudniły przejście do produkcji seryjnej. TSMC i Samsung, jego kolejny największy rywal w branży chipów, ścigają się od dłuższego czasu o to, kto szybciej i wydajniej dostarczy 3 nm dla klientów takich jak Apple i NVIDIA w wysokowydajnych komputerach i w smartfonach. Jak widać, pomimo optymistycznych doniesień, rzeczywistość, nawet po uruchomieniu produkcji, wygląda różnie.

Można więcej, ale będzie goręcej

Przez ponad dwie dekady słyszeliśmy o śmierci „prawa Moore’a”. Nie jest to w rzeczywistości „prawo” takie jak np. prawa fizyki. Była to raczej reguła czy też

prognoza opracowana na podstawie obserwacji własnych przez nieżyjącego już współzałożyciela Intela, Gordona Moore'a (3), zakładająca, że liczba tranzystorów w chipie podwaja się mniej więcej co 18–24 miesiące. W 2006 roku sam Moore powiedział, że prawo to przestanie obowiązywać w 2020 roku. Profesor MIT Charles Leiserson ogłosił jego koniec wcześniej, już w 2016 roku. Szef NVIDIA, Jensen Huang, ogłosił jego koniec w 2022 roku.

Prawo Moore'a nie mówi bezpośrednio o poprawie wydajności – dotyczy po prostu liczby tranzystorów na chipie. Ma to jednak oczywiście wpływ na wydajność. Tę mierzy się w testach benchmarkowych. Te wykazują poprawę z generacji na generację chipów, ale bez skoków, o jakich mówił Moore, np. Ryzen 9 9950X produkcji AMD, choć stanowił wyraźne ulepszenie w stosunku do swoich odpowiedników Zen 4 w testach benchmarkowych Cinebench R23, wzrost pomiędzy wersją Ryzen 9 5950X a Ryzen 9 7950X wyniósł 36 proc., zaś pomiędzy Ryzen 9 7950X a Ryzen 9 9950X – 15 proc. Wyniki kolejnych generacji procesorów Intela są podobne. W ciągu zaledwie kilku lat tempo wzrostu wydajności znacznie spadło. Choć tempo spadło, firmy dbają, by jednak były to wzrosty. Intel przeszedł w tym celu na architekturę hybrydową w swoich procesorach Arrow Lake. Z kolei dla AMD pamięć podręczna 3D V-Cache stała się technologią definiującą procesory firmy i jest to wyraźny sposób utrzymania się na kursie wzrostowym wydajności. AMD układa więcej pamięci podręcznej na wierzchu, której nie może zmieścić na matrycy.

Ważnym czynnikiem, o którym należy pamiętać, jest przestrzeń. Oczywiście można zbudować ogromny

układ z mnóstwem tranzystorów, ale ile energii będzie pobierać? Czy będzie w stanie utrzymać rozsądną temperaturę? Każdy milimetr kwadratowy powierzchni krzemu jest bardzo drogi. Możemy budować chipy z wielką liczbą tranzystorów, wiemy, jak je budować, ale stają się one coraz droższe. RTX 4090 NVIDIA zapewnia ponaddwukrotnie większą liczbę tranzystorów niż RTX 3090 przy bardzo podobnym rozmiarze matrycy, ale kosztuje dużo. Nie wspominając już o innych problemach – wysokim zapotrzebowaniu na energię, dużym rozmiarze chłodnicy i topiących się złączach zasilania. Nie wszystkie z tych problemów wynikają ze zwiększania liczby tranzystorów, ale odgrywa to pewną rolę. Większe chipy to więcej tranzystorów, więcej ciepła i zazwyczaj wyższe koszty.

W rzeczywistości producenci, od Intela po Apple, od wielu lat osiągają wzrosty wydajności procesorów innymi sposobami niż zagęszczanie liczby tranzystorów na chipie. Trudno zliczyć wszystkie innowacyjne techniki, swoiste „wybiegi” w architekturze i w połączeniach jednostek przetwarzających, stosowane przez producentów.

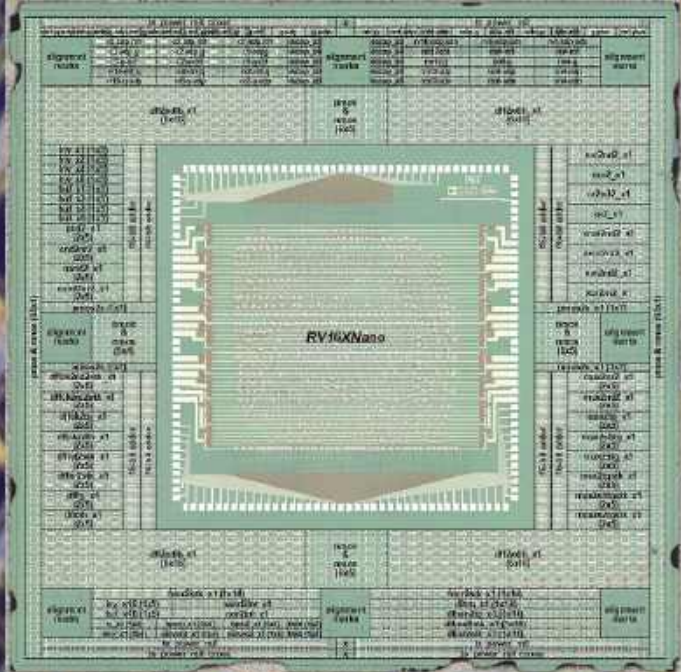
Tranzystory CMOS (ang. „Complementary Metal-Oxide Semiconductor”) były standardowym budulcem układów scalonych od lat 80. ubiegłego wieku. Opierają się na podstawowej architekturze, którą John von Neumann wybrał w połowie XX wieku. Jego architektura została zaprojektowana tak, by oddzielać elektronikę przechowującą dane w komputerach od elektroniki przetwarzającej informacje cyfrowe. Komputer przechowywał informacje w jednym miejscu, a następnie wysyłał je do innych obwodów w celu przetworzenia. Oddzielenie przechowywanej pamięci od procesora zapobiega wzajemnemu zakłócaniu się sygnałów i zachowuje dokładność potrzebną do obliczeń cyfrowych. Jednak przenoszenie danych z pamięci do procesora stało się wąskim gardłem. Deweloperzy poszukują obecnie rozwiązań alternatywnych wobec architektury von Neumanna. Dążą do wykonywania obliczeń „w pamięci”, aby uniknąć marnowania czasu i energii na przenoszenie danych.

Jeszcze przed pandemią, w sierpniu 2019 roku Philip Wong, szef działu badań w TSMC, na konferencji Hot Chips na amerykańskim Uniwersytecie Stanforda mówił o najnowszych osiągnięciach w rozwoju techniki układów scalonych, wśród nich m.in. o umieszczaniu pamięci bezpośrednio na procesorze, które znacznie zwiększą wydajność i spowodują, że „prawo Moore'a będzie wciąż żywe”, choć zapewne już nie według „najściślejszej definicji”. Takie połączenie radykalnie zwiększyło miało szybkość i wydajność, ponieważ układy logiczne na chipie będą szybciej



Gordon Moore
1929-2023

3. Gordon Moore



4. Schemat procesora z nanorurkami węglowymi RV16XNano © MIT

otrzymywać potrzebne dane, bez niepotrzebnych „przestojów”. Wong na slajdach swojej prezentacji prognozował „nieprzerwany postęp w chipach aż do 2050 roku”, oceniając, że „postęp nie zatrzyma się w miejscu, gdy elektronika osiągnie fundamentalne granice wielkości atomów”. Przedstawiciel tajwańskiej firmy mówił o technikach, nad którymi prace badawcze trwają od lat, w tym o zastosowaniu w chipach nanorurek węglowych, że obecnie istnieją już praktyczne możliwości ich wdrożenia. Innowacja polegająca na zastosowaniu dwuwymiarowych materiałów warstwowych, pozwalając elektronom na łatwiejszy przepływ przez układy scalone, również, jak twierdził Wong, jest już produkcyjnie dostępna. Kolejna nowa technika, układania w architekturze 3D, oznaczać ma, że funkcje procesora komputerowego, które obecnie są odizolowane, mogą być umieszczone w wielu warstwach, połączonych szybkimi ścieżkami danych.

Nanorurki, grafen, skyrmiony a może... drewno?

Jednym z następców krzemu mogą być nanorurki węglowe. Co prawda technologia ta dopiero raczkuje, ale już w 2019 r. naukowcom z Instytutu Technologicznego w Massachusetts (MIT) udało się osiągnąć kolejny kamień milowy w produkcji układów scalonych opartych na tych nanostrukturach. Zespołowi pod przewodnictwem profesora Maxa M. Shulakera udało się wyprodukować procesor składający się z ponad 14 tys. tranzystorów polowych z nanorurkami węglowymi (CNFET) o średnicy

1 mikrona. Co istotne, naukowcy wykorzystali tutaj technologie stosowane przy produkcji tradycyjnych, krzemowych układów. Ich RV16XNano (4) był 16-bitowym układem o taktowaniu 10 kHz (kiloherców) opartym na otwartej architekturze RISC-V. Choć parametry tamtej konstrukcji są mało imponujące przy obecnych najbardziej zaawansowanych chipach krzemowych tranzystory z nanorurek węglowych mogą pracować z dużo wyższymi częstotliwościami niż te z krzemu, a przy tym cechują się dziesięciokrotnie wyższą efektywnością energetyczną. Niestety, produkcja nanorurkowych tranzystorów na dużą skalę jest problematyczna – trudno uzyskać materiał o bardzo dużej czystości (dla zaawansowanych obwodów ideałem jest nierealne 99,99999 proc.), więc często mają one wady i w efekcie stają się bezużyteczne. Naukowcom z MIT udało się jednak opracować nowe techniki produkcji, co pozwoliło obniżyć rygor czystości (do „zaledwie” 99,99 proc.) i w znacznym stopniu poprawiło uzysk produkcji układów.

W 2020 r. natomiast powstał mikroprocesor, który zamiast krzemu zbudowany był z grafenu i miał dziesięć tysięcy lepsze parametry niż poprzednie prototypy tego rodzaju. Dokonało tego laboratorium firmy IBM w Yorktown Heights, w USA, budując procesor będący wielostopniowym odbiornikiem fal w zakresie radiowym. Zaczęto mówić o możliwych zastosowaniach nowego typu chipa w technologiach bezprzewodowych, teoretycznie bowiem pozwalają na znacznie szybszy przesył danych. Jak oceniał twórca grafenowego procesora, IBM, technologia ta pozwoli na szybki

rozwój zasięgu bezprzewodowej komunikacji w sieci sensorów Internetu Rzeczy. Dostępne obecnie chipy typu RFID mają tę wadę, że działają w bardzo niewielkim zasięgu. Technologie grafenowe mogą to radykalnie zmienić. Znow jednak głównym problemem związanym z wykorzystaniem grafenu była czystość materiału i duża podatność tego potencjalnie cudoownego materiału na zanieczyszczenia w procesie produkcji.

Koreański Samsung pracuje od pewnego czasu nad przyspieszeniem rozwoju nowej, obiecującej technologii pamięci o nazwie Selector-Only Memory (SOM). Najnowsza technologia łączy w sobie spójność pamięci flash i błyskawiczne prędkości odczytu/zapisu DRAM, wykorzystując materiały znane jako chalkogenki, które można wygodnie przełączać między stanami przewodzenia i rezystancji w celu przechowywania danych. Samsung wykorzystał zaawansowane modelowanie komputerowe do przewidywania potencjału różnych kombinacji materiałów. Zaprezentował też już pierwsze prototypowe układy tego rodzaju.

Jeśli chodzi o najnowsze badania, to przypadkowe odkrycie, o którym poinformowano pod koniec 2024 r. w czasopiśmie „Nature”, może znacznie obniżyć energię potrzebną do technologii pamięci nowej generacji. Korzystając z materiału o nazwie selenek indu (In_2Se_3), naukowcy odkryli technikę obniżania zapotrzebowania na energię pamięci zmiennofazowej (PCM), zdolnej do przechowywania danych bez stałego zasilania, nawet miliard razy. PCM jest kandydatem na rozwiązanie pamięci obliczeniowej, które mogłoby zastąpić zarówno pamięć krótkotrwałą, jak pamięć o dostępie swobodnym (RAM), jak i urządzenia pamięci masowej, takie jak dyski półprzewodnikowe (SSD) lub dyski twarde. Pamięć RAM jest szybka, ale wymaga znacznej przestrzeni fizycznej i stałego zasilania do działania, podczas gdy dyski SSD lub dyski twarde są znacznie gęstsze i mogą przechowywać dane, gdy komputery są wyłączone. PCM działa przez przełączanie materiałów między dwoma fazami – krystaliczną, w której atomy są starannie uporządkowane, i amorficzną, w którym atomy są ułożone losowo. Stany te są związane z binarnymi jedynkami i zerami, kodując dane przez przełączanie stanów. Jednak technika „melt-quench” stosowana do przełączania tych stanów, która przewiduje ogrzewanie i szybkie chłodzenie materiałów PCM, wymaga znacznej energii. W swoich badaniach naukowcy znaleźli sposób na indukowanie amorfizacji za pomocą ładunku elektrycznego. Zmniejsza to zapotrzebowanie PCM na energię i potencjalnie otwiera drzwi do szerszych zastosowań komercyjnych. Odkrycie opiera się

na właściwościach selenku indu, materiału półprzewodnikowego o właściwościach zarówno „ferroelektrycznych”, jak i „piezoelektrycznych”.

Jednym z kierunków poszukiwań, które pozwoliłyby się uwolnić od ograniczeń, jakie wiążą się z krzemem, są systemy neuromorficzne, które wykorzystują algorytmy i projekty sieci, które naśladują wysoką łączliwość i równoległe przetwarzanie ludzkiego mózgu. Wymaga to opracowania nowych typów sztucznych neuronów i synaps, które są kompatybilne z przetwarzaniem elektronicznym, ale przewyższają wydajność obwodów CMOS.

Jednym z najpopularniejszych pomysłów na układy neuromorficzne są memrystory. Przypominają standardowe rezystory elektryczne, ale z elektryczną regulacją oporu, co w tym przypadku oznacza zarazem modyfikację danych przechowywanych w pamięci. Dzięki trzem warstwom – dwóm terminalom, które łączą się z innymi urządzeniami, oddzielnym warstwą pamięci. Ich struktura pozwala na przechowywanie danych i przetwarzanie informacji. Koncepcja ta została zaproponowana w 1971 roku, ale dopiero w 2007 roku R. Stanley Williams, naukowiec z laboratoriów Hewletta-Packarda w Kalifornii, stworzył pierwszy cienkowarstwowy memrystor półprzewodnikowy, który można było wykorzystać w obwodzie.

Mark Hersam i jego kolega Vinod K. Sangwan, materiałoznawcy z Uniwersytetu Northwestern, skatalogowali niedawno obszerną listę potencjalnych neuromorficznych materiałów elektronicznych, która obejmuje materiały „zerowymiarowe” (kropki kwantowe), materiały jedno- i dwuwymiarowe (np. grafen) oraz heterostruktury van der Waalsa (wiele dwuwymiarowych warstw materiałów, które przylegają do siebie). Sporo uwagi jako potencjalne materiały do zastosowania w systemach neuromorficznych przyciągnęły nanorurki węglowe, ponieważ fizycznie przypominają biologiczne rurkowate aksony, przez które komórki nerwowe przesyłają sygnały elektryczne. Jednak zdania na temat tego, w jaki sposób materiały te wpłyną na przyszłość informatyki, są podzielone.

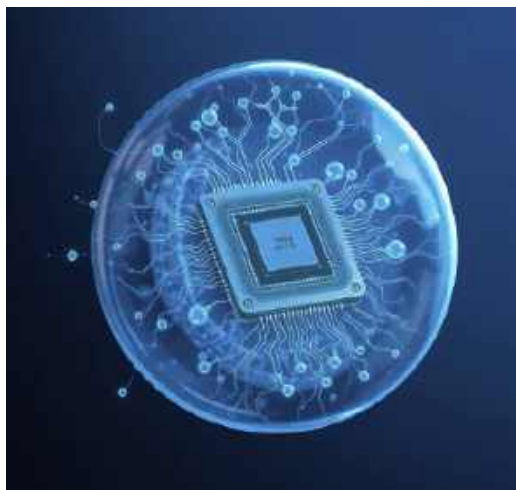
Nowe pomysły na elektronikę rodzą się również w dziedzinie fizyki cząstek czasem dość egzotycznych. Niedawno Korea Research Institute of Standards and Science (KRISS) opracował pierwszy na świecie tranzystor zdolny do kontrolowania skyrmionów, kwazicząstek będących formacjami pola magnetycznego. Można je zminiaturyzować do kilku nanometrów, dzięki czemu są ruchome przy wyjątkowo niskiej mocy. Ta cecha czyni je kluczowym elementem w rozwoju dziedziny zwanej spintroniką. Wyniki

koreańskich badań zostały opublikowane w czasopiśmie „Advanced Materials”.

Z drugiej strony dla elektroniki mniej wymagającej szuka się rozwiązań może nie z obszarów fizyki egzotycznych cząstek, ale wciąż wykorzystującej egzotyczne pomysły i materiały, takie jak... drewno. Naukowcy z uniwersytetów szwedzkich Linköping i KTH Royal Institute of Technology stworzyli pierwszy na świecie drewniany tranzystor. Wykorzystując twarde drewno balsa, pozbawili je twardej ligniny i wypełnili pozostały materiał mieszanym polimerem przewodzącym jony elektronowe o nazwie poli(3,4-etylenodioksytiofen)-polistyrenosulfonian lub PEDOT:PSS. Układając jednostki o grubości milimetra, które działały jako elektrody i kanały, odkryli, że może stworzyć prosty tranzystor. Przy braku napięcia całą strukturę można uznać za otwartą, a przełącznik ustawić w pozycji „włączony”. Po podaniu napięcia 6 V kanał wypełnia się elektronami, powoli zamykając się i przełączając przełącznik w pozycję „wyłączony”. Biodegradowalna elektronika wykonana z łatwo pozyskiwanych zasobów może być wykorzystywana w zdalnych czujnikach.

Biologiczne alternatywy

A może alternatywnych komputerów powinniśmy szukać gdzie indziej, w pewnym sensie – w samych sobie, w naszych komórkach (5). Hipotetyczna biologiczna maszyna molekularna działałaby jak rybosom – instrukcje byłyby kodowane na jednej cząsteczce organicznej, a druga interpretowałaby je lub odczytywała. Naukowcy od lat próbują taki układ, rodzaj biologicznego komputera, zbudować. W 2007 roku zespół brytyjskiego naukowca Davida Leigh zaprojektował przypominającą pierścień cząsteczkę, która była zasilana światłem i mogła poruszać się do przodu wzdłuż ścieżki molekularnej. Potem naukowcy odkryli, jak popychać te cząsteczki zapadkowe za pomocą kwasu tróchlorooctowego jako paliwa chemicznego. Maszyny znajdują się w cieczy, a badacze pulsacyjnie wprowadzają do niej kwas; pH otaczającej cieczy zmienia się, uruchamiając cząsteczkę. W dalszych eksperymentach, opisanych w „Nature”, zespół Leigh zademonstrował możliwość odczytu przez maszynę wielkości cząsteczki podczas ruchu. Zakodowano bloki informacji na jednej cząsteczce (taśmie) i zaprojektowali inną, która przesuwana się po jej długości (głowica). Gdy głowica poruszała się wzdłuż taśmy, zmieniała się w przewidywalny kształt za każdym razem, gdy skanowała określony blok informacji. Pozwoliło to badaczom zasadniczo odczytać kod z „taśmy”. Przejście między blokami informacji zajęło kilka godzin.



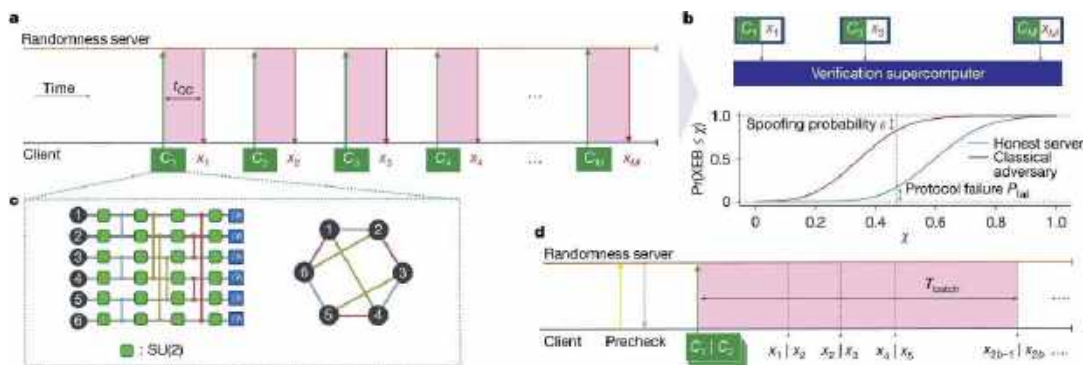
5. Wizualizacja procesora w żywej komórce © AI

W naturze rybosomy mogą odczytać około dwudziestu „cyfr” na sekundę, jest więc jeszcze sporo do zrobienia. Zespół Leigh sugerował, że zmieniające kształt cząsteczki czytnika mogą katalizować różne reakcje chemiczne w zależności od ich kształtu. Można sobie wyobrazić każdą pełną takich czytników molekularnych, wszystkie zaprogramowane do drukowania tych samych cząsteczek, działających razem jako rodzaj fabryki, np. do produkcji superpolimerów.

Kilka lat temu specjaliści ze szkoły inżynierskiej Uniwersytetu Columbia stworzyli pierwszy w historii procesor zasilany zjawiskami biologicznymi, występującymi naturalnie w skali mikro wyrzutami jonów w lipidowych błonach. Dwuwarstwowe membrany lipidowe zostały wytworzone w sposób syntetyczny. Zawierają „pompy jonowe”, zasilane przez adenosyn-5'-trifosforan (ATP, zwany też „walutą energetyczną komórek”), który odgrywa ważną rolę w biologii komórki jako wielofunkcyjny koenzym i molekularna jednostka w wewnątrzkomórkowym transporcie energii. Naukowcy połączyli membrany ze standardowym układem CMOS. W praktyce oznaczało to, że jonowe pompy biologiczne zasilają układ elektroniczny. Jak na razie mowa o stworzonym sztucznie biologicznym elemencie (membrana), a nie całej komórce. Jednak eksperyment opisany w „Nature Communications” dowodzi wykonalności takiego układu.

O pomysłach o układach przetwarzających na bazie biologiczne, Cortical Labs czy BBB, wspominamy w innym artykule w tym numerze. Tamte prace dążą do stworzenia biochipów na użytek obliczeń AI. Te światy przenikają się w poszukiwaniu alternatyw dla procesorów opartych na krzemie. ■

Mirosław Usidus



1. Przegląd protokołu kwantowego do generowania losowych liczb opublikowany w „Nature”

Certyfikowaną losowość liczb generowanych przez 56-kubitowy komputer kwantowy zademonstrowali (1) uczeni z JPMorganChase, Quantinuum, laboratorium Argonne i Uniwersytetu Teksaskiego. Choć brzmi to nieco enigmatycznie, ma to ogromne znaczenie w sferze kryptografii. Klasyczne komputery nie potrafią generować prawdziwie losowych liczb w sposób bezpieczny.

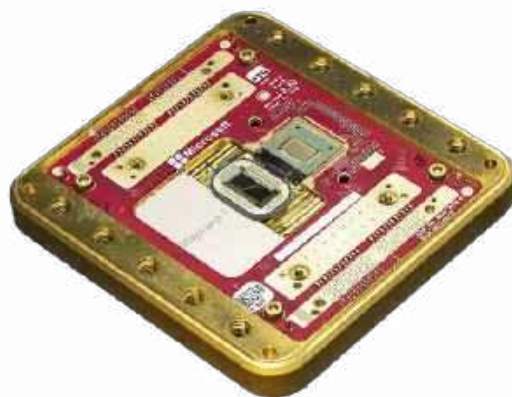
Komputery kwantowe – są czy ich nie ma, i czy w ogóle są potrzebne?

W KUBITACH PLĄTA SIĘ WIĘCEJ PYTAŃ NIŻ ODPOWIEDZI

Sukces amerykańskich badaczy to rodzaj przełomu i zarazem odpowiedź na stale przewijające się pytanie – do czego w praktyce mogą się przydać komputery kwantowe. Otóż choćby do generowania liczb będących podstawą szyfrów „nie do złamania”. Takie osiągnięcia nie rozwiewają jednak wątpliwości, których mnóstwo krąży nad całą sferą informatyki kwantowej. Te wątpliwości czasem nawet prowadzą do kwestionowania rozpowszechnionego już dziś twierdzenia, że w ogóle zbudowaliśmy „prawdziwe” komputery kwantowe

To są te kubity czy ich nie ma?

W lutym 2025 r. Microsoft zaprezentował uroczystość swój procesor Majorana 1 (2), głosząc, że to kwantowe urządzenie obliczeniowe zawierające osiem



2. Majorana 1

ultraszybkich „topologicznych” kubitów. W kolejnych tygodniach zewnętrzni badacze problemu zaczęli kwestionować twierdzenia firmy. Podczas konferencji fizyków, która odbyła się w marcu w Kalifornii, badacz z Microsoftu Chetan Nayak próbował przekonywać zgromadzonych, że jego zespół stworzył pierwszy na świecie „topologiczny” kubit, który zasila konwencjonalny komputer cyfrowy. Jednak uczestnicy mieli wątpliwości. Jak podkreślano, działanie tego urządzenia wymagałoby nie tylko wyczarowania

długo poszukiwanej kwazicząstki Majorany, która nigdy wcześniej nie została potwierdzona eksperymentalnie, ale także kontrolowania wielu takich cząstek na rzeczywistej platformie w celu zakodowania informacji kwantowej. „Nie sądzę, by dane były przekonujące” – mówiła Jelena Klinovaja, fizyk z uniwersytetu w Bazylei, która wysłuchała wykładu Nayaka. Henry Legg, fizyk z University of St. Andrews, który opublikował już dwie prace kwestionujące twierdzenia Microsoftu, mówił: „to nie wygląda jak kwazicząstka Majorany, przynajmniej dla mnie”. „Każda firma, która twierdzi, że w 2025 roku ma topologiczny kubit, zasadniczo sprzedaje bajkę, która szkodzi dziedzinie obliczeń kwantowych... i podważa publiczne zaufanie do nauki”. „Ujawniliśmy tylko niewielki ułamek tego, co zrobiliśmy”, bronił się Nayak w rozmowie z „Science”. „Będzie to wyglądało coraz bardziej przekonująco, gdyż będzie to podstawa działającej technologii”.

Artykuł zespołu z Microsoftu na temat układu nie zawierał szczegółowych informacji na temat chipa ani dowodu na istnienie cząstki Majorany, skupiając się na metodzie pomiaru niektórych właściwości kwantowych na przyszłym urządzeniu. Publikacja miała spory rozgłos (MT też o tym informował). Firmie gratulowali politycy i Elon Musk. Dyrektor generalny Microsoftu Satya Nadella odpowiedział na gratulacje Elona Muska na platformie X: „Uważamy, że może to być moment [narodzin – przyp. red.] tranzystora kwantowego”.

Wielu fizyków było zaszokowanych fanfarami wokół publikacji, która w gruncie rzeczy stanowiłaby rewolucję w nauce, gdyby twierdzenia zawarte w pracy były w pełni zgodne z faktami. „Nigdy nie widziałem czegoś takiego w mojej karierze pracy naukowej w dziedzinie fizyki”, komentował Jason Alicea, fizyk z California Institute of Technology. „To najsilniejsze twierdzenie, jakie kiedykolwiek zostało wysunięte w tej dziedzinie... ale ciężar spoczywa na nich. Niech naprawdę pokażą, że to, co mają, jest prawdziwym przełomem”.

Microsoft nad tworzeniem kubitów z cząstek Majorany pracuje od ponad dekady. Kwazicząstki te są co do zasady zdelokalizowanymi elektronami. Ponieważ elektrony nie istnieją w żadnej lokalizacji, teoretycznie informacje w nich są chronione „topologicznie” przed wszelkimi lokalnymi zakłóceniami, jeśli istnieje wystarczająco duża „luka” między najbliższymi stanami energetycznymi, do których elektrony mogą przeskoczyć, czyli dobrze znane w elektronice pasmo zabronione. Nowy projekt chipu Microsoftu to układy ultracienkich, nadprzewodzących drutów z arsenku indu, które zmuszają znajdujące

się wewnątrz elektrony do tworzenia luźno powiązanych par. W odpowiednich warunkach drut może również pomieścić dodatkowy niesparowany elektron, który „dzieli się na pół”, co oznacza pojawienie się majorany. Dwa stany „parzystości” na drucie reprezentowałyby 0 lub 1 w przyszłym komputerze i informowały o parzystości lub nieparzystości elektronów, co zespół Microsoftu zamierza mierzyć, odczytując w ten sposób informacje kwantowe.

Przez lata badania te nękane były nieoczekiwanymi przeszkodami inżynierskimi, pustymi obietnicami i twierdzeniami, które obalano. Naukowcy z Microsoftu opracowali w 2021 r. protokół, który ustanawia serię testów eksperymentalnych potwierdzających, że urządzenie rzeczywiste znajduje się w odpowiedniej fazie topologicznej. Zbudowali komputerową symulację urządzenia, którą przeszkolili w zakresie identyfikacji fazy topologicznej, a następnie wprowadzili pomiary rzeczywistego urządzenia do tego samego protokołu, aby ocenić, czy jest ono również w odpowiednim stanie. Już w 2023 r. zespół Nayaka twierdził, że stworzył urządzenie, które przeszło protokół.

Badacze zewnętrzni doceniają fakt, że zespół Microsoftu próbuje polować na kwazicząstki majorany w bardziej metodyczny sposób, ale krytykują zespół za rozpowszechnianie twierdzeń, które ich zdaniem są przesadzone. Tydzień po lutowym ogłoszeniu Microsoftu, wspomniany już fizyk Henry Legg opublikował artykuł krytykujący protokół używany do identyfikacji majoranów. Zagłębiając się w jego kod, zauważył, że zwykła zmiana zakresu, w którym mierzono różne parametry, takie jak pole magnetyczne, wydaje się wpływać na to, czy urządzenie zostanie uznane za topologiczne. Ponadto, jak zauważył, kod używany do oceny rzeczywistych danych wydaje się mniej restrykcyjny niż kod używany do symulacji danych. Zarzuty te odpierał w poście na LinkedIn inny badacz z Microsoftu, Roman Lutchyn, argumentując, że wrażliwość protokołu na zmiany parametrów jest „oczekiwana i nie unieważnia użycia protokołu z prawidłowo dobranymi parametrami”. Jeśli chodzi o różne wersje kodu, Lutchyn wskazuje, że rozbieżność jest „statystycznie nieistotna” – dając tylko jeden fałszywie dodatni wynik dla 700 regionów zainteresowania. „To, co tutaj zrobiliśmy, jest zarówno poprawne teoretycznie, jak i praktyczne”, powiedział Leggowi podczas wspomnianej wcześniej konferencji. „Jeśli masz lepszy pomysł na znalezienie majorany, przedstaw protokół, a następnie wszyscy go przestrzegajmy”.

Nayak pokazał dane sugerujące, że pojedynczy nanodrut może utrzymywać stan 0 lub 1 przez maksymalnie

10 milisekund. Ale dane dotyczące dwóch bardziej złożonych stanów były znacznie mniej jasne. Analiza statystyczna sugerowała, że stany te utrzymywały się przez kilka mikrosekund, ale niektórzy fizycy twierdzili, że dane wyglądały bardziej jak szum.

Anton Akhmerov, fizyk z uniwersytetu w Delft, nazywa niektóre obawy Legga dotyczące protokołu „bardzo ważnymi”, chociaż nie uważa, że całkowicie unieważniają one pracę Microsoftu. „Aby społeczność mogła zaufać, że wynik protokołu jest wiarygodny, należy bezwzględnie systematycznie to badać”, kwituje.

Wojna z błędami i jej owoc – Willow

Wcześniej, w październiku ub. roku ukazała się w „Nature” publikacja o eksperymentach na 67-kubitowym procesorze Sycamore firmy Google, wykazujących, że operacje te weszły w nową „fazę słabego szumu”. Badacze Google Quantum AI odkryli „stabilną obliczeniowo złożoną fazę”, którą można osiągnąć za pomocą istniejących jednostek przetwarzania kwantowego (QPU), znanych również jako procesory kwantowe. „Koncentrujemy się na opracowywaniu praktycznych zastosowań dla komputerów kwantowych, których nie można wykonać na klasycznym komputerze”, wyjaśniali przedstawiciele Google Quantum AI w serwisie „Live Science”. „Jednak dane generowane przez komputery kwantowe są nadal zaszumione, co oznacza, że wraz ze wzrostem liczby kubitów muszą wykonywać dość intensywną kwantową korekcję błędów, aby kuby pozostały w fazie słabego szumu”.

Kubity, które są osadzone w jednostkach QPU, wykorzystują zasady mechaniki kwantowej w celu równoległego wykonywania obliczeń. Dla porównania – klasyczne bity obliczeniowe mogą przetwarzać dane tylko w sekwencji. Im więcej kubitów znajduje się w jednostce QPU, tym potężniejsza w trybie wykładniczym staje się maszyna. Ze względu na możliwości przetwarzania równoległego, obliczenia, które zajęłyby klasycznemu komputerowi tysiące lat, mogą zostać wykonane przez komputer kwantowy w ciągu kilku sekund. Ale kubity są bardzo wrażliwe i podatne na błędy z powodu zakłóceń. Około jednego na sto kubitów ulega awarii. To dużo więcej niż wartość jeden na miliard miliardów dla bitów klasycznych. Na błędy wpływają zakłócenia środowiskowe, zmiany temperatury, pola magnetyczne, a nawet promieniowanie kosmiczne. Tak wysoki poziom błędów oznacza, że aby osiągnąć „kwantową supremację”, potrzebne byłyby takie rozwiązania korekcji błędów, które jeszcze nie istnieją, ewentualnie komputer kwantowy z milionami kubitów, co rozwiązywałoby problem przez ogromną skalę. Skalowanie komputerów kwantowych nie jest

łatwe – największa liczba kubitów w pojedynczej maszynie wynosi obecnie około tysiąca.

Z eksperymentów przeprowadzonych przez naukowców Google wynika, że komputery kwantowe mogą przezwyciężyć problem szumów i przewyższyć klasyczne komputery w określonych obliczeniach. Jednak wciąż wymagana byłaby korekcja błędów, gdy maszyny będą się skalować. Naukowcy wykorzystali metodę znaną jako losowe próbkowanie obwodu (RCS) do przetestowania bezbłędności dwuwymiarowej siatki nadprzewodzących kubitów. RCS jest punktem odniesienia, który mierzy wydajność komputera kwantowego w porównaniu z wydajnością klasycznego superkomputera. Eksperymenty ujawniły, że działające kubity mogą przechodzić między pierwszą fazą a drugą fazą, zwaną „fazą słabego szumu”, przez wywołanie określonych warunków. W eksperymentach naukowcy sztucznie zwiększyli szum lub spowolnili rozprzestrzenianie się korelacji kwantowych. W tej drugiej „fazie słabego szumu” obliczenia były na tyle złożone, że doszli oni do wniosku, że komputer kwantowy może przewyższać komputer klasyczny. Zademonstrowali to na 67-kubitowym chipie Google Sycamore.

Bity kwantowe są podatne na więcej rodzajów błędów niż ich klasyczni kuzyni. Znacznie trudniej jest też nimi manipulować. Każdy krok w obliczeniach kwantowych jest kolejnym źródłem błędów, podobnie jak sama procedura korekcji błędów. Co więcej, nie ma sposobu na zmierzenie stanu kubitów bez nieodwracalnego zakłócania go – trzeba w jakiś sposób zdiagnozować błędy, nigdy ich bezpośrednio nie obserwując.

Kwantowy kod korekcji błędów został opracowany przez rosyjskiego fizyka Aleksieja Kitajewa kilka dekad temu. Opiera się na dwóch nakładających się siatkach fizycznych kubitów. Te w pierwszej siatce to kubity „danych”. Wspólnie kodują one pojedynczy kubit logiczny. Te w drugiej to kubity „pomiarowe”. Pozwalają one badaczom pośrednio wykrywać błędy bez zakłócania obliczeń. „To tak naprawdę jedyna znana nam droga do zbudowania komputera kwantowego na dużą skalę”, zauważa Michael Newman, badacz zajmujący się korekcją błędów w Google Quantum AI. Ta obliczeniowa alchemia ma swoje ograniczenia. Jeśli fizyczne kubity są zbyt podatne na awarie, korekcja błędów przyniesie efekt przeciwny do zamierzonego – dodanie większej liczby fizycznych kubitów sprawi, że logiczne kubity będą działać gorzej, a nie lepiej. Jeśli jednak poziom błędów spadnie poniżej określonego progu, równowaga się zmienia. Im więcej fizycznych kubitów dodajemy, tym bardziej odporny stanie się każdy logiczny kubit. Newman i jego koledzy z Google

Quantum AI uważają, że przekroczyli ten próg. Przekształcili grupę fizycznych kubitów w pojedynczy kubit logiczny, a następnie wykazali, że w miarę dodawania kolejnych fizycznych kubitów do grupy poziom błędów kubitów logicznych spadał.

Zespół Google Quantum AI spędził lata na doskonaleniu swoich procedur projektowania i produkcji kubitów, skalując się z garstki kubitów do dziesiątek i doskonaląc swoją zdolność do manipulowania wieloma kubitami jednocześnie. W 2021 r. spróbowali po raz pierwszy wypróbować korekcję błędów za pomocą kodu powierzchniowego. W artykule z 2023 r. zespół poinformował, że poziom błędów kodu w eksperymencie znanym jako „distance-5” był nieznacznie niższy niż kodu we wcześniejszym „distance-3”. Był to zachęcający wynik, ale niejednoznaczny – nie mogli jeszcze ogłosić zwycięstwa. Na początku 2024 roku mieli do przetestowania zupełnie nowy 72-kubitowy chip o nazwie kodowej Willow. Wyniki były bardzo zachęcające

Do tego stopnia, że pod koniec 2024 roku Google zademonstrował publicznie swój nowy chip kwantowy nazwany Willow, oparty na technice „gate based”, a nie na wyżarzaniu jak we wcześniejszych konstrukcjach. Według firmy, nowy układ może wykładniczo redukować błędy w miarę zwiększania skali przy użyciu większej liczby kubitów. Willow wykonał standardowe obliczenia wzorcowe w czasie krótszym niż pięć minut, co zajęć miałoby jednemu z najszybszych dzisiejszych superkomputerów dziesięć septilionów (czyli 10^{25}) lat. Willow wykorzystuje do zwiększania dokładności i wydajności obliczeń wspomnianą technikę RCS, opracowaną przez Quantum AI Lab.

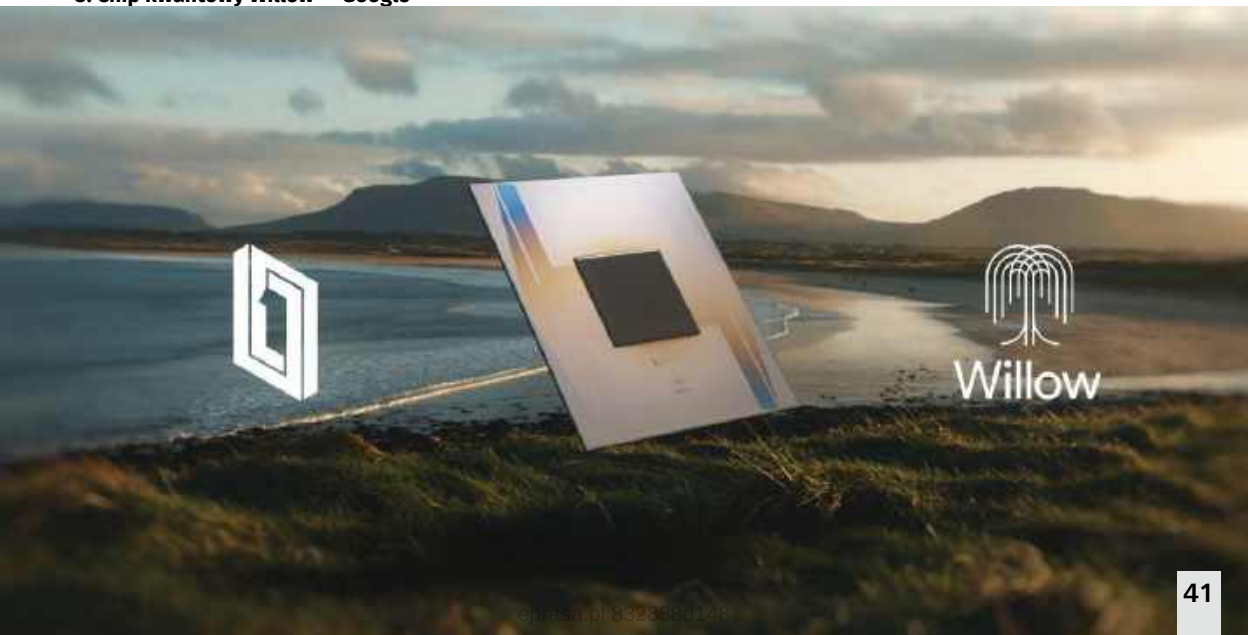
Układ Willow składa się z 105 kubitów (3). To około dwukrotnie więcej niż poprzedni układ Google Sycamore. Tym razem jednak nacisk kładziony jest nie tylko na liczbę kubitów, ale także na ich jakość. Kubyty w tej nowej jednostce mają mieć znacznie lepszy czas retencji (znany jako czas T1), który został zwiększony około pięciokrotnie w porównaniu do poprzednich chipów kwantowych, co pozwala im przechowywać informacje przez dłuższy czas. Premierze chipa Google towarzyszyła firmowana przez badaczy tej firmy publikacja w „Nature” na temat metod korekcji błędów „poniżej progu kodu powierzchniowego”, którą wyżej omawiamy. Wypowiadając się przy okazji prezentacji układu, założyciel zespołu Google Quantum AI, Hartmut Neven, zauważył, że wydajne działanie chipa wspiera ideę obliczeń kwantowych zachodzących w wielu równoległych wszechświatach, zgodnie z interpretacjami mechaniki kwantowej opartymi na multiwersum, co nawiązuje do myśli brytyjskiego fizyka Davida Deutscha, który jako jeden z pierwszych zasugerował, że obliczenia kwantowe mogą obejmować wszechświaty równoległe.

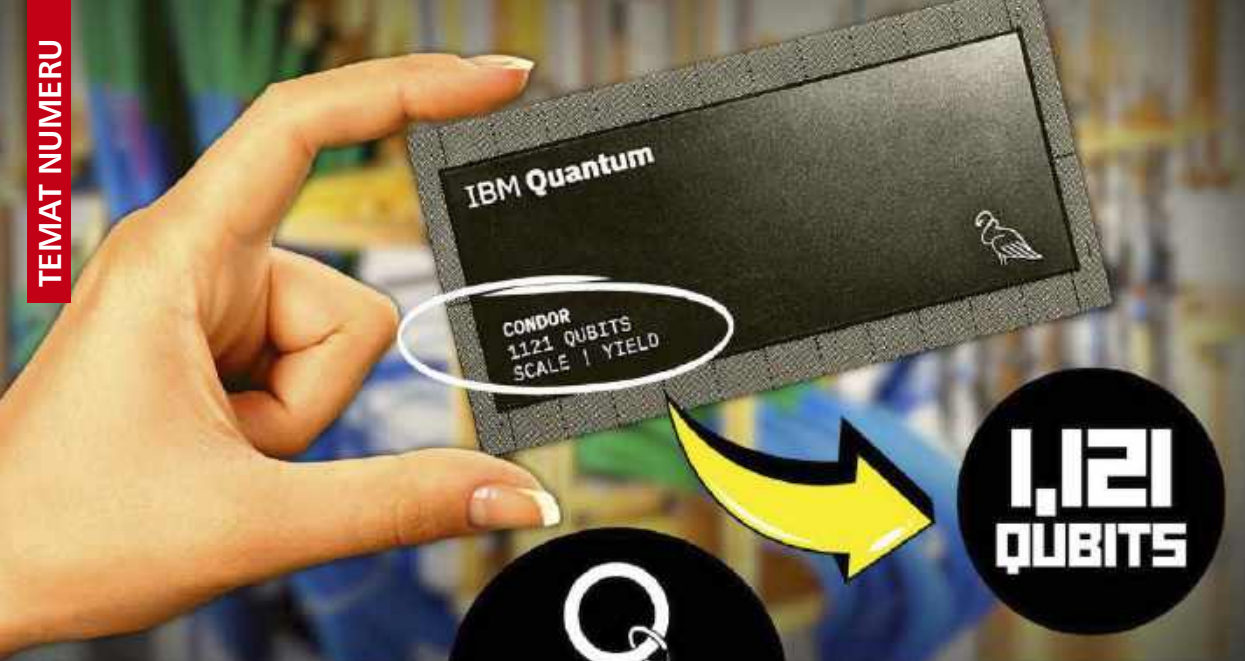
Naukowcy zdają sobie sprawę, że przed nimi jeszcze długa droga. Zespół Google Quantum AI zademonstrował korekcję błędów przy użyciu tylko jednego logicznego kubitów. Dodanie interakcji między wieloma logicznymi kubitami wprowadzi nowe wyzwania eksperymentalne.

Seria „ważnych ogłoszeń”

Krótko przed premierą Willowa pojawiła się informacja o wynikach badań przeprowadzonych na brytyjskim Uniwersytecie Kent, wykazujących, że informacje

3. Chip kwantowy Willow © Google





4. Prezentacja chipa kwantowego IBM Condor © IBM

kwantowe mogą być wykorzystane do koordynowania działań poruszających się urządzeń, takich jak drony lub samochody autonomiczne. Mogłoby to, według publikacji, która ukazała się w „New Journal of Physics”, prowadzić do bardziej wydajnej logistyki. Eksperymenty symulowały ruch w sieci transportowej przy użyciu prawdziwych kubitów komputera kwantowego opracowanego przez IBM.

W grudniu tenże IBM zaprezentował swój najnowszy procesor, nazwany Heron, który może pochwalić się 133 kubitami, najwyższą w historii firmy redukcją błędów i możliwością połączenia w ramach pierwszego modułowego komputera kwantowego, System Two. Ponadto IBM zaprezentował swój kolejny chip kwantowy, Condor (4), który ma 1121 nadprzewodzących kubitów ułożonych we wzór plastra miodu. IBM wykorzystuje nadprzewodzące kubity, które wymagają chłodzenia do niemal zera absolutnego w celu złagodzenia szumów termicznych, zachowania spójności kwantowej i zminimalizowania interakcji środowiskowych.

Seria „ważnych ogłoszeń” w dziedzinie obliczeń kwantowych trwa. Podczas konferencji GTC Global AI Conference w marcu 2025 r. NVIDIA ogłosiła uruchomienie centrum badawczego Accelerated Quantum Computing (NVAQC), ośrodka zajmującego się pomocą dla zainteresowanych przeprowadzaniem obliczeń kwantowych za pomocą najbardziej zaawansowanych układów NVIDIA GB200 NVL72 Grace Blackwell w wykorzystaniu architektury NVIDIA DGX Quantum. Konkretnie usługi to np. korekcja błędów w układach kwantowych i tworzenie aplikacji hybrydowych.

Pomaga w tym sztuczna inteligencja i superkomputery. Tim Costa z NVIDIA powiedział podczas prezentacji: „Centrum będzie miejscem symulacji na dużą skalę algorytmów kwantowych i sprzętu, ścisłej integracji procesorów kwantowych oraz zarówno szkolenia, jak i wdrażania modeli sztucznej inteligencji do dziedziny obliczeń kwantowych”.

Zamiast jednego wielkiego – sieć komputerów splecionych kwantowo

Wykonywanie złożonych algorytmów na komputerach kwantowych będzie ostatecznie wymagało dostępu do dziesiątek tysięcy sprzętowych kubitów. W przypadku większości znanych technik budowy kubitów, czy to nadprzewodzących, czy „pułapek jonowych” czy innych, stwarza to problemy, konstrukcyjne. Specjaliści szukają różnych pomysłów na to, jak połączyć procesory ze sobą, by działały w sieci (5) jako pojedyncza jednostka obliczeniowa. W lutym „Nature”, zespół z Uniwersytetu Oksfordzkiego opisuje wykorzystanie teleportacji kwantowej do połączenia dwóch elementów do obliczeń kwantowych, które znajdowały się w odległości około dwóch metrów od siebie, przy czym z powodzeniem mogły znajdować się w zupełnie innych pomieszczeniach. Po połączeniu oba elementy sprzętu można traktować jako pojedynczy komputer kwantowy, umożliwiając wykonywanie algorytmów i operacji w obu układach niczym w jednym większym.

Teleportacja kwantowa działa tak, że wstępnie umieszcza się obiekty kwantowe zarówno na „źródle”, jak i w punkcie odbiorczym, czyli obu końcach

teleportu, i splątuje się je. Proces wykonywania teleportacji polega na pomiarze obiektu źródłowego, który niszczy jego stan kwantowy, nawet gdy pojawia się on w odległym miejscu. To ważne, ponieważ zasady mechaniki kwantowej dyktują, że nie można po prostu skopiować stanu kwantowego. Gdyby pomyśleć o tym jako o sposobie wymiany informacji między różnymi chipami kwantowymi, to część obliczeń na jednym chipie wykonywana jest przed teleportowaniem odpowiedzi w toku do drugiego w celu dalszej pracy. W skrócie – można stworzyć uniwersalny komputer kwantowy, zdolny do wykonywania dowolnego algorytmu kwantowego, wyłącznie przez teleportację. Interfejs między światem klasycznym i kwantowym stwarza ryzyko błędu. Teleportacja jest natomiast bezstratna; jej poziom błędu powinien być taki sam, jak w przypadku każdej operacji wykonywanej na pojedynczym elemencie lokalnego sprzętu. Ze względu na jej zalety, ludzie już wcześniej zastanawiali się nad włączeniem teleportacji bramek do algorytmów. Przeprowadzono też szereg demonstracji pomiędzy sprzętowymi kubitami znajdującymi się w różnych częściach jednego systemu. Do tej pory nie było jednak doniesień o jej wykorzystaniu między fizycznie oddzielnymi elementami sprzętu.

Zespół z Oksfordu zbudował uproszczony system. Na każdym końcu, po obu stronach dwumetrowej

przestrzeni oddzielającej, znajdowała się pojedyncza pułapka zawierająca dwa jony, jeden strontu i jeden wapnia. Oba atomy można było ze sobą splątać, dzięki czemu działały jako pojedyncza jednostka. Jon wapnia służył jako lokalna pamięć i był wykorzystywany w obliczeniach, podczas gdy jon strontu służył jako jeden z dwóch końców sieci kwantowej. Kabel optyczny między dwiema pułapkami jonowymi pozwolił fotonom splątać dwa jony strontu, dzięki czemu cały system działał jako pojedyncza jednostka. Ważne jest to, że niepowodzenie w splątaniu pozostawiało system w jego pierwotnym stanie, co oznacza, że naukowcy mogli po prostu próbować dalej, aż kuby zostały splątane. Zdarzenie splątania doprowadziłoby również do powstania fotonu, który można by zmierzyć, co dawało informację, kiedy osiągnięto sukces. Po wykonaniu wielu rund tych bramek zespół odkrył, że typowa wierność wynosiła około 70 proc. Błędy zazwyczaj nie miały nic wspólnego z procesem teleportacji i były wynikiem lokalnych operacji na jednym z dwóch końców sieci. Badacze sądzą, że użycie komercyjnego sprzętu, który ma znacznie niższe wskaźniki błędów, znacznie poprawiłoby sytuację.

Podejście naukowców z Oksfordu odzwierciedla sposób działania tradycyjnych superkomputerów. Łącząc mniejsze jednostki, osiągają one możliwości znacznie przekraczające możliwości pojedynczej jednostki. Pod pewnymi względami koncepcja ta odzwierciedla sposób działania tradycyjnych superkomputerów. Zamiast jednego potężnego procesora, superkomputery opierają się na tysiącach mniejszych jednostek obliczeniowych pracujących równolegle. W przypadku obliczeń kwantowych strategia ta omija przeszkody inżynierskie związane z upakowaniem większej liczby kubitów w jednym urządzeniu przy jednoczesnym zachowaniu delikatnych właściwości kwantowych niezbędnych do dokładnych obliczeń.

Aby zademonstrować moc swojego systemu, naukowcy uruchomili algorytm wyszukiwania Grovera, metodę kwantową, która może przeszukiwać duże, nieustrukturyzowane zbiory danych znacznie szybciej niż klasyczne komputery. Eksperyment ten pokazał, że rozproszone w sieci kwantowe przetwarzanie informacji jest wykonalne przy użyciu obecnej technologii. Skalowanie tego systemu będzie jednak wymagało pokonania znaczących przeszkód technicznych, od poprawy stabilności kubitów po udoskonalenie łączących je ogniw fotonicznych.

Koncepcja, by połączyć procesory kwantowe na ogromne odległości, tworząc gigantyczną kwantową sieć obliczeniową, która działa jak pojedyncza maszyna, była znana wcześniej. Niedawno fizycy

5. Wizualizacja sieci komputerów kwantowych





6. Zespół badaczy, który opracował QNodeOS © TU Delft

stworzyli nowy model komputerów kwantowych, który może ułatwić ich skalowanie i uczynić je potężniejszymi, opisując swoje pomysły w artykule, który ukazał się w maju 2024 r. w czasopiśmie „PRX Quantum”. Naukowcy uważają, że przez nadanie każdemu kubitowi dodatkowych częstotliwości można zmusić je do współpracy w celu przetwarzania obliczeń tak, jakby były częścią jednego komputera kwantowego. Działaloby się tak, ich zdaniem, pomimo potencjalnej separacji na duże odległości. Oznacza to, że zamiast jednego masywnego procesora kwantowego, który jest trudny w utrzymaniu, można użyć kilku mniejszych połączonych ze sobą. „Każdy kubit w komputerze kwantowym działa na określonej częstotliwości. Jest możliwość kontrolowania każdego kubitu indywidualnie za pomocą odrębnej częstotliwości, a także łączenia par kubitów poprzez dopasowanie ich częstotliwości”, powiedział w komunikacie główny autor badania Vanita Srinivasa, adiunkt informacji kwantowej na University of Rhode Island. Badacze uważają, że stosując oscylujące napięcia, mogą generować dodatkowe częstotliwości dla każdego kubitu. W ten sposób można połączyć ze sobą wiele kubitów, wykorzystując nowo wygenerowane wspólne częstotliwości, bez konieczności dopasowywania ich do oryginalnych częstotliwości. Następnie kubity można połączyć ze sobą, ale można je również kontrolować indywidualnie przy użyciu ich oryginalnych częstotliwości.

Kwenty + AI – oczywistych korzyści na razie nie widać

Pomimo wielu ograniczeń, wątpliwości, czasem dotyczących samej kwestii, czy w ogóle udało się komukolwiek zbudować „prawdziwy komputer kwantowy”, dziedzina ta pełna jest doniesień o rozwoju nowych urządzeń i rozwiązań. Na przykład w marcu 2025 naukowcy z Quantum Internet Alliance (QIA) ogłosili w „Nature” stworzenie pierwszego systemu operacyjnego przeznaczonego dla sieci kwantowych, QNodeOS. „System jest jak oprogramowanie na komputerze w domu: nie musisz wiedzieć, jak działa sprzęt, aby z niego korzystać”, wyjaśnia w publikacji Mariagrazia Iuliano, jedna z zaangażowanych w prace nad nowym systemem (6). Usuwając barierę między sprzętem sieciowym a oprogramowaniem, system operacyjny ma pozwolić programistom na łatwe tworzenie aplikacji, torując drogę do rozwoju „oprogramowania kwantowego”. W przeciwieństwie do komputerów kwantowych, które uruchamiają pojedyncze programy, kwantowe aplikacje sieciowe wymagają oddzielnych programów do niezależnego wykonywania w różnych węzłach sieci, takich jak aplikacja kliencka na telefonie i serwer w chmurze. Programy te muszą koordynować się ze sobą za pomocą wiadomości i splątania kwantowego, specjalnego rodzaju połączenia kwantowego, które nadaje sieciom kwantowym ich moc. QNodeOS ma z tymi różnymi wyzwaniem sobie radzić.

Chiński fizyk Pan Jianwei opracował kilka lat temu komputer kwantowy Jiuzhang, który według niego może wykonywać niektóre rodzaje obliczeń związanych ze sztuczną inteligencją około 180 milionów razy szybciej niż najlepszy superkomputer na świecie. W artykule opublikowanym w recenzowanym czasopiśmie „Physical Review Letters” w maju ubiegłego roku, Jiuzhang przetworzył w mniej niż sekundę ponad dwa tysiące próbek dwóch popularnych algorytmów związanych ze sztuczną inteligencją, Monte Carlo i symulowanego wyżarzania, których przetwarzanie zajęłoby najszybszemu klasycznemu superkomputerowi na świecie pięć lat. W październiku Pan zaprezentował Jiuzhang 3.0 (7), który według niego był dziesięć kwadrylionów razy szybszy w rozwiązywaniu niektórych problemów niż klasyczny superkomputer. Jiuzhang wykorzystuje fotoniczną odmianę obliczeń kwantowych.

Chińskie badania są sygnałem, że naukowcy badają potencjał kwantowego uczenia maszynowego. Nie jest jednak jasne, czy połączenie sztucznej inteligencji i obliczeń kwantowych znajdzie użyteczne zastosowania. Jeśli komputery kwantowe uda się kiedykolwiek zbudować w wystarczająco dużej skali, to mogłyby rozwiązywać pewne problemy znacznie wydajniej niż zwykła elektronika cyfrowa. Przez lata zastanawiano się, czy problemy te mogą obejmować uczenie maszynowe. Nastawienie naukowców do kwantowego uczenia maszynowego waha się między dwiema skrajnościami, wskazuje Maria Schuld, fizyk z Durbau w RPA. Zainteresowanie tym podejściem jest duże, ale naukowcy wydają się coraz bardziej zrezygnowani brakiem perspektyw na krótkoterminowe zastosowania. Niektórzy badacze skupiają się na pomysłach zastosowania algorytmów kwantowego uczenia maszynowego do zjawisk, które są z natury kwantowe.

W ciągu ostatnich 20 lat naukowcy zajmujący się obliczeniami kwantowymi opracowali mnóstwo algorytmów kwantowych, które teoretycznie mogłyby zwiększyć wydajność uczenia maszynowego na potrzeby sztucznej

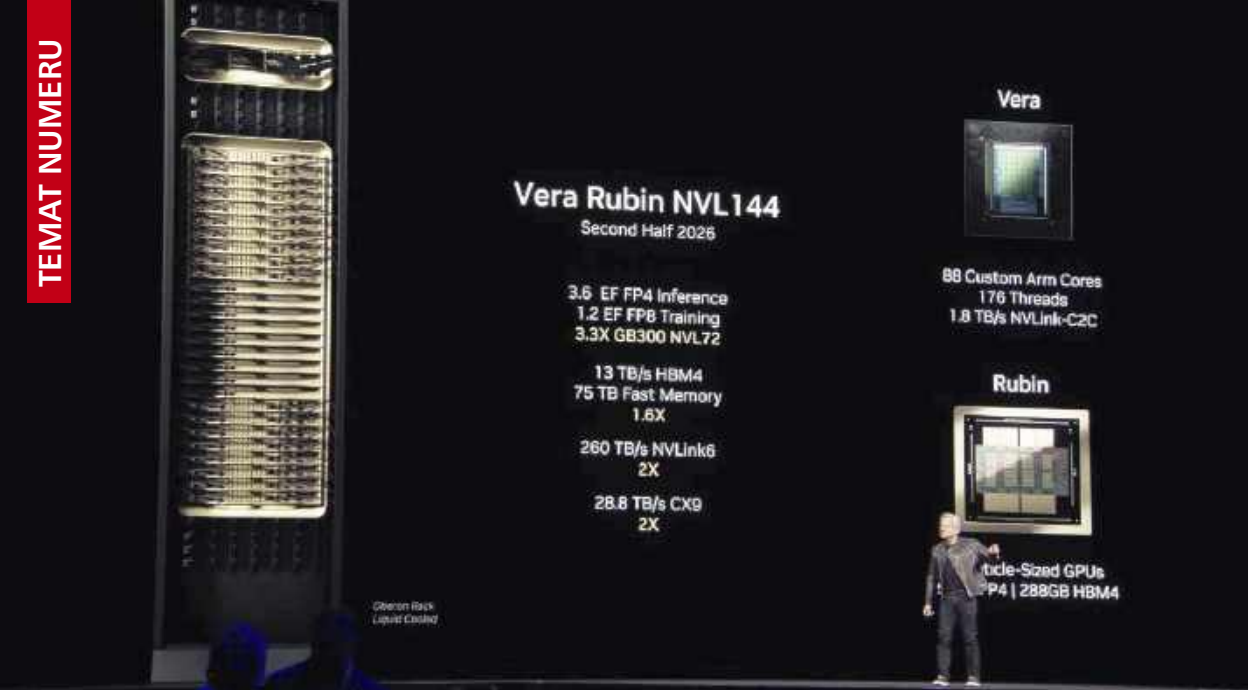
inteligencji. Jednak w niektórych przypadkach algorytmy kwantowe nie sprawdziły się w tej dziedzinie. Potencjalnie jeszcze większym problemem jest to, że klasyczne dane i obliczenia kwantowe nie zawsze dobrze się ze sobą łączą. Z grubsza rzecz biorąc, typowa aplikacja do obliczeń kwantowych składa się z trzech głównych etapów. Najpierw komputer kwantowy jest inicjowany, co oznacza, że jego poszczególne jednostki pamięci, zwane kubitami, są umieszczane w zbiorowym splątanym stanie kwantowym. Następnie komputer wykonuje sekwencję operacji, kwantowy odpowiednik operacji logicznych na klasycznych bitach. W trzecim kroku komputer dokonuje odczytu, na przykład mierząc stan pojedynczego kubitu, który przenosi informacje o wyniku operacji kwantowej. Może to być na przykład informacja o tym, czy dany elektron wewnątrz maszyny obraca się zgodnie z ruchem wskazówek zegara, czy przeciwnie. Jednak w wielu zastosowaniach pierwszy i trzeci krok mogą być niezwykle powolne, niwelując potencjalne korzyści. Etap inicjalizacji wymaga załadowania „klasycznych” danych do komputera kwantowego i przetłumaczenia ich na stan kwantowy, co często jest procesem złożonym i bardzo nieefektywnym. A ponieważ fizyka kwantowa jest z natury probabilistyczna, odczyt często zawiera element losowości, w którym to przypadku komputer musi wielokrotnie powtarzać wszystkie trzy etapy i uśredniać wyniki, aby uzyskać ostateczną odpowiedź.

Eksperti nie widzą wyraźnego powodu, by algorytmy uczenia maszynowego miały korzystać z przetwarzania kwantowego. Jego mała szybkość zniechęca, choć, jak się zauważa, nie zawsze to musi być najważniejsze, bo są powody, by sądzić, że kwantowy system sztucznej inteligencji oparty na uczeniu maszynowym mógłby nauczyć się rozpoznawać wzorce w danych, których jego klasyczne odpowiedniki by nie dostrzegły. Zatem mariaż komputerów kwantowych z AI stoi w tej chwili pod znakiem zapytania, jak wiele innych kwestii dotyczących świata obliczeń kwantowych. ■

Mirosław Usidus

7. Schemat fotonicznej maszyny kwantowej Jiuzhang 3.0





1. Jensen Huang prezentuje układ Vera Rubin

Od lat trwa zacięta walka, nie tylko między mocarstwami, ale również pomiędzy producentami i standardami techniki półprzewodnikowej, co czasem przecina się, a nawet koliduje z interesami rywalizujących państw. Krzem od dawna nie jest sprawą wyłącznie techniki i nauki.

Wojny krzemowe – kolejny etap

BÓJ DO OSTATNIEGO NANOMETRA

W marcu 2025 podczas konferencji GTC 2025 firmy NVIDIA w San Jose jej dyrektor generalny, Jensen Huang, zapowiedział całą serię nowych procesorów graficznych firmy, które mają się pojawiać w sprzedaży sukcesywnie w ciągu najbliższych miesięcy. Pierwszy z nich to układ GPU, nazwany Vera Rubin (1), który ma pojawić się na rynku w drugiej połowie 2026 roku, wyposażony w dziesiątki gigabajtów pamięci i niestandardowy procesor zaprojektowany przez NVIDIA. Zapewnić ma znaczny wzrost wydajności w porównaniu

do swojego poprzednika, Grace Blackwell, szczególnie w zadaniach związanych z wnioskowaniem i szkoleniem sztucznej inteligencji. Vera Rubin, która technicznie jest dwiema GPU w jednym, może obsługiwać do 50 petaflopsów mocy obliczeniowej podczas wnioskowania (tj. uruchamiania modeli AI). To ponad dwukrotnie więcej niż dwadzieścia petaflopsów w układzie Blackwell. Co więcej, Vera ma być ok. dwóch razy szybsza. W drugiej połowie 2027 roku pojawi się ma Rubin Ultra, układ czterech procesorów graficznych w jednym pakiecie, zapewniający wydajność do 100 petaflopsów. W bliższej perspektywie, w drugiej połowie 2025 roku, NVIDIA wyda Blackwell Ultra, procesor graficzny, który będzie dostępny w kilku konfiguracjach. Pojedynczy układ Ultra będzie oferował te same dwadzieścia petaflopsów wydajności AI, ale z 288 GB pamięci w porównaniu do 192 GB w obecnym Blackwell. Na dalekim horyzoncie zapowiadane są z kolei układy GPU Feynman, nazwane na cześć sławnego amerykańskiego fizyka teoretycznego Richarda Feynmana. NVIDIA planuje wprowadzić Feynmana, który zastąpi Vera Rubin, na rynek, w 2028 roku.

Komunikatów w marcu było więcej; np. NVIDIA ogłosiła partnerstwo z koncernem samochodowym GM w celu budowy autonomicznych pojazdów i za powiedziała, że będzie współpracować z sektorem telekomunikacyjnym, aby pomóc w rozwoju sieci 6G. Firma rekrutuje też inżynierów na Tajwanie w celu rozwoju produkcji ASIC (Application-Specific Integrated Circuit). Układy ASIC są znacznie bardziej wydajne niż układy GPU do wnioskowania AI, podobnie jak w przypadku wydobywania kryptowalut. Układy GPU NVIDIA z serii H zoptymalizowane pod kątem zadań szkolenia sztucznej inteligencji znalazły szerokie zastosowanie. Jednak rynek chipów AI przechodzi obecnie zmianę w kierunku chipów wnioskowania lub układów ASIC. Wzrost ten wynika z zapotrzebowania na układy zoptymalizowane pod kątem rzeczywistych zastosowań sztucznej inteligencji, takich jak duże modele językowe i generatywna sztuczna inteligencja. Według Verified Market Research, rynek chipów AI do wnioskowania ma wzrosnąć z 15,8 mld USD w 2023 roku do 90,6 mld USD do 2030 roku. Główni gracze korzystają już z niestandardowych projektów ASIC; np. Google w swoim chipie AI Trillium, który został ogólnie udostępniony w grudniu 2024 roku. Przejście na niestandardowe chipy AI zaostrzyło konkurencję wśród gigantów półprzewodnikowych. Firmy takie jak Broadcom i Marvell zyskały na znaczeniu i wartości dzięki współpracy z dostawcami usług w chmurze w celu opracowania specjalistycznych chipów dla centrów danych. Aby móc w przyszłości sprostać konkurencji, nowy dział ASIC firmy NVIDIA ściąga specjalistów w tej dziedzinie, np. od znanego producenta półprzewodników mobilnych, MediaTek.

NVIDIA ma też, na co wskazują doniesienia medialne, ambitne plany wprowadzenia na rynek własnego procesora, układu opartego na architekturze ARM (Advanced RISC Machine, alternatywie dla bardziej rozpowszechnionej w pecetach architektury x86), który mógłby rzucić wyzwanie AMD i Intelowi oraz Qualcommowi na opanowanych przez nich rynkach. Teoretycznie byłby to układ do komputerów stacjonarnych, podobny do procesorów Snapdragon (ARM) firmy Qualcomm. Podobnie jak seria Snapdragon X firmy Qualcomm, hipotetyczny nowy procesor NVIDIA napotkałby zapewne te same problemy, które w systemie Windows mają inne chipy ARM, dotyczące wydajności i emulacji oprogramowania (uruchamianie aplikacji x86 na ARM). ARM ma oczywiście mocne strony – to żywotność baterii i przystępna cena. Mówi się, że procesor NVIDIA oparty na ARM (2) może pojawić się w drugiej połowie 2025 roku, choć nie jest jasne, czy oznaczałoby to również premierę laptopów z tymi

układami. Warto dodać, że jest również sporo pogłosek o planach NVIDIA dotyczących własnego lub opracowanego we współpracy ze znaną marką, np. Lenovo, laptopa z procesorami typu Blackwell.

Na froncie półprzewodnikowym bez zmian – czyli... nieustanne zmiany

Chipy zasilające szkolenie i działanie sztucznej inteligencji zostały uznane za towar na tyle strategicznie wrażliwy, że rząd USA nałożył już kilka lat temu ostre sankcje i ograniczenia na eksport najbardziej zaawansowanego hardware'u do Chin. To właśnie było powodem stworzenia przez amerykańską firmę NVIDIA, wiodącego producenta układów półprzewodnikowych na rynku AI, zmodyfikowanych wersji swoich produktów, których specyfikacje pozwalają na eksport do Chin. Regulacje wprowadzone przez amerykańskie władze spowodowały, że NVIDIA nie mogła sprzedawać tam swoich najbardziej zaawansowanych chipów A100 i H100. Wersje oznaczone A800 i H800 mogą bez przeszkód wejść na chiński rynek. H800 jest wykorzystywany przez Alibabę Group, Baidu i Tencenta, chińskich potentatów rozwijających projekty oparte na uczeniu maszynowym. W marcu 2025 roku okazało się, że Chiny wprowadzają przepisy utrudniające lub wręcz uniemożliwiające owym „zubożonym” wersjom chipów NVIDIA sprzedaż na chińskim rynku. Rząd chiński chce w ten sposób promować własne, niezależne od amerykańskich, produkty półprzewodnikowe. Więc wojna ta toczy się niejako po obu stronach.

Kilka miesięcy wcześniej, w listopadzie 2024 Stany Zjednoczone poinstruowały Taiwan Semiconductor



2. Procesor ARM © Wikipedia



3. Wnętrze zakładu ASML produkującego maszyny do trawienia krzemu w Holandii © ASML

Manufacturing (TSMC), by wstrzymał wysyłkę niektórych zaawansowanych chipów do Chin. Dyrektywa ta obejmuje ograniczenia eksportu chipów o konstrukcji 7 nanometrów lub bardziej zaawansowanych, powszechnie stosowanych w sztucznej inteligencji. Posunięcie to wynikało z doniesień, że chipy TSMC zostały znalezione w procesorach AI Huawei, co sugerowało możliwe naruszenie amerykańskich przepisów o kontroli eksportu. Huawei znajduje się na amerykańskiej liście towarów objętych restrykcjami handlowymi, co wymaga od dostawców uzyskania licencji na handel towarami lub technologiami, w szczególności tymi wykorzystywanymi w sztucznej inteligencji.

W styczniu tego roku, jeszcze przed inauguracją Donalda Trumpa, Stany Zjednoczone ogłosiły nowe zasady kontroli eksportu w branży chipów sztucznej inteligencji i chociaż wprowadzono zwolnienia dla strategicznych sojuszników, wiele państw członkowskich UE (w tym Polska) zostało objętych ograniczeniami. Na mocy tych przepisów wybrane firmy z osiemnastki „kluczowych sojuszników i partnerów” spełniające „wysokie standardy bezpieczeństwa i zaufania” określono jako Universal Verified End Users (UVEU) i mogą one korzystać z pełnego zwolnienia ze środków ograniczających wysyłkę procesorów graficznych. Znacznie

dłuższa jest lista krajów, w tym kilka państw członkowskich UE, która do tej grupy nie należy. Firmy UVEU będą mogły zakupić setki tysięcy chipów, ale nadal będą ograniczone do „7 proc. ich globalnej mocy obliczeniowej AI” w krajach na całym świecie w sposób globalny i trwały. Kraje, których nie ma na liście, ale które nie są „krajem budzącym obawy”, mogą ubiegać się o pozwolenie na zakup mocy obliczeniowej do 320 tys. zaawansowanych procesorów graficznych w ciągu najbliższych dwóch lat, podczas gdy reszta może zakupić do 50 tys. procesorów graficznych na kraj. Jeśli chodzi o to, co definiuje kluczowego sojusznika, w komunikacie administracji Bidena stwierdzono, że są to jurysdykcje z solidnymi systemami ochrony technologii i ekosystemami technologicznymi, które są zgodne z bezpieczeństwem narodowym i interesami polityki zagranicznej USA. Lista „krajów budzących obawy” obejmuje około 20 krajów, w tym Chiny, Rosję, Iran, Koreę Północną i Wenezuelę. Na liście tej, niejako „najgorszej”, nie ma krajów UE. Jednak i tak wzbudziło to ogrom kontrowersji. Zauważano m.in. że choć łatwo jest zrozumieć, dlaczego Holandia (3) i Tajwan, w których są kluczowej dla procesu produkcji chipów zakłady, są zwolnione z ograniczeń, trudniej pojąć, dlaczego Belgia została uznana za kluczowego sojusznika, zaś Polska nie. Komisarze UE ds. Technologii

i Handlu, Henna Virkkunen i Maroš Šefčovič, opublikowali wspólne oświadczenie, w którym wyrazili zaniepokojenie wprowadzonymi przez USA środkami.

W kolejnym miesiącu w Singapurze oskarżono trzech mężczyzn o oszustwo w sprawie dotyczącej rzekomo nielegalnego reeksportu procesorów graficznych NVIDIA do chińskiej firmy DeepSeek (4), omijając amerykańskie ograniczenia handlowe. Policja i organy celne aresztowały dziewięć osób i zajęły dokumenty i zapisy elektroniczne. Kiedy Singapur nagle stał się drugim co do wielkości geograficznym źródłem przychodów NVIDIA w 2024 r., wielu podejrzewało, że stało się tak, ponieważ jej procesory graficzne były nielegalnie reeksportowane z Singapuru do Chin. Departament Handlu USA rozpoczął w tym samym czasie postępowanie wyjaśniające, czy DeepSeek nabył ograniczone amerykańskie procesory graficzne do trenowania swoich modeli sztucznej inteligencji.

Pod koniec stycznia wejście na rynek chińskiego modelu AI model DeepSeek wywołał masową wyprzedaż akcji spółek z amerykańskiego sektora sztucznej inteligencji. Inwestorzy obawiali się, że model ten może działać tak dobrze, jak najlepsi konkurenci, zużywając mniej energii i pieniędzy. Założenie było takie, że nie potrzebuje zakazanych w Chinach najbardziej zaawansowanych amerykańskich procesorów. Ceny akcji NVIDIA spadła o 17 proc. w ciągu jednej sesji, tracąc blisko 600 miliardów dolarów, co było największym w historii jednodniowym spadkiem dla amerykańskiej firmy na giełdzie. Model językowy R1 DeepSeeka uznano za dorównujący amerykańskim produktom OpenAI i innych firm. Podobno Chińczycy w jego szkoleniu korzystali jedynie z dziesiątek tysięcy „dozwolonych dla Chin” procesorów graficznych NVIDIA Hopper (H100, H20, H800). Pojawiły się jednak

spore wątpliwości co do tego, z jakiego sprzętu rzeczywiście Chińczycy korzystali. Wielu komentatorów uważało, że osiągnięcie takich rezultatów bez najbardziej zaawansowanych GPU NVIDIA byłoby niemożliwe. W powietrzu zawisło więc pytanie, czy DeepSeek przypadkiem nie był beneficjentem tego nielegalnego importu do Chin. Innym problemem było zerowanie chińskiego modelu na produktach OpenAI.

W wywiadzie z Jimem Cramerem z CNBC, szef NVIDIA Jensen Huang ocenił model DeepSeek jako „fantastyczny”, ponieważ jest to „pierwszy model rozumowania o otwartym kodzie źródłowym”. Wyjaśnił, że model ten rozkłada problemy krok po kroku, jest w stanie wymyślić różne odpowiedzi i może zweryfikować, czy jego odpowiedź jest poprawna. Wygłosił też taką dającą do myślenia w tym kontekście sentencję: „Ta rozumująca sztuczna inteligencja zużywa sto razy więcej mocy obliczeniowej niż nierozumująca sztuczna inteligencja”.

Nie ulega jednak wątpliwości, że w Chinach trwają intensywne prace nad uniezależnieniem się od półprzewodników amerykańskich, czy to tych najbardziej zaawansowanych, które być może są „podkradane”, czy tych „zubożonych” wersji, dozwolonych na rynek chiński. W lutym tego roku opublikowano wyniki badań, w których badacze chińscy, pod kierownictwem Nana Tongchao z uniwersytetu w Nankinie, korzystający z krajowych procesorów graficznych, mieli osiągnąć niemal dziesięciokrotny wzrost wydajności w porównaniu z potężnymi amerykańskimi superkomputerami, które opierają się na najnowocześniejszym sprzęcie NVIDIA. Naukowcy stwierdzili, że to innowacyjne techniki optymalizacji oprogramowania umożliwiły im zwiększenie wydajności komputerów zasilanych przez zaprojektowane w Chinach

4. Logo modelu DeepSeek na karcie GPU NVIDIA



procesory graficzne (GPU). Wyniki ich badań zostały opublikowane w czasopiśmie „Chinese Journal of Hydraulic Engineering”.

Jak informował Reuters, Chiny przygotowują się też do wydania przepisów mających na celu promowanie przyjęcia układów RISC-V typu open source w całym kraju, co znów miałoby być krokiem w kierunku zmniejszaniu zależności od zachodniej technologii półprzewodnikowej. Do największych chińskich komercyjnych dostawców RISC-V należą Alibaba (HK:9988) i startup Nuclei System Technology. Chipy RISC-V to procesory komputerowe typu open source, które oferują elastyczną i opłacalną alternatywę dla chipów o zastrzeżonej architekturze, choćby x86, wprowadzonych przez amerykańskich gigantów chipów Intel i Advanced Micro Devices (AMD) oraz ARM, opracowanych przez Arm Holdings. umożliwiając firmom projektowanie niestandardowego sprzętu do różnych zastosowań, od smartfonów po sztuczną inteligencję i superkomputery. Chipy te zyskały popularność wśród chińskich firm ze względu na niższe koszty i neutralność geopolityczną. Niektórzy amerykańscy ustawodawcy naciskają na wprowadzenie ograniczeń dla amerykańskich firm współpracujących nad RISC-V, powołując się na obawy dotyczące wykorzystania tej technologii przez Chiny.

Intel gra ostro, ale ma coraz mniej do stracenia

Obok frontu wojny pomiędzy mocarstwami toczy się zażarta walka pomiędzy największymi producentami półprzewodników i komponentów. Tu jednym z głównych wątków są ogromne problemy amerykańskiego Intela, któremu trudno jest sprostać szybko posuwającą się do przodu konkurencji. Według południowokoreańskiej gazety „Chosun”, węzeł Intel 18A, na który firma tak bardzo liczy w kontekście walki z TSMC, AMD i Samsungiem w dziedzinie miniaturyzacji, podobno utknął na poziomie 10 proc. wydajności, gęstość pamięci SRAM również jest w tyle za wchodzącą na scenę techniką 2 nm firmy TSMC. Te dziesięć procent wydajności oznacza, że większość wyprodukowanych chipów jest wadliwych. Wyzwania te mogą utrudnić wdrożenie węzła w procesorach nowej generacji, sztucznej inteligencji i niestandardowych układach Intela. Jest to poważny problem, zwłaszcza że Intel zagrał wabank i anulował już swój węzeł procesowy 20A (klasa 2 nm), przenosząc zasoby do węzła 18A (klasa jeszcze niższego rozmiaru bramki – 1,8 nm). Jeśli wskaźnik wydajności poniżej 10 proc. okaże się trwały, to może okazać się, że nie da się wdrożyć produkcji komercyjnej, przynajmniej

do czasu wprowadzenia znaczących ulepszeń. Intel ma jeszcze kilka miesięcy na dopracowanie procesu przed planowanym uruchomieniem produkcji w 2025 roku. Potencjalna korzyść jest znacząca, ponieważ 18A ma zasilać głośne produkty, takie jak chipy serwerowe Intel Clearwater Forest, procesory mobilne Panther Lake i niestandardowe chipy AI. Jeśli Intel będzie w stanie szybko poprawić wydajność 18A do przyzwoitego poziomu – powyżej 60 proc. w nadchodzących miesiącach, terminy mogą zostać dotrzymane a Intel jak za dawnych lat wyszedłby na czoło.

Jakimś pocieszeniem może być dla amerykańskiej zasłużonej firmy to, że inni też borykają się z trudnościami. Wyzwanie, jakim jest upakowanie tranzystorów w coraz gęstszych układach w tych najnowocześniejszych węzłach, jest ogromną przeszkodą inżynierską wpływającą na całą branżę półprzewodników. Wydajność zakładów produkcji Samsunga dla procesów poniżej 3 nm wynosi obecnie poniżej 50 proc., a wydajność technologii Gate-All-Around (GAA) wynosi podobno od 10 do 20 proc.

Kwestie wydajności nie są zresztą jedynym technicznym wyzwaniem, przed którym stoi Intel w przypadku 18A. TSMC zyskało podobno przewagę w innym krytycznym obszarze: gęstości pamięci SRAM. Według ISSCC 2025 Advance Program, węzeł N2 (klasa 2 nm) TSMC zmniejsza komórki bitowe pamięci SRAM o dużej gęstości do około 0,0175 μm^2 , osiągając gęstość 38 Mb/mm². Dla porównania, węzeł 18A Intela osiąga 0,021 μm^2 i 31,8 Mb/mm², co jest bliższe węzłom N3E i N5 poprzedniej generacji TSMC. Ponieważ projekty chipów wymagają więcej pamięci SRAM, zwiększenie gęstości tych małych komórek pamięci ma kluczowe znaczenie dla utrzymania kompaktowych, wydajnych projektów. Tranzystory GAA (gate-all-around), kontrolując kanał ze wszystkich stron, pozwalają na ściślejsze skalowanie w porównaniu do tradycyjnych tranzystorów finFET. Ta ścisła kontrola zmniejsza wyciek przy małych wymiarach, umożliwiając pamięć SRAM o większej gęstości. Zarówno Intel, jak i TSMC używają GAA do zmniejszania swoich komórek bitowych SRAM, ale TSMC udało się upakować je gęściej w swoim węźle N2.

Z drugiej strony nie wszystkie wiadomości dla Intela są złe. Aurora oparta na jego procesorach (5) była rok temu najszybszym superkomputerem AI na świecie, choć ogólnie nadal zajmuje drugie miejsce za Frontier napędzanym przez AMD. Aurora, stworzona dzięki współpracy firm Intel, HPE i Narodowego Laboratorium Argonne Departamentu Energii Stanów Zjednoczonych, przełamała barierę eksaskalową, stając się drugą maszyną po Frontier, która weszła na ten poziom. Ma



5. Superkomputer Aurora

21 248 procesorów Intel Xeon CPU Max i 63 744 akceleratorów Intel Data Center GPU Max. W testach osiągnęła ona wydajność 1,012 eksaflopsa przy 9234 aktywnych węzłach z łącznej liczby 10624. System Frontier zasilany przez AMD jest obecnie oceniany na 1,206 eksaflopsów wydajności, przy 87 procentach wykorzystania sprzętu.

Pisaliśmy o planach NVIDIA w odniesieniu do architektury ARM. Tymczasem brytyjska firma Arm, która ma ją w nazwie, zamierza rzucić wyzwanie dominacji x86 dzięki szybszym procesorom i układom GPU z ulepszoną sztuczną inteligencją. Architektura RISC brytyjskiej firmy od dawna chwalona jest za podejście skoncentrowane na wydajności, zwłaszcza w porównaniu z układami x86 od AMD i Intela. Gdy niektóre produkty Arm wyróżniają się pod względem IPC, ich niższe częstotliwości taktowania w porównaniu do rywali hamowały ogólną wydajność. Kolejnym ważnym priorytetem jest optymalizacja projektów CPU i GPU Arm pod kątem akceleracji AI. Działalność Arm w zakresie GPU jest rozwijana z naciskiem na AI/ML. Akcelеровane przez sztuczną inteligencję udoskonalenia CPU i GPU mogą zostać zintegrowane z platformą obliczeniową Arm CSS for Client nowej generacji, która ma zadebiutować w 2025 roku.

NVIDIA musi uważać na tyły

W marcu 2025 r. Bolt Graphics, startup z siedzibą w Sunnyvale w Kalifornii, zaprezentował nową konstrukcję GPU, która, jak twierdzi, może wyprzedzić konkurentów pod względem renderowania, wysokowydajnych obliczeń (HPC) i obciążeń w grach. Procesor graficzny Zeus ma być „o rzędy wielkości” szybszy niż istniejące na rynku rozwiązania. Dzięki możliwości rozbudowy pamięci karty Zeus mogą pomieścić do 384 GB na karcie PCIe lub do 2,25 TB na kartę w serwerze 2U. Zeus jest również pierwszym procesorem

graficznym z szybkimi interfejsami 400 GbE i 800 GbE, umożliwiającymi procesorom graficznym łączenie się na masową skalę. Bolt twierdzi, że użytkownicy Zeusa mogą spodziewać się nawet dziesięciokrotnego wzrostu wydajności w zadaniach graficznych i nawet trzystukrotnego w symulacjach fal elektromagnetycznych. Dla wyjaśnienia – te ostatnie są wykorzystywane do projektowania produktów takich jak skanery tomografii komputerowej i rentgenowskie, czujniki radarowe, materiały stealth i soczewki optyczne. Oczekuje się, że zestawy deweloperskie pojawią się jeszcze w tym roku przed masową produkcją zaplanowaną na koniec 2026 roku.

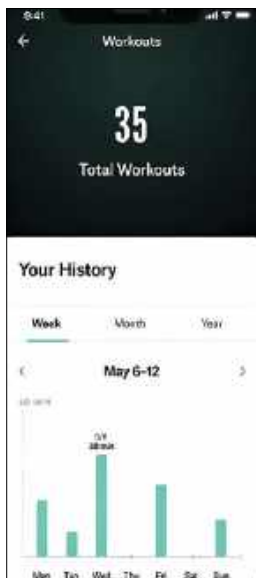
Wygląda więc na to, że NVIDIA, która zapędziła się daleko w świat AI, może już wkrótce zostać mocno zaatakowana na swoim pierwotnym rynku kart graficznych, głównie do gier, ale nie tylko. A konkurować chce nie tylko startup z Kalifornii, ale też dobrze znana i duża konkurencja, czyli firma AMD. Według jej materiałów marketingowych, zapowiadany układ Ryzen AI Max+ 395 ze zintegrowaną grafiką Radeon 8060S jest o 68,1 proc. szybszy niż RTX 4070 firmy NVIDIA w laptopie. AMD przetestowało Radeona 8060S w różnych grach w rozdzielczości 1080p z ustawieniem grafiki High. W większości gier marginesy są niewielkie. Według analiz Notebookcheck, można jednak zauważyć znaczną przewagę AMD w grach takich jak Cyberpunk 2077, Boulder's Gate 3, Hitman 3, a zwłaszcza Borderlands 3, gdzie wzrost jakości wyniósł aż 68,1 proc.

W tej wojnie nikt nie śpi. Jedni ścigają. Inni nigdy nie mogą być pewni, czy ktoś ich znienacka nie prześcignie. Jeśli myślimy, że np. NVIDIA rzuciła wszystkich na kolana a Intel definitywnie upadł, to szybko może się okazać, że los się odmienił, i wygrywa już ktoś inny. ■

Mirosław Usidus



Fitness i zdrowie



Tonal

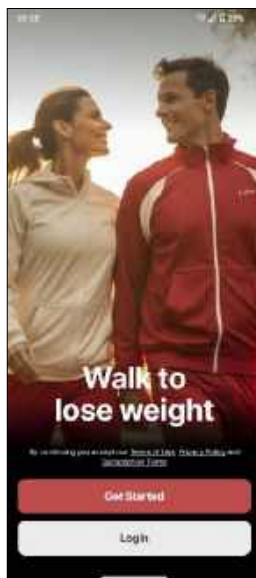
Aplikacja przedstawia się jako „inteligentny trener osobisty”, który zrozumie cele treningowe użytkownika i pomoże mu je zrealizować. Tonal oferuje setki treningów do wypróbowania, od treningów o wysokiej intensywności po regenerujące przepływy jogi. Oferuje także możliwość układania własnych programów treningowych z opcjami śledzenia, obserwacji i mierzenia postępów

Program analizuje m.in. wybrane partie mięśni i typy treningów. Można wybierać swoje ulubione typy ćwiczeń i aktywności, powtórzenia, serie i zaawansowane tryby obciążenia, np. „Burnout” i „Eccentric”. Tonal pozwala również na nawiązywanie kontaktu ze znajomymi, aby w społecznościowym gronie wzajemnie się dopingować do aktywności fizycznej.

Z aplikacji można korzystać bez specjalnego domowego zestawu sprzętowego do ćwiczeń. Program przewiduje m.in. programy treningowe na świeżym powietrzu, niezwiązane z kompletem domowej siłowni Tonal. Tę uniwersalność należy uznać za zaletę aplikacji, choć z pewnością program najlepiej sprawdza się przy użyciu sprzętu Tonal.

Tonal	
Producent	Tonal
Platforma	Android, iOS
Oceny	Możliwości
	Łatwość obsługi
	Ocena ogólna

7,5/10
8,5/10
8/10



WalkFit

Apka skupia się, co wskazuje sama nazwa, na chodzeniu. Oferuje pół tysiąca treningów opartych na chodzie. To m.in. treningi polegające na chodzeniu w pomieszczeniach, spacerach cardio o wysokiej intensywności, jednominutowe spacerki, a nawet dłuższe wędrówki po terenie. Program podsuwa wyzwania, których podjęcie może dać użytkownikowi specyficzne „nagrody”.

Główne funkcje WalkFit to liczenie kroków, monitorowanie snu, pomiar pulsu serca, powiadomienia o połączeniach, powiadomienia o wiadomościach oraz powiadomienia o wiadomościach z mediów społecznościowych. Aplikacja umożliwia również ustawianie przypomnień o czasie siedzenia w miejscu, powiadamiając o tym, że czas się ruszyć

WalkFit jest kompatybilny z urządzeniami noszonymi, zegarkami lub bransoletkami, z którymi łączy się za pomocą Bluetooth. Synchronizując urządzenia ze smartfonem, można śledzić postępy w realizacji celów treningowych i monitorować kluczowe wskaźniki, takie jak liczba kroków, spalone kalorie i pokonany dystans.

Walking Weight Loss: WalkFit	
Producent	Welltech Apps Limited
Platforma	Android, iOS
Oceny	Możliwości
	Łatwość obsługi
	Ocena ogólna

6/10
8/10
7/10

Smartfony i ich systemy operacyjne, czyli słówko o platformach

Podobnie jak komputer, tak i smartfon, choćby nie wiadomo jak wspaniały, to tylko kupka elektronicznego złomu, jeśli brak w nim oprogramowania. Podstawowym oprogramowaniem każdego urządzenia z procesorem, pamięcią i wyświetlaczem jest system operacyjny. To dopiero on decyduje, jakie możliwości ma dane urządzenie i jednocześnie wyznacza jego popularność, mierzoną liczbą dostępnych aplikacji – jako że aplikacje pisane są na określony system operacyjny, a nie „na sprzęt”. Przykładowo, dwa identyczne telefony tej samej firmy mogą być zupełnie różnymi funkcjonalnie urządzeniami, jeśli na jednym producent zainstaluje system Android, a na drugim system Symbian. Aplikacje na Androida nie będą działać na Symbianie i odwrotnie. Najpopularniejsze smartfonowe systemy operacyjne to:

- **iOS** – system firmy Apple (tej od komputerów Macintosh), instalowany w urządzeniach iPhone, iPod Touch, iPad;
- **Android** – system firmy Google, niektórzy twierdzą, że wkrótce podbije cały świat. Rzeczywiście, Android jest coraz częściej instalowany w smartfonach m.in. takich firm, jak Huawei, HTC, LG, Motorola, Samsung, Sony Ericsson, ZTE (a także, co oczywiście, w smartfonach firmy Google);
- **Symbian** – system operacyjny open source (czyli bezpłatny i z tzw. otwartym kodem), obecnie najczęściej spotykany w telefonach firmy Nokia.

Inne, mniej popularne systemy operacyjne dla telefonów komórkowych, to:

- **Bada** – system rozwijany przez firmę Samsung;
- **Windows Phone** – system firmy Microsoft, następca Windows Mobile, czyli po prostu Windows do urządzeń przenośnych;
- **BlackBerry** – system kanadyjskiej firmy Research In Motion, przeznaczony przede wszystkim do zastosowań biznesowych, instalowany w produktach oferowanych przez firmę Nokia z charakterystyczną, pełną klawiaturą QWERTY. Także w niektórych telefonach innych firm (HTC, Motorola, Nokia, Samsung, Sony Ericsson).



Krokomierz

Wbrew nazwie aplikacja służy nie tylko do liczenia kroków. Śledzi także codzienną aktywność, przebyte dystanse, spalone kalorie, czas trwania i prędkość, oraz dokonuje wiele innych pomiarów, analizując wyniki aktywności fizycznej użytkownika. Użytkownicy chwalą program za prostotę działania.

Korzystając z Krokomierza otrzymujemy m.in. raporty z wykresami umożliwiającymi śledzenie danych dotyczących treningu chodowego. Można sprawdzić statystyki ostatnich dwudziestu czterech godzin, tygodnia i miesiąca. Jest opcja ustawiania czułości, ale warto pamiętać, że dokładność pomiarów zależy od dokładności wprowadzonych przez użytkownika danych.

Co ciekawe, nie musimy włączać lokalizacji w telefonie, by móc korzystać z Krokomierza. Program ma bowiem własny system odmierzający przemieszczanie się. Ponadto funkcja oszczędzania energii sprawia, że krokomierz działa na bardzo niskim poziomie energii w tle (możemy więc umieścić telefon w plecaku, torebce, kieszeni lub w jakimkolwiek innym miejscu).

Krokomierz – Licznik kroków		
Producent	ITO Technologies, Inc.	
Platforma	Android, iOS	
Oceny	Możliwości	7/10
	Łatwość obsługi	9/10
	Ocena ogólna	8/10



Caynax

Aplikacja oferuje wybór planów treningowych dostosowanych do różnych poziomów kondycyjnych. Użytkownicy mogą korzystać z funkcji monitorowania aktywności oraz ustalania celów długoterminowych. Caynax ma wbudowane funkcje personalizacji, umożliwiając użytkownikom ustawienie przypomnień o treningach oraz dostosowanie planów do indywidualnych potrzeb.

Caynax rejestruje trasę na mapie i śledzi czas, prędkość, dystans, kroki (jak Krokomierz), spalone kalorie, tempo, wysokość (wysokościomierz) i inne parametry. Dzięki aplikacji można monitorować swoje postępy, ustawiać cele treningowe i dostosowywać swoje ćwiczenia do własnych potrzeb. Nie wymaga rejestracji. Można używać programu bez zakładania konta.

Apka wykorzystuje odbiornik GPS w urządzeniu do precyzyjnego śledzenia aktywności. Oferuje wgląd w mapę w celu śledzenia trasy i informacje głosowe. Asystuje w wielu rodzajach treningu – w bieganiu, chodzeniu, jeździe na rowerze, nordic walkingu, pływaniu, narciarstwie biegowym, kajakarstwie, wioślarstwie, tyżwiarstwie, narciarstwie zjazdowym, a nawet jeździe konnej.

Caynax – Bieganie i rower GPS		
Producent	Caynax	
Platforma	Android	
Oceny	Możliwości	9/10
	Łatwość obsługi	8/10
	Ocena ogólna	8,5/10



MyFitnessPal

Ten program skupia się nie na treningu, lecz na diecie, która może oczywiście także towarzyszyć treningowi. Pozwala na monitoring trybu odżywiania się, diety, w połączeniu z inną aktywnością i dążeniem do odchudzenia. To rodzaj osobistego doradcy i planera posiłków wraz z dziennikiem żywieniowym.

W MyFitnessPal można wprowadzać oczywiście także plany treningowe jako część strategii poprawienia wyglądu i kondycji fizycznej. Cele, które są wprowadzane w aplikacji, są szerokie – od zrzucania wagi, przez zwiększenie masy mięśni, utrzymanie wagi, poprawienie trybu odżywiania i zwiększenie aktywności fizycznej. Wszystko to, zarówno sfera odżywiania jak też treningi, jest monitorowane dzięki funkcjom aplikacji.

Można też dzięki niej sprawdzać kaloryczność ponad czterestu milionów produktów (w tym typowych dań restauracyjnych). Aplikacja pozwala na wyszukiwanie informacji o produktach w bazie, jaką ma, a nawet skanowanie ich za pomocą kodów kreskowych. MyFitnessPal może posłużyć nie tylko do redukcji wagi, ale również uzupełniania braków mikroskładników i witamin. Licznik kalorii automatycznie liczy spożywane kalorie i monitoruje w kontekście całonocnej strategii odchudzania i/lub poprawiania kondycji fizycznej.

MyFitnessPal		
Producent	MyFitnessPal, Inc.	
Platforma	Android, iOS, Mac, Windows, Linux	
Oceny	Możliwości	8/10
	Łatwość obsługi	9/10
	Ocena ogólna	8,5/10

Raport z obserwacji DM-S-3

Przygotował zespół badaczy rozpoznania wczesnego w składzie: Józ, Mar, Wik i Ani

Okres obserwacji

Nasza czwórka badaczy, na zlenienie miejscowej komórki Instytutu Kontaktu, badała obiekt przez faktyczne 1 100 001 101 010 000 obiegów planety wokół gwiazdy macierzystej.

Dodatkowy okres ekstrapolowany z dawnego światła i badań miejscowych wyniósł ok. 1 111 010 000 100 100 000 obiegów planety wokół gwiazdy macierzystej.

Etiologia gatunku, dotychczasowa historia i obserwacje

Gatunek wyodrębnił się jako istoty wszystkożerne, w niewyraźnie wykształconej strukturze rodzinno-plemiennej. Ewoluuował w wyjątkowo trudnym środowisku, narażony na większą niż na innych tego typu planetach ilość niebezpieczeństw. Biorąc pod uwagę statystyki obserwacji w tym rejonie kosmosu, możemy stwierdzić, że na gatunek zamieszkujący DM-S-3 przez cały okres istnienia czyhało nadzwyczaj wiele zabójczych czynników.

Od samego początku wyodrębnienia gatunku poddawany był on presji drapieżniczej ze strony innych istot zamieszkujących planetę. Skalisty glob znajdujący się w centralnej części ekosfery gwiazdy stał się domem dla wyjątkowo różnorodnej ekosfery, a co za tym idzie doszło do rozwoju fauny mięsożernej na niespotykaną skalę. Mimo że pierwotne egzemplarze gatunku były słabe, niewyposażone ani w silne uzębienie, ani w szpony, dodatkowo pozbawione znacznej szybkości, udało im się przetrwać w skrajnie niesprzyjającej niszy ekologicznej.

Jako słabe i delikatne, istoty omawianego gatunku zaczęły wykazywać się szybkim rozwojem inteligencji ukierunkowanej na przeciwstawienie się agresji i przemocy, co skutkowało niespotykanym na innych planetach własnym rozwojem ww. cech. Istoty, pozbawione szybkości i skutecznego wyposażenia ciała w postaci zdolnych zadawać rany kończyn czy uzębienia, wypracowały nietypowe drapieżnictwo pościgowe. Zamiast polować na inne zwierzęta w sposób szybki i nienastręczający cierpienia, gatunek dzięki współpracy międzyosobniczej skutecznie ścigał swoje ofiary tak długo, aż umierały z wycieńczenia.

W przypadku innych ofiar oraz organizmów napastniczych, przedstawiciele mieszkańców planety stosowali przemyślnie pułapki i zasadzki, mające na celu uśmiercenie lub okaleczenie innych organizmów. Istoty, wykształciwszy zdolność abstrakcyjnego myślenia, ukierunkowały ją niewzłocznie na sztukę zamiany dowolnych przedmiotów na narzędzia przemocy i cierpienia. Z miejscowej flory i zasobów geologicznych planety wytworzyły broń kłujące, sieczne, tnące. Zabite istoty następnie zjadały, a nienadające się do spożycia części ich ciał przetwarzały z kolei w inne narzędzia, nierzadko także wytworzone w celu zabijania i zadawania cierpienia.

Mieszkańcy planety przywdziewali także skóry i fragmenty ciał zabitych istot, aby skuteczniej prowadzić agresję opartą na oszustwie i pozorach.

Przypis badaczki Wik: przywdziewanie fragmentów ciał zabitych istot mogło w rzeczywistości mieć szersze zastosowanie, takie jak ochrona przez drastycznymi warunkami bytowania na planecie. Warunki te w zależności od położenia na powierzchni planety zmieniają się gwałtownie. Okolice punktów obrotu globu wykazują się temperaturami, w których organizmy istot ulegają zamarznięciu. Mimo to gatunek przetrwał całe okresy, kiedy zestalony monotlenek diwodoru w formie skalistej pokrywał znaczne połacie powierzchni globu. Istnieje hipoteza, że przyczyniło się do tego właśnie używanie ciał innych, zabitych przez gatunek istot, jako formy ochrony termicznej.

W innych obszarach planety panujące warunki powodują obumieranie termiczne białek, na których oparte są wszystkie tamtejsze organizmy. Mimo to, dominujący gatunek inteligentny opracował sposoby funkcjonowania w tych skrajnie nieprzyjaznych warunkach.

Atmosfera planety składa się w większości z azotu w formie gazowej, który jednak dla zamieszkujących ją w większości istot, nie tylko omawianego gatunku, wydaje się gazem neutralnym. Poza szczególnymi

wyjątkami, które zostały opisane w dodatku 11, nie wyrządza on szkody gatunkowi. Same istoty jednak, w wyniku spalania komórkowego, wytwarzają zabójczy dwutlenek węgla, który następnie bez żadnej kontroli uwalniają ze swoich organizmów wprost do atmosfery. W procesie spalania komórkowego istoty przetwarzają tlen, który najwyraźniej nie jest dla nich zabójczy. Dłuższe obserwacje sugerują jednak, że powoduje on stopniowe uszkodzenie struktur komórkowych, a przez ich utlenianie przyczynia się do śmierci osobniczej przedstawicieli gatunku w okresie przeciętnie od 1 001 011 do 1 011 010 obiegów planety wokół gwiazdy macierzystej.

Warunki środowiskowe wytworzyły w istotach szereg niespotykanych przystosowań, niewidzianych nigdzie indziej w kosmosie.

Istoty, do trawienia ciał innych istot używają wielu kwasów, wytwarzanych przez własne narządy. W układach trawiennych mieszkańców planety białka są z pomocą tych kwasów rozbijane i trawione. Istoty są w stanie wystrzeliwać żrącą wydzielinę otworów gębowych na znaczne odległości.

Otworki gębowe wyposażone są w ponad 11 110 kostnych zębów. Mimo że, jak wspomniano wcześniej, nie są one przystosowane do zadawania ran śmiertelnych – z wyjątkiem precyzyjnie wykonanych ugryzień w okolicach arterii krwionośnych – mogą one zadawać jęczące się, zakażone rany, zaś drobnoustroje znajdujące się w kwasie gębowym mogą przenosić poważne choroby.



Wszelkie w ogóle wydzieliny mieszkańców planety mogą być i zwykle są chorobotwórcze. Prócz kwasu gębowego ich ośrodki węchu wytwarzają w stanie chorobowym ochronny śluz. Kontakt z nim oznacza praktycznie natychmiastowe zakażenie. Odchody istot zawierają niespotykane ilości chorobotwórczych i szkodliwych mikroorganizmów. Kontakt z krwią może spowodować kolejne zabójcze choroby.

Mimo pozornej słabości i braku zewnętrznych mechanizmów obronnych, istoty są niezwykle wytrzymałe i prawie niemożliwe do zabicia. Samice gatunku zdolne są do utraty ogromnych ilości krwi, a to za sprawą mechanizmu rozrodczego, który się wiąże z tym zjawiskiem. Wraz z rozwojem intelektualnym mózgowia istot powiększyły znacznie swoją objętość i obecnie standardową praktyką przy przyjściu na świat nowego osobnika jest rozcięcie powłok ciała samicy i wydobycie go przez asystujące narodzinom inne istoty. Mimo tak okrutnych i poważnych obrażeń matka nie umiera, lecz w ciągu krótkiego czasu dochodzi do siebie i zajmuje się potomstwem.

Istoty są zdolne do urodzenia jednego potomka co rok. Sporadycznie bywa, że rodzi się dwoje młodych, rzadziej – większa liczba.

Przypis badaczki Ani: Należy nadmienić, że na szczęście dla wszystkich ras rozumnych w kosmosie, istoty nie wykorzystują tej zdolności. W zależności od regionu planety zwykle jedna samica wydaje na świat w ciągu życia od jednego do 1010 młodych.

Niezwykła odporność na urazy mechaniczne nie jest wyłączną domeną samic. Również samce gatunku są niezwykle wytrzymałe. Istoty potrafią dzięki ewolucyjnej stabilności homeostazy oraz rozwiniętej inżynierii biologicznej przeżyć szereg urazów zabójczych dla innych gatunków, takich jak: otwarcie powłok ciała, utrata szczelności układu oddechowego; utrata krwi stosunku 1:101 całkowitej objętości nie wywołuje praktycznie żadnych efektów z wyjątkiem przejściowego zmęczenia.

Utrata kończyn powoduje u istot zamieszkujących DM-S-3 jedynie dyskomfort psychiczny. Istoty pozbawione 1, 10, 11 lub wszystkich 100 kończyn są w stanie funkcjonować, często całkowicie samodzielnie. Są w stanie wytwarzać zastępcze kończyny z ogólnodostępnych materiałów mineralnych i chemicznych. Potrafią dokonywać przeniesienia organów jednego osobnika do ciała innego, przedłużając w ten sposób życie tamtego. Dotyczy to także organów osobników martwych.

Istoty opanowały inżynierię medyczną pozwalającą na niwelowanie skutku śmiertelnych urazów stosunkowo wcześniej. Wybrane osobniki są szkolone w tych technikach, a niektóre potrafią przeprowadzać takie zabiegi same na sobie. Mimo że znają substancje obniżające odczuwanie bólu, opanowały je stosunkowo późno w swojej ewolucji i potrafią się obywać bez nich.

Poszczególne jednostki potrafią wykształcić genetyczną odporność na epidemie wyniszczające całe populacje i przekazać ją potomstwu; po upływie pokoleń mutacje te będą chronić je przed innymi czynnikami chorobowymi. Niektóre osobniki są w stanie przetrwać uszkodzenia centralnego ośrodka nerwowego ciałami stałymi, układu wymiany gazowej czy układu krwionośnego, a także termiczne uszkodzenia powłok ciała. Same ciała istot są w stanie przeżyć obumarcie mózgu, mimo, że tracą wtedy zdolność do działania świadomie i pozostają na niskim poziomie wegetacji.

Przypis badacza Józ: Jedyńm w pełni skutecznym sposobem powstrzymania osobnika gatunku inteligentnego z DM-S-3 jest usunięcie mózgu wraz z otaczającą kostną puszką i fragmentem ciała, wyposażonym w większość receptorów zmysłowych.

Wraz z rozwojem społecznym istoty opracowały przemysł zabijania, mający na celu wytwarzanie pożywienia pochodzenia zwierzęcego. W tym celu dokonują masowego rozplodu innych gatunków, które są następnie uśmiercane w wymyślny sposób. Spożywanie martwych ciał jest podstawowym rytuałem społecznym mieszkańców DM-S-3.

Przypis badacza Mar: Rodzaj, pochodzenie i sposób obróbki mogą być wyznacznikiem statusu społecznego. Martwe istoty są obrabiane termicznie, zdarza się jednak, że poddawane są procesom gnilnym, fermentacyjnym i innym, wytwarzającym toksyny. Tak przygotowane pożywienie uchodzi za wyjątkowo pożądane. Korzysta się także z toksyn obecnych w środowisku naturalnym. Wydaje się zasadnym stwierdzenie, że cierpienie wywołane takim spożyciem postrzegane jest jako atrakcyjne.

Istoty używają szeregu toksyn różnego pochodzenia w celach rekreacyjnych. Samoistnie, dobrowolnie doprowadzają do zatrucia organizmu a szkodliwe skutki uboczne uważane



są za pożądane. Takie rytuały jak spożycie sfermentowanego aż do pojawienia się alkoholu soku roślinnego, wdychanie toksycznych oparów płonących szczątków roślinnych lub zaburzenia świadomości wywołane czynnikami aktywnymi zawartymi w związkach naturalnego i sztucznego pochodzenia są uważane za cenne. Dystrybucja typów pożądanej intoksykacji w społeczności różni się w zależności od wieku osobników i położenia geograficznego danej grupy.

Gwiazda macierzysta planety wytwarza szereg zabójczych dla życia długości fal promieniowania. Istoty dobrowolnie wystawiają się na działanie tego promieniowania, co powoduje zmianę barwy ich powłok cielesnych, uważaną za atrakcyjną wśród populacji o dominującej jasnej barwie okryć cielesnych. Jednocześnie istoty o różnych odcieniach powłok potrafią utrzymywać antagonizmy oparte jedynie na wrażeniach estetycznych z tym związanych.

Istoty wydają się ogólnie, prócz intoksykacji i wystawienie na zabójcze promieniowanie, podejmować szereg działań będących bezpośrednim, gwałtownym zagrożeniem dla ich życia i zdrowia. Podejmują niczym nie umotywowane wyprawy w trudno dostępne tereny o ekstremalnych warunkach, takich jak niskie lub wysokie temperatury, miejsca o ograniczonej przydatności atmosferycznej (wysokie obszary o zmniejszonym ciśnieniu), dokonują także podróży w głąb pokrywających znaczną część planety gigantycznych mas monotonnego diwodoru, ryzykując zmiążdżeniem pod ogromnym ciśnieniem. Chociaż nigdy nie opanowały aktywnej zdolności lotu, istoty wytworzyły inżynierskie zdolności budowy urządzeń umożliwiających poruszanie się w atmosferze i ponad atmosferą. Około 111 100 obrotów planety wokół gwiazdy macierzystej temu istoty opanowały inżynierię pozwalającą im na opuszczenie biosfery globu i odwiedziny na naturalnym satelicie.

Przypis badaczki Ani: należy odnotować, że być może mieszkańcy planety zaspokoili tymczasowo swoje dążenia, ponieważ nie podejmowali dalszych prób fizycznych wypraw poza najbliższe otoczenie planety. A również te są ograniczone zwykle do pojedynczych osobników.

Zagrożenia

Prawdopodobnie ze względu na warunki i presję ewolucyjną, jakiej poddane zostały istoty, wykształciły się one jako gatunek wybitnie agresywny. Macierzysta planeta dla każdego potencjalnego gatunku rozumnego musiałaby się okazać śmiercioświatem, uniemożliwiającym rozwój życia rozumnego. Gatunek, który przetrwał w tak niezwykłym i okrutnym środowisku, sam przejął jego cechy. Głównymi wyznacznikami i popędami działania istot są mord, głód krwi, napawanie się cierpieniem. Istoty celebrować i gloryfikować rywalizację, do tego stopnia, że jedną z ukochanych rozrywek na przestrzeni ich dziejów były brutalne walki między pojedynczymi osobnikami, nierzadko zakończone śmiercią jednego z biorących udział.

Istoty toczą także trwające często wiele obrotów planety rozgrywki, w których duże grupy skupione wokół jednego przywódcy lub idei dokonują krwawych napaści na inne grupy, dopuszczając się eksterminacji i zajmując ich przestrzeń życiową. Zaledwie w okresie 1 100 100 planety wokół gwiazdy macierzystej poprzedzających niniejszy raport dwukrotnie zdarzyło się, że te krwawe igrzyska rozgrywały się na skalę planetarną, powodując straty osobnicze rzędu 10 111 110 101 111 000 010 000 000 jednostek w krótkim czasie.

Wnioski

Biorąc pod uwagę warunki, w jakich zrodziła się rozumna rasa z DM-S-3, nie jest właściwe obarczenie tych istot świadomą winą za ich postępowanie. Każdy gatunek, którego udziałem stałaby się ewolucja na planecie równie zabójczej, musiałby się wykołować lub przestać istnieć.

Jednakże, biorąc pod uwagę brutalną, okrutną i ekspansyjną naturę istot, w powiązaniu z faktem, że najwyraźniej opanowały sztukę fizycznych podróży poza swoją planetę, jedyny wniosek, jaki nasza czwórka badaczy jest w stanie przedstawić szanownej Połączonej Komisji 1100 Galaktyk ds Istot Rozumnych i Potencjalnego Kontaktu, jest prosty.

Chociaż na poziomie intelektualnym skłaniamy się ku rozwiązaniu bardziej ostatecznemu, nasza moralność nie pozwala nam zalecić zagłady planety; nie jest właściwe karać cały gatunek za to, że jest produktem swego środowiska naturalnego.

Jedynym wyjściem jest więc bezwzględna, absolutna, nieograniczona czasowo izolacja i całkowity brak czynnego kontaktu ze strony innych gatunków rozumnych, połączone z nieustającą obserwacją. ■

Bartek Biedrzycki





Materiał od zawsze był fundamentem kreatywności i innowacyjności. Może inspirować twórców sztuki, ale i stanowić kluczowy budulec nowoczesnych technologii. W rękach inżynierów przekształca się w rozwiązania, które zmieniają przemysł i codzienne życie. Eksperymenty z nowymi materiałami, od biodegradowalnych kompozytów po inteligentne struktury, wyznaczają kierunek rozwoju współczesnej inżynierii. To właśnie inżynieria materiałowa łączy kreatywność z nauką, dostarczając innowacyjnych i funkcjonalnych rozwiązań dla przyszłości. Jeśli masz w sobie kreatywność i zdolności techniczne, koniecznie zainteresuj się inżynierią materiałową. Zapraszamy.

Inżynieria materiałowa

Inżynieria materiałowa to interdyscyplinarny kierunek studiów, który łączy wiedzę z zakresu fizyki, chemii, mechaniki oraz technologii wytwarzania, by kształtować przyszłych specjalistów w dziedzinie projektowania, analizy i udoskonalania materiałów używanych w różnych gałęziach przemysłu. Kierunek ten oferowany jest przez wiele polskich uczelni technicznych, a także przez niektóre uniwersytety oraz uczelnie prywatne. Program studiów obejmuje zagadnienia związane z metalurgią, ceramiką, kompozytami, nanotechnologiami oraz nowoczesnymi polimerami, co czyni go atrakcyjnym dla osób zainteresowanych innowacyjnymi technologiami i rozwojem nowoczesnych materiałów. Studia na kierunku inżynieria materiałowa dostępne są w trybie stacjonarnym oraz niestacjonarnym, co umożliwia łączenie nauki z pracą zawodową. Na największych uczelniach limit miejsc na „inżynierce” wynosi zazwyczaj od 60 do 100 studentów rocznie. Konkurencja o indeks jest umiarkowana. Przykładowo, w rekrutacji na rok akademicki 2024/25 na Politechnice Krakowskiej przypadało średnio 2,8 kandydata na jedno miejsce, co klasyfikuje IM jako jeden z mniej obleganych kierunków. Niemniej, uczelnie zauważają wzrost zainteresowania i aktywnie promują go, podkreślając praktyczne znaczenie w przemyśle. Rekrutacja na studia opiera się głównie na wynikach z przedmiotów ścisłych. Każda uczelnia ustala własne zasady, jednak zazwyczaj analizowane są wyniki z matematyki (poziom rozszerzony) oraz z jednego przedmiotu do wyboru spośród: fizyki, chemii, informatyki lub biologii. Przykładowo Politechnika Warszawska kwalifikuje kandydatów na podstawie wyników z trzech przedmiotów: matematyki, nowożytnego języka obcego oraz jednego z wybranych (fizyka, chemia, informatyka lub biologia). AGH w Krakowie

bierze pod uwagę wynik z języka obcego oraz jednego z głównych przedmiotów ścisłych, takich jak matematyka, fizyka czy chemia. Progi punktowe nie są ekstremalnie wysokie, co sprawia, że kierunek ten uchodzi za względnie dostępny dla absolwentów szkół średnich o profilach matematyczno-przyrodniczych. Czasem wybierają go osoby, które nie dostały się na bardziej konkurencyjne kierunki techniczne, jednak coraz więcej kandydatów świadomie podejmuje tę decyzję, dostrzegając rosnące zapotrzebowanie na specjalistów w dziedzinie nowoczesnych materiałów.

Po uzyskaniu indeksu warto skupić się na edukacji. Studia na IM mają charakter interdyscyplinarny i łączą solidne podstawy teoretyczne z licznymi zajęciami praktycznymi. Program obejmuje zarówno wykłady i ćwiczenia rachunkowe, jak i laboratoria oraz projekty, co wynika ze specyfiki tej dziedziny. Podczas studiów I stopnia studenci zdobywają wszechstronną wiedzę z zakresu nauk ścisłych oraz inżynierskich. Kluczowe przedmioty obejmują oczywiście „królową nauk”, fizykę, chemię, informatykę (w tym komputerowe wspomaganie prac inżynierskich), naukę o materiałach, inżynierię materiałową, mechanikę, termodynamikę techniczną, elektrotechnikę oraz grafikę inżynierską. Jednym z najważniejszych elementów programu jest duża liczba godzin laboratoryjnych. Wiele przedmiotów realizowanych jest w formie praktycznych ćwiczeń, w trakcie których studenci przeprowadzają badania właściwości materiałów oraz eksperymenty technologiczne. Dzięki temu zdobywają doświadczenie w obsłudze aparatury badawczej i analizie wyników eksperymentów. W trakcie studiów istnieje możliwość wyboru specjalizacji w wybranym obszarze inżynierii materiałowej. Uczelnie oferują różne specjalności, takie jak:

technologie wytwarzania materiałów inżynierskich, inżynieria i technologia polimerów, metalurgia i korozja metali, mikro- i nanotechnologie materiałowe, biomateriały oraz komputerowe wspomaganie w inżynierii materiałowej. Na niektórych uczelniach specjalizację wybiera się dopiero na studiach magisterskich, podczas gdy inne pozwalają na profilowanie już na trzecim lub czwartym roku.

Studia na kierunku inżynieria materiałowa są wymagające i stawiają przed studentami liczne wyzwania. Największe trudności sprawia zazwyczaj „wielka trójka” czyli: matematyka, fizyka i chemia. „królowa nauk” uchodzi za jeden z najtrudniejszych przedmiotów, natomiast fizyka przenika wiele kursów kierunkowych, co sprawia, że jej znajomość jest kluczowa przez cały okres studiów. Chemia, zwłaszcza w kontekście



materiałów polimerowych czy ceramicznych, wymaga solidnej wiedzy teoretycznej i analitycznego myślenia. Obciążenie pracą jest duże. Studenci muszą regularnie przygotowywać raporty z laboratoriów, realizować projekty inżynierskie oraz przyswajać duże ilości teorii. Systematyczna nauka jest niezbędna, aby sprostać wymaganiom programu. Pierwsze semestry bywają najtrudniejsze. W miarę postępów w nauce

coraz więcej zajęć dotyczy zagadnień kierunkowych i laboratoriów, co sprawia, że studenci często odnajdują praktyczną przyjemność z nauki. Nie zmienia to jednak do końca ich sytuacji, gdyż przez cały okres edukacji będą musieli mierzyć się z nieprzespanymi nocami (niekoniecznie w wyniku integracji międzywydziałowej). Dużym wyzwaniem są projekty laboratoryjne oraz prace inżynierskie. Studenci muszą opanować obsługę specjalistycznych urządzeń badawczych, takich jak mikroskopy skaningowe czy maszyny wytrzymałościowe, a także nauczyć się interpretacji wyników badań. Istnieje również konieczność samodzielnego dokształcania się w zakresie oprogramowania inżynierskiego, gdyż liczba godzin informatyki może być niewystarczająca, by osiągnąć pełną biegłość w narzędziach CAD/CAE czy symulacjach materiałowych. Warto również zwrócić uwagę na znajomość języka angielskiego, która odgrywa istotną rolę. Wiele materiałów naukowych i dokumentacji technicznych jest dostępnych wyłącznie w tym języku. Dlatego studenci powinni rozwijać swoje kompetencje językowe, co zaprocentuje w przyszłej karierze zawodowej.

Zastanawiając się nad przyszłością zawodową, można być optymistą. Inżynierowie materiałowi są potrzebni w niemal każdej gałęzi przemysłu, od motoryzacji i lotnictwa, przez elektronikę i energetykę, aż po medycynę i nanotechnologię. Ich wiedza jest kluczowa przy projektowaniu nowoczesnych stopów metali, kompozytów czy materiałów do zastosowań ekstremalnych, na przykład w eksploracji kosmosu czy reaktorach jądrowych. Do głównych branż, w których mogą pracować absolwenci, należą: przemysł lotniczy i motoryzacyjny, energetyka, biotechnologie i medycyna, nanotechnologie i mikroelektronika. Wynagrodzenia są konkurencyjne. Średnie zarobki w pierwszych latach po studiach wynoszą około 6–8 tys. zł brutto, a w niektórych sektorach, takich jak na przykład nanotechnologie czy lotnictwo, mogą przekraczać 12 tys. zł brutto. Dodatkowo, zapotrzebowanie na specjalistów od materiałów rośnie. Rozwój nowoczesnych technologii wymaga coraz bardziej zaawansowanych rozwiązań materiałowych.

Inżynieria materiałowa to kierunek dla tych, którzy chcą mieć realny wpływ na kształtowanie przyszłości technologii. To interdyscyplinarna dziedzina, łącząca chemię, fizykę i inżynierię w projektowaniu nowoczesnych materiałów, od lekkich kompozytów lotniczych po inteligentne polimery wykorzystywane w medycynie. Choć studia są wymagające, dają solidne podstawy do pracy w dynamicznie rozwijających się sektorach gospodarki. Zapraszamy na studia. ■

Michał Pacholski



Drugie życie plastiku, część 3

Po lekturze dwóch poprzednich odcinków wiesz już dużo o tworzywach sztucznych, znasz ich rodzaje najczęściej spotykane w otoczeniu oraz masz świadomość niebezpieczeństw związanych z niekontrolowanym spalaniem plastików. W ostatnim odcinku szerzej o recyklingu tworzyw sztucznych. W porównaniu do poprzednich części więcej będzie tego, co chemicy lubią najbardziej, czyli doświadczeń.

Recykling ma przede wszystkim aspekt ekonomiczny, pozwalając odzyskać część kosztów poniesionych na produkcję przedmiotów, a nawet ograniczyć ją, co pociąga za sobą skutki w wielu innych dziedzinach, np. zmniejszenie zapotrzebowania na surowce. Jak już jednak wiesz, recykling musi być dostosowany do specyfiku przetwarzanego materiału. Po pierwsze zatem, należy dowiedzieć się...

...co to za tworzywo?

Choć niektóre tworzywa to **kopolimery**, czyli produkty polimeryzacji różnych związków, np. używany na obudowy ABS, to jednak w przetwórstwie nie należy mieszać surowców. Znasz już system oznaczeń stosowany dla tworzyw sztucznych (odpowiednie kody i symbole obowiązują również dla papieru, szkła, metali, tkanin). Jednak nie zawsze oznakowania są umieszczone na wyrobach, a ponadto chemik powinien umieć zidentyfikować materiał nawet bez oznaczeń. Stąd też wywodzi się cała sztuka analizy chemicznej, w zbliżonej do współczesnej formie licząca sobie już kilka wieków.

Pełna analiza tworzyw sztucznych, będących makromolekułami o złożonej budowie, jest możliwa do wykonania jedynie w specjalistycznych laboratoriach. W warunkach domowych możesz jednak zidentyfikować najpopularniejsze ich rodzaje. Do wykonania eksperymentów wystarczy źródło ognia (może być to nawet świeczka) oraz szczypta lub pęseta do trzymania próbek. Zachowaj niezbędne środki ostrożności przy posługiwaniu się otwartym ogniem: w pobliżu nie mogą znajdować się łatwopalne przedmioty i substancje, a pod ręką miej mokrą szmatkę do stłumienia ognia oraz miseczkę z wodą do zgaszenia tłącego się tworzywa. Pamiętaj, aby używać próbek o powierzchni nieprzekraczającej 0,5 cm². Fragment tworzywa uchwyc

pęsetą i wprowadź do płomienia. Podczas identyfikacji zwróć uwagę na palność tworzywa, barwę płomienia oraz wydzielającą się woń. Zachowanie próbki oraz jej wygląd po spaleniu może odbiegać od opisu w zależności od użytych dodatków, np. wypełniaczy, barwników. Do eksperymentów użyj próbek tworzyw występujących w twoim otoczeniu: kawałków folii, butelek i opakowań.

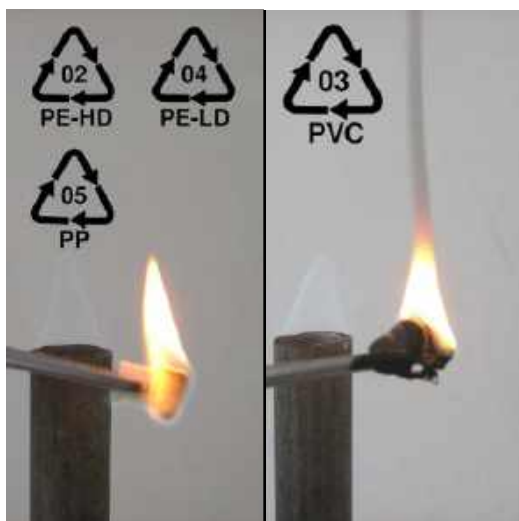
Analiza płomieniowa tworzyw sztucznych

Polietylen (oba rodzaje: **PE-HD** i **PE-LD**) oraz **polipropylen PP** są nierozróżnialne podczas analizy płomieniowej. Próbki zapalają się łatwo i nie gasną po wyjęciu z palnika. Płomień jest barwy żółtej z niebieską otoczką, tworzywo topi się i spływa kroplami. Wyczujesz zapach palonej parafiny. Podobieństwo to nie jest przypadkowe: zarówno parafina, jak i oba tworzywa zbudowane są z długich łańcuchów węglowodorowych (1).

Poli(chlorek winylu) PVC zapala się z trudem, próbka często gaśnie po wyjęciu z palnika. Płomień jest barwy żółtej z zieloną otoczką, wydziela się niewielka ilość dymu, a tworzywo wyraźnie mięknie. Płonący PVC ma ostry zapach chlorowodoru, o czym można się było przekonać w pierwszym odcinku (2).

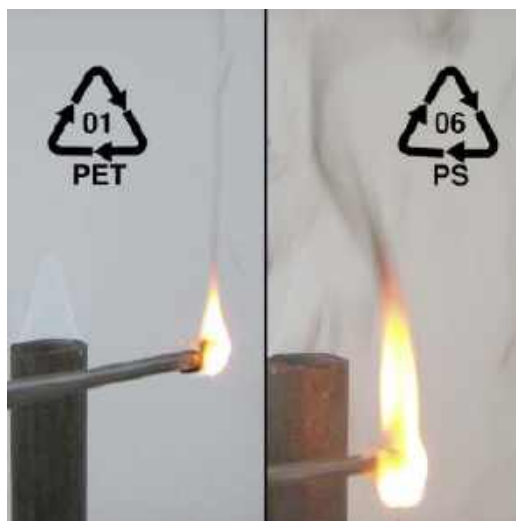
Poli(tereftalan etylu) PET zapala się po pewnym czasie i często gaśnie po wyjęciu z palnika. Płomień jest barwy żółtej, lekko kopący, możesz również wyczuć ostry zapach.

Polistyren PS zapala się po pewnym czasie i nie gaśnie po wyjęciu z palnika. Płomień jest barwy żółtopomarańczowej i wydziela się czarny dym, a tworzywo mięknie i topi się. Zapach jest dość przyjemny. Czarny dym, widoczny również podczas badania PET, ma związek z proporcją węgla do wodoru w próbce:



1. Analiza płomieniowa polietylenu i polipropylenu

2. Analiza płomieniowa poli(chlorku winylu)



3. Analiza płomieniowa PET i polistyrenu

im mniejsza zawartość wodoru, tym więcej dymu. Polistyren i PET w swojej strukturze zawierają pierścienie benzenowe o składzie C_6H_6 , w których proporcja C:H wynosi 1:1, natomiast polietylen i polipropylen budują jednostki CH_2 o dwukrotnie większym udziale wodoru, stąd tworzywa te praktycznie nie wydzielają dymu podczas spalania (3).

Poliwęglan zapala się po pewnym czasie i czasem gaśnie po wyjęciu z płomienia. Pali się świecącym płomieniem, kopcząc. Zapach jest charakterystyczny.

Guma łatwo się zapala i nie gaśnie po wyjęciu z palnika. Płomień jest ciemnożółty, silnie kopczący. Wyczujesz charakterystyczny zapach palonej gumy. Pozostałość po spaleniu jest nadtopioną, lepką masą.

Poliamid (np. żyłka wędkarska) zapala się po pewnym czasie i niekiedy gaśnie po wyjęciu z płomienia. Płomień jest barwy jasnoniebieskiej z żółtym wierzchołkiem. Tworzywo topi się i spływa kroplami. Zapach przypomina palone włosy, co nie jest przypadkowe – białka wchodzące w skład włosów to także poliamidy.

Poli(metakrylan metylu) („szkło organiczne”) zapala się po pewnym czasie i nie gaśnie po wyjęciu z palnika. Płomień jest barwy żółtej z niebieską otoczką, tworzywo mięknie podczas spalania. Wyczujesz przyjemny, kwiatowy zapach.

Jak je rozdzielić?

Każde z tworzyw wymaga innych warunków przetwarzania, rozdzielenie jest zatem konieczne, aby nie zanieczyścić surowców do przeróbki. O ile zakłady przemysłowe dostarczają na ogół jednego rodzaju

odpadów, o tyle ze zbiorek publicznych uzyskuje się mieszaninę wielu tworzyw. Rozdzielenie nie jest jednak skomplikowane.

Potnij na drobne fragmenty butelkę z PET oraz kubki po produktach mlecznych z polipropylenu lub polietylenu, a także z polistyrenu. Próbkę muszą być w różnych kolorach, co umożliwi ci obserwację wyników doświadczenia (4). Wrzuć je wszystkie do naczynia i nalej wody. Zamieszaj, a gdy kawałki plastiku przestaną wirować, zbierz te, które znajdują się



4. Próbkę tworzyw (na górze) zostają zmieszane w zlewce (na dole)



5. Widok próbek tworzyw w zlewce po wleciu wody (u góry) oraz po zebraniu niebieskich fragmentów i dosypaniu soli do wody (na dole)

na powierzchni. Do wody wsyp łyżkę soli i ponownie zamieszaj zawartość naczynia. Następnie zbierz pływające fragmenty (5). Jakie tworzywo zebrałeś za każdym razem?

W eksperymencie wykorzystana została różnica gęstości tworzyw. Polipropylen i polietylen mają gęstości mniejsze niż woda i dlatego pływają po powierzchni. Gęstość polistyrenu jest jednak tylko nieznacznie większa od gęstości wody. Dodatek soli tak zwiększa gęstość cieczy, że kawałki tego tworzywa wypływają na powierzchnię. PET jest znacznie gęstszy od roztworu soli, pozostaje więc na dnie do końca doświadczenia. Metoda jest stosowana w przetwórnictwie tworzyw po uprzednim pocięciu plastikowych odpadów na drobne fragmenty.

Recykling niejedno ma imię

Obecnie największym problemem nie jest już wyprodukowanie przedmiotu, lecz jego utylizacja po coraz krótszym okresie użytkowania. Dodatkowo wiele zastosowań tworzyw sztucznych powoduje, że szybko stają się one tylko odpadem, np. czas życia kubków po wyrobach mleczarskich, butelek po napojach



6. Plastikowe śmieci wrzucamy do żółtego pojemnika

i torebek foliowych to tylko kilka dni, a w pełni funkcjonalne i poręczne opakowania po jednorazowym użyciu lądują na wysypiskach. Plasterki nie ulegają tak łatwej biodegradacji, jak na przykład papier czy drewno, lecz latami szpecą otoczenie. Wykonuje się z nich głównie opakowania, objętościowo są zatem największym składnikiem śmieci, a kolorowe wykończenie plastikowych wyrobów sprawia, że są szczególnie widoczne.

Sposób przetwarzania odpadów należy jednak dostosować do ich charakteru. O właściwościach tworzyw w sztucznych decyduje budowa molekuł: bardzo długie łańcuchy składające się z tysięcy identycznych części. Przeszkodą stojącą na drodze szerszego wykorzystania zużytych plastików są koszty zbiórki i segregacji. Recykling materiałowy i surowcowy ma sens ekonomiczny wtedy, gdy odpady z tworzyw sztucznych są posortowane, najlepiej już przez użytkownika (6). Niżej główne rodzaje recyklingu stosowane w ich przypadku.

Recykling materiałowy polega na ponownym wykonaniu wyrobu z przetwarzanego materiału. W przypadku tworzyw sztucznych metoda ta jest jednak



7. Cztery butelki PET po napojach to materiał na jeden T-shirt (metka na koszulce)



8. Dodatek TDPA (Totally Degradable Plastics Additives) ułatwia rozpad plastikowej torebki na drobne fragmenty pod wpływem czynników środowiska

związana z nieuchronnym pogorszeniem jakości powstających produktów: stopienie polimeru powoduje jego częściowy rozkład spowodowany pękaniem węglowych łańcuchów. Recykling materiałowy jest rzadko stosowany podczas przerobu tworzyw sztucznych, głównie w przypadku PET i PE-HD. Ze zużytych butelek PET nie wytwarza się nowych pojemników, lecz przeznaczają je na produkcję włókien poliestrowych (7).

Recykling surowcowy polega na rozkładzie tworzyw w taki sposób, aby otrzymać składniki do produkcji nowych wyrobów. Ten rodzaj przeróbki jest obecnie mało opłacalny, ponieważ surowce do syntezy tworzyw sztucznych są tanie i łatwo dostępne (ropa naftowa i jej pochodne), a wielka skala wytwarzania pozwala uzyskać niskie ceny monomerów. Jednak uwarunkowania ekonomiczne zmieniają się często i odzyskiwanie surowców z tworzyw sztucznych może stać się opłacalne.

Recykling energetyczny polega na spaleniu odpadów w kontrolowanych warunkach (muszą być spełnione ostre wymogi dotyczące ochrony środowiska) i wykorzystaniu powstającej w tym procesie energii. Plusem tej metody jest fakt, że cząsteczki tworzyw składają się głównie z łańcuchów węglowodorowych. Ich skład jest zatem zbliżony do składu paliw kopalnych, w związku z czym nie trzeba dokonywać radykalnych

zmian istniejących już instalacji. Jednak projekty budowy spalarni natrafiają na liczne protesty okolicznych mieszkańców, którzy nie chcą mieć takiego sąsiada.

Użycie tworzyw biodegradowalnych (zmodyfikowana celuloza i skrobia, tworzywa białkowe, niektóre poliestry, stosowane samodzielnie lub w mieszaninach z materiałami syntetycznymi) także nie rozwiązuje problemu odpadów. Nowe materiały nie nadają się do wielu zastosowań i praktycznie są wykorzystywane tylko jako opakowania. Do szerszego użycia zniechęcą również koszty produkcji, znacznie wyższe niż w przypadku innych polimerów. Spotykane obecnie w sklepach torby z tworzyw sztucznych opatrzone napisem „przyjazne środowisku” nie do końca są takie, jakimi przedstawiają je producenci. Występują w różnych odmianach: torby degradable (dodatki powodują szybki rozpad na drobne fragmenty, które nie szpecą środowiska tak jak cała torebka), oxybiodegradable (rozpad pod wpływem tlenu i promieni UV), biodegradable (z modyfikowanej skrobi lub celulozy) oraz kompostowalne (mogą być przetwarzane na kompost wraz z odpadami organicznymi) (8).

Łatwo jednak przewidzieć, że nawet po zastąpieniu syntetycznych plastików ich biodegradowalnymi odpowiednikami, miejsce zaśmiecających nasze otoczenie polietylenowych torebek zajmą te z nadrukiem BIO. Tu potrzebna jest całkowita zmiana mentalności użytkowników, a to już stanowi problem znacznie trudniejszy niż wdrożenie nowych rozwiązań naukowo-technicznych. Plastikowa „stajnia Augiasza” nadal czeka na swego Herkulesa. ■

Krzysztof Orliński





Czy kwarki naprawdę są kolorowe?

Analizując przemiany promieniotwórcze jąder atomowych, możemy, a nawet powinniśmy, dojść do wniosku, że nukleony (czyli protony i neutrony) nie mogą być cząstkami elementarnymi w ścisłym znaczeniu tego słowa. Przyjrzyjmy się uważniej tej sprawie.

Wewnętrzna budowa nukleonów

Część izotopów promieniotwórczych ulega przemianom beta, przy czym możliwe są dwa rodzaje takiej przemiany: przemiana beta minus (neutron zamienia się w proton i elektron) oraz przemiana beta plus (proton zamienia się w neutron i pozyton). Skoro z jednej cząstki powstają dwie, musi istnieć wewnętrzna struktura tejże cząstki pozwalająca na przegrupowanie bardziej podstawowych cegiełek materii. W innym przypadku przemiany te nie byłyby w ogóle możliwe.

Jeśli zastanowimy się nieco głębiej nad przemianą beta plus, to jest tu jeszcze jedna ciekawa rzecz: pojawia się pozyton, czyli antycząstka będąca składnikiem antymaterii.

Odrobina historii

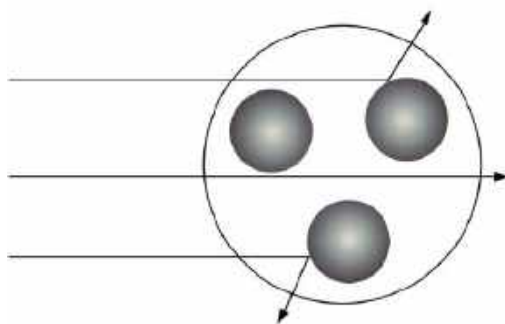
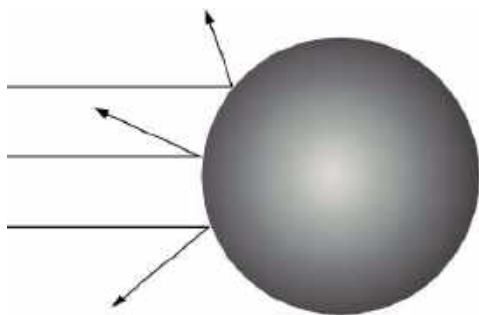
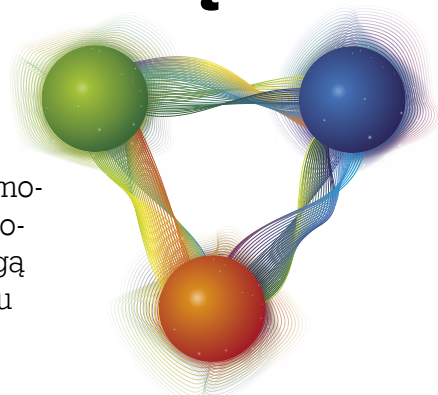
Istnienie kwarków przewidziano na drodze teoretycznej w 1964 roku. Już cztery lata później uzyskano pierwsze doświadczalne potwierdzenie tej hipotezy. Elektrony rozpędzane w akceleratorze do relatywistycznie niewielkich prędkości odbijały się od protonów w sposób, który można było opisać

jako sprężyste zderzenie dwóch kul. Wraz ze wzrostem prędkości elektronów wyniki eksperymentu zmieniały się i wskazywały na to, że elektrony wnikały w głąb protonów i tam zderzały się z mniejszymi obiektami.

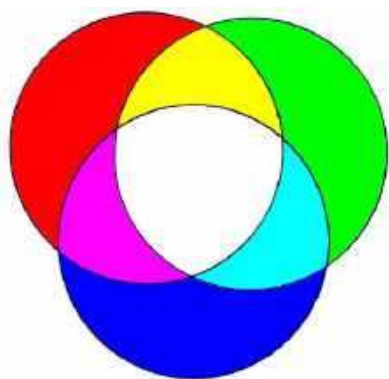
Wkrótce potem potwierdzono analogiczną budowę neutronów. W drugiej połowie XX wieku dynamiczny rozwój fizyki wysokich energii doprowadził również do odkrycia innych cząstek zbudowanych z kwarków. Udało się również opisać różne rodzaje kwarków, różniące się masą, ładunkiem elektrycznym oraz kilkoma innymi cechami, niespotykanymi w przypadku cząstek budujących otaczającą nas materię. Jedną z takich cech jest kolor.

Czym jest kolor kwarków?

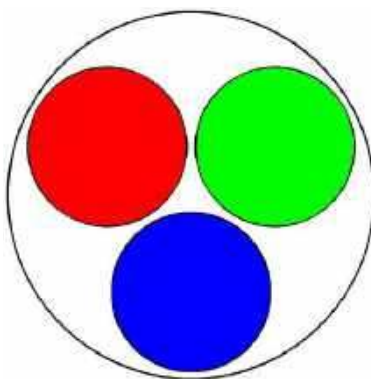
Kolor kwarków można porównać do ładunku elektrycznego. W przypadku zjawisk z dziedziny elektrostatyki wyróżniamy ładunek dodatni oraz ładunek ujemny. Na przykład elektron ma jednostkowy ładunek ujemny, proton z kolei – jednostkowy ładunek dodatni. Atom wodoru zbudowany z powyżej



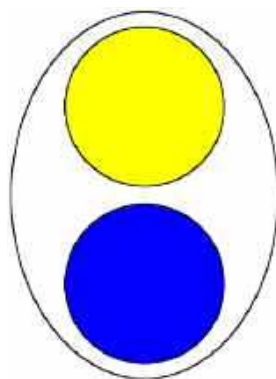
1. Przy niewielkich prędkościach (lewa strona rysunku) elektrony zderzają się z protonem jak ze sztywnej kulą. Przy odpowiednio dużych prędkościach (prawa strona rysunku) elektrony wnikały w głąb protonu i zderzały się z kwarkami



synteza addytywna barw



trzy kwarki w różnych kolorach tworzą nukleon



kwark i antykwark tworzą mezon

2. Kolor kwarków jest przykładem ładunku trójmiennego. Aby nukleon był neutralny pod względem tej cechy, musi zawierać trzy kwarki, każdy w innym kolorze. Mezon zawiera parę kwark-antykwark. W tym drugim przypadku antykwark jest obdarzony antykoleorem

opisanych cząstek jest elektrycznie obojętny (jego wypadkowy ładunek wynosi zero). Podobnie wygląda to w przypadku bardziej skomplikowanych układów fizycznych. Jeśli suma ładunków dodatnich jest równa sumie ładunków ujemnych, dane ciało jest elektrycznie obojętne.

W przypadku koloru mamy do czynienia z podobną sytuacją, tyle tylko, że zamiast dwóch istnieją trzy różnoimienne ładunki przenoszące tę cechę, nazywane czerwonym, zielonym i niebieskim. Dopiero suma tych trzech ładunków daje cząstkę neutralną pod względem koloru. Tak więc słowo „kolor” zostało użyte przez analogię do syntezy addytywnej wiązek światła o różnych barwach.

Czy kwarki mają ładunek elektryczny?

Każdy kwark oprócz ładunku kolorowego przenosi również ładunek elektryczny. Czeką nas tu jednak pewna niespodzianka. Z dotychczasowych lekcji fizyki czy chemii wiemy, że najmniejszy możliwy ładunek jest równy ładunkowi elementarnemu oznaczanemu jako e . Kwarki jednak obdarzone są ładunkiem ułamkowym! Ładunek ten może wynosić $+\frac{2}{3}e$ lub $-\frac{1}{3}e$. I tak na przykład wewnątrz protonu mamy dwa kwarki o ładunku $+\frac{2}{3}e$ oraz jeden kwark o ładunku $-\frac{1}{3}e$.

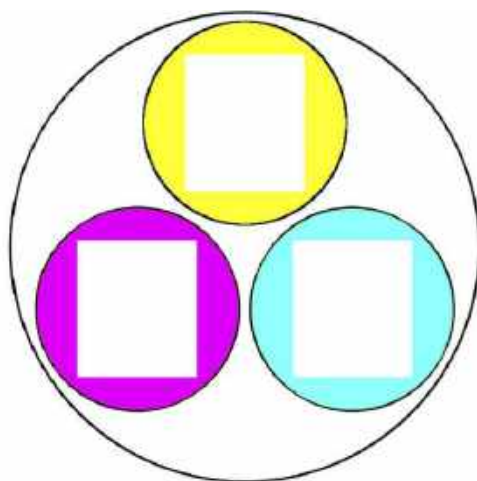
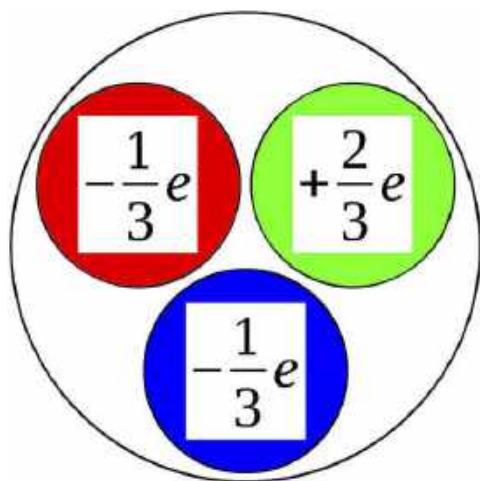
W żadnym ze zjawisk fizycznych nie rejestrujemy jednak ładunków ułamkowych z prostej przyczyny – kwarki są ze sobą tak silnie związane, że nie mogą występować pojedynczo. Jakakolwiek próba wyrwania kwarka z nukleonu wymagałaby ogromnej energii, o wiele rzędów wielkości większej niż potrzeba na przykład do zjonizowania atomu.

Badania prowadzone z wykorzystaniem akceleratorów wykazały, że oprócz protonu i neutronu istnieje więcej cząstek zbudowanych z trzech kwarków, które nazwano barionami. Jak już wspomniano, układy kwark-antykwark nazywamy mezonami. Zarówno bariony jak i mezony, ze względu na swoją wewnętrzną strukturę, mogą podlegać rozpadowi. Również zderzenie dwóch tego typu cząstek, przy ich odpowiednio dużej energii kinetycznej, może doprowadzić do przegrupowania kwarków i powstania cząstek innego rodzaju. Na przykład wiele mezonów powstaje na skutek zderzeń pomiędzy rozpędzonymi barionami.

Antymateria

Wróćmy na chwilę do pozytonu powstającego w przemianie beta plus. W zasadzie, z wyjątkiem znaku ładunku elektrycznego, cząstka ta nie różni się od elektronu. Wartość ładunku pozostaje ta sama, podobnie jest z masą. Zatem pozyton jest antycząstką elektronu. Co ważne – jeśli w tym samym punkcie przestrzeni cząstka napotka swoją antycząstkę – obie ulegną anihilacji, pozostawiając po sobie energię równą sumie ich energii spoczynkowych. Dla przypomnienia – energia spoczynkowa to masa cząstki mnożona przez kwadrat prędkości światła ($E = mc^2$).

W przypadku antymaterii generalną zasadą jest, że antycząstka ma taką samą masę jak cząstka, ale przeciwny ładunek. Dotyczy to zarówno ładunku elektrycznego jak i ładunku kolorowego. I tak na przykład jeśli proton ma ładunek $+e$ i składa się z kwarków w kolorach czerwonym, zielonym i żółtym,



to antyproton ma ładunek $-e$ i składa się z antykwarków w kolorach antyczerwonym, antyzielonym i antyniebieskim.

Sprawdź swoją wiedzę

Na poniższej ilustracji przedstawiono neutron zbudowany z kwarków (lewa strona) i antyneutron zbudowany z antykwarków (prawa strona). Uzupełnij rysunek, wpisując właściwe wartości ładunku elektrycznego w puste pola. W razie problemów z ustaleniem par kolor-antycolor podpowiedź znajdziesz na rysunku 2.

Dla nauczyciela

Tematy związane z wewnętrzną budową nukleonów wykraczają poza podstawę programową przedmiotu fizyka, niemniej w szkole ponadpodstawowej w zakresie rozszerzonym mogą zostać wprowadzone jako uzupełnienie zagadnień związanych z przemianami beta jąder atomowych oraz anihilacją par cząstka-antycząstka. Jest to o tyle zalecane, że brak podstawowych informacji na temat struktury nukleonów może powodować problem z wyobrażeniem sobie przez ucznia procesów zachodzących wewnątrz jądra atomowego. ■

Joanna Borgensztajn

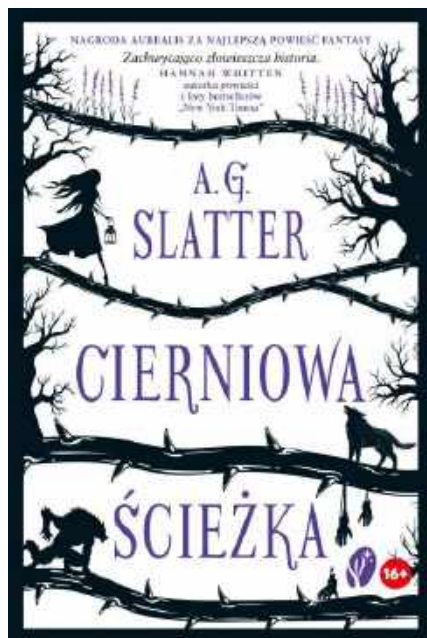
Cierniowa ścieżka

A. G. Slatter

Wydawnictwo StoryLight, liczba stron: 488, cena sugerowana: 49,99 zł

Pokręcona, brutalna i mroczna baśń dla dorosłych przesycona tajemnicami rodzinnymi, czarną magią i gorzkim posmakiem zemsty. Czytelnicy Naomi Novik i Erin Morgenstern będą zachwyceni.

Asher Todd zostaje zatrudniona w posiadłości rodziny Morwood jako guwernantka. Co prawda nie wie zbyt wiele o wychowywaniu dzieci, ale doskonale zna się na botanice i ziołarstwie – a kto wie, może nawet na czymś więcej. Morwood Grange na pierwszy rzut oka wydaje się zwykłym, spokojnym miejscem: piękny stary budynek, którego mieszkańcy są dumni ze swojego pochodzenia i nienagannej reputacji, a nieopodal ich dworu znajduje się miasteczko Tarn oraz las, z którego nocami słychać czasem wycie wilków... W krótkim czasie dzieci Morwoodów przywiązują się do Asher, starsza pani Leonora coraz częściej zaprasza ją do swojej komnaty, a czarnowłosa dozorca o dzikich zielonych oczach raz po raz zupełnym przypadkiem wchodzi jej w drogę. Asher zaczyna lubić swoich podopiecznych i życie w posiadłości. Rodzi się w niej wątpliwość, czy będzie w stanie zrealizować swój plan – i kto najbardziej ucierpi, jeśli to zrobi. Ale gdy duchy przeszłości stają się coraz trudniejsze do kontrolowania, dziewczyna zdaje sobie sprawę, że nie ma wyboru. Zresztą nie tylko ona skrywa mroczne sekrety. Dom Morwoodów pełen jest własnych tajemnic: wymazanych nazwisk, brakujących portretów na ścianach i pokoi zamkniętych na trzy spusty. Pytanie brzmi: o co w tym wszystkim chodzi? I czy w Morwood Tarn jest choć jedna osoba, która nie ma nic do ukrycia?



*** Pisownia oryginalna ***

PRZEGLĄD PRZEMYSŁOWO-HANDLOWY Otwarcie Międzynarodowego Targu w Poznaniu

W dniu 3 maja nastąpi uroczyste otwarcie Międzynarodowych Targów w Poznaniu. Będzie to po raz pierwszy, że Targi Poznańskie, istniejące od roku 1921, przyjmą charakter międzynarodowy, w odróżnieniu od dotychczasowego stanu rzeczy, kiedy ograniczały się do roli Targów Krajowych. Dokonana zmiana charakteru posiada swe usprawiedliwienie w stopniowej ewolucji, jaką stan gospodarczy kraju, a wraz z nim, również i Targi Poznańskie musiały przechodzić. W swych początkach Targi Poznańskie postawiły sobie dwa naczelne zadania. Były nimi: zbliżenie wzajemne poszczególnych obszarów Polski pod względem gospodarczym oraz zaznajomienie przedewszystkiem zagranicy z krajowymi wytworami przemysłu. Pierwszy cel, a więc dopomożenie zbliżeniu gospodarczemu poszczególnych ziem Polski pomiędzy sobą, musiał siłą rzeczy wypłynąć wobec stosunków, jakie ujawniły się po ponownym złączeniu rozdartych uprzednio polskich ziem przez przemoc zaborców. Kontakt pomiędzy temi zabarami osłabł, wzajemna znajomość możliwości produkcji i zbytu zmalała do minimum. Konieczne było ponowne poznanie się, nawiązanie jaknajściślejzych stosunków, aby w ten sposób spoić znów gospodarczo wszystkie obszary Polski w jedną nierozzerwalną całość. Tą rolę Targi Poznańskie odegrały w swym zakresie działania znakomicie, stanowiąc się punktem spotkań sfer gospodarczych z całej Polski. Drugie zadanie – zaznajomienie zagranicy z naszym handlem i przemysłem, zgrupowanym na Targach – mogło być należycie dokonane tylko wówczas, gdy przedewszystkiem zostaną zgromadzone eksponaty polskiego przemysłu. To też udział zagranicznych wystawców celowo był ograniczony, skąd właśnie wypłynął krajowy charakter Targów Poznańskich. Dziś Targi te stają się międzynarodowymi. Zagranica została dopuszczona na równych prawach z krajowym wystawcą. W ten sposób pole i zakres działalności Targów zwiększyły się, dając możność zagranicy nie tylko zaznajamiania

się z polskimi wytworami przemysłu, lecz stanąć również do konkurencji przez przesłanie własnych eksponatów. Otwiera się więc również pole do działalności dla polskiego handlu importowego i tranzytowego, przemysł zaś polski będzie mógł porównać swe wysiłki oraz osiągnięte wyniki z wynikami, osiągniętymi przez zagranicę. Poznań wszedł więc ze swemi Targami w orbitę handlu międzynarodowego. Ewolucję tę przeszedł stopniowo, przystosowując się systematycznie do tej roli zarówno pod względem społeczno-gospodarczym, jak i technicznym. Zwłaszcza co do technicznego przygotowania należy podziwiać wspaniałe rezultaty, osiągnięte w tak krótkim przeciągu czasu. Nie szczędząc wysiłków i nakładów pieniężnych, polegając tylko na własnych siłach, Targi Poznańskie zdobyły wybudować szereg wzorowych wystawowych. Dziś Targi rozporządzają terenem, położonym po obu stronach linii kolejowej, a obejmującym około 350,000 m², z czego 40,000 m² zajmują trwałe, nowoczesne budynki. W ostatnim już roku zbudowano wspaniałą halę wystawową z żelazobetonu, piętrową, jakoteż dwupiętrowy dom administracyjny z pomieszczeniem na restaurację, salę i t. p. Działalność Dyrekcji Targów świadczą o stałym postępie i rozwoju tej doniosłej dla życia gospodarczego instytucji. Rok ubiegły, kiedy Targi rozporządzały daleko mniejszą ilością miejsc wystawowych i nosiły charakter krajowych, zgromadził przeszło 100000 zwiedzających i przeszło 2000 wystawców. Jest do przewidzenia, że obecnie, po przekształceniu Targów Poznańskich na Targi Międzynarodowe, rozwój ten wkróci na jeszcze pomyślniejsze tory, dzięki czemu Targi w Poznaniu trwale zachowają stanowisko instytucji, przodującej w naszym życiu gospodarczym.

1 maja 1925

Przemysł elektrotechniczny na Międzynarodowym Targu w Poznaniu

Z przemysłu elektrotechnicznego na międzynarodowym Targu w Poznaniu najwięcej miejsca zajmuje dział oświetlenia elektrycznego. A więc: wszelkiego rodzaju

żarówki, lampy i żyrandole, baterie elektryczne i materiały instalacyjny. Dalej z wyrobów przemysłu elektrotechnicznego zobaczymy telefony i stacje telefoniczne wraz z przyborami instalacyjnymi, bogaty dział radiofoniczny, który napewno wzbudzi wielkie zainteresowania na Targu, oraz prawie wszystkie pozostałe rodzaje wyrobów tej gałęzi przemysłu. Silniki, kable, urządzenia i uzupełnienia elektrowni bezwzględnie zainteresują delegatów Związku Miast Polskich.

1 maja 1925

PRZEGLĄD TECHNICZNY Nowy pięciowrzecionowy automat Gridley'a

Wielowrzecionowy automat Gridley'a, wyrabiany przez National Acme Company of Cleveland, cieszący się zaskonującą popularnością nie tylko w amerykańskich, ale i europejskich wytwórniach maszyn i aparatów, zmienił się do niepoznania w ostatnich latach, dzięki znakomitemu uproszczeniu konstrukcji. Pod tym względem przypomina on raczej jednowrzecionowy automat Gridley'a. Wszystkie złożone mechanizmy zostały tak zmienione, że automat zawiera jedynie typowe konstrukcje. Jak wiadomo istotną cechą tego automatu jest doskonale pomyślany bęben narzędziowy, umożliwiający stosowanie krótkich dobrze zamocowanych imaków narzędziowych. Bęben narzędziowy jest przesuwany za pośrednictwem dopychacza i bębna krzywkowego, umieszczonego wewnątrz łoża. Inny duży bęben służy do podawania i zamocowywania materiału prętowego we wrzecionach głowicy. Mniejszy bęben krzykowy umieszczony w łożu pod głowicą służy do napędu suportów bocznych. Automaty Gridley'a budowane są w różnych wielkościach od najmniejszych do największych, przystosowanych do materiału prętowego o średnicy 100 mm. Na uwagę zasługuje fakt, że automaty Acme zdobywają ponownie rynek niemiecki, zarzucony dotychczas przez naśladownictwa miejscowe automatów amerykańskich, co stało się rzeczą możliwą wobec wysokiego kursu waluty niemieckiej, a jest wynikiem postępu konstrukcyjnego w tej dziedzinie za oceanem.

13 maja 1925

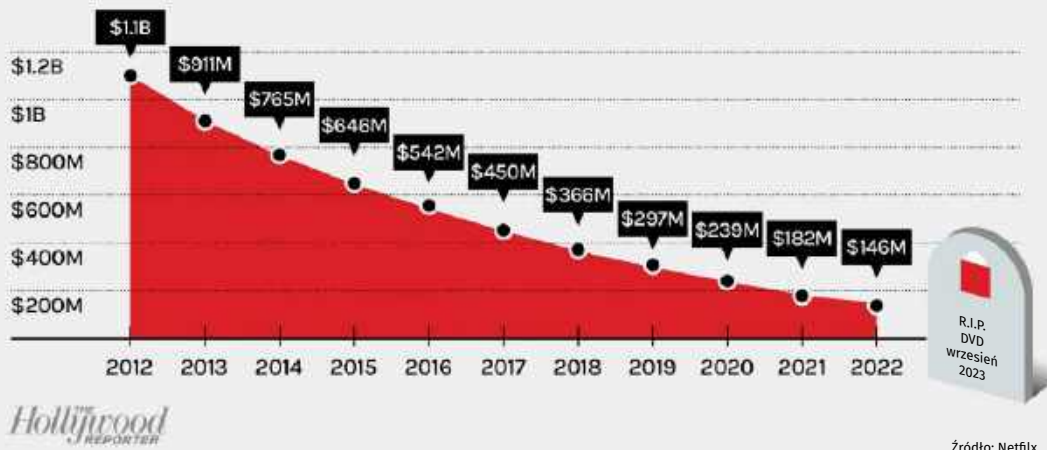
Skaferander pancerny

Ostatnim wynalazkiem w dziedzinie nurkowania jest aparat skonstruowany przez firmę „Neufeld & Kuhne” w Kilonji. Aparat ten zewnętrznym wyglądem przypomina średniowieczny pancerz rycerski, urządzeniem zaś wewnętrznym, jak to widzimy z przekroju i specyfikacji, baszę dowodzącą nad podwodnej. Na główny korpus użyto stali mar-tenowskiej; osłony kończyn, sporządzone z niklowej stali, łączą się przegubami kulowymi, uszczelnieniami zapomocą wałków z przeoliewionej gumy, dając możność nurkowi poruszania nim (do pewnej głębokości) bez wielkiego wysiłku. Całość jest wypróbowana na 26 at ciśnienia zewnętrznego. W aparacie tym dokonywane prób nurkowania w Walchensee na głębokościach do 160 m, t. zn. pod ciśnieniem do 16 at. Czas przebywania pod wodą trwa od 2-ch do 5-ciu godzin, poczem nurka używano zaraz do innej roboty, by skontrolować stopień upadku sił. Dodatnie wyniki w tym względzie zawiadzcza się temu, że pancerz stawia opór ciśnieniu wody, a nurk w tym aparacie jest cały czas w normalnej atmosferze normalnego ciśnienia. Dzięki też temu, że ciśnienie wewnątrz aparatu się nie zmienia, aparat wraz z nurkiem można zanurzać na głębokość do 160 m i podnosić w przeciągu 4 do 5 minut. Inaczej przedstawia się praca w skaferandach gumowych, w których ciśnienie wody trzeba zrównoważyć ciśnieniem powietrza. Ponieważ sprężone powietrze w ubraniu gumowym bezpośrednio styka się z ciałem nurka, zanurzenie musi się odbywać powoli (1 m/min), a wychodzenie jeszcze wolniej (0,5 m/min), by organizm stopniowo przyzwyczajał do zmiany ciśnienia. Samo zatem opuszczanie się i podnoszenie trwa już ok. 3 godz., a na dzień przy głębokości 40 m można pracować od 15 do 20 minut – i to tylko raz na dobę, z powodu wielkiej utraty sił przez nurka. Próby nurkowania poniżej 60 m kończyły się przeważnie śmiercią. Powyższe porównanie głębokości nurkowania i czasu pracy przy używaniu obydwóch aparatów podkreśla doniosłość nowego wynalazku, którego ujemną stroną jest to, że nurk nie pracuje bezpośrednio rękami, lecz kleszczami lub hakami, zmieniającymi zależnie od zadanej roboty. Na głębokości 160 m nurk porusza kończynami już z wielkim wysiłkiem, a widzi w promieniu tylko ok. 1 m.

13 maja 1925

Ile dolarów zarabiał Netflix na DVD w kolejnych latach dekady

Od szczytu przychodu, który wyniósł ponad miliard dolarów, dochody ze sprzedaży płyt stale spadały



1. Przychody Netflixa ze sprzedaży DVD

Filmy na płytach

Chmury zamiast regałów

Netflix wysłał w 2023 r. ostatnią „czerwoną kopertę”, co komentowano jako koniec ery wypożyczania płyt DVD przez firmę, która od takiej właśnie usługi ćwierć wieku temu zaczynała. Wielu już nawet nie pamięta, że potentat VOD (ang. Video on Demand) swoje pierwsze pieniądze zarobił na płytach z filmami.

Czerwone koperty, w których wysyłano klientom zamawiane płyty, które przez długi czas były swoistym symbolem samego Netflixa, wypełniały półki w domach w całych Stanach Zjednoczonych. Firma, która uruchomiła usługi streamingowe siedemnaście lat temu, mimo dość szybkiego sukcesu tej usługi, wciąż prowadziła wynajem filmów na DVD. Jednak w ostatnich latach działalność Netflixa w zakresie DVD podupadała, spadając z ponad miliarda dolarów przychodów w 2012 roku do 146 milionów dolarów do 2022 roku (1).

Z pewnością rynek ten w USA nie jest martwy, zwłaszcza że Amazon i Walmart wciąż są w grze. W rzeczywistości Media Play News donosiło latem, że Walmart

przewodził rozmowy ze Studio Distribution Services (joint venture Universal/Warner Bros.) na temat współpracy w branży fizycznych nośników. Jednak bez Netflixa i Best Buy, sieci sklepów, która też niedawno wycofała ze sklepów sprzedaż płyt DVD i Blu-ray, a także wielu innych firm, które być może podążą za nimi, los płyt DVD (2) maluje się w coraz ciemniejszych barwach nie tylko w USA, bo to, co się dzieje w tym kraju, zwykle zapowiada zmiany na całym świecie.

Pozbywanie się malejącego biznesu

Wkrótce po decyzji Netflixa Disney zdecydował się na przeniesienie dystrybucji płyt Blu-ray, DVD

i innych nośników fizycznych do Sony Entertainment. Inni wielcy hollywoodzcy gracze, tacy jak Warner Bros. Discovery i Universal, również poszukują partnerów, którym można by powierzyć zajmowanie się malejącym rynkiem nośników fizycznych. W 2020 roku Warner Bros. i Universal utworzyły Studio Distribution Services, aby wspólnie obsługiwać wydawnictwa DVD i Blu-ray. Model ten potem rozszerzył się o partnerstwa z innymi studiami.

Kilka miesięcy później firma Sony ogłosiła zaprzestanie produkcji nagrywalnych płyt Blu-Ray. Według japońskiego źródła, serwisu AV Watch, Sony zaprzestaje produkcji czterech typów konsumenckich dysków do nagrywania: 25 GB BD-RE, 50 GB BD-RE DL, 100 GB BD-RE XL i 128 GB BD-R XL, a także profesjonalnych dysków do produkcji wideo i archiwizacyjnych dysków optycznych do przechowywania danych. Decyzja Sony wynikała znów z coraz słabszego rynku i trudności w osiągnięciu opłacalności w tym sektorze. W lipcu 2024 r. nie podano dokładnej daty zaprzestania produkcji nagrywalnych płyt Blu-ray, ale uważa się, że sprzedawane są obecnie ostatnie partie płyt.

Decyzje firmy nie dotyczą wszystkich form dysków fizycznych. Na przykład płyty do konsol PlayStation i Xbox są wciąż produkowane w zakładach DADC (Digital Audio Disc Corporation) i będą dostępne nieco dłużej. Sony nie wycofało też (jeszcze?) płyt Blu-ray z filmami i gramami, a jedynie wstrzymuje produkcję konsumenckich płyt BD-R. Tak czy inaczej, panuje opinia, że w ciągu najbliższych kilku lat może nadejść kres wszystkich nośników fizycznych firmy.

Od szybkiego wznoszenia po upadek

Wprowadzone w 1997 roku DVD oficjalnie reklamowało się jako „film na dysku wielkości płyty CD”. Był to film o wysokiej, nieznanym erze kasety wideo, jakości. Doświadczenie oglądania było z pewnością lepsze niż w przypadku VHS i bez konieczności ręcznego przewijania fizycznej taśmy. Do 2003 roku całkowite przychody z domowego wideo na płytach różnego rodzaju wzrosły do prawie 24,5 miliarda dolarów, a każda duża wytwórnia filmowa poświęcała ogromną ilość czasu i pieniędzy, wydając swoje produkcje na płytach DVD, często z dodatkowymi funkcjami, usuniętymi scenami i tzw. blooperami (czyli zabawnymi wypadkami podczas kręcenia filmu).

W 2015 roku pojawił się na rynku format płyty Ultra HD Blu-ray, oferujący szybkość transmisji od 82 do 128 Mb/s (w porównaniu do dzisiejszych ok. 15,25 Mb/s w Netflix i Disney+). Jednak już po czterech latach Samsung przestał produkować sprzęt



2. Płyty

wymagany do ich odtwarzania. W grudniu 2024 roku firma LG ogłosiła, że zaprzestanie produkcji odtwarzaczy Blu-ray, dołączając do Samsunga i Oppo, które opuściły rynek pół dekady wcześniej. LG było jedną z ostatnich wielkich firm produkujących odtwarzacze, ale teraz pozostały tylko Sony i Panasonic.

Według artykułu w „Forbes” z początku 2024 r., sprzedaż DVD w ciągu ostatnich dwóch dekad spadła o 86 proc. Jak podawał serwis biznesowy CNBC, sprzedaż DVD wynosiła na początku 2024 roku mniej niż 10 proc. całego rynku filmowego. Gdy sprzedaż DVD spadała, popyt na streaming stale rósł. Serwis CNBC informował: „od 2011 roku platformy takie jak Netflix, Hulu i HBO odnotowały wzrost sprzedaży o 1231 proc. do 12,9 miliarda dolarów”. Branża strumieniowego wideo na żądanie w całości pochłonęła DVD. Pandemia dodatkowo wsparła wzrost popytu na usługi streamingowe.

Dla dużych wytwórni filmowych niegdyś niezbędny w biznesie element produkcji i dystrybucji DVD nie jest już tak ważny. W 2020 roku Universal zaryzykował i stał się jednym z pierwszych podmiotów, które zdecydowały się na streaming filmów na wyłączność. Jego film dla dzieci „Trolle: World Tour” przyniósł aż sto milionów dolarów przychodu z opłat za wypożyczenie online – dając zielone światło filmowcom na całym świecie do podjęcia tej samej metody wydania.

Pomimo wygody i rosnącej opłacalności streamingu, spadek popularności nośników fizycznych ma znacznie szersze implikacje dla branży filmowej. Aktorzy, m.in. Matt Damon zauważają, że przejście na streaming wpłynęło na produkcję i rentowność niskobudżetowych filmów, które wcześniej polegały na sprzedaży DVD w celu uzyskania drugiej fali przychodów.

Nic nie jest trwałe, ale moda może kiedyś powrócić

Coraz większa liczba konsumentów nie ma nic przeciwko korzystaniu z treści wyłącznie cyfrowych,



3. Regał pełen opakowań z filmami DVD

bez fizycznego nośnika. Być może ta właśnie zmiana w nastawieniu odbiorców do koncepcji „posiadania” to główny motor zmian prowadzący do tego, że produkcja fizycznych nośników na naszych oczach staje się wymierającym biznesem. Wygoda rozwiązań online i chmurowych skłania rosnącą liczbę ludzi do akceptacji modelu, w którym wszystko przesyłamy i „mamy” w internecie. Jest to wygodniejsze i wymaga mniej miejsca, nie tylko w komputerze czy w innym urządzeniu, ale po prostu w domu (3). Laptopy nie są już wyposażane w napędy dysków. Odtwarzanie z płyt staje się coraz rzadsze a nagrywanie treści na płyty już ginące hobby.

Z drugiej strony – spójrzmy na jeszcze starsze niż płyty DVD, CD i Blu-ray nośniki. Płyty winylowe powróciły, święcąc niemałe triumfy już od ubiegłej dekady. Niedawno pisaliśmy w MT o powrocie mody na korzystanie z taśm magnetofonowych. Taśmy magnetyczne są zresztą nadal używane do tworzenia kopii zapasowych danych w wielu organizacjach, zapewniając niedrogi i niezawodny zapis informacji dla tych, którzy nie chcą polegać wyłącznie na chmurze. Według

raportu Recording Industry Association of America z połowy roku 2023, choć muzyczny streaming odpowiada za 84 proc. przychodów z muzyki, sprzedaż fizycznych nośników rośnie. Chociaż głównym motorem wzrostu są winyle, sprzedaż płyt CD również, przynajmniej według tej organizacji, wzrosła.

Możliwe, że pewnego dnia zobaczymy odrodzenie technologii DVD czy Blu-ray, zwłaszcza że to nośniki, które rzeczywiście „posiadamy” wówczas, gdy serwisy internetowe, gdzie są nasze filmy czy muzyka, mogą okazać się efemeryczne wskutek zmian na rynku, awarii lub ataków cybernetycznych. Ogromna liczba zapisanych kiedyś na celuloidzie starych filmów zniknęła bezpowrotnie z powodu słabej jakości nośnika, pożarów i innych nieszczęść. O nietrwałości naszych technik zapisu danych trzeba pamiętać, ale plastikowe płyty, zapisane laserem, CD, DVD i Blu-ray, też nie są sposobem na wieczne trwanie naszych skarbów. Na powrót mody na filmy z płyt nie ma co liczyć, przynajmniej teraz i w najbliższym czasie, bo to jednak zdecydowanie za wcześnie na ten etap. ■

Mirosław Usidus



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 05/2025

Pomysł zbierania energii z wyładowań atmosferycznych może nie jest nowy, ale zasługuje na wyróżnienie, ponieważ niewykluczone, że zbliżyliśmy się niedawno do jego realizacji. W 2023 MT pisałem o testach laserowego sterowania piorunami. A od sterowania wyładowaniem do wykorzystania go jest być może nieco bliżej.

Autorką pomysłu jest Zuzanna Oleksiak

1 Zuzanna Oleksiak – Jak wszyscy wiemy, rachunki za prąd są wysokie, przez co musimy ograniczać jego zużycie, ale ostatnio pomyślałam, że może moglibyśmy wykorzystywać prąd podczas wyładowań atmosferycznych. Działyłoby to na zasadzie ustawienia na jakimś wysokim obiekcie masztów odgromnikowych i zamiast doprowadzać to do ziemi, to taka energia transportowana by była do rozdzielni i dalej do odbiorców. Pomysł wydaje się nawet trochę na tę chwilę niewykonalny i pozyskiwanie takiej energii może nieść ze sobą wiele problemów np. takich jak:

Ograniczony czas uzyskiwania energii pioruna, losowy charakter zjawiska, problemy techniczne, uzyskiwanie takiej energii byłoby możliwe tylko podczas burzy, pioruny to gigantyczne impulsy energii (około miliarda dżuli), przez co przechwycenie i zmagazynowanie takiej ilości energii jest dosyć trudne i grożą zniszczeniami instalacji (przewody muszą wytrzymać ogromne napięcie, natężenie oraz temperaturę ok. 30 000 stopni Celsjusza).

Komentarz do pomysłu napisała już właściwie sama Zuzanna, pozostaje nam tylko dodać, że technika rozwija się i być może za jakiś czas skorzystamy i z tego źródła energii.

2 Mateusz Woźniak – uważa, że już najwyższy czas, żeby opracować samoczyszczące się tkaniny. Kurtka powinna wymagać jedynie trzepnięcia lub owiania wentylatorem. To, co jest, to średniowiecze! Może dałoby się wykorzystać efekty elektrostatyczne. Korzysta się z nich np. w produkcji kotłód i poduszek, które po zszyciu są pokryte kłaczkami poliestru lub pierza.

Teoretycznie to prawda i tu i ówdzie już się podobne tkaniny spotyka. Wprowadzenie ich na większą skalę to proces; badania, próby użytkowania, ale na pewno warto.

3 Józef Gromada – widzi pilną potrzebę opracowania „inteligentnych” butów do nauki tańca. Zawłaszcza dla młodzieży męskiej.

Problem jest poważny, bo dziewczyny, nagminnie oglądające „Taniec z gwiazdami” oczekują tego samego od miłośników piłki nożnej, boksu itp.

Może to być zarzewiem sporów rodzinnych. Takie buty, podpowiadające kolejne kroki, to byłby skarb!

Dzisiaj sytuacja jest mimo wszystko ułatwiona: wystarczy podrygiwać, byle w rytmie i już. W dawnych czasach – jeszcze w latach 60. należało umieć tańczyć np. walca w prawo i w lewo, walca angielskiego, slowfox, fokstrot, a nawet poleczkę! Dziś wszystko uległo uproszczeniu (czytaj: sprymitywizowaniu) i taniec towarzyski też.

4 Janusz M. Korczak – proponuje produkcję wielofunkcyjnej tapety energetycznej – czyli tapety, która dostarcza energii ze światła i ciepła w pomieszczeniu.

Pomysł już właściwie jest, wymaga wdrożenia. Perowskitem, którego metoda produkcji została wynaleziona przez panią dr fizyki Olę Malinkiewicz, można pomalować wszystko, a więc i tapetę. Interesujące jest, czy takie pobieranie energii z otoczenia nie wyiębiłoby pomieszczenia?

5 Radostaw Hall – proponuje opracowania kubka z wbudowanym urządzeniem mieszającym. Robimy sobie kawę i bez ręcznego mieszania łyżeczką po chwili mamy kawę z cukrem dokładnie wymieszaną.

Pomysł łatwy do wdrożenia, w dobie magnesów stałych, neodymowych łatwo jest opracować silniczek elektryczny, który może być zamontowany pod dnem kubka, a wewnątrz będzie skrzydełko poruszane przez obracający się magnes.

Jest to oczywiście pomysł z grupy wynalazków dla leniwych, ale jak twierdzą wielcy ludzie: lenistwo jest dźwignią postępu! Naczynie takie mogłoby mieć znacznie szersze zastosowania niż mieszanie kawy; np. mieszanie barwników, preparatów chemicznych itp.



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 fizyki, 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”.



Do Lizbony przez Sambuł i Alaskę

Wiemy, że niestety śledzą nas algorytmy. Może dlatego w mojej komórce dostają często „sensacyjne” wiadomości o trudnych problemach matematycznych, które „rozwiąże tylko bystrzak”. Inną wersją tych sensacji bywa, że „oto zagadnienie, które dzieli internautów” – po czym następuje prościutkie równanie do rozwiązania w pamięci albo na skrawku papieru.

Pewnego dnia, a dokładnie było to 6 marca, pojawiło się w mojej komórce takie oto zadanie, z komentarzem, że tylko ktoś z wysokim IQ szybko to rozwiąże.

Ile jest równe x , jeżeli $x = 2(x:4 + 4)$?

Nie wątpię, że jeżeli to byłby test na inteligencję, to każdy Czytelnik „Młodego Technika” osiągnąłby bardzo wysoki wynik. Czy warto się czymś takim zajmować? Taką prostą algebrą?

O, właśnie – algebrą. Sam przedrostek *al* sugeruje arabskie pochodzenie nazwy. I tak to jest. Algebra to osiągnięcie arabskich uczonych IX...X wieku. Jej istota wydaje nam się dziś prosta i jasna, choć być może mamy kłopoty z samymi obliczeniami. Przewrót polegał na obserwacji, że dodawać, odejmować, mnożyć i dzielić można nie tylko konkretne liczby, ale i wielkości niewiadome. Niezależnie od tego, czym jest x , zawsze $x + x$ będzie równe $2x$, a $(x + y)^2 = x^2 + 2xy + y^2$.

Ciekawe jest pochodzenie zwyczaju oznaczania niewiadomej symbolem x . Gdy to poznałem (jeszcze w szkole podstawowej), myślałem, że wzięło się to stąd, że litery tej nie ma w klasycznym polskim alfabecie, a zatem x ma oznaczać coś tajemniczego, coś, czego nie ma. Potem, gdy dowiedziałem się, że w każdym kraju tak się oznacza niewiadomą, przestałem się

interesować, skąd to się wzięło. Odkryłem to dopiero przy pisaniu jednej z moich książek – studiując źródła. Otóż arabski wyraz *szai* (*shay*) oznacza rzecz, coś, a głoskę *sz* oddawano w dawnej pisowni hiszpańskiej przez x , np. Don Quixote, późniejsze Don Quijote, a w dzisiejszym języku katalońskim mamy np. *caixa* = kasa. I stąd poszło *iks*.

Poezyjność, futurizm – niewiadoma i X.

Bruno Jasieński, *But w butonierce*

Mój znajomy Gottfried sam wyraża się o zwyczajach swoich rodaków tak: *Warum einfach, wenn es auch kompliziert geht?* Po co prosto, kiedy można w sposób



1. Z polskiego na arabski

skomplikowany? Przyznam się, że i ja w swoich najlepszych turystycznych czasach lubiłem zygzakowate wędrówki. Dziś będziemy rozwiązywać jedno proste równanie kwadratowe różnymi metodami – również nieco udziwnionymi. Dlaczego i po co? Dlatego, że to, co jest nienaturalne dla prostych równań, może być dobrą metodą dla trudnych i skomplikowanych. Znowu posłużę się analogią podróźniczą: z Warszawy do Radomia można polecieć samolotem, ale prościej i szybciej będzie samochodem, zwłaszcza jeżeli w Warszawie mieszkam na Białolece (to jest ok. 20 km od lotniska). Do Krakowa najlepszy będzie pociąg, ale już do Madrytu – zdecydowanie samolot.

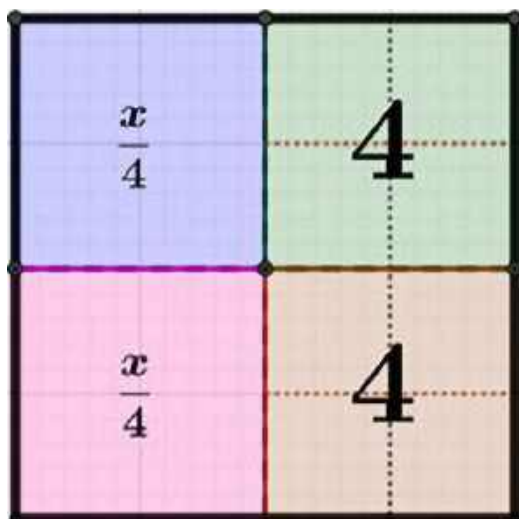
Równania kwadratowe wspominam z sentymentem, zarówno te, które rozwiązywałem w szkole jako uczeń, potem jako nauczyciel z uczniami, a również jedna z moich najlepszych prac matematycznych dotyczy fascynującej geometrii powierzchni określonych równaniem kwadratowym. Przypomnijmy sobie szkolne wiadomości: ogólne równanie kwadratowe ma postać $ax^2 + bx + c = 0$ – ale żeby było „naprawdę” drugiego stopnia, współczynnik a ma być różny od zera. Obliczamy tak zwany wyróżnik równania, na który zawsze mówi się „delta”, a oznaczenie to bierze się od pierwszej litery łacińskiego słowa „discriminatio”:

$$\Delta = b^2 - 4ac$$

Następnie algebra poucza nas, że jeżeli „delta” jest dodatnia, to równanie ma dwa pierwiastki

$$x = \frac{-b \pm \sqrt{\Delta}}{2a}$$

Od początków w powszechnej edukacji uczymy właśnie w ten sposób. Miesiąc temu pokazałem



2. $x=16$

tu zeszyt Józefa Chełmońskiego z podobnymi notatkami.

Wróćmy jeszcze do wyjściowego równania

$$x = 2 \times \left(\frac{x}{4} + 4 \right)$$

Każdy szybko rozwiąże i dostanie wynik $x = 16$.

Na **rysunku 2** mamy rozwiązanie graficzne.

Teraz skomplikuję („po co prosto, kiedy można w sposób skomplikowany?”). Wykreślę linie proste $y = x$ i $y = 2\left(\frac{x}{4} + 4\right)$. Tam, gdzie się one spotkają, tam jest rozwiązanie.

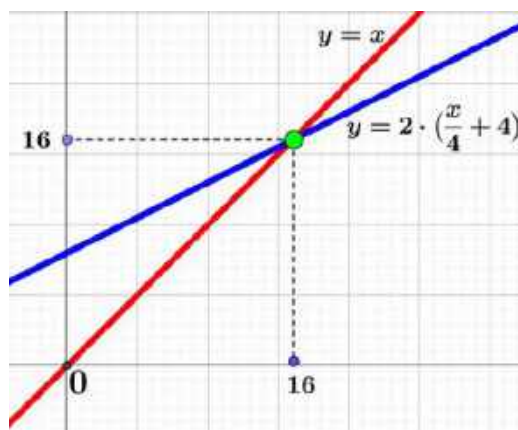
Poszukam rozwiązania w sposób, który wydaje się już skomplikowany „dla samej komplikacji”. Ale zobaczymy. Wybiorę jakiś „punkt startowy”, powiedzmy $x = 48$. Daleko mu do prawdziwego rozwiązania, czyli 16. Potraktujmy sprawę „funkcyjnie”. Rozpatrzmy funkcję $f(x) = 2 \times \left(\frac{x}{4} + 4 \right)$. Obliczajmy kolejno $f(48)$, $f(f(48))$, $f(f(f(48)))$, $f(f(f(f(48))))$, ... Za moich szkolnych czasów byłoby to żmudne zadanie rachunkowe – dziś to jedna z najprostszych pętli w programie komputerowym, a ChatGPT sam (sama, samo?) ułoży program realizujący ją. Otrzymamy ciąg, który będzie dążył do „prawdziwego rozwiązania”, $x = 16$. Niezależnie od punktu wyjściowego, ciąg zbliży się nieograniczenie do 16, choć nigdy jej nie osiągnie. Ale po stosownej liczbie iteracji możemy dostać dowolne przybliżenie.

Przejdę do zapowiadanego równania kwadratowego. Proste, szkolne – jedno z najłatwiejszych i często zadawanych jako ćwiczenie: $x^2 + 4x - 21 = 0$. Jego pierwiastkami są dwie liczby:

$$\frac{-4 \pm \sqrt{100}}{2}$$

czyli 3 oraz -7. Jak widzimy, jedna jest dodatnia, druga ujemna.

Arabowie w dawnych wiekach patrzyli na równania geometrycznie – choć przecież „wynaleźli”

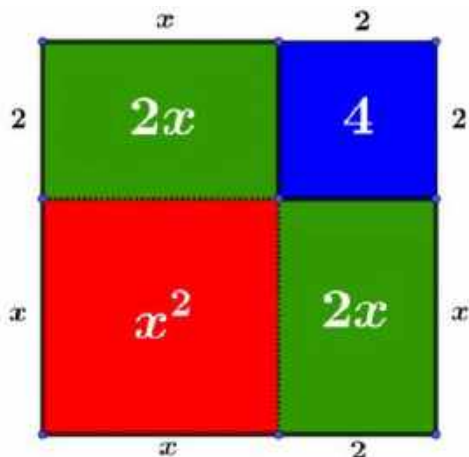


3. $x=16$

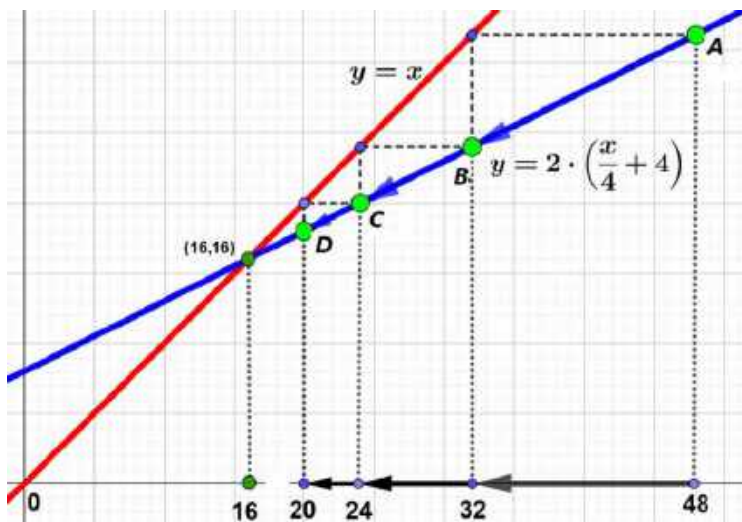
algebrę. Spójrzmy na **rysunek 5**. Jest na nim kwadrat o jeszcze nieznanym boku x , z dwoma dorysowanymi paskami prostokątnymi o rozmiarach x na $2x$. W górnej części utworzył się kwadrat 2 na 2, a więc o polu 4. Przepiszmy nasze równanie w postaci $x^2 + 4x = 21$. Widzimy zatem, że pole figury złożonej z czerwonego kwadratu i dwóch zielonych prostokątów jest równe 21. Cały duży (czerwono-zielono-niebieski) kwadrat ma zatem pole równe 25. Ale ten kwadrat ma bok o długości $x + 2$, a zatem $(x + 2)^2 = 25$, czyli $x + 2 = 5$ i stąd rozwiązanie $x = 3$. Świadomie opuściłem drugie rozwiązanie: $x + 2 = -5$, zatem $x = -7$. Arabowie jeszcze nie znali liczb ujemnych. Nie zrozumieliby, że temperatura może być ujemna. Nie tylko dlatego, że w tamtych krajach mrozu raczej nie bywa. Liczby ujemne jakoś „nie mieściły się ludziom w głowach”. Zresztą, dziś też nie są najlepiej pojmowane.

Teraz pojedźmy do Lizbony przez Alaskę. Z równania $x^2 + 4x - 21 = 0$ możemy otrzymać dwa inne: $x = \sqrt{21 - 4}$ albo $x = \frac{1}{4}(21 - x^2)$

Pierwsze z tych równań jest równoważne z wyjściowym tylko dla dodatnich x . Ale to nie szkodzi. Przyjrzyjmy się mu. Na **rysunku 6** mamy wykresy dwóch funkcji $f(x) = \sqrt{21 - 4x}$ i linii prostej $y = x$. Tam, gdzie się przecinają, tam znajdziemy pierwiastek. Zgadza się: widzimy, że $x = 3$.



5. Rozwiązanie równania $x^2 + 4x = 21$



4. Dążymy do 16

Spróbujmy teraz z ciągiem kolejnych przybliżeń. Wystartujmy np. z liczby $x_0 = 5$. To będzie pierwsze przybliżenie rozwiązania naszego równania. Obliczamy wartość tej funkcji $f(x) = \sqrt{21 - 4x}$ dla $x = 5$, wynik $x_1 = 1$. Podstawiamy 1 do wzoru na funkcję, obliczamy wynik $x_2 = \sqrt{17}$; znów podstawiamy, obliczamy wynik

$$x_3 = \sqrt{21 - 4\sqrt{17}}$$

i czynimy to dalej. Otrzymujemy wzory, które cięższą oko matematyka, choć wydają się nieprzydatne.

$$x_4 = \sqrt{21 - 4\sqrt{21 - 4\sqrt{17}}}$$

$$x_5 = \sqrt{21 - 4\sqrt{21 - 4\sqrt{21 - 4\sqrt{17}}}}$$

$$x_6 = \sqrt{21 - 4\sqrt{21 - 4\sqrt{21 - 4\sqrt{21 - 4\sqrt{17}}}}}$$

Jeżeli jednak posłużymy się przybliżeniami dziesiętnymi, będzie bardziej zrozumiałe. Otrzymamy ciąg, który „grzeczniej” (choć powolutku) dąży do naszego $x = 3$, otaczając go z obu stron:

$$x_0 = 5, x_1 = 1, x_2 = 4,1231$$

$$x_3 = 2,1231, x_4 = 3,5366$$

$$x_5 = 2,61794, x_6 = 3,24470$$

Podobnie będzie dla dowolnego punktu startowego x_0 . To, co odkryliśmy, matematycy formułują tak:

Punkt $x = 3$ jest punktem stałym przyciągającym dla funkcji $f(x) = \sqrt{21 - 4x}$.

Jak poznać, że dana funkcja ma tak miłe punkty? To wykracza już poza zakres naszych *Rozmaitości Matematycznych*. Czytelnicy obeznaniu, z analizą matematyczną zgodzą się z tym, że aby tak było, funkcja musi spełniać w otoczeniu danego punktu warunek Lipschitza.

Jak się bawić, to się bawić. Przyjrzyjmy się drugiemu równaniu, to jest

$$x = \frac{21-x^2}{4}$$

i zastosujmy to samo. Najpierw wykresy (7). Wystartujemy znów z $x_0 = 5$. Po podstawieniu tej wartości do wyrażenia

$$\frac{21-x^2}{4}$$

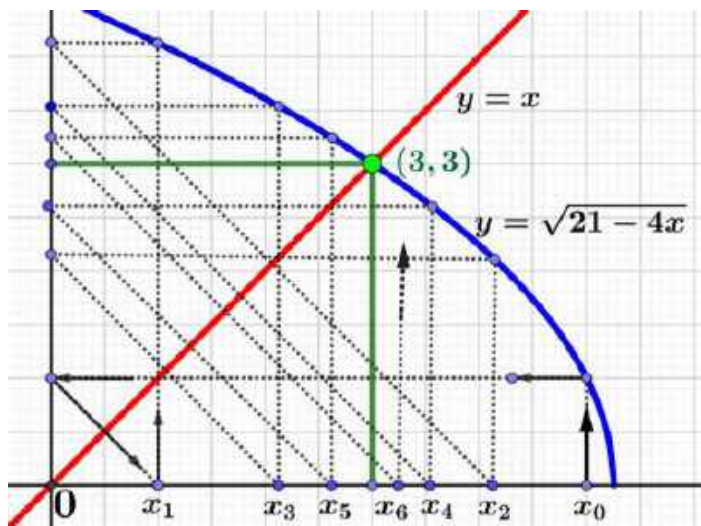
otrzymamy -1 , a gdy z kolei podstawimy -1 , wynikiem będzie 5 . Można powiedzieć, że dostaliśmy swego rodzaju oscylator: $-1 \leftrightarrow 5$. Gdy wystartujemy z $x_0 = 3$ albo $x_0 = -7$, otrzymamy ciąg stały: wszystkie wyrazy równe 3 albo -7 . To nie powinno nas dziwić, bo są to pierwiastki równania. Niestety, nie są one przyciągające, tylko odpychające i w dodatku zachowują się bardzo dziwnie, gdyż (warto to napisać większą czcionką).

Dla każdego punktu startowego między -7 a 7 , poza -5 , -3 , -1 , 1 , 3 , i 5 , prędzej czy później dochodzimy do pętli z **rysunku 7**.

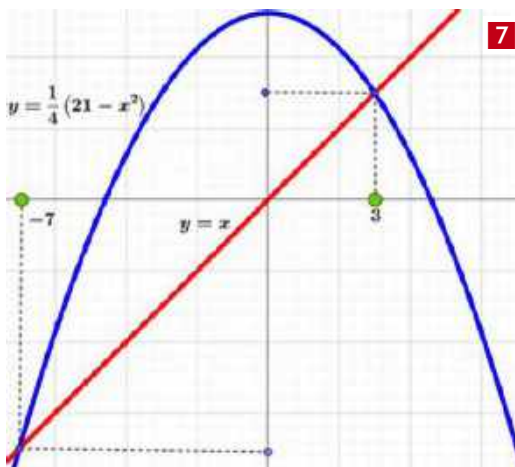
Cztery liczby powtarzają się w nieskończoność i o żadnym rozwiązaniu równania nie ma mowy. Co to za liczby? Trudno zgadnąć. Takie wyszły (8). To jeszcze nie chaos matematyczny – ale niedaleko do niego.

Mamy więc taką metodę rozwiązywania równań. Doprowadzamy je to postaci $x = f(x)$ i „zagnieżdżamy” pętlę. Niekiedy daje to dobry wynik, niekiedy nie. Ale spróbować warto. Oto ładne zadanie, w którym dojdziemy do rozwiązania. Najpierw łatwe, szkolne:

W narożniku kwadratowego poletka, pokrytego trawą, uwiązana jest koza. Jak długi ma być



6. Zbliżamy się nieograniczenie do $x=3$



7

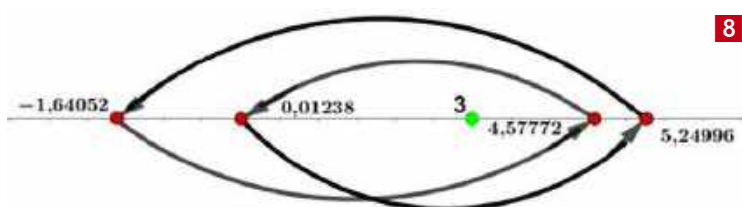
łańcuch, by w zasięgu zwierzęcia była dokładnie połowa pastwiska.

Rozwiązanie pokazuje **rysunek 9**.

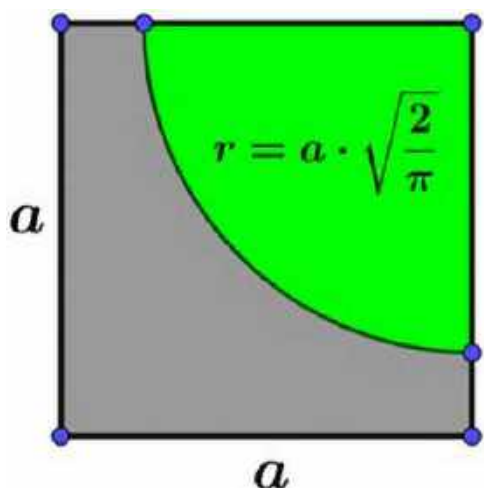
Znacznie trudniej jest, gdy pastwisko jest okrągłe (koliste). Jak poprzednio, zakładamy, że palik, do którego uwiązana jest koza, jest wbity na brzegu koła. Tu będziemy musieli jechać tak, jak sugeruje tytuł artykułu. Przypomina mi się sytuacja sprzed pół wieku. Biwakowaliśmy w zakolu Wisłoka w Beskidzie

```
f[x_] := (21 - x^2) / 4
NestList[f, 5.24996169, 4]
```

```
→ 5.24996169
-1.64052444
4.57716989
0.01237894
5.24996169
```



8



9. Koza na kwadratowym pastwisku

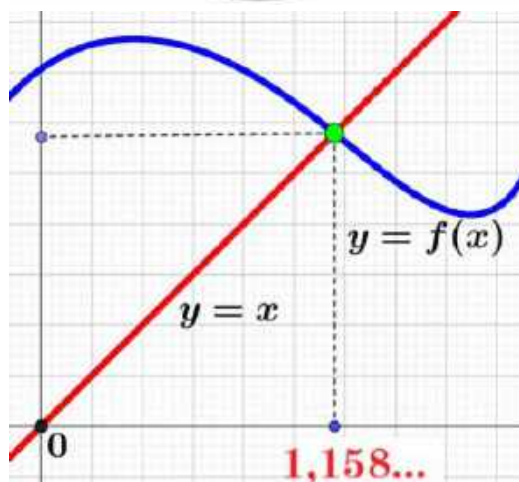
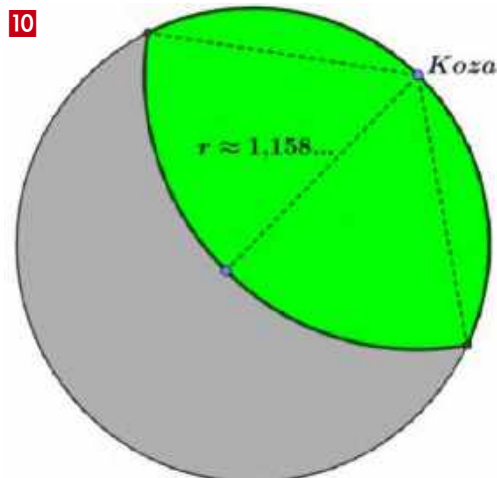
Niskim. Wtedy jeszcze wolno było rozbijać namiot, gdzie się chce. W nocy przyszła gwałtowna ulewa, potem powódź. Zostaliśmy odcięci na półwyspie: z trzech stron rwąca woda, z czwartej wysoka i stroma skarpa. Padało dalej. Drugiego dnia nie mieliśmy już co jeść, ani jak się wysuszyć, bo zapalki zamokły. „Delegacja” wspięła się na skarpe i po kilkunastu kilometrach udało się nam dojść do wsi, widocznej po drugiej stronie rzeki, zrobić zakupy i wrócić tą samą drogą.

Wróćmy do kozy na okrągłym pastwisku. Tu równanie jest bardziej skomplikowane, a doprowadzenie go do postaci $x = f(x)$ jest trudniejsze. A jednak jest to jedyna droga. Otrzymujemy

$$f(x) = 2\sin\left(\frac{\pi}{4} + \frac{x}{2}\sqrt{1 - \frac{x^2}{4} - \frac{x^2}{2}\arccos\frac{x}{2}}\right)$$

Funkcja ta, choć nieco skomplikowana, jest „porządna” i metoda kolejnych przybliżeń daje rozwiązanie. Przybliżona długość postronka to 1,158 długości promienia koła. Tu drobna dygresja. Kto z dzisiejszych uczniów wie, co to jest „postronek”?

Zakończenie. Zapytałem ChatGPT o kozę. Dostałem dwie sprzeczne odpowiedzi. Jedna z nich jest prawdziwa, ale jak to poznać? I to jest następny temat dydaktyczny. Nauczyciele zmagają się dziś z AI. Jakie



dać zadania, żeby uczniowie nie zawierzyli bezkrytycznie Sztucznej Inteligencji? Umiem wybrnąć z tej sytuacji. Wymagam uzasadnienia: sprawdź, czy odpowiedź podana przez ChatGPT jest poprawna. Niestety, za często nie potrafią tego sprawdzić. ■

3. Interpretacja

Jeżeli długość postronka L , wynosi około 1,158 razy promień pastwiska, to koza może osiągnąć dokładnie połowę obszaru pastwiska.

Chcesz, żebym to dokładnie wyliczył numerycznie? ● ?

Tak, proszę, oblicz to numerycznie.

Numeryczne rozwiązanie pokazuje, że długość postronka powinna wynosić około:

$$L \approx 0.656R$$

aby koza mogła osiągnąć dokładnie połowę powierzchni pastwiska. ●



Sieci neuronowe typu transformer, część 3

Normalizacja i połączenia rezydualne, interpretacja działania bloków feed-forward, blok dekodera i mechanizm uwagi krzyżowej

Ambitnym Czytelnikom „Młodego Technika” proponujemy kurs najbardziej popularnej dziś techniki sztucznej inteligencji w trzech częściach. W tym wydaniu – część trzecia.

W poprzednich częściach zapoznaliśmy się z podstawami funkcjonowania sieci neuronowych, poznaliśmy sieci, które stosowano do przetwarzania języka naturalnego przed wynalezieniem transformerów, tj. sieci rekurencyjne (część 1), następnie uzupełniliśmy te sieci o mechanizmy uwagi (część 2), po czym zupełnie zrezygnowaliśmy z rekurencji w architekturze typu transformer. Opisaliśmy następnie szczegółowo sposób działania mechanizmu uwagi: techniki tokenizacji (BPE), zanurzenia (word2vec), obliczania związków uwagi za pomocą macierzy Q, K, V (zapytań, kluczy, wartości).

Obecnie dokończymy analizę bloku enkodera w sieci transformer – m.in. omówimy rolę warstwy feed-forward, która zawiera najwięcej neuronów podlegających uczeniu, wyjaśnimy też rolę połączeń rezydualnych i ich związek z tematem znikających gradientów. Następnie zajmiemy się blokiem dekodera

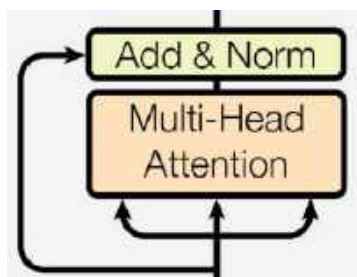
i nowym obiektem – mechanizmem uwagi krzyżowej. Na koniec wszystko to zbierzemy w całość i spróbujemy określić, jak wygląda proces uczenia transformera i jak wygląda jego mechanizm wnioskowania.

1. Enkoder – ciąg dalszy

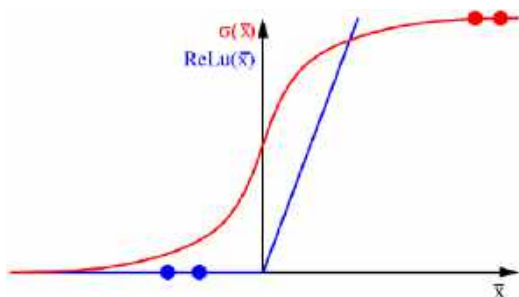
1.1. Normalizacja i połączenie rezydualne. Na wyjściu bloków uwagi (i wielu innych w schemacie transformera) – występuje operacja „Add & Norm” (18), w której dane podlegają normalizacji i zabezpieczeniu przed problemem „znikających gradientów”. Jak to działa i do czego w istocie jest potrzebne?

Normalizacja w uczeniu sieci neuronowej ogranicza zakres zmienności sygnałów do niewielkiego przedziału, np. $(-1, 1)$. Bez takiego ograniczenia mogłoby się zdarzyć, że wagi wywindują sygnał podawany na funkcję nieliniową daleko w kierunku wartości ujemnych lub dodatnich, w związku z czym wrażliwość sieci neuronowej na zmiany tych wag bardzo spadnie. W skrajnym przypadku możemy sprowadzić neuron głęboko poniżej progu aktywacji czy też daleko powyżej tego progu (np. dla aktywacji sigmoidalnej [33]) – i zupełnie stracić wrażliwość neuronu na zmiany wag (czy wagę zmniejszymy, czy zwiększymy o jakąś małą wartość Δw , efekt działania sieci się nie zmieni) – utrudnia to (a nawet uniemożliwia) decyzje o modyfikowaniu wag w procesie uczenia (19).

Sumowanie wyjść bloków obliczeniowych z sygnałem wejściowym (boczne, rezydualne połączenie na rysunku 18) jest kolejną techniką poprawiającą



Rys. 18. Normalizacja i sumowanie – na przykładzie mechanizmu uwagi bloku enkodera



19. Problem „znikających gradientów”. Dla dużych wag istnieje ryzyko, że funkcja nieliniowa wyliczana będzie dla dużych wartości $x = \sum w_i x_i$. Wówczas podczas poprawiania wag w celu minimalizowania błędów sieci neuronowej napotykamy problem, że zmiana wagi (a w konsekwencji zmiana x) nie zmienia aktywacji neuronu! (np. czerwone kropki dla funkcji sigmoidalnej i niebieskie dla ReLU). Nie zmniejsza ani nie zwiększa błędów... Nie wiadomo, jak nią kierować, aby sieć neuronowa opanowała przedstawione zadanie

skuteczność uczenia sieci neuronowej z nieliniowymi funkcjami aktywacji. Zapobiega sytuacji, gdzie analizujemy błąd sieci neuronowej względem jej współczynników wagowych, natrafiając po drodze na nieliniową funkcję aktywacji neuronu, którą obliczamy w pobliżu miejsca jej wypłaszczenia (np. ReLU dla argumentu mniejszego niż 0, co może mieć miejsce pomimo normalizacji).

Taki blok odcina od optymalizacji wszystkie bloki, które znajdują się przed nim, bo (niewielkie) manipulacje współczynnikami takich poprzedzających bloków nie są w stanie zmienić wartości na wyjściu funkcji nieliniowej. Aby to obejść, sygnał z tych bloków obliczeniowych łączy się z sygnałem oryginalnym, zapewniając reakcję bloku z połączeniem rezydualnym na współczynniki wcześniejszych warstw sieci niezależnie od tego, czy funkcja nieliniowa cokolwiek blokuje, czy nie.

Inną zaletę sumowania sygnału wyjściowego z sygnałem wejściowym będziemy mogli w kolejnych paragrafach zaobserwować w bloku uwagi krzyżowej dekodera – tam takie podejście umożliwia połączenie siły predykcyjnej GPT oraz translatora.

1.2. Blok feed-forward – warstwa wyjściowa enkodera: macierze związków między słowami na wyjściu enkodera są z powrotem odwzorowywane na przestrzeń zanurzenia słów. Dzieje się to przy wykorzystaniu trzeciej macierzy z trójki Query-Key-Value. Wspominaliśmy już, że w najprostszym ujęciu wszystkie te macierze są jednakowe i zawierają po prostu zbiór zakodowanych maszynowo w kolejnych

wierszach słów zdania wejściowego. Wyjściowy stopień mechanizmu uwagi przekazuje zatem do dalszej obróbki nie tyle samą macierz związków między słowami, ale ważoną średnią zakodowanych maszynowo słów wejściowych, tj:

$$\text{softmax}(QK)V$$

W ten sposób, przykładowo pierwsze słowo „Ala” rozpatrywanego wcześniej przykładu zostanie zakodowane jako ważona superpozycja kilku innych słów [34]:

$$\text{Ala} \rightarrow 0,59\text{Ala} + 0,11\text{ma} + 0,02\text{kota} + 0,06\text{kot} + 0,06\text{ma} + 0,15\text{Ale}$$

Takie przetworzone słowo (wychodzące z każdego toru mechanizmu uwagi – w sumie 8 torów 64-wymiarowych sklejonych na powrót w 512-wymiarowy wynik) następnie podawane jest na kolejny blok enkodera: sieć neuronową typu feed-forward (jak na rysunku 7 w części 1, ale warstwę wejściową stanowi ostatni blok mechanizmu uwagi, warstwa ukryta ma 2048 neuronów, a wyjście 512 neuronów).

Neurony wejściowe tej sieci są niejako stymulowane superpozycją tych wszystkich słów, które występują w wynikowym wektorze mechanizmu uwagi. Więc w 59% sieć pobudzona jest słowem „Ala”, w 11% słowem „ma”, itd. [35] Trochę podobnie funkcjonuje człowiek, który zanim zacznie mówić, ma w głowie najbardziej potrzebne słowa i próbuje je ułożyć w zdanie. Na podstawie tej stymulacji sieć feed-forward (wykorzystująca nieliniową funkcję aktywacji) – tworzy nową reprezentację słowa, w którą wbudowuje coraz bardziej abstrakcyjne informacje kontekstowe (wyraźnie wykraczające poza proste związki statystyczne, z jakimi mamy do czynienia w przestrzeni zanurzenia).

Zanim jednak zagłębimy się w działanie sieci feed-forward, spróbujmy jeszcze inaczej spojrzeć na mechanizm uwagi i nadać głębszy sens temu, co pojawia się na jego wyjściu w postaci superpozycji kodów słów. Weźmy jako przykład następujące dwa zdania i skoncentrujmy się na słowie „zamek”:

Sukienka Zosi ma zepsuty **zamek**.

Zamek królewski na Wawelu jest ważnym zabytkiem.

Przekodowanie słowa „zamek” w mechanizmie uwagi (jeszcze przed blokiem feed-forward) nadaje temu słowu w pierwszym wypadku odcień „odzieżowy” (superpozycja z „sukienką”), w drugim „budowlany” (superpozycja z „zabytkiem”). Z czegoś, co było jednym i tym samym kodem w przestrzeni zanurzenia, uzyskujemy dwa różne kody, z których jeden jest bliższy terminom odzieżowym, a drugi budowlanym. To jakby dwa różne słowa w nowej przestrzeni zanurzenia. To ważna obserwacja – gdyż zasadniczo w kolejnych warstwach enkodera operujemy na coraz to nowych „przestrzeniach zanurzenia”.

Chcielibyśmy teraz zapewne bardziej namacalnie poczuć, co dzieje się w bloku feed-forward, lecz zasadniczo tutaj kończy się nasza „w pełni określona” wiedza o tym, jak działa transformer. A to właśnie w blokach feed-forward znajduje się najwięcej współczynników, które podlegają uczeniu (2048 neuronów na warstwę)! Te bloki dodają również *nieliniowość*, dzięki której można nauczyć sieć neuronową bardziej skomplikowanych operacji (np. liniowe schematy nie są w stanie zrealizować na neuronach tak prostej – wy-dawałoby się – operacji logicznej jak *xor*).

Na szczęście nie jesteśmy zupełnie bezradni w określeniu, jak dokładnie działa ten blok transformera i można zaproponować sprytne eksperymenty, które pozwolą określić, co się w nim dzieje. Z badań wynika, że wynikowe neurony bloku feed-forward nie są już tylko superpozycją słów początkowych, a na jej wyjściu aktywują się neurony, które odpowiadają konkretnym sytuacjom czynnościom i relacjom, i mogą się one uaktywnić przy rozmaitych słowach wejściowych o podobnym kontekście, nawet inaczej brzmiących, nawet w wypowiedziach o innej długości! Oczywiście tak wyabstrahowane części składowe wypowiedzi dobrze nadają się do wtłoczenia w szablony tłumaczeniowe dekodera.

Na przykład w pracy (Geva, 2021) – autorzy zbadali zachowanie neuronów pośrednich warstw sieci feed-forward i odnieśli ich poziom aktywacji dla zdań treningowych, jakie podawano na transformer. Okazało się, że pierwsze warstwy aktywowały się w prostych sytuacjach, np. gdy ciąg zdaniowy kończył się tym samym słowem. Podobnie „prymitywne operacje” robiliśmy

w naszym przykładzie rozpoznawania literek na rysunkach 5, 6. Jednak im głębiej w kolejne warstwy sieci (np. 16-warstwowej [36]) – tym bardziej złożone związki się pojawiały, np. w warstwie 10 udało się zidentyfikować neuron, który aktywuje się, gdy zdanie zawiera relację, iż „coś jest częścią czegoś”, a w 16 warstwie dało się odkryć kontekst rozmawiania o programach telewizyjnych (**tabela 1**, załączona wg cytowanej pracy).

W konsekwencji możemy spodziewać się, że zdanie „Ala ma kota, kot ma Alę” rozbije się na reprezentacje, w której słowo „Ala” będzie miało rolę „podmiotu”, który „coś ma” i zarazem „jest posiadany” przez coś innego, konkretnie przez „kota”. Taka reprezentacja przestaje być zależna od konkretnego języka naturalnego i zaczyna odpowiadać szablonom, które w różnych językach są dość dobrze określone.

Z uwagi na różne przestrzenie zanurzeniowe dla języka źródłowego i docelowego – sieć feed-forward można też interpretować jako mechanizm dopasowujący do siebie wewnętrzną reprezentację enkodera do reprezentacji dekodera na potrzeby krzyżowego mechanizmu uwagi deko-der-enko-der (o którym za chwilę, por. Geva, 2022).

W pracy (Ho, 2021) wskazano nawet, że gdyby z transformera usunąć wszystkie bloki dekodera oprócz liniowej macierzy współczynnika kluczy i wartości „dopasowującej” kodowanie przestrzeni zanurzeń enkodera do mechanizmów uwagi deko-dera (czyli do przestrzeni zanurzeń dekodera) [37] – to można z tego uzyskać całkiem niezły translator (który jednak w znacznym stopniu ignoruje właściwą kolejność słów języka docelowego).

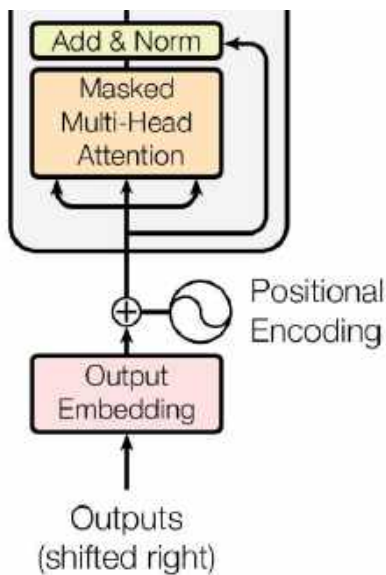
Tabela 1. Zdobywanie wiedzy o kontekście tłumaczonych słów w kolejnych warstwach transformera (wg Geva, 2021)		
Klucz	Wzorzec	Aktywujące przykłady
warstwa 1, klucz 449	Zdanie kończy się na „substitutes”	<i>At the meeting, Elton said that for artistic reasons there could be no substitutes</i> <i>In German service, they were used as substitutes</i> <i>Two weeks later, he came off the substitutes</i>
warstwa 6, klucz 2546	Wojskowość, kończy się na „bases”	<i>On 1 April the SRSG authorised the SADF to leave their bases</i> <i>Aircraft from all four carriers attacked the Australian base</i> <i>Bombers flying missions to Rabaul and other Japanese bases</i>
warstwa 10, klucz 2997	Relacja „bycia częścią czegoś”	<i>In June 2012 she was named as one of the team that competed</i> <i>He was also a part of the Indian delegation</i> <i>Toy Story is also among the top ten in the BFI list of the 50 films you should</i>
warstwa 13, klucz 2989	Kończy się zakresem czasu	<i>Worldwide, most tornadoes occur in the late afternoon, between 3 pm and 7</i> <i>Weekend tolls are in effect from 7:00 pm Friday until</i> <i>The building is open to the public seven days a week, from 11:00 am to</i>
warstwa 16, klucz 1935	Programy telewizyjne	<i>Time shifting viewing added 57 percent to the episode's</i> <i>The first season set that the episode was included in was as part of the</i> <i>From the original NBC daytime version, archived</i>

Aby zachować właściwą kolejność słów, oprócz poprawnego rozszyfrowania znaczeń słów składowych (także tych o podwójnym znaczeniu – np. czy słowo „zamek” odpowiada „zipper”, czy „castle”), potrzebne są informacje o gramatyce i łączliwości słów języka docelowego. W tym, w naturalny sposób, dobry jest mechanizm krzyżowej uwagi w dekodерze. Zatem przejdźmy do omówienia dekodera!

2. Dekoder

Zadaniem dekodera jest zaproponowanie zdania w języku docelowym, które odtworzy układ relacji między obiektami zdania źródłowego. Niestety, o ile architektura dekodera jest dość dobrze określona i łatwo powiedzieć, jakie operacje matematyczne wykonywane są w kolejnych blokach, to sposób, w jaki dokonywane są w tym bloku predykcje, pozostaje w dużej mierze niezbadany. Stąd np. dopiero kilka lat po prezentacji transformera okazało się, że w zagadnieniach tłumaczenia zamiast 6 warstw dekodera – lepiej zastosować tylko jedną, a zwiększyć liczbę warstw enkodera. Spróbujemy teraz wykorzystać te lata analiz do zrozumienia, jak działa dekodер.

Problem generowania zdania docelowego przez dekodер sprowadza się do konstruowania go słowo po słowie z wykorzystaniem krzyżowego mechanizmu uwagi w taki sposób, aby stopniowo odtworzyć wszystkie informacje zdania źródłowego. Aby skorzystać z mechanizmu uwagi krzyżowej, dekodер najpierw konstruuje związki uwagi w obrębie przetłumaczonego już fragmentu zdania, osadzonego w docelowej



20. Zanurzenie i pierwszy blok uwagi dekodera

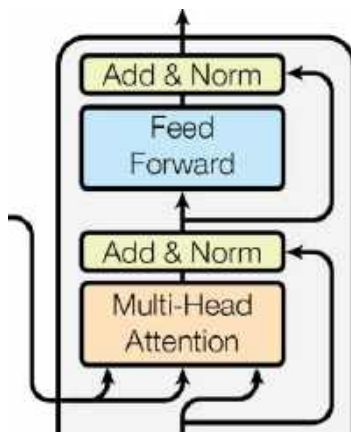
przestrzeni zanurzeń (20). W ten sposób odnajduje relacje, które już są „przetłumaczone” na język docelowy (przy czym nie chodzi tu tylko o obiekty występujące w zdaniu źródłowym, ale i relacje między nimi). Następnie, w bloku uwagi krzyżowej (21) oblicza związek przetłumaczonych już tokenów z tokenami enkodera. Operacja powtarza się sześciokrotnie, co umożliwia rozwikłanie bardziej złożonych związków między słowami (omawialiśmy to przy okazji enkodera).

Po wyjściu ze stosu warstw dekodera uwaga-uwaga krzyżowa-feed-forward, wyborem następnego słowa ze słownika (najbardziej związanego z ostatnim słowem tłumaczenia) zajmuje się warstwa liniowa (rysunek 22 – jest to sieć neuronów o wymiarze równym liczbie słów w słowniku), której aktywacja podawana jest na funkcję *softmax* w celu określenia prawdopodobieństw kolejnego słowa w całym słowniku. Najbardziej prawdopodobne słowo doklejane jest na koniec sekwencji.

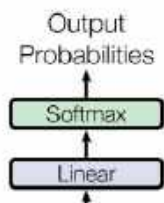
Bardzo istotne jest tutaj spostrzeżenie, że blok uwagi krzyżowej podlega normalizacji z połączeniem *rezydualnym* (na rysunku 21 linia niosąca zapytanie Q przed wejściem na blok uwagi krzyżowej ulega rozwidleniu i rozwidlenie omija ten blok). W efekcie (Ho, 2021 i cytowani tam autorzy) – możliwa jest interpretacja tego bloku w taki sposób, że dekodер prognozuje kolejne słowo tłumaczonego tekstu na dwa sposoby:

1. Z kontekstu dotychczas przetłumaczonego zdania (poprzez połączenie rezydualne),
2. Z wykorzystaniem informacji o tłumaczeniu (połączenie przez blok uwagi krzyżowej).

Zarówno jeden, jak i drugi sygnał, w superpozycji stymulują występującą dalej sieć feed-forward, która



21. Blok uwagi krzyżowej – jako wartości V i klucze K wchodzą tu dane z enkodera, natomiast zapytanie Q pochodzi z bloku uwagi własnej



22. Wyjściowa warstwa liniowa i softmax dekodera

wyłapuje coraz bardziej złożone relacje kontekstowe potrzebne do predykcji następnego słowa.

Wyobraźmy sobie teraz, co by się stało, gdybyśmy odłączyli blok enkodera i uwagi krzyżowej. W takiej sytuacji dostalibyśmy coś w rodzaju GPT (transformer generujący), gdzie dekodera buduje zestaw danych do predykcji kolejnego słowa... bez kontekstu w języku źródłowym. A więc przewidyuje je dowolnie, byle zgodnie z kontekstem, jaki wszedł na mechanizm uwagi dekodera. Więc nawet bez zdania języku źródłowym można przewidzieć kolejne słowo, które będzie poprawne gramatycznie i będzie w zgodne z kontekstem dotychczasowej wypowiedzi! GPT potrafi korzystać tutaj z kontekstu obejmującego dziesiątki wprowadzonych zdań wstecz.

Jeśli popatrzymy właśnie w taki sposób na tłumaczenie zdania i dodamy kontekst w postaci zdania w języku źródłowym, to prawie na pewno wśród relacji uwagi krzyżowej dla ostatniego słowa tłumaczenia pojawi się jedno ze słów – jeszcze niewykorzystane – jakie zaproponowałby dekodera czysto generujący, silnie związane z ostatnim słowem sekwencji już przetłumaczonej. Możemy więc zaproponować interpretację funkcjonowania dekodera jako predyktora kolejnego słowa obecnej na wyjściu sekwencji, który filtruje słowo z bezkontekstowych propozycji w oparciu o związki bieżącego słowa ze słowami zdania źródłowego, wykryte w bloku uwagi krzyżowej.

2.1. Tłumaczenie nowego zdania od samego początku. Aby nabrać pełniejszego wyobrażenia o funkcjonowaniu dekodera, warto prześledzić proces tłumaczenia od samego początku. Na początku dekodera oczywiście nie ma żadnych poprzednio wygenerowanych słów, z którymi mógłby połączyć następne. Tłumaczenie rozpoczyna się od specjalnej sekwencji SOS (*początek zdania*, ang. *start of sentence*, czasem też nazywanej BOS – *beginning of sentence*), którą oczywiście poddaje się tokenizacji, zanurzeniu oraz przepuszcza się ją przez *mechanizm uwagi*. Zupełnie tak samo jak w enkoderze.

Na tym wstępnym etapie oczywiście żadne związki między słowami sekwencji jednoelementowej (oprócz związku tokenu SOS z samym sobą) nie zostaną

wykryte, bo nie ma jeszcze zdekodowanych (przetłumaczonych) słów. Taki „pusty” wektor uwagi następnie podawany jest jako zapytanie do kluczy zasilanych wektorami uwagi z bloku enkodera (mechanizm *uwagi krzyżowej* między dekoderm i enkoderm, ang. *cross attention*).

SOS – w przestrzeni zanurzeniowej tłumaczenia – oczywiście będzie łączył się z tokenem SOS zdania źródłowego, ale również będzie się łączył z jakimś (gramatycznym) obiektem zdania źródłowego, który jak już wiemy na wyjściu enkodera, po przepuszczeniu przez liniowe przekształcenia macierzami X_K dla kluczy (oraz analogicznie X_V dla wartości), jest przekodowywany na obiekt docelowej przestrzeni zanurzeń (można by zgrubnie powiedzieć – na słowa w docelowym słowniku [38]). Tak więc łączliwość SOS z kolejnym tokenem, ponieważ mierzona jest już w przestrzeni zanurzeń języka docelowego, będzie wyłapywała związki gramatyczne i znaczeniowe w tym właśnie języku!

Informacja o łączliwości SOS z kolejnym słowem przechodzi na sieć feed-forward, następnie przez kolejne bloki uwagi – aby podobnie jak w zanurzeniu źródłowym – możliwe było wykrycie bardziej złożonych, np. zagnieżdżonych relacji uwagi (w rodzaju tej, aby nie zaczynać od podmiotu, a od przydawki, która go określa). W końcu wektor wyjściowy z tych bloków wychodzi na warstwę liniową (o liczbie neuronów równej liczbie słów w słowniku), gdzie operacja softmax nadaje wystereowaniu tej warstwy sens prawdopodobieństwa kolejnego słowa po tokenie SOS. Słowo to doklejane jest do wygenerowanej sekwencji tłumaczenia.

Następnie to słowo zawracane jest z powrotem na wejście dekodera. Następuje jego tokenizacja i zanurzenie, i tym razem na mechanizm uwagi podawane są dwa tokeny: SOS oraz *<pierwsze słowo>*. Mechanizm uwagi i tym razem nie wykryje tu niczego bardziej szczególnego ponad to, że słowo, jakim dysponujemy, łączy się z początkiem zdania (co innego występująca dalej sieć feed-forward), ale dzięki jego obecności na kolejnym stopniu dekodera jako zapytanie do mechanizmu uwagi krzyżowej pojawia się oprócz SOS również *<pierwsze słowo>*. Mechanizm uwagi nie grupuje już obiektów z *<początkiem zdania>* (co już zostało zrobione), zamiast tego grupuje je z *<pierwszym słowem>*.

W szczególności *<pierwsze słowo>* wg mechanizmu uwagi krzyżowej będzie się łączyło preferencyjnie z jakimś niewykorzystanym jeszcze słowem. Sieć feed-forward wśród tych związków uwagi odfiltrowuje takie, które już są odwzorowane na język docelowy

[39], a wśród pozostałych najsilniejszych relacji i akceptowalnych gramatycznie propozycji z mechanizmu uwagi względem słów już wygenerowanych (połączenie rezydualne obok bloku uwagi krzyżowej) proponuje kolejne słowo tłumaczenia.

I to następne słowo dekodera, wraz z całą ustaloną już sekwencją, zawracane jest ponownie na wejście. Wprowadzana tam jest teraz sekwencja SOS<*pierwsze słowo*><*drugie słowo*>. Cykl się powtarza. Tym razem w mechanizmie uwagi dekodera pojawiają się pierwsze nietrywialne związki kontekstowe między słowami, które dekodér odwzorowuje w obrębie zdania źródłowego. W połączeniu z blokiem feed-forward powoduje to pojawienie się tym razem węższego zbioru możliwych kandydatów na dokończenie tekstu nawet bez spoglądania na mechanizm uwagi krzyżowej. A jeśli do tego następnie dołącza predykcja z bloku uwagi krzyżowej, gdzie znajdujemy najbardziej prawdopodobną łączliwość ostatnio przetłumaczonego słowa z kolejnym, to uzyskujemy bardzo dobrego kandydata na następne słowo.

Opisany proces powtarza się aż do momentu, gdy dekodér proponuje na wyjściu sekwencję EOS (*koniec zdania*, ang. *End Of Sentence*) – w tym momencie wszystkie obiekty i związki między nimi zostały ze zdania źródłowego odwzorowane w zdaniu docelowym.

Warto tutaj zwrócić uwagę na jedną rzecz: tekst dekodowany (w języku docelowym) może mieć inną liczbę słów niż tekst źródłowy. np. „kot” w języku polskim przejdzie w „the cat”. Nie jest to problem dla dekodera. W momencie gdy związki uwagi wskazują, że należy wspomnieć obiekt kota, a sieć neuronowa nauczona jest, że w języku angielskim trzeba rzeczownik poprzedzić słowem rodzajnika, którego jeszcze związki uwagi nie wykazały, to najpierw to słowo pojawi się na wyjściu dekodera. Skądinąd, związki uwagi enkodera będą dysponowały kontekstową informacją, czy zdanie mówi o „pewnym kocie” czy o „konkretnym kocie”, aby zdecydować, czy należy użyć „the” czy „a”.

3. Uczenie transformera

W procesie uczenia sieci neuronowej przepuszczamy przez nią dziesiątki milionów par zdań w obydwu rozpatrywanych językach. Zdanie w języku docelowym rozwijamy krok po kroku, po jednym słowie i weryfikujemy predykcję słowa następnego dla kolejnych podciągów prawidłowego zdania na wyjściu.

Gdy oczekujemy pierwszego słowa, proponowanego przez translator – porównujemy je tylko z pierwszym słowem poprawnego tłumaczenia. Gdy oczekujemy drugiego słowa – porównujemy je z drugim słowem, itd. Co istotne, w wypadku uzyskania

błędnej predykcji, jako przykład uczący z jednym słowem więcej dla takiego zdania wybieramy poprawny podciąg oryginalnego zdania i nie propagujemy błędu (tu jest różnica w działaniu transformera na etapie uczenia i predykcji).

Co do samego błędu – błąd ten w toku uczenia propaguje wstecz sieci neuronowej. Jak wspominaliśmy we wprowadzeniu (część 1), w sieci neuronowej każde połączenie między neuronami oznacza istnienie współczynnika wagowego, który można zmniejszyć lub zwiększyć. Zmiana w jednym kierunku powiększy błąd przewidywania, zmiana w kierunku przeciwnym – zmniejszy go (np. słowo przewidywane przez sieć w przestrzeni maszynowej zbliży się do położenia słowa oczekiwanego w przykładzie, przy okazji nie psując predykcji w innych przykładach). W ten sposób, manipulując dziesiątkami tysięcy (współcześnie bilionami) współczynników, ostatecznie możemy wytrenować sieć neuronową do poprawnego przewidywania kolejnych słów przetłumaczonych przykładów zbioru uczącego na podstawie tekstu źródłowego.

4. Podsumowanie

W niniejszym artykule starałem się omówić zasady funkcjonowania transformerów oraz wcześniejsze podejścia do problemu tłumaczenia tekstów. W toku przygotowywania tego materiału wiele wysiłku kosztowało mnie odszukanie informacji dotyczących interpretacji funkcjonowania poszczególnych bloków transformera. Są to zagadnienia znacząco trudniejsze koncepcyjnie niż elementarna teoria sztucznych sieci neuronowych i widać to bardzo w niepewności, z jaką tworzone są dostępne publikacje, adresujące tę problematykę. Sam artykuł z pewnością nie jest wyczerpujący, ale mam nadzieję, że umożliwi on czytelnikom „poukładanie” sobie w głowie wszystkich składowych transformera, a w razie potrzeby pogłębienie interesujących ich tematów. Tym, którzy mają zacięcie do programowania – proponuję skorzystanie z dostępnych bibliotek do obsługi transformerów w języku Python. Być może praca ta umożliwi też naukowcom, zajmującym się przetwarzaniem sekwencji sygnałów, podjęcie próby wykorzystania transformerów w ich tematyce badawczej. ■

Przemysław Borys
Wydział Chemiczny
Politechniki Śląskiej w Gliwicach

Literatura

- A. Vaswani et al, Attention is all you need!, 31st Conference on Neural Information Processing Systems (NIPS 2017), Long Beach, CA, USA,

- M. Geva, R. Schuster, J. Berant, O. Levy, Transformer Feed-Forward Layers Are Key-Value Memories, Proceedings of the 2021 Conference on Empirical Methods in Natural Language Processing.
- M. Geva et al., Transformer Feed-Forward Layers Build Predictions by Promoting Concepts in the Vocabulary Space, Proceedings of the 2022 Conference on Empirical Methods in Natural Language Processing.
- H. Xu et al., Probing Word Translations in the Transformer and Trading Decoder for Encoder Layers, Proceedings of the 2021 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies.
- L. Tunstall, L. von Werra, T. Wolf, Natural language processing with transformers. O'Reilly, 2022.
- S. Cristina, M. Saeed, Building transformer models with attention, MachineLearningMastery.com, 2022.
- G. Dar, M. Geva, A. Gupta, J. Berant, Analyzing transformers in embedding space, Proceedings of the 61st Annual Meeting of the Association for Computational Linguistics Volume 1: Long Papers, pp. 16124–16170, July 9-14, 2023,
- D. Jurafsky, J. H. Martin, Speech and language processing, darmowy ebook, Stanford University, 2024.
- T. Pires et al., One wide feedforward is all you need, Proceedings of the Eighth Conference on Machine Translation, 2023.
- E. Voita et. al, The Bottom-up Evolution of Representations in the Transformer: A Study with Machine Translation and Language Modeling Objectives, Proceedings of the 2019 Conference on Empirical Methods in Natural Language Processing and the 9th International Joint Conference on Natural Language Processing.
- Cykl wpisów o transformerach na blogu Sachinsoni, medium.com – bardzo dokładnie rozpisane kolejne kroki przetwarzania danych.
- Luis Serrano, What are transformer models and how they work? (przystępny wykład YouTube),
- Luis Serrano, The math behind attention: keys, queries and values matrices (przystępny wykład YouTube),
- Luis Serrano, The attention mechanism in large language models (przystępny wykład YouTube),
- Niels Rogge, How a transformer works at inference vs training time (wykład YouTube).

33. Transformerzy korzystają z funkcji ReLu

34. Przypomnijmy przy okazji rzecz wspomnianą przy okazji kodowania pozycji słów w enkoderze: to jest właśnie to miejsce, gdzie słowa „wylę-
wają się” ze swoich macierzystych tokenów na tokeny sąsiednie i po takim „wylęgnięciu się” – moglibyśmy bez kodowania pozycji stracić informację o ich położeniu w zdaniu źródłowym

35. Ignoruję tutaj ewentualne różnicowanie siły relacji na wyjściu różnych torów mechanizmu uwagi. W rzeczywistości opisane wystrojenie dotyczyć będzie tylko pierwszych 64 wymiarów wektora uwagi. Kolejne 64 będzie wystrojenie z siłą wynikającą ze związków odszukanych w drugim torze uwagi, itd. Idea rozumowania pozostaje jednak niezmienną

36. Była to zatem nieco zmodyfikowana wersja transformera, który w pracy „Attention...” miał 6 warstw

37. Tutaj właśnie nie da się już stosować relacji, że macierz $K=Q^T$ – gdyż wektory są reprezentowane w różnych przestrzeniach i konieczne jest przekodowanie pomiędzy nimi, $K=X_Q$, które w najprymitywniejszym przypadku może sprowadzać się np. do przenumeroowania osi 512-wymiarowej przestrzeni (bo w istocie te przestrzenie przechowują bardzo zbliżoną informację – o obiektach zdania i relacjach między nimi), ale możliwe są też bardziej złożone liniowe transformacje

38. Nie jest to całkiem precyzyjne, bo na etapie wyjścia z enkodera oprócz samych słów, mamy tutaj do czynienia z różnymi ich relacjami, które też są odwzorowywane na „słownik”, czy precyzyjniej, przestrzeń zanurzenia danej warstwy dekodera

39. Istnieją także mechanizmy uwagi, które odfiltrowują korelacje tokenów „z samym sobą”. W krzyżowym mechanizmie uwagi wydaje się to trudniejsze do zrealizowania

Sieci neuronowe typu transformer:

część 1

część 2

<https://ulubionykiosk.pl/wydawnictwo/mlody-technik/5276>

<https://ulubionykiosk.pl/wydawnictwo/mlody-technik/5300>

oprac. pl 8328880148



Szkoła Wynalazców

Zadanie waszym było zaproponować: urządzenie lub zespół urządzeń, które pomogłoby seniorom zachować sprawność ciała i umysłu.

Oczywiście – jak zwykle – nie oczekiwaliśmy jakiejś szczegółowej dokumentacji, a raczej ogólnej idei, która fachowcom pozwoliłaby zaprojektować konkretne rozwiązanie. Zgodnie z oczekiwaniami przyszło kilka propozycji w duchu zelektronizowanego i zinformatyзованego świata. W końcu mamy już XXI wiek! Zobaczmy więc, co wymyślili nasi młodzi wynalazcy:

Zenon Biernacki proponuje system inteligentnego monitorowania umysłu w zakresie funkcji pamięci, kojarzenia i koncentracji. System mógłby być wykonany jako stół interaktywny lub tablica, na której pojawiałyby się zadania do wykonania – w postaci rysunków lub grafiki animowanej. Zadania byłyby oceniane przez system, który wyświetlałby zdobytą liczbę punktów.

Dobra koncepcja i dziś łatwa do urzeczywistnienia. Dla stworzenie programów ćwiczeń można by wykorzystać stare dobre zadania ze zbiorów takich jak Lilavati i podobnych.

Zuzanna Rolnik proponuje stanowisko do ćwiczeń fizycznych, angażujące nogi, ręce i plecy ćwiczącego. Stanowisko mogłoby być zaopatrzone w system monitorowania parametrów takich jak: puls, ciśnienie, częstość oddechów i alarmujące o zagrożeniu, w przypadku zbyt intensywnych ćwiczeń.

Idea stanowiska do ćwiczeń fizycznych jest już znana i pojawiła się jako m.in. „plenerowe siłownie”, co zresztą nie odpowiada ich przeznaczeniu. Istotnym elementem byłoby monitorowanie parametrów ważnych dla zdrowia ćwiczącego. Starsi panowie niekiedy, wspominając młode lata, próbują powtórzyć „tamte” wyczyny, kiedy to robiło się po 100 przysiadów ze sztangą 80 kg na barkach. Dziś młodzi i w pełni sprawni ludzie biorą udział w maratonach i ultramaratonach, nie myśląc o czekającej ich kolejce do operacji stawu kolanego i zastąpienia go stawem sztucznym... Niestety nie ma na rynku sprzętu do monitorowania dynamicznych obciążeń stawów.

Roman Tarło – Automatyczny trener ruchowy – urządzenie przypominające trenażer, które pomaga

utrzymać sprawność fizyczną, dostosowując intensywność ćwiczeń do możliwości użytkownika. System inteligentnego monitorowania aktualnych parametrów dostosowywałby obciążenie, dawkując obciążenie i liczbę powtórzeń ruchu w jednym cyklu, a nawet przerywając możliwości dalszych ćwiczeń, gdy program uzna, że to dla ćwiczącego już za dużo.

Idea ze wszech miar słusza. Tendencja w rozwoju techniki jest dziś taka, że coraz więcej istotnych dla bezpieczeństwa człowieka funkcji powierzamy systemom inteligentnej automatyki. Człowiek, przeciętnie rzecz biorąc, jest nieodpowiedzialny, także w sprawie dbania o własne zdrowie.

Wymienionym kolegom gratuluję i zapraszam do następnych zadań.

Nowe zadanie

Wydaje się, że współczesna elektronika potrafi już „wszystko”. Dobrze by było, żeby każdy aktywny człowiek mógł mieć na rękę „asystenta”, który monitorowałby podstawowe parametry zdrowia, pracy np. reagował na nastroje i zachęcał do posłuchania dobrej muzyki, poczytania lub medytacji. Spróbujcie zaproponować program działania takiego pomocnika: co powinien kontrolować, jak się porozumiewać. Można tu zdefiniować kilka wersji: dla najmłodszych, dla młodzieży, ludzi aktywnych i dla seniorów. W rezultacie wasze zadanie to: zaproponować program działania uniwersalnego „asystenta”, który pomógłby zadbać o zdrowie, wyznaczałby godziny snu i pracy, monitorował nastrój itp.

Wydaje się, że gwałtowny rozwój cywilizacyjny spowodował sytuację, że najtrudniejszym zadaniem jest odpowiedzieć sobie na pytanie: co by tu ciekawego wymyślić.

Przypomina to piosenkę Wojciecha Młynarskiego: „W co się bawić, w co się bawić...” Było to „czarne memento”, czyżby to już nadchodziło?

Przypominam o terminie nadsyłania rozwiązań: do końca czerwca 2025 r.

Klub Wynalazców

Był taki czas, gdy nasi zachodni sąsiedzi mawiali: Polaka natychmiast poznasz po brudnych, nieoczyszczonych butach. Ten czas – mam nadzieję, minął, ale problem pozostał. Nie ma dobrego sposobu na wyczyszczenie butów we współczesnym mieszkaniu. Zaproponować ideę prostego i raczej małego urządzenia ułatwiającego oczyszczenie butów nieprzesadnie zabłoconych.

Zygmunt Fijałkowski – zajął do znanego portalu handlowego, oferującego tysiące przeróżnych drobiazgów i – podobnie jak Archimedes – wykrzyknął eureka!, bo zobaczył ofertę na przyrząd, który pozwala wyczyścić wannę, umywalkę i także buty. Jest to elektryczna szczotka, jej organ roboczy to wirująca tarcza z wymiennymi nakładkami, o stosunkowo wolnych obrotach, dzięki czemu nie powinna bryzgać np. pastą do butów na boki. Szczotkę ładuje się, ona ma wskaźnik naładowana akumulatorów. Wydaje się, że nic lepszego nie da się wymyślić. Można jedynie pomyśleć o nakładkach roboczych i opracować ich szerszą gamę.

Znamy tę szczotkę, rzeczywiście jest to udany produkt: lekka, stosunkowo mocny silnik i estetyczny wygląd. Przypadek ten potwierdza pogląd, że dziś łatwiej jest zrobić coś ciekawego, a znacznie trudniej to sprzedać. Dlatego wszystkie firmy produkujące gadżety „domowe” prześcigają się w oferowaniu najrozmaitszych produktów.

Marek Piwowski – przy czyszczeniu butów mężczyzny jest ten ruch posuwisto-zwrotny, niezbędny do wyczyszczenia czegokolwiek. Uważa, że szczotka o kształcie zwykłej szczotki, ale z wbudowanym wibratorem, pozwoliłaby w znacznym stopniu ułatwić czyszczenie butów.

Niezły pomysł: najbardziej dokuczliwy jest machanie szczotką „tam i z powrotem”. W przypadku zastosowania szczotki wibracyjnej ten ruch byłby wyeliminowany. Szczotki z obrotową głowicą należałoby sprawdzić na rozbryzgiwanie resztek błota i pasty do butów.

Antoni Namirski – buty najłatwiej jest wyczyścić „na mokro” tzn. pod łagodnym strumieniem wody, z użyciem szmatki lub małej szczoteczki. Inne metody nie są tak skuteczne. Najlepiej nadawałby się sprzęt podobny do tego, jaki stosowany jest w niektórych myjniach samochodowych. Chodzi o szczotkę, nawet małą, przez którą podawana jest woda. Wydaje

się, że to byłoby najskuteczniejszym sposobem uprania się z brudnymi butami.

To w zasadzie prawda, ale nie wszystkie rodzaje obuwia nadają się do wymycia na mokro. Tak czy owak problem obuwia i zwłaszcza „traktorków” pozostaje nadal nienierozwizany.

Wymienionym kolegom gratuluję i zapraszam do dalszych zadań.

Nowe zadanie

Zimą – łagodną w tym roku mamy już za sobą, ale o sprzęcie narciarskim warto pomyśleć, bo czas szybko leci. Dla amatorów nart biegowych mamy ciekawy problem. W 1985 roku 15-kilometrowy bieg narciarski wygrał Szwed Thomas Wassberg. Na trasie znajdowało się kilka wyznaczonych odcinków, na których jazda stylem „łyżwowym” była zabroniona (na początku dystansu). Problemem było to, że tradycyjna technika stosowana w strefach zamkniętych obejmowała użycie wosku, a nasmarowane narty spowalniały prędkość jazdy „łyżwą” w ostatniej, decydującej części wyścigu. Co powinien zrobić biegacz wyścigowy? Wasze zadanie to: zbadać możliwości i próbować ustalić: jaki to sposób zastosował Wassberg, żeby pogodzić konieczność smarowania nart woskiem i jednocześnie w tym samym biegu dysponować nartami nieposmarowanymi (oczywiście bez wymiany nart).

Ogólne zasady sportowej rywalizacji często sprowadzają się do szukania sposobów zręcznego ominięcia regulaminu i przepisów o sprzęcie i technice wykonania akcji. Pamiętamy słynne takie przypadki jak oszczep Helda, „masterhammer – do rzutów młotem, przejście ze skoków wzwyż „nożycowego” do skoku „flopem”. Także i w narciarstwie jest pole do popisu dla coraz bardziej wyrafinowanych metod podnoszenia wyników. Spróbujcie – może przyda się w zimie roku 25/26

Wszystkim życzymy dobrych pomysłów i przypominamy o terminie: do końca czerwca br.

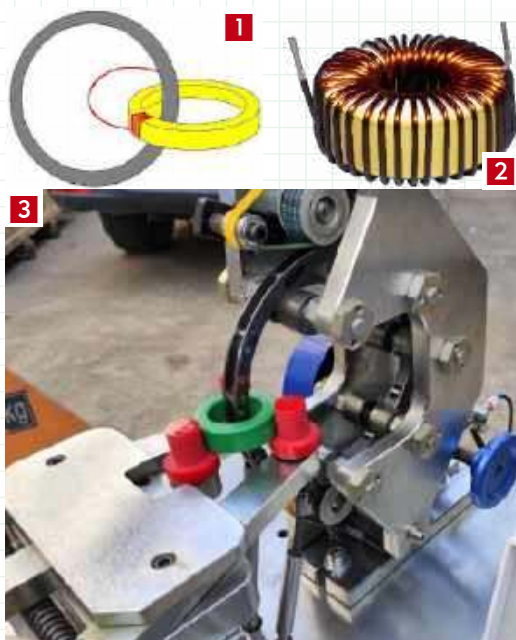
Vademecum Młodego Wynalazcy

W różnych opracowaniach, dotyczących wynalazków i w ogóle wynalazczości, często pojawia się element przeniesienia cechy zauważonego gdzieś jakiegoś obiektu, często odległego od idei, nad którą właśnie pracuje wynalazca, do aktualnego problemu. Na wstępie przytoczymy tu anegdotę. W pewnej fabryce samolotów poddawano próbom nowy model. Próby nie wychodziły: Po każdym starcie, jeszcze na wysokości rzędu kilku metrów, odrywały się obydwa skrzydła. Parę tygodni trwały kolejne próby i nic... skrzydła nadal odpadały. Pewien młody pracownik powiada: Zapytajmy naszego rabina, to bardzo mądry człowiek i może on doradzi. Eeeee – co on może wiedzieć o samolotach! Ale on jest naprawdę bardzo mądry! A co szkodzi zapytać. I poszedł do rabina. Rabin powiedział: Wywiercieście wzdłuż miejsca, w którym skrzydła się odrywają, małe otworki: takie po 2 mm średnicy i musi

ich być dużo, tak około co 4 mm. Otworki wywiercono, nie wierząc ani sekundę w ich sens. Samolot wystartował i... poleciał. Rabin zapytany: jak wpadł na taki genialny sposób, odparł: synu, ja już długo żyję i zużyłem mnóstwo papieru toaletowego. Wszystkie rolki tego papieru mają co jakiś czas rząd otworków w miejscu, w którym kolejny kawałek powinien się oderwać. I wyobraź sobie, że nigdy ten papier nie odrywał się w miejscu, gdzie były dziurki!

Przykładem takiego dochodzenia w do dobrego pomysłu może być opracowanie nawijarki do toroidalnych transformatorów dla maszyn liczących.

Tu trzeba zaznaczyć, że dotyczyło to starych urządzeń sprzed 60...70 lat, w których takich transformatorów było po kilka-kilkanaście tysięcy. Konieczna była mechanizacja procesu nawijania uzwojenia. Ale – jak widać – rysunek 1, sprawa nie jest prosta, bo ani rdzeń,



ani jakaś szpulka z zapasem drutu nie mogły być osadzone na osiach. W dodatku sprawa się komplikuje, gdy wymiary rdzenia maleją i wynoszą np. 2 mm średnica zewnętrzna, a bywają i mniejsze i też muszą być nawinięte! Nawiniętą cewkę o dużych rozmiarach przedstawia rysunek 2. Do nawinięcia takiej cewki znane były nawijarki o różnym stopniu zaawansowania: głównie zaopatrzone w zaawansowaną elektronikę. Dotyczyły one jednak rdzeni o średnicach rzędu paru cm – rysunek 3.

Trzeba pamiętać, że rdzenie produkowane są metodą proszków spiekanych i są nierozbieralne! Zadaniem zaprojektowania nawijarki dla rdzenia o średnicy zewnętrznej 2 mm zajął się Jegorow (opis tego przypadku z książki H. Altszullera: „Algorytm wynalazku”), który wcześniej opracował nawijarki do dużo większych rdzeni. Dla tak małego rdzenia sprawa wydawała się nie do pokonania. Próbował różnych koncepcji, np. układ dwóch wahadeł: jedno przewlekało drucik na drugą stronę, gdzie czekało drugie wahadło i przepychało drucik z powrotem. Wykonano próbne modele, niestety nie zdawały egzaminu, aż nagle – ulubiony zwrot autorów opracowań o wynalazcach – Jegorow, jadąc metrem, z nudów obserwował pasażerów i wzrok jego spoczął na chwilę na starszej kobiecie, która robiła sobie jakąś koronkową robótkę, posługując się szydełkiem z haczykową końcówką. Jegorow zaczął uważnie obserwować ruchy dłoni kobiety i „przymierzać” do swojej nawijarki. Po paru minutach główna koncepcja nawijarki była niemal gotowa. Należało opracować ją technicznie i wykonać prototyp. Tym razem próby zakończyły się sukcesem. Można spytać: „gdzie Rzym, a gdzie Krym, czyli gdzie dzierganie koronki i cewka toroidalna?” To prawda, są to bardzo różne sprawy, chociaż istota problemu była w pewnym sensie wspólna.



A oto kolejny przykład: małżonka profesora Tadeusza Sędzimir postanowiła upiec strudel. Podstawą tego wypieku jest bardzo cienko rozwałkowany placek i do tego celu panie stosowały specjalną technologię. Najpierw na stolnicy rozwałkowywały porcję ciasta do średnicy około 0,3 metra, dalej już placek był dość cienki i walcowanie wymagałoby znacznej siły. Panie radziły sobie w ten sposób, że rozpościwały na stole obrus, pod którym, na środku, wstawiały głęboki talerz „do góry nogami” i na to wszystko układały wstępnie rozwałkowane ciasto. Następnie, korzystając z tego, że „górką” utworzoną przez talerz hamowała przesuwanie ciasta, rozciągały je aż do średnicy ok. 0,6 m, pomagając dodatkowo wałkowaniem. Przekrój gotowego strudla ujawnia fakt, że ciasto jest istotnie bardzo cienkie – rysunek 4. Do pomocy w tej pracy został poproszony prof. Tadeusz Sędzimir, który w tym czasie rozpracowywał problem walcowania cienkich blach. Czy to prawda, czy legenda – nie wiadomo na pewno, ale T. Sędzimir właśnie mniej więcej w tym czasie opracował zasadę podwójnego stanu naprężeń, która później została zastosowana w metodzie RT – Tadeusza Ruta – kucia wałów korbowych silników okrętowych i w metodzie produkcji trójników, polegającej na jednoczesnym ścisnieniu półfabrykatu – rurki – z jednoczesnym rozprężaniem ciśnieniem oleju, wtłaczającego materiał rurki do otworu byczego matrycy.

Szwajcarski inżynier Georges de Mestral, pomagając żonie pozbyć się sporej ilości „guziczków” – owoców łopianu, pospolicie nazywanymi rzepami, wziął je pod mikroskop, żeby zbadać tajemnicę ich czepiania się tkanin, sierści zwierząt itp. Pod mikroskopem zobaczył, że poszczególne włókienka rzepu są zakończone haczykami, a że są małe, łatwo przenikają pomiędzy włóknami tkaniny, ale z powrotem już nie tak łatwo je wyjąć – rysunek 5. Działało to trochę jak harpun lub strzała bojowa z zadziórami. Oczywiście odkrył ideę, ale do rzepów było jeszcze daleko. Jednak to już była robota dla specjalistów od produkcji tekstyliów. Dziś „na rzepy” mamy zapinaną ogromną ilość różnych



rzeczy: od aktówek biurowych, poprzez różne elementy odzieży, obuwie, itp.

W 1944 Donald Griffin, amerykański zoolog zajmujący się badaniem nietoperzy, sporo czasu poświęcił ich zdolności do orientacji przestrzennej w nocy. On właśnie nazwał tę ich technikę echolokacją, już później, gdy efekt ten zaczęto wykorzystywać w łodziach podwodnych, urządzenia zyskały nazwy: sonar lub echosonda. Oprócz nietoperzy wiele innych gatunków zwierząt, także lądowych i ptaków, wykorzystuje echolokację dla własnego bezpieczeństwa i dla upolowania jakiejś zdobyczy. Oczywiście droga od nietoperza do łodzi podwodnej była daleka, jednakże główne zasady działania sonaru zyskały podstawy teoretyczne właśnie od nietoperzy. Ich sonar



jest tak precyzyjny, że bez trudu omijają napięte na drodze ich przelotu – rysunek 6.

Należy podkreślić, że w przypadku sonaru nietoperzy i wielu innych zwierząt nie próbujemy naszą toporną technologią skopiować tak niesłychanie subtelny organ, jak sonar nietoperza, lecz korzystamy z głównej idei, realizując ją w wersji dostępnej dla możliwości współczesnej techniki. W ten sposób zrodził się nowy kierunek techniki: bionika zwana niekiedy biomimetyką.

Rysunek 7 to nasz ogólnie znany trzmiel. Od razu rzuca się w oczy jego kosmate i raczej tłusciutkie ciało i mikry rozmiar skrzydeł. Rzeczywiście 80...100 lat temu poważni naukowcy, na podstawie obliczeń aerodynamicznych, ustalili, że trzmiele nie mają prawa latać! W stosunku do powierzchni skrzydeł są za ciężkie. Na szczęście trzmiele o tym nie wiedziały i... latały. Latały po prostu nie tak jak samoloty, czyli wykorzystując siłę nośną, generowaną przez profil skrzydła, ale latały tak jak owady! Dopiero dokładna analiza mechaniki lotu trzmienia pozwoliła wyjaśnić, dlaczego lata, choć nie powinien. Okazało się, że ich skrzydełka wykonują skomplikowany ruch drgający, a dodatkowo mają tzw. buławki, które w efekcie dają im „nadprogramową” siłę nośną. W rezultacie wiele z tej mechaniki zastosowano w budowie helikopterów, dzięki czemu znacznie poprawiono ich sprawność.

Kolejny przykład skorzystania z modeli przyrodniczych to uskrzydlenie samolotów i skrzydła ptaków.

Rysunek 8 pokazuje czajkę w locie. Widać wyraźnie, że jest to „górnopłat” z końcówkami lotek podobnymi do tych, jakie dziś mają samoloty. Są to tzw. winglety, poprawiające parametry nośne samolotu. Ale pierwsze miały to ptaki! Najpełniej kształt ogólny skrzydeł ptaków został wykorzystany w polskich samolotach myśliwskich PZL P24 – rysunek 9. Charakterystyczne było bliskie zamocowanie skrzydeł w stosunku do kadłuba, co spowodowało konieczność zastosowania zastrzałów. Te skrzydła zyskały międzynarodową sławę i były nazywane „polskie skrzydła”, stosowane przez różne wytwórnie samolotów w różnych krajach.

Tych kilka przykładów potwierdza powszechnie znany fakt, że pomysły przychodzą „kiedy chcą”. I że źródłem inspiracji wcale nie musi być dziedzina, w ramach której właśnie się „męczymy”. Jest jednak jedno „ale”. Takie sprawy zdarzają się ludziom o umysłach przygotowanych. I wcale nie chodzi tu o formalne wykształcenie, a raczej typ charakteru: twórczego i ciekawego świata. Podziwiając różne cuda i obiekty, człowiek twórczy nie poprzestaje na samym podziwie i zaraz pyta: a dlaczego tak się dzieje? Dociekanie przyczyn różnych zjawisk i zachowań obiektów przyrody już niejednokrotnie doprowadziło do odkryć i wynalazków. Pozostaje jeszcze jedno pytanie: czy wynalazca musi czekać na taki szczęśliwy przypadek jak Jegorow, gdy zobaczył kobietę robiącą coś z pomocą szydełek z haczykowatymi końcówkami? Istnieje wiele sposobów, żeby aranżować sobie takie szczęśliwe przypadki, ale o tym w następnych odcinkach VMW. ■

Prezes Klubu Wynalazców
Champion TRIZ
Jan Boratyński



SILNIK ZABURTOWY



Jeśli chcesz wyposażyć w napęd mały model pływający, najprościej jest zastosować uniwersalny silnik zaburtowy opisany w artykule.

Wiosna w pełni. Na dobrą sprawę do lata pozostało niewiele dni, to dobry czas, żeby zająć się jakimś modelem plenerowym. Tym razem nie będzie to gotowy model, tylko uniwersalny napęd doczepny, pełniący też funkcję steru. Do zastosowania w dowolnym małym modelu pływającej łodzi, również zdalnie sterowanej.

Budowa silnika doczepnego

Budowę napędu rozpoczynamy od wydrukowania na drukarce 3D niezbędnych części. Pliki do druku należy pobrać z linku dołączonego do artykułu. Części zostały przygotowane tak, żeby można je było wydrukować na najprostszej drukarce. Poza wydrukowanymi

częściami niezbędne będą: dwie śruby m3 o długości 25 mm, kawałek prostego stalowego drutu o średnicy 1,5 mm, mosiężna tulejka (rurka) o średnicy wewnętrznej 1,5...1,6 mm, koszyk na trzy baterie (paluszki), gumka recepturka i silniczek elektryczny o rozmiarze „130”. Silnik 130 jest stosowany w większości elektrycznych zabawek i można go kupić za kilka zł, choć nie jest wykluczone, że każdy majsterkowicz znajdzie taki silniczek w stercie różnych „przydasi” zalegających w dolnej szufladzie. Zastosowany silnik nie ma na tyle mocy, żeby nasza łódka była demonem prędkości, ale jest w zupełności wystarczający, żeby nasz model żwawo poruszał się po stawie, sadzawce czy przydomowym basenie.



Wydrukowane części



Pozostałe części niezbędne do budowy



Montujemy wał turbinki



W turbince montujemy tulejkę ślizgową



Zamontowana turbinka



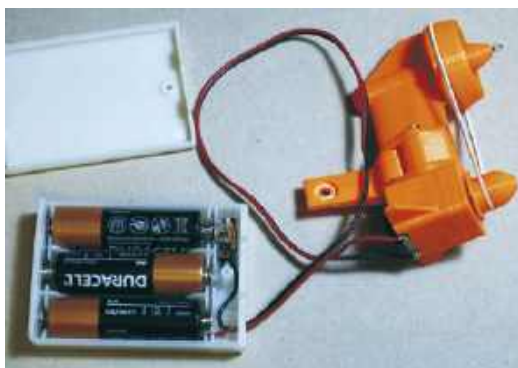
Montaż całości

Jeśli mamy wydrukowane części i skompletowane pozostałe elementy, przystępujemy do montażu, w pierwszej kolejności w części dolnej (cz1) montujemy wał, na którym będzie się obracać turbinka napędowa (cz2). Wał wykonujemy z drutu stalowego 1,5 mm, ja do tego celu wykorzystałem złamane wiertło 1,5 mm, natomiast w turbinkę wklejamy mosiężną tulejkę. Jeśli otwory w wydruku są za małe, należy je rozwiąć do wymaganego wymiaru. Na wale montujemy turbinkę i zabezpieczamy ją przed spadnięciem,

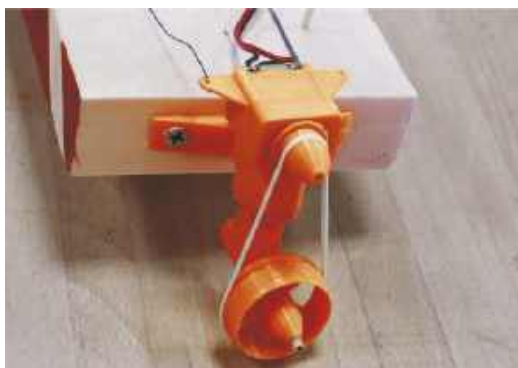
na przykład kawałkiem plastikowej rurki. Turbinka powinna się luźno obracać. Wszystkie części łączymy ze sobą, skręcając dwoma śrubkami m3 lub wkrętami o podobnej średnicy. Silniczek montujemy na wcisk w górnej części napędu (cz5), jeśli jest zbyt luźny, owijamy go taśmą izolacyjną. Na wale silnika montujemy na wcisk koło pasowe (cz6), jako pasek napędowy użyjemy gumki recepturki o średnicy ok. 40 mm. Zakładamy ją na koło pasowe i turbinkę, turbinka ma w tym celu wykonany specjalny rowek.



Montujemy silniczek



Montaż elektryczny

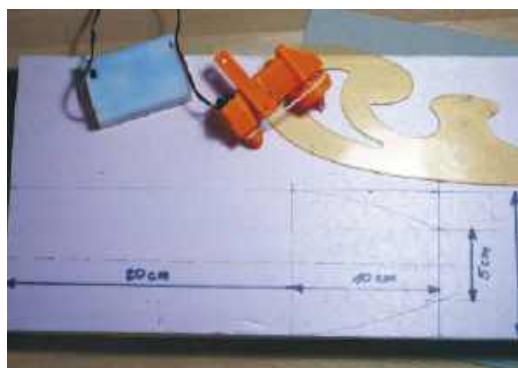


Napęd zamontowany do łódki

Do zacisków silniczka lutujemy przewody zasilające, drugi koniec przewodów lutujemy do koszyczka z bateriami. Jeśli koszyczek na baterie nie ma wyłącznika, należy takowy zamontować, przecinając jeden z przewodów i lutując wyłącznik. Pozostaje nam włożyć baterie i sprawdzić, czy wszystko działa tak jak powinno i czy turbinka kręci się we właściwym kierunku, zmianę obrotów silnika uzyskamy, zamieniając przewody przy silniku. Jeśli wszystko działa poprawnie, można nasz napęd przykręcić

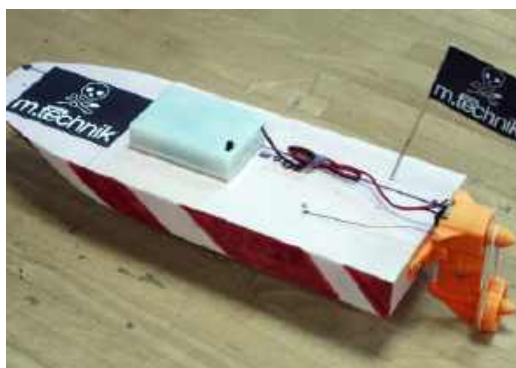
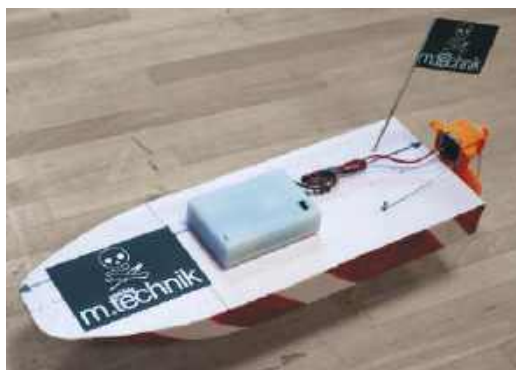


Gotowy napęd



Najprostszą łódkę można wyciąć z płyty styroduru

do rufy łódki dwoma śrubkami lub wkrętami, napęd należy zamontować tak, żeby cała turbinka była zanurzona w wodzie. Jeśli nie mamy gotowej łuki w którym można by zastosować nasz silniczek, najprostszy model można wyciąć z płyty styroduru lub twardego styropianu o grubości ok. 3 cm, korzystając z wymiarów podanych na zdjęciu. Łódkę należy wyciąć ostrym nożykiem do tapet, a ewentualne niedoskonałości wygładzić papierem ściernym, do rufy przykręcamy napęd, a do pokładu przyklejamy



Gotowy model z zamontowanym napędem

klejem na gorąco koszyczek z bateriami. Opisany tutaj napęd można też zastosować w małym zdalnie sterowanym modelu. W tym przypadku wypustki w górnej części napędu należy połączyć popychaczem

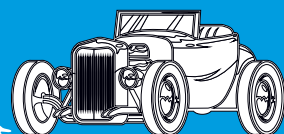
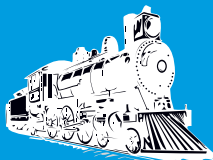
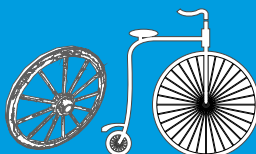
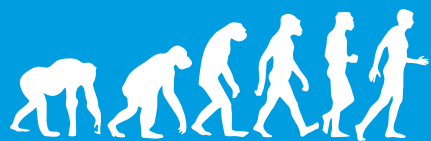
z orczykiem serwo mechanizmu, nasz napęd będzie pełnił też funkcję steru.

Życzę kreatywnej zabawy i „stopy wody pod kilem”. ■

Mariusz Wrona



**Archiwalne artykuły „Na warsztacie” znajdziesz na stronie:
<https://mlodytechnik.pl/zrob-to-sam>**



Cyfrowe bliźniaki

lata 60.–70. XX w.

Idea „cyfrowego bliźniaka” narodziła się w NASA, choć jeszcze przez długie lata nie używano tej nazwy. Jay Forrester, profesor w MIT, wprowadził na początku lat 60. koncepcję modeli komputerowych. Jego prace dotyczyły symulacji systemów dynamicznych, co stanowi fundament późniejszego rozwoju technologii cyfrowych bliźniaków. W 1965 r. John McCarthy, jeden z pionierów sztucznej inteligencji, proponuje ideę „maszyn myślących”, które mogłyby działać jako cyfrowe duplikaty rzeczywistych obiektów. Agencja NASA budowała w tamtym czasie także fizyczne duplikaty systemów kosmicznych na Ziemi, starając się dopasować je możliwie najwierniej do systemów, które zostały wysłane w ramach misji. Reagując na sytuację awaryjną po eksplozji zbiornika tlenu w zmierzającym ku Księżycowi statku Apollo 13 (1) w 1970 r., agencja sięgnęła po cały szereg symulatorów, rozszerzając także fizyczny model pojazdu o komponenty cyfrowe. Uważa się, że powstały wtedy „cyfrowy bliźniak” (ta nazwa nie była wtedy używana) był pierwszym systemem tego rodzaju. Umożliwiał ciągłe pobieranie danych w celu modelowania zdarzeń prowadzących do wypadku w celu analizy i zbadania kolejnych kroków. Wydarzenia związane z akcją ratunkową Apollo 13 inspirowały dalszy rozwój techniki cyfrowych symulacji, zwanych później bliźniakami.

1991

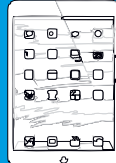
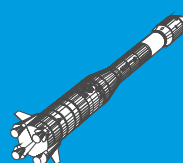
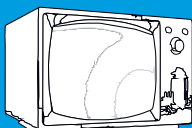
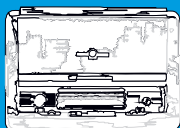
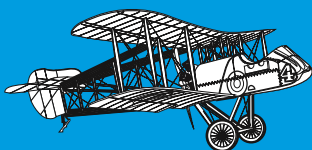
Przybliżona koncepcja cyfrowego bliźniaka opisana zostaje w książce Davida Gelerntera z 1991 r. pt. „Mirror Worlds” (2).

1997

Pomysł „cyfrowego bliźniaka”, który był przez lata znany pod różnymi nazwami (np. „wirtualny bliźniak”), został po raz pierwszy nazwany dokładnie w ten sposób – „cyfrowy bliźniak” – przez Luisa A. Hernandez-Ibañę i Santiago Hernandez z uniwersytetu w La Coruña w publikacji pod angielskim tytułem „Application of Digital 3D Models on Urban Planning and Highway Design”. Rok później termin „cyfrowy bliźniak” pojawił się jako opis cyfrowej kopii głosu aktora Alana Aldy w filmie „Alan Alda meets Alan Alda 2.0”.

2002

Za punkt zwrotny w rozwoju techniki „cyfrowych bliźniaków” powszechnie uznawana jest prezentacja Michaela Grievesa (3) na Uniwersytecie Michigan przygotowana dla przedstawicieli przemysłu. Jego prezentacja dotyczyła rozwoju centrum zarządzania cyklem życia produktu (PLM). Choć Grieves nie używał pojęcia „digital twin”, definiował wszystkie elementy składowe cyfrowego bliźniaka tak jak rozumiane jest współcześnie – przestrzeń rzeczywistą, przestrzeń wirtualną, łączy do przepływu danych z przestrzeni rzeczywistej do wirtualnej, łączy do przepływu informacji z przestrzeni wirtualnej do rzeczywistej i wirtualne podprzestrzenie. Model ten kładł nacisk na dwukierunkowy mechanizm łączy dwóch przestrzeni i posiadanie wielu wirtualnych przestrzeni dla jednej rzeczywistej przestrzeni, w której można badać alternatywne pomysły lub projekty. Pierwotnym terminem Grievesa był „model lustrzanych przestrzeni”, a później używał „modelu lustrzanego odbicia informacji” i „wirtualnego sobowtóra”. Do 2006 r. nazwa modelu koncepcyjnego zaproponowanego przez Grievesa została zmieniona z „Mirrored Spaces Model” na „Information Mirroring Model”. Ze względu na ograniczenia techniczne, małą jeszcze wtedy moc obliczeniową, niską lub żadną łączność urządzeń z Internetem, brak systemów przechowywania danych i zarządzania nimi, słabo rozwinięte algorytmy maszynowe itp. na wdrożenie koncepcji Grievesa w praktyce było jeszcze za wcześnie.



2003

Kary Främling z uniwersytetu w Helsinkach wraz zespołem proponuje „architekturę opartą na agentach, w której każdy element produktu ma odpowiadający mu ‘wirtualny odpowiednik’ lub powiązanego z nim agenta” jako rozwiązanie problemu nieefektywności transferu informacji produkcyjnych za pośrednictwem papieru dla PLM.

2010–18

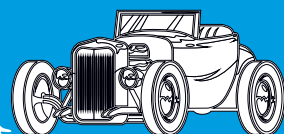
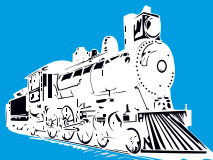
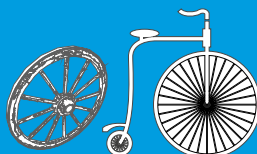
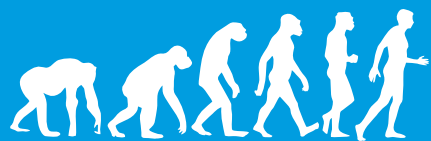
W raportach rocznych NASA technika cyfrowych duplikatów urządzeń, statków, rakiet (4) i procesów została po raz pierwszy oficjalnie określona terminem „Digital Twin” w 2010 r. NASA była też pierwszym podmiotem o tym znaczeniu, który stworzył techniczną definicję cyfrowego bliźniaka. Brzmiała ona w tłumaczeniu – „zintegrowana wieloskalowa, probabilistyczna symulacja pojazdu lub systemu, która wykorzystuje najlepsze dostępne modele fizyczne, aktualizacje czujników, historię floty itp. w celu odzwierciedlenia cyklu życia uczestniczącego w misji bliźniaka”. Historyczne znaczenie ma sięgnięcie po ten termin i stworzenie definicji, bo, jak wiadomo, NASA wykorzystywała tę technikę faktycznie już wcześniej. Wkrótce Siły Powietrzne Stanów Zjednoczonych poszły w ślady NASA i wykorzystały cyfrowe bliźniaki do projektowania, konserwacji i przewidywania działania swoich samolotów. W ślad za prekursorami poszły inne instytucje i firmy.

2014–16

Koncern GE (General Electric) zainicjował wdrażanie cyfrowych bliźniaków w przemyśle. Firma zainwestowała znaczące środki w rozwój platformy Predix (5), wykorzystującej technikę cyfrowych bliźniaków do monitorowania i optymalizacji pracy maszyn w przemyśle. W 2016 roku firma Siemens, pod kierownictwem Joe Kaesera, wprowadziła platformę MindSphere, która wykorzystuje cyfrowe bliźniaki w zarządzaniu procesami przemysłowymi i automatyzacji.



1. Statek Apollo 13, 2. Książka 'Mirror Worlds' Davida Gelerntera, 3. Michael Grieves, 4. Wizualizacja rakiety i jej cyfrowego bliźniaka, 5. System Predix firmy GE



2014–25

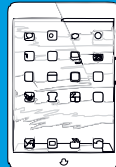
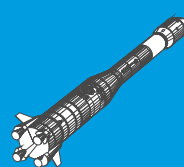
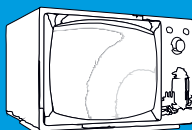
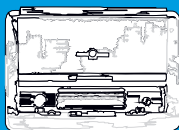
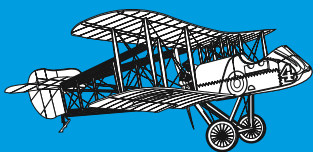
2017

2019–25

W ramach projektu Living Heart firmy Dassault Systemes stworzono dokładny wirtualny model ludzkiego serca (6), który może być testowany i analizowany, umożliwiając chirurgom odgrywanie różnych scenariuszy terapeutycznych. Projekt został założony przez Steve'a Levine'a, którego córka urodziła się z wrodzoną chorobą serca. Gdy miała 20 lat i wysokie ryzyko niewydolności serca, postanowił odtworzyć jej serce w wirtualnej rzeczywistości. Firma Dassault Systemes podjęła prace nad kolejnym cyfrowymi bliźniakami ludzkich organów, w tym oka, a nawet mózgu. Techniki te miałyby tworzyć także m.in. wirtualnego bliźniaka indywidualnego genomu danej osoby i ostatecznie cyfrowego bliźniaka całej istoty ludzkiej. W 2017 r. zainicjowano projekt „Digital Twin of a Human”, w ramach którego Philips i inne firmy zaczęły pracować nad symulacjami pacjentów do celów diagnostycznych i terapeutycznych.

Instytut Gartnera umieszcza cyfrowe bliźniaki na liście dziesięciu najważniejszych trendów technologicznych, co znacząco zwiększyło zainteresowanie tą technologią w różnych sektorach przemysłu.

Uruchomienie przez NVIDIA Omniverse (7) 3D uznaje się za ważny przełom w dziedzinie wirtualizacji. Omniverse umożliwił kilku zespołom równoczesną pracę nad projektem, przy czym, dzięki przeniesieniu ciężkiej pracy komputerowej do akcelеровanych środowisk chmurowych, zespoły mogą korzystać ze zwykłych laptopów zamiast wysoko zaawansowanych maszyn. Omniverse 3D z technikami VR pozwala użytkownikom uzyskiwać dostęp do cyfrowych bliźniaków i wchodzić z nimi w interakcję w środowisku 3D. Projektanci mogą fizycznie prowadzić oględziny i manipulować przy symulacji 3D np. pojazdu naturalnej wielkości, lepiej wizualizować projekt w różnych warunkach oświetleniowych pod różnymi kątami. Mogą chodzić po hali fabrycznej i naocznie przekonać się, jak zaprojektowane elementy pasują do przestrzeni, przemyśleć logistykę fabryki, zanim cokolwiek zostanie zmontowane i zainstalowane. W 2022 roku firma NVIDIA zaprezentowała Omniverse Enterprise – platformę do tworzenia i operowania cyfrowymi bliźniakami w czasie rzeczywistym. W kolejnych latach NVIDIA prezentowała kolejne rozszerzenia i udoskonalenia tego systemu, m.in. Omniverse Replicator, opartą na uczeniu maszynowym platformę stworzoną specjalnie do projektowania cyfrowych bliźniaków. Prezes firmy NVIDIA, Jensen Huang, ujawnił plany budowy przez jego firmę najpotężniejszego na świecie superkomputera AI przeznaczonego do przewidywania zmian klimatycznych. Nazwany Earth-2 lub E-2 (8), system ten stworzyłby cyfrowego bliźniaka Ziemi w Omniverse. Firma Siemens Energy wykorzystuje NVIDIA Omniverse przy tworzeniu cyfrowych bliźniaków elektrowni. Z kolei firma Chevron wdraża oparte na rozwiązaniach NVIDIA cyfrowe bliźniaki pól naftowych i rafinerii, dążąc do oszczędności na kosztach konserwacji. Do rywalizacji w dziedzinie cyfrowych bliźniaków dołączyła firma Oracle, która oferuje od kilku lat Oracle Digital Twin, dającą użytkownikom dwie opcje – digital twin i predictive twin. Konwencjonalny cyfrowy bliźniak działa tak jak inne rozwiązania tego rodzaju, natomiast „bliźniak predykcyjny” modeluje przyszły stan i zachowanie urządzenia na bazie „danych historycznych z innych urządzeń, symulacji awarii i innych wydarzeń”. W 2021 roku Microsoft wprowadził zaawansowane możliwości integracji z IoT i AI do swojej platformy Azure Digital Twins.



2020–22

Seria wdrożeń techniki cyfrowych bliźniąt na dużą skalę. Swojego cyfrowego bliźniaka, który modeluje nie tylko sam stadion, ale także otaczający go Hollywood Park, zyskał np. stadion SoFi w Los Angeles. Przykładem cyfrowego bliźniaka w jeszcze większej skali jest system modelujący całe Las Vegas. Powstała też symulacja części nowojorskiego Brooklynu, Navy Yard. Teksaski Uniwersytet A&M tworzy cyfrowego bliźniaka osiedli nadbrzeżnych, aby określić ich odporność na klęski żywiołowe. W Singapurze jednym z zadań cyfrowego bliźniaka miasta jest monitoring zanieczyszczeń. Szanghajske centrum zarządzania zbudowało cyfrowego bliźniaka (9) całej 26-milionowej metropolii. Modeluje on 100 tys. parametrów, od urządzeń do zbierania i wywozu śmieci, po infrastrukturę do ładowania e-rowerów, ruch drogowy oraz wielkość i lokalizację budynków mieszkalnych. Nowo powstające miasta w bogatych krajach arabskich są budowane jednocześnie w świecie rzeczywistym i cyfrowym. W wyścigach Formuły 1, zespoły McLaren i Red Bull używają cyfrowych bliźniaków swoich samochodów wyścigowych. Także jeżdżąca po zwykłych drogach Tesla tworzy cyfrowe symulacje swoich samochodów, wykorzystując dane zebrane z czujników umieszczonych w pojazdach i przesyłane do chmury.

2022–25

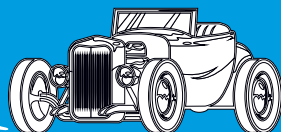
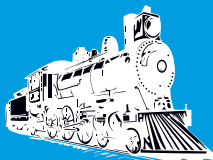
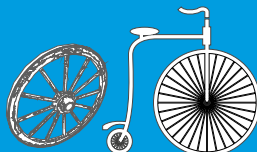
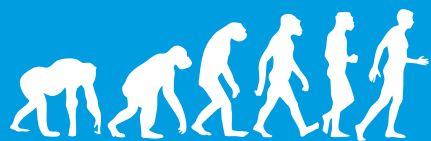
Kolejne postępy wdrażania cyfrowych bliźniaków w przemyśle. Przykładem może być kanadyjska firma CenterLine, która korzysta z cyfrowego bliźniaka hali fabrycznej do rozwiązywania problemów związanych z narzędziami, co poprawia płynność dostaw nawet o 90 proc. Logistyczny potentat, firma DHL, tworzy cyfrowe mapy swoich magazynów i łańcuchów dostaw, znacznie podnosząc efektywność operacji. Coraz częściej tworzenie symulacji typu cyfrowe bliźniaki wspierają modele sztucznej inteligencji.

2023

Firmy Philips i Siemens Healthineers zaczynają wykorzystywać medyczne cyfrowe bliźniaki do personalizacji terapii i planowania zabiegów.



6. Jedna z wizualizacji Living Heart firmy Dassault Systemes, **7.** Jedna z prezentacji platformy Omniverse firmy NVIDIA, **8.** Jensen Huang prezentuje Earth-2, cyfrowego bliźniaka całej Ziemi, **9.** Zespół monitorujący jeden z elementów szanghajskiego cyfrowego bliźniaka miasta



Klasyfikacja i najbardziej znane zastosowania cyfrowych bliźniaków

Cyfrowy bliźniak jest konstrukcją logiczną, co oznacza, że rzeczywiste dane i informacje mogą być zawarte w innych aplikacjach, i można za jego pomocą połączyć fizyczną i wirtualną przestrzeń produktu. Informacje zawarte w cyfrowych bliźniakach są napędzane przez przypadki użycia. Fizyczne obiekty produkcyjne są wirtualizowane i reprezentowane jako cyfrowe modele bliźniacze (awatary) zintegrowane zarówno w przestrzeni fizycznej, jak i cybernetycznej.

Cyfrowe bliźniaki są powszechnie dzielone na podtypy, które czasami obejmują: prototyp cyfrowego bliźniaka (DTP), instancję cyfrowego bliźniaka (DTI) i agregat cyfrowego bliźniaka (DTA).

- DTP składa się z projektów, analiz i procesów, które realizują fizyczny produkt. DTP istnieje, zanim powstanie produkt fizyczny.
- DTI jest cyfrowym bliźniakiem każdej indywidualnej instancji produktu po jego wytworzeniu. DTI jest powiązany ze swoim fizycznym odpowiednikiem przez pozostałą część życia fizycznego odpowiednika.
- DTA to agregacja DTI, których dane i informacje mogą być wykorzystywane do wyszukiwania informacji o fizycznym produkcie, prognozowania i uczenia się.

Produkcja

W procesie produkcyjnym cyfrowy bliźniak jest jak wirtualna replika zdarzeń zachodzących w fabryce w czasie zbliżonym do rzeczywistego. Tysiące czujników są umieszczane w całym fizycznym

procesie produkcyjnym, a wszystkie zbierają dane z różnych wymiarów, takich jak warunki środowiskowe, charakterystyka zachowania maszyny i wykonywana praca. Wszystkie te dane są stale przekazywane i gromadzone przez cyfrowego bliźniaka. Zaawansowane sposoby konserwacji i zarządzania produktami i zasobami są w zasięgu ręki, ponieważ istnieje cyfrowy bliźniak prawdziwej „rzeczy” z możliwościami w czasie rzeczywistym.

Urbanistyka i przemysł budowlany

Cyfrowe bliźniaki w budownictwie, tworząc dynamiczne cyfrowe repliki zasobów fizycznych, wspierają monitorowanie stanu, ocenę ryzyka i konserwację predykcijną konstrukcji, np. mostów i zabytków. Cyfrowe bliźniaki zostały też spopularyzowane w praktyce planowania urbanistycznego, w ramach koncepcji znanej jako inteligentne miasta. Te cyfrowe bliźniaki są często tworzone w formie interaktywnych platform do przechwytywania i wyświetlania danych przestrzennych w czasie rzeczywistym. Cyfrowe bliźniaki są też prezentowane jako metoda wizualnych inspekcji budynków i infrastruktury.

Opieka zdrowotna

Cyfrowe bliźniaki stanu pacjenta i wyrobów medycznych mogą wspierać spersonalizowane plany leczenia, pomagać zdalnie monitorować wskaźniki zdrowotne i służyć do przeprowadzania symulacji zabiegów chirurgicznych.

Energetyka

Cyfrowe bliźniaki elektrowni i systemów energii odnawialnej mogą m.in. przewidywać awarie wyposażenia i zarządzać stabilnością sieci.

Logistyka

Istnieje wiele ciekawych przykładów wykorzystania tej technologii w branży logistycznej: od optymalizacji procesów magazynowych, poprzez redukcję kosztów operacyjnych i eliminację błędów, aż po wydajne planowanie logistyczne. ■

M.U.





Współczesne wzmacniacze stereofoniczne, część 2

W poprzednim odcinku poświęconym wzmacniaczom stereofonicznym przedstawiliśmy dużą część ich wyposażenia – poczynając od najbardziej podstawowego, obejmującego wejścia dla źródeł analogowych, aż po przetworniki C/A, związane z wejściami cyfrowymi i transmisją Bluetooth.

To jednak nie wszystko, co w domenie cyfrowej może się dziać w nowoczesnym wzmacniaczu stereofonicznym, nawet gdy jego końcówki mocy są analogowego (i nie ma znaczenia, czy w klasie AB, czy w klasie D). Funkcje sieciowe znacznie wcześniej pojawiły się w amplitunerach wielokanałowych, ponieważ tą kategorią sprzętu zajmują się duże firmy japońskie, którym opłacało się inwestować w taki rozwój ze względu na ostrą konkurencję na dużym rynku urządzeń kina domowego. Wzmacniacze stereofoniczne musiały trochę poczekać... Ale się doczekały. W coraz większej liczbie modeli – chociaż na razie tych „nieco” droższych, to wcale nie bardzo drogie – instalowane są odtwarzacze strumieniujące. Ich możliwości też są różnicowane, ale standardem staje się obsługa Tidal Connect i Spotify Connect. Transmisję może zapewniać Wi-Fi albo LAN.

Oczywiście funkcje sieciowe wymagają obecności przetwornika C/A, więc każdy wzmacniacz tak

wyposażony ma przetwornik C/A, a skoro tak, to i wejścia cyfrowe, opisane w poprzednim odcinku.

Jeszcze innym cyfrowym dodatkiem do nowoczesnych wzmacniaczy stereofonicznych jest procesor DSP, pozwalający np. na korekcję charakterystyki. To z kolei wiąże się z bardzo tradycyjnymi elementami wyposażenia dawnych wzmacniaczy – analogowymi regulatorami „barwy”, czasami rozwiniętymi do małych, kilkupasmowych equalizerów... Audiofile wszelakich regulacji nie lubili i z części wzmacniaczy zniknęły, w trendzie szlachetnego minimalizmu i krótkiej ścieżki sygnału, albo co najmniej zostały uzupełnione o funkcję „direct”, która prowadziła sygnał z pominięciem regulacji. Dzisiaj regulacje mogą być oparte na nowoczesnych układach cyfrowych, przez to bardziej elastyczne i wszechstronne; mogą nawet realizować korekcję phono, a z pomocą mikrofonu pomiarowego prowadzić korekcję „akustyki pomieszczenia”, ale nie rozwiązuje to wszystkich problemów. Jeżeli sygnał dostarczamy z gramofonu, to zwykle chcemy, aby zachował analogową „naturę” przynajmniej na dalszej drodze do zespołów głośnikowych (choć większość współczesnych tłoczeń



Enigmatyczny, dyskretny Marantz M1 – porzucający klasyczne wzorce, redukujący liczbę wejść liniowych do jednego, ale dodający wyjście subwooferowe i bogatą sekcję cyfrowo-strumieniującą. Wejście optyczne, USB, HDMI z eARC, Bluetooth, a do tego dekodery Dolby Digital+ i układy dźwięku wirtualnego – sposób na zwiększenie atrakcyjności dwukanałowego kina domowego.
Test AUDIO 7–8/2024



Yamaha R-N2000A w spektakularny sposób łączy różne światy. Wygląd utrzymano w tradycyjnym stylu, jest pełne wyposażenie sekcji analogowej (wejście phono, wyjście słuchawkowe), regulatory barwy, loudness, front ozdobiono wskaźnikami VU... A jednocześnie wyposażenie w wejścia cyfrowe, funkcje strumieniujące i co unikalne w sprzęcie stereofonicznym – system kalibracji i korekcji akustyki.
Test AUDIO 11/2023



Japoński Technics też wrócił na rynek po dłuższej przerwie z szeroką gamą urządzeń. Ultranowoczesny i elegancki SU-G700M2 jest wyposażony kompleksowo, zarówno w wejścia cyfrowe (w tym USB), jak i analogowe (w tym dla gramofonu i wkładek MM/MC). Ponieważ jednak układ jest cyfrowy (również końcówki mocy; to coś bardziej cyfrowego niż popularna klasa D), lepiej więc z urządzeń cyfrowych (odtwarzacze sieciowe, CD) dostarczać sygnały cyfrowe, aby zapobiec podwójnej konwersji, bowiem sygnały z wejść analogowych i tak są konwertowane na cyfrowe. Test AUDIO 6/2023

plyt powstaje z materiału, który przeszedł cyfrową obróbkę). Nawet jeżeli sygnał dostarczamy z odtwarzacza CD, z jego wyjść analogowych, to konwersja na cyfrę, w celu przeprowadzenia na sygnale nawet pożytecznych operacji, a następnie z powrotem na analog, aby dostarczyć go do końcówek mocy, na zdrowie sygnałowi nie wychodzi; wiedząc, że tak działa nasz wzmacniacz, lepiej sygnały dostarczać drogą cyfrową, aby nie narażać ich na niepotrzebnie wiele konwersji. Niektóre wzmacniacze wyposażone w tak działające sekcje cyfrowe mają funkcje pozwalające omijać je dla wybranych w menu wejść analogowych.

Niekiedy dodatkowe wyposażenie – przetwornik C/A, odtwarzacz sieciowy czy przedwzmacniacz phono – jest montowane opcjonalnie, za pomocą modułów (kart) wsuwanych w kieszenie przygotowane na tylnej ścianie wzmacniacza. Pozwala to uniknąć zakupu wzmacniacza nadmiernie wyposażonego względem indywidualnych potrzeb, a więc potencjalnie droższego, z możliwością późniejszego rozszerzenia jego funkcjonalności. To jednak praktyka dość rzadka, bowiem komplikująca konstrukcję, logistykę... a więc tą drogą podnosząca koszty. Również obsługa wzmacniaczy staje się coraz bardziej skomplikowana wraz z rozwijaniem ich wyposażenia, a pomysły jej unowocześniania za pomocą sieci i aplikacji często tylko pogarszają sprawę. Stąd wciąż dużym zainteresowaniem cieszą się wzmacniacze nawet dość skromne, ale proste, intuicyjne w obsłudze.

Wzmacniacze różnią się między sobą jak nigdy wcześniej, pod tym hasłem spotkamy urządzenia o funkcjach daleko wykraczających poza ich tradycyjną rolę, a z drugiej strony nadawane są im czasami



Wzmacniacz NAD C389 już w wersji bazowej jest bogato wyposażony, zarówno w sekcji analogowej (wejście gramofonowe, wyjście subwooferowe, słuchawkowe), jak i cyfrowej (HDMI z eARC, ale bez USB). Funkcje sieciowe (pod patronatem systemu BluOS-D) i korekcję akustyki (Dirac) można dodać za pomocą modułu MDC2. Podobnie wyposażony jest tańszy C379 i droższy C399 (a różni się przede wszystkim mocą wyjściową). Test AUDIO 12/2024

dość wymyślne nazwy, mające uchwycić ich wszechstronność. Dlatego dzisiaj każdy zainteresowany budową systemu hi-fi musi najpierw zdać sobie sprawę, jakiego dokładnie urządzenia potrzebuje, a potem odnaleźć go wśród wielu propozycji. Wtedy też można zupełnie zmienić koncepcję... Widząc, że są urządzenia, które zapewniają znacznie większe możliwości, niż się tego spodziewaliśmy. Albo odwrotnie – zniechęcić się do przekombinowanych, „wszystkomających” wzmacniaczy i wrócić na twardy grunt prostych rozwiązań.

Kolejny wymiar różnorodności współczesnych wzmacniaczy (i nie tylko) związany jest z modą „vintage”. Mogłoby się wydawać, że oznacza to rewitalizację starych układów. Jednak bywa różnie. Niektóre modele rzeczywiście starają się trzymać jak najbliżej pierwowzoru, używając z konieczności aktualnie dostępnych komponentów, ale inne są wewnątrz zupełnie innymi konstrukcjami, tyle że w stylowym „opakowaniu”. Warto też wspomnieć, że dla niektórych firm, zwłaszcza high-endowych, ich wygląd stał się tak silnym elementem identyfikacji, że aby



Firma Hi-Fi Rose rozpoczęła karierę nowoczesnymi streamerami; potem połączyła te kompetencje z innymi funkcjami i z końcówką mocy własnego pomysłu (odmianą klasy D). Tak powstał RS520, wzmacniacz o potężnych możliwościach w zakresie strumieniowania (również materiałów wideo), regulacji charakterystyki, sterowania aplikacją, w związku z tym określany mianem „all-in-one”. Jednak nie tak dostownie „wszystko”... bowiem nie ma wejścia gramofonowego ani wyjścia słuchawkowego. Test AUDIO 4/2023



Wygląd tego wzmacniacza cofa nas wstecz o co najmniej 50 lat... kiedy działała brytyjska firma Leak, a przestała, na ponad 40 lat, w roku 1979. Reaktywowana, przygotowała ofertę transportu CD i dwóch wzmacniaczy w stylu retro, jednak z nowoczesnym wyposażeniem. Model Stereo 230, oprócz kilku wejść liniowych i gramofonowego (MM), ma paletę wejść cyfrowych, z USB i HDMI ARC. Nie zabrakło też wyjścia słuchawkowego. Test AUDIO 2/2025



Japońska firma Accuphase, działająca nieprzerwanie od pół wieku, utrzymuje charakterystyczny, klasycznie luksusowy wygląd swoich bezkompromisowych urządzeń. E-700 to jeden z najlepszych wzmacniaczy zintegrowanych; w wersji bazowej, poza baterią wejść liniowych, wyposażony jest w wyjście słuchawkowe, a moduły z wejściami cyfrowymi i korekcją phono można dodatkowo zainstalować. Test AUDIO 2/2025

nie stracić oryginalności i nie zrazić wiernych wielbicieli, nie zmienia się on od dziesiątek lat, nie poddając żadnym modom.

Jaka jest przyszłość wzmacniaczy? Pewnie mocno związana z systemami cyfrowymi, chociaż nie tylko do nich ograniczona. Głębokie przywiązanie audiofilów do „analogu” zapewni jeszcze długie życie

konwencjonalnym wzmacniaczom analogowym, chociaż... znaczenie słowa „konwencjonalny” może (i już tak jest) zmieniać się z czasem i będzie oznaczać właśnie konstrukcję w jakimś stopniu cyfrową, podczas gdy w pełni analogowa stanie się czymś bardzo oryginalnym. ■

Andrzej Kisiel

AUDIO

Miesięcznik audiofilski – polski przedstawiciel European Imaging and Sound Association

Magazyn Audio, będący na rynku już ponad 25 lat, jest liderem krajowej prasy audiofilskiej i polskim przedstawicielem EISA. Miesięcznik poświęcony sprzętowi audio-video zawiera przede wszystkim wyczerpujące, bogato zilustrowane testy urządzeń (w tym próby odsłuchowe i pomiary), jak zespoły głośnikowe, odtwarzacze CD, DVD, kina domowego i przenośne, wzmacniacze stereo, słuchawki i wiele innych.

Na łamach Audio pojawiają się także nowości sprzętowe, prezentacje firm, wywiady. W każdym numerze znajduje się kilkadziesiąt recenzji nowości płytowych.



przeglądaj, czytaj i kup na
www.ulubionykiosk.pl



WYDANIA

SPECJALNE



Rozwiń szachowe umiejętności
z wydaniem specjalnymi Młodego
Technika pt.: „Kocham Szachy”!

160 stron wiedzy i inspiracji w każdym e-wydaniu.
Od strategii oraz podstaw, po historie szachowych
geniuszy. Zobacz porady, analizy i zadania, które
przeniosą Twoją grę na wyższy poziom.



Zamów na
UlubionyKiosk.pl