



nr 9. wrzesień 2023

e-suplement www.mt.com.pl



Tu przejrzysz
i kupisz ten numer

NEWS 24/7
przeglądaj codziennie
na swoim smartfonie

mlody **m.technik**

Ciekawi świata są zawsze młodzi

WOJNA O CHIPY

Krzemowa władza nad światem

KONIEC I CO DALEJ
Piloty do TV i innych urządzeń

**Zaprenumeruj „Młodego Technika”,
a zawsze dostaniesz najnowszy numer
wprost do Twojej skrzynki!**



**do 6* wydań
gratis!**

* Cena prenumeraty rocznej wynosi 163,90 zł.

Przy zamówieniu prenumeraty dwuletniej w cenie 268,20 zł

oszczędność wynosi równowartość sześciu wydań „Młodego Technika”

Wszystkie opcje prenumeraty i e-prenumeraty znajdziesz na stronie

www.UlubionyKiosk.pl

prenumerata@avt.pl

AVT-Korporacja sp. z o.o., ul. Leszczynowa 11, 03-197 Warszawa

konto 18 1050 1012 1000 0024 3173 1013

eprasa.pl 86a0ce6243



Temat okładkowy

Wprowadzane przez USA sankcje na dostawy do Chin najbardziej zaawansowanych komponentów elektronicznych, zwłaszcza sprzętu na którym pracuje i rozwija się AI, zmieniają reguły gry, której wynik nie jest pewny i wcale nie oczywisty.

Prenumerata
dla szkół i placówek
oświatowych
30% taniej!

Prenumerata roczna
(12 wydań) MT
w wersji drukowanej
kosztuje 125,20 zł

Roczny dostęp online
kosztuje 100,00 zł

Zamów na
[www.UlubionyKiosk.pl/](http://www.UlubionyKiosk.pl/prenumerata/szkolna)
prenumerata/szkolna

Wojny krzemowe nabierają tempa

Sankcje nałożone przez USA na Chiny, którym rząd amerykańskim chce odciąć dostęp do najbardziej zaawansowanych rozwiązań krzemowych, spotkały się działaniami odwetowymi ze strony Pekinu, budząc przy okazji wiele innych krajów i regionów, które przypomniwały sobie, że warto mieć, i to najlepiej na swoim terenie, przemysł procesorowy. Przypomniwała sobie o swoich krzemowych tradycjach także Europa. My mieliśmy niedawno chwilę radości, gdy okazało się, że zakład produkcyjny wybuduje w Polsce Intel.

„Wojna o chipy” nakłada się na szerszy kontekst rywalizacji mocarstw i groźby gorącej wojny o Tajwan, który tak się

*Okazuje się,
że najlepiej mieć
własną produkcję
procesorów*

składa jest światowym centrum produkcyjnym komponentów elektronicznych. Teoretycznie, gdyby Chiny dokonały inwazji na wyspę, którą uważają za swoją zbuntowaną prowincję, mogłyby przejąć za-

kłady TSMC, wytwarzające większość komponentów elektronicznych na potrzeby świata, ale każdy chyba rozumie, że to tak nie działa.

Tak czy inaczej, nie tylko Stany Zjednoczone, ale wiele innych krajów podjęło w ostatnim czasie działania, aby źródła dostaw strategicznego towaru, jakim są dziś komputerowe komponenty krzemowe, „zdywersyfikować” a najlepiej mieć produkcję i know-how na własnym terenie.

Obok gry geopolitycznej toczy się wyścig technologiczny, np. o to, kto ostatecznie najwięcej skorzysta na eksplozji AI. Na razie głównym dostawcą bazy sprzętowej jest Nvidia, której wartość rynkowa przekroczyła niedawno bilion dolarów. Jednak karty graficzne tej firmy nie są tanie. Niewykluczone więc, że, jeśli komuś uda się opracować inne, odpowiednio wydajne a przy tym tańsze układy obsługujące uczenie maszynowe, to sytuacja może się zmienić.

Choć dużo się wydarzyło, rozgrywka, jak się wydaje, dopiero się zaczęła. Wojny krzemowe mogą być jednym z głównych wydarzeń obecnej dekady.

Mirosław Usidus

Spis treści

Temat numeru:

Wojna o chipy. Krzemowa władza nad światem

- 24 • Czy AI jest skazane na karty graficzne? Nvidia i długo nic... ale nadzieja w rywalach nie gaśnie
- 28 • Dokąd zmierzamy idąc krzemową doliną i poza nią? Egzotyka nieograniczonych możliwości
- 33 • Komputery jakie znamy mogą nie przetrwać rewolucji w technice i nauce. Czy to już kres świata zer i jedynek?
- 38 • W świecie półprzewodników trwa bezwzględna rywalizacja mocarstw i firm. Wojna o chipy

Technika

- 8 Info Zoom
- 16 Dodaj do obserwowanych Horyzonty mgłą spowite
- 17 • Kiedy wreszcie nadejdą dostawy dronami? Marzenia sprowadzone na ziemię
- 20 • Czy elektroluminescencyjne kropki kwantowe zdobędą ekrany i wyświetlacze? Kwantowe maleństwo zamiast LED i LCD
- 22 • AI uczona na danych wygenerowanych przez AI. Wąż zjadający własny ogon
- 45 Raport MT: Luki z historii Ziemi. Planeta dotknięta amnezją

m.technik

- 56 Mobilne aplikacje. Test aplikacji: Programy do obsługi wideo i multimedialnych

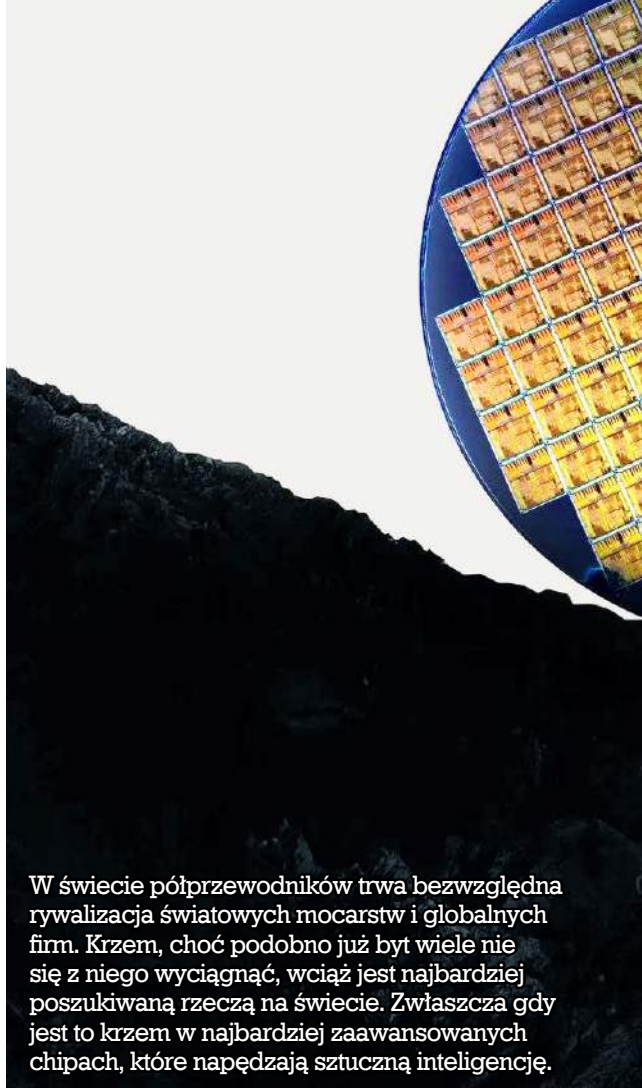
Szkoła

- 58 Chemia inna niż w szkole: Wcale nie takie rzadkie (1)
- 62 Fizyka bez granic: Mikroskop sił atomowych (1). Zobaczyć atomy
- 64 MT studiuj: Informatyka
- 66 Matematyka z ludzką twarzą: Felga z Bielska i powrót do przeszłości
- 71 Edukacja przez szachy: Szachowa dyplomacja Benjamina Franklina
- 77 Koniec i co dalej: Piloty do TV i innych urządzeń. Czy zmienimy władzę?
- Klub i Szkoła Wynalazców
- 80 • Szkoła Wynalazców, dozwolone do lat 15
- 81 • Klub Wynalazców, bez ograniczeń wieku
- 82 • Vademecum Młodego Wynalazcy
- 85 Pomysły genialne, zwirowane i takie sobie
- 86 Na warsztacie. Elektronika dla Ciebie: Płynny regulator prędkości i kierunku obrotów silnika 12 V
- 89 Konkurs „As Majsterkowania”
- Odkryj historię wynalazców
- 90 • Śruby i gwinty
- 94 • Rodzaje śrub i gwintów

Hobby

- 95 Akademia audio: Gorące czary klasy A. Musical Fidelity A1
- 97 Fotografia: Kreatywne fotografowanie

- 2 Prenumerata
- 3 Od wydawcy
- 6 Listy, Facebook
- 99 Sędziwy Technik – 100 lat temu prasa pisała



W świecie półprzewodników trwa bezwzględna rywalizacja światowych mocarstw i globalnych firm. Krzem, choć podobno już być wiele nie się z niego wyciągnąć, wciąż jest najbardziej poszukiwaną rzeczą na świecie. Zwłaszcza gdy jest to krzem w najbardziej zaawansowanych chipach, które napędzają sztuczną inteligencję.

W tym wydaniu MT m.in.:

- **Horyzonty mgłą spowite: Kiedy wreszcie nadejdą dostawy dronami?**
Amazon nie wydaje się obecnie bliski uruchomienia usługi transportu zamówionych zakupów i paczek na większą skalę. Może uda się to innym firmom?
- **Koniec i co dalej: Piloty do TV i innych urządzeń**
Kto je miał – ten miał władzę. Teraz nadchodzi zmiana władzy.
- **Test aplikacji: Programy do obsługi wideo i multimedialnych**

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

• Adres wydawnictwa:
03-197 Warszawa, ul. Leszczyńska 11,
tel. 22 257 84 98, 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.
e-mail: dtp@mad.media.pl

• Konsultacja graficzna:
Małgorzata Jabłońska

• Dział Reklam:
e-mail: reklama@mt.com.pl

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

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

• Prenumerata w RUCH S.A.
www.prenumerata.ruch.com.pl
lub tel. 801 800 803, 22 717 59 59
e-mail: prenumerata@ruch.com.pl

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



Krzemowa władza nad światem 24



Planeta dotknięta amnezją **45**

List miesiąca



Poszukiwanie życia pozaziemskiego

Piszę do Państwa pod wpływem tego, co niedawno przeczytałem w „Młodym Techniku”, chcąc przedstawić garść moich przemyśleń na temat kwestii poszukiwania życia poza Ziemią.

Jak wiemy, poszukiwania SETI prowadzone są od ponad 50 lat, ale jak dotąd nie udało się znaleźć żadnych jednoznacznych dowodów na istnienie życia poza Ziemią. Poszukiwania SETI mogą pomóc nam zrozumieć, czy jesteśmy sami we Wszechświecie i czy istnieje inne życie, które jest podobne do naszego. Poszukiwania SETI mogą również pomóc nam zrozumieć, jak powstało życie na Ziemi i jakie są warunki niezbędne do powstania życia.

Istnieje wiele różnych metod poszukiwania życia poza Ziemią. Najczęściej wykorzystywaną metodą jest poszukiwanie sygnałów radiowych wysyłanych przez inteligentne cywilizacje. Inne metody poszukiwania życia poza Ziemią obejmują poszukiwanie mikroorganizmów na Marsie i innych planetach, poszukiwanie śladów życia w kometach i meteoroidach oraz poszukiwanie biomolekuł w atmosferze planet pozasłonecznych.

Mimo wielu wysiłków naukowców, jak dotąd nie udało się znaleźć żadnych jednoznacznych dowodów na istnienie życia poza Ziemią. Jednak brak dowodów nie oznacza braku życia. Możliwe jest, że życie poza Ziemią istnieje, ale jest na tyle słabo zauważalne, że nie jesteśmy w stanie go wykryć.

W ostatnich latach odkryto wiele planet pozasłonecznych, czyli planet krążących wokół gwiazd innych niż Słońce. Niektóre z tych planet znajdują się w strefie zamieszkalnej, czyli w obszarze wokół gwiazdy, w którym może występować woda w stanie ciekłym. Woda jest niezbędna do powstania życia, jak znamy je na Ziemi.

Oznacza to, że istnieje wiele potencjalnych miejsc, w których może istnieć życie poza Ziemią. Jednak jak dotąd nie udało się znaleźć żadnych jednoznacznych dowodów na istnienie życia poza Ziemią.

Istnieje kilka powodów, dla których tak trudno znaleźć dowody na istnienie życia poza Ziemią. Po pierwsze, życie może być bardzo rzadkie we Wszechświecie. Po drugie, , nawet jeśli jesteśmy w stanie wykryć życie, może ono być tak inne od życia na Ziemi, że nie będziemy w stanie go rozpoznać.

Mimo tych trudności, wierzę, że poszukiwania życia poza Ziemią powinny być kontynuowane. Poszukiwania SETI to bardzo ważne zadanie, które może mieć znaczący wpływ na nasze zrozumienie wszechświata i naszego miejsca w nim.

Jeśli kiedykolwiek uda nam się znaleźć dowody na istnienie życia poza Ziemią, będzie to jedno z najważniejszych odkryć w historii ludzkości. To odkrycie zmieniłoby nasze postrzeganie Wszechświata i naszego miejsca w nim.

Edward Turczynowski, Jurata

Tajemnice protonu

Dziękuję za zajęcie się na łamach waszego pisma tak fascynującym obiektem jakim jest proton, podstawowy składnik materii, o którym im więcej wiemy tym mniej rozumiemy.

Choć od dawna znany, do dzisiaj stanowi zagadkę dla fizyków. Proton składa się bowiem z kwarków – cząstek elementarnych o ułamkowym ładunku elektrycznym. W protonie znajdują się dwa kwarki górne i jeden kwark dolny. Kwarki te oddziałują ze sobą za pośrednictwem gluonów, cząstek przenoszących oddziaływanie silne. To właśnie silne oddziaływania sprawiają, że kwarki są na stałe uwięzione w protonie.

Choć znany jest skład protonu, wciąż trwają badania nad poznaniem dokładnej struktury tej cząstki. Najnowsze eksperymenty pokazują, że kwarki nie są rozmieszczone symetrycznie, a ich ruchy są skorelowane w złożony sposób. Co więcej, rozkład ładunku elektrycznego wewnątrz protonu też jest asymetryczny. Wciąż wiele pytań pozostaje bez odpowiedzi.

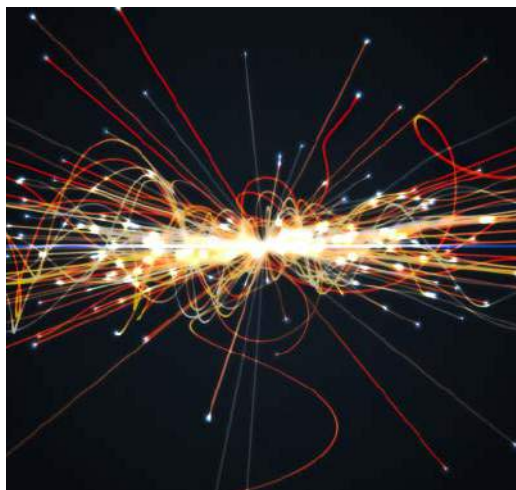
Oprócz składu i struktury, kolejnym fascynującym zagadnieniem związanym z protonem jest jego spin, czyli moment pędu. Wiadomo, że proton składa się z cząstek o spinie $1/2$ (kwarków i gluonów), ale całkowity spin protonu wynosi $1/2$. Nie jest jasne, w jaki sposób spiny cząstek składowych sumują się, by dać spin całkowity protonu.

Aby to wyjaśnić, fizycy wprowadzili pojęcie spinu kwarków i gluonów. Okazuje się, że tylko niewielka część całkowitego spinu protonu pochodzi od spinów kwarków. Pozostała część jest wynikiem ruchów kwarków i gluonów wewnątrz protonu. Obecnie trwają eksperymenty mające na celu dokładne zmierzenie wkładu poszczególnych składników do całkowitego spinu.

Oprócz zagadnień związanych ze strukturą i spinem protonu, istotnym problemem pozostaje kwestia jego masy. Zgodnie z teorią, proton powinien ważyć znacznie mniej niż wynika z pomiarów. Rozbieżność pomiędzy przewidywaniami teoretycznymi a wartością eksperymentalną nazywana jest problemem masy protonu.

Skąd bierze się nadmiar masy protonu? Otóż, obliczając masę z pierwszych zasad, uwzględnia się jedynie masy kwarków i energii wiązania gluonów. Jednak w rzeczywistości, duży wkład pochodzi od energii kinetycznej kwarków oraz wirtualnych par cząstka-antycząstka, które powstają w wyniku fluktuacji kwantowych. Te dodatkowe, dynamiczne czynniki są trudne do oszacowania teoretycznie.

Aby dokładniej zrozumieć skąd bierze się masa protonu, fizycy opracowują coraz doskonalsze modele



oraz przeprowadzają precyzyjne obliczenia z zakresu chromodynamiki kwantowej. Dzięki gigantycznej mocy obliczeniowej superkomputerów, jesteśmy w stanie coraz lepiej symulować złożone procesy zachodzące wewnątrz protonu.

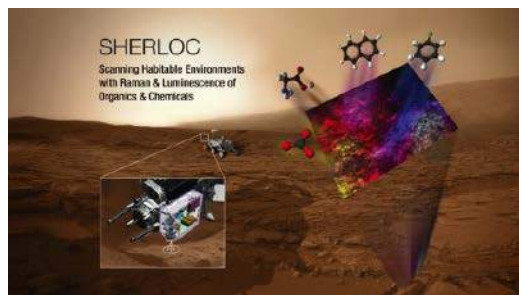
Oprócz zagadnień teoretycznych, istotnym elementem badań struktury protonu są również eksperymenty. W ciągu ostatnich dekad fizycy opracowali wiele pomysłów metod badania wnętrza protonu.

Jedną z nich jest rozpraszanie głęboko nieelastyczne elektronów na protonach przy bardzo wysokich energiach. Pozwala to badać rozkład kwarków i gluonów w protonie. Inną metodą są zderzenia protonów z antyprotonami, w wyniku których powstają liczne cząstki umożliwiające rekonstrukcję struktury protonu. Nowe możliwości dają też kolizje protonów przyspieszanych w Wielkim Zderzaczu Hadronów (LHC).

Eksperymenty dostarczyły wielu fundamentalnych odkryć, jak asymetria rozkładu kwarków czy niewielki udział spinu kwarków w całkowitym spinie protonu. Jednak wciąż pozostaje wiele niewiadomych. Plany budowy bardziej zaawansowanych akceleratorów oraz detektorów pozwolą jeszcze głębiej zajrzeć w strukturę protonu.

Jak widać, struktura i własności protonu skrywają jeszcze wiele tajemnic. Zrozumienie budowy tej podstawowej cząstki ma kluczowe znaczenie dla całej fizyki cząstek elementarnych. Dalsze badania z pewnością przyniosą jeszcze wiele odkryć i pozwolą lepiej poznać fundamentalne prawa natury rządzące mikroświatem.

Tomasz Nowak, Świecko



PLANETY

Różnorodność substancji organicznych na Marsie

Cząsteczki organiczne o niespotykanej dotychczas różnorodności zostały znalezione na Marsie przez łazik Perseverance w kraterze Jezero na Marsie. Badacze nie potrafili wykluczyć biologicznego pochodzenia wykrytych substancji. Nie mogą też wykluczyć, że dowodzą one istnienia życia na Marsie, choć uznają za prawdopodobne, że trafiły na czerwoną planetę z zewnątrz.

Odkrycia, opisanego w „Nature”, dokonał instrument o nazwie SHERLOC (skrót od angielskiego terminu – Scanning Habitable Environments with Raman and Luminescence for Organics and Chemicals). Pochodzenie organicznych związków chemicznych opartych na węglu nie jest w tej chwili w pełni wyjaśnione, stąd mnogość hipotez.

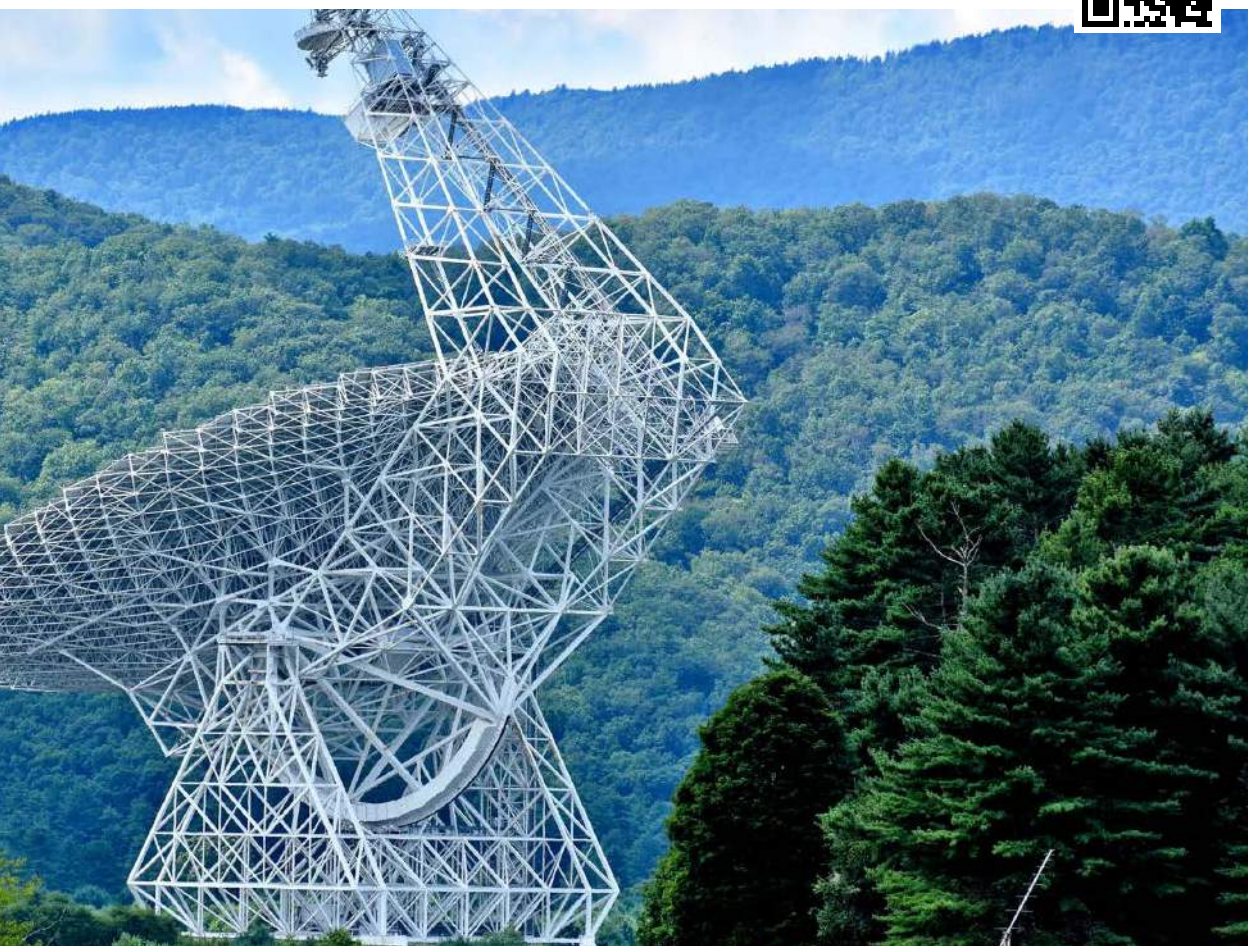
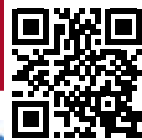
Do tej pory organiczne związki węglowe były wykrywane wcześniej przez lądownik Mars Phoenix i łazik Mars Curiosity przy użyciu zaawansowanych technik, takich jak analiza gazów ewolucyjnych i chromatografia gazowa ze spektrometrią mas. Nowe badanie wprowadza inną technikę, która identyfikuje związki organiczne na Marsie. ■

20 km wysokości
liczy sobie Verona Rupes,
najwyższy ze znanych klifów
w Układzie Słonecznym,
znajdujący się na Mirandzie,
księżycu Urana.



W ramach projektu badawczego trwającego prawie od dwu dekad, kilka zespołów naukowców na całym świecie niezależnie od siebie wykryło ten sam sygnał pochodzący od gwiazd zwanych pulsarami, który wskazuje na istnienie gigantycznych fal grawitacyjnych, przetaczających się przez galaktykę. Uczeni zakładają, że tworzą one tak zwane tło fal grawitacyjnych, wypełniający czasoprzestrzeń „szum” powstały w wyniku skumulowanych zderzeń zachodzących w całym kosmosie. Według badaczy, każda z tych fal powinna mieć w przybliżeniu długość całego naszego Układu Słonecznego, co, wbrew pozorom, czyni je znacznie trudniejszymi do wykrycia niż dotychczas znane fale grawitacyjne.

Ogłoszony wynik badań to rezultat prac w ramach projektu znanego jako North American Nanohertz Observatory for Gravitational Waves (NANOGrav). Od 2004 roku zaangażowani w NANOGrav naukowcy z różnych stron świata monitorują regularne błyski światła pochodzące z rozciągającej się w Drodze Mlecznej sieci pulsarów. Wszystkie grupy korzystały z czułych radioteleskopów do monitorowania pulsarów „milisekundowych”. Za każdym razem, gdy pulsar obraca się wokół własnej osi, jego wiązka

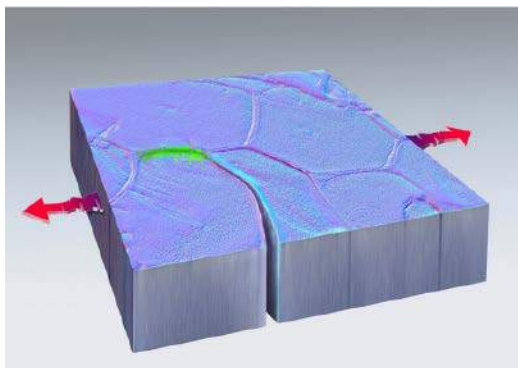


ASTRONOMIA

Fale grawitacyjne jako „czasoprzestrzenne tło” Wszechświata

radiowa wędruje w kierunku Ziemi i dalej. Są to emisje w regularnych odstępach czasu. Pulsary milisekundowe obracają się najszybciej, do kilkuset razy na sekundę. Jako że „fale tła”, o których mowa różnią się od tych, które zostały wcześniej wykryte przez Laser Interferometer Gravitational-Wave Observatory (LIGO) i inne ziemskie detektory, potrzeba było wielu lat żmudnych obserwacji drobnych ruchów pulsarów, aby je wykryć.

„Szum” fal grawitacyjnych przenikający Wszechświat pochodzi z wielu wydarzeń, tych o których wiemy z ze znanych detekcji „mniejszych” fal dokonanych przez LIGO i docenionych już Nagrodą Nobla, ale również wielu innych wydarzeń, których natury nie znamy i być może nawet nie podejrzewamy, że kosmosie zachodzą. Naukowcy określają nowe odkrycie jako kolejny krok na drodze rozwoju tzw. astronomii fal grawitacyjnych. ■



TECHNIKA MATERIAŁOWA

Metal regeneruje się sam – pierwsza tego rodzaju obserwacja

Proces „samo-regeneracji” uszkodzenia w metalu zaobserwowali naukowcy z amerykańskich laboratoriów Sandia i teksaskiego Uniwersytetu A&M. Podczas testowania odporności metalu, z wykorzystaniem specjalistycznej techniki transmisyjnego mikroskopu elektronowego zauważyli, że w bardzo małych skalach, w kawałku platyny o grubości 40 nanometrów zawieszonym w próżni, dochodzi do samo-naprawy uszkodzeń.

W próbie pojawiły się wywołane powtarzającymi się naprężeniami pęknięcia, znane jako uszkodzenia zmęczeniowe. Po około 40 minutach obserwacji, w miejscu pęknięcia w platynie zaczęło się łączenie struktury. I choć jest to obserwacja niezwykle, dla której nie ma w tej chwili naukowego wyjaśnienia, nie była całkowicie nieoczekiwana dla uczestników eksperymentu. W 2013 r. Michael Demkowicz z Uniwersytetu A&M, przewidywał w swojej pracy, że tego rodzaju „samo-leczenie” nano-pęknięć może zachodzić, napędzane przez drobne ziarna krystaliczne metali, które przesuwają się w odpowiedzi na nacisk.

Badacze, którzy opisali swoje wyniki na łamach „Nature”, sugerują, że mogą w tym przypadku zachodzić zjawiska znane z procesów określanych jako „spawanie na zimno”. Zazwyczaj cienkie warstwy powietrza lub zanieczyszczeń zakłócają ten proces. Jednak w warunkach próżni, czyste powierzchnie metali mogą znaleźć się na tyle blisko siebie, że dochodzi do „sklejania”. ■



ROBOTYKA

Robot w sferycznej klatce ma toczyć się, latać i eksplorować niedostępne zakamarki

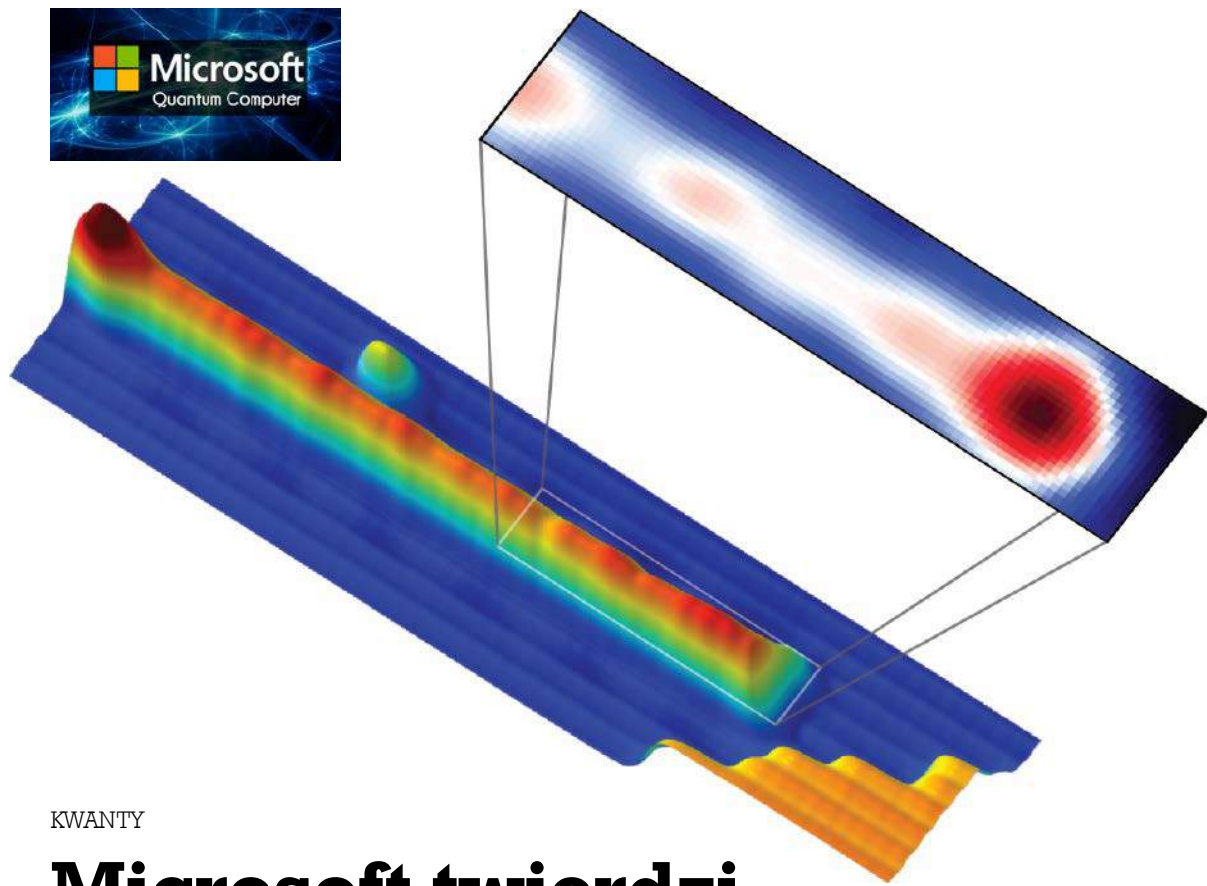
Firma Revolute Robotics z Arizony zaprezentowała konstrukcję autonomicznego robota Hybrid Mobility Robot (HMR), którego główna konstrukcja to wirująca, sferyczna klatka, która może latać jak multicopter lub toczyć się w dowolnym kierunku z wykorzystaniem dwóch żyroskopowo mocowanych pierścieni.



Reportaż o robocie firmy Revolute Robotics:
<https://tiny.pl/cw2z1>

HMR cechuje możliwość odkształcania się do pewnego stopnia, co pozwala amortyzować uderzenia podczas lądowania, chroniąc elektronikę przez izolowanie jej od nadmiernych wibracji. Konstrukcja klatkowa zapewnia również przestrzeń dla układu napędowego służącego do lotu. Jednocześnie osłania i izoluje ruchome elementy, takie jak śmigła, od otoczenia, co zwiększa bezpieczeństwo np. dla przebywających w pobliżu ludzi.

Revolute, założona przez absolwentów Uniwersytetu w Arizonie chciałaby sprzedawać w pełni rozwiniętego HMR jako autonomicznego robota inspekcyjnego do ograniczonych przestrzeni, np. rurociągów. Zakłada się, iż tryb lotu pozwoliłby mu łatwo wspinać się po pionowych odcinkach rur w sposób, w jaki nie mogą tego robić roboty na kołach. Mówi się też o możliwym wykorzystaniu robota w warunkach kopalni lub w wojsku. ■



KWANTY

Microsoft twierdzi, że jest bliski zbudowania stabilnego komputera kwantowego

Zespół naukowców z zespołu badawczego Microsoft Quantum w artykule opublikowanym w czasopiśmie „Physical Review B” opisał, jak to określa „kamień milowy” w kierunku budowy niezawodnego i praktycznego komputera kwantowego. Chodzi o budowę kubitów o dużej stabilności z wykorzystaniem tzw. fermionów Majorany, które są same swoimi antycząstkami. Po raz pierwszy wykryli je trzy lata temu naukowcy z MIT na powierzchni złota.

W opisie badań zespołu czytamy, że jego urządzenia „są wytwarzane z dwuwymiarowych gazów elektronowych o wysokiej ruchliwości, w których funkcjonują elektrostatyczne bramki. Urządzenia te umożliwiają pomiary lokalnych i nielokalnych właściwości transportowych”. Ogólnie rzecz biorąc, nie wchodząc w dość skomplikowaną fizykę, działanie układu opiera się tu na zmianach

fazy kwantowej topologicznej nadprzewodnictwo, co jest podstawą zapisu i sterowania kubitami. Naukowcy z Microsoftu stworzyli topologiczny nadprzewodnik, wykorzystując nadprzewodnik aluminiowy i półprzewodnik z arsenku indu. Ich urządzenie przeszło rygorystyczny protokół, sugerując wysokie prawdopodobieństwo występowania zerowych stanów Majorany.

W swoim komunikacie grupa badaczy z Microsoft informuje również o stworzeniu nowej miary do pomiaru wydajności superkomputera kwantowego – liczby niezawodnych operacji kwantowych na sekundę (rQOPS). Według nich, aby maszyna kwalifikowała się jako superkomputer kwantowy, jej rQOPS musi wynosić co najmniej milion. Osiągnięcie wartości miliardowych czyni, ich zdaniem, układ kwantowy prawdziwie praktycznym komputerem. ■



LASERY

Mały laser z wielką mocą

Laser półprzewodnikowy zdolny do cięcia metali to przełom, który udało się osiągnąć, według publikacji w „Nature”, zespołowi japońskich badaczy z Uniwersytetu w Kioto. Lasery półprzewodnikowe, w porównaniu z laserami gazowymi i światłowodowymi, są niewielkie, energooszczędne i wysoce sterowalne. Jednak dotychczas nie osiągały naprawdę dużych mocy.

Wygląda na to, że japońska innowacja może to zmienić. Zespół pod kierownictwem Susumu Nody znalazł sposób na przezwycięzenie ograniczeń laserów półprzewodnikowych przez zmianę struktury fotoniczno-krystalicznych laserów emitujących powierzchniowo (PCSEL). Kryształ fotoniczny składa się z arkusza półprzewodnika przebitego regularnie ułożonymi, nanometrowymi otworami wypełnionymi powietrzem. Lasery fotoniczno-krystaliczne nie są całkowitą nowością, ale do tej pory inżynierowie nie byli w stanie skalować ich tak, aby dostarczały wiązki o mocy wystarczającej do praktycznego cięcia i obróbki metali.

Nowy laser ma moc wyjściową 50 watów. Jego jasność, wyrażana w jednostce 1 GW/cm²/str, jest, jak głosi publikacja, wystarczająco wysoka do zastosowań obecnie zdominowanych przez nieporęczne lasery gazowe i światłowodowe, w tym precyzyjnej inteligentnej produkcji w przemyśle elektronicznym i motoryzacyjnym. Jest również wystarczająco wysoka dla bardziej egzotycznych zastosowań, takich jak komunikacja satelitarna i napędy. ■



Prezentacja skuteczności lasera PCSEL w cięciu stali nierdzewnej: <https://tiny.pl/cw235>



SIEĆ

Standard internetu świetlnego

Organizacja IEEE (skrót od ang. Institute of Electrical and Electronics Engineers) opublikowała standard 802.11bb dla LiFi, opartego na falach świetlnych bezprzewodowego systemu transmisji. Nowy standard toruje drogę różnym producentom do tworzenia kompatybilnych ze sobą urządzeń LiFi. Technika transmisji oparta na świetle rozwijana była od lat przez wiele ośrodków badawczych i firm a „Młody Technik” pisał o tym wiele razy. Teraz doczekała się standaryzacji.

W miejsce fal radiowych, technika Light Fidelity (LiFi) wykorzystuje do przesyłu danych światło ze zwykłych lamp LED, które trafia do odbiorników rejestrujących fotony i konwertujących je na bity informacji. Dane zawarte są w modulacji sygnału świetlnego czyli migotaniu, które jednak nie jest widoczne dla użytkowników, ponieważ występuje ono przy częstotliwościach powyżej 60 Hz – zbyt szybko, aby ludzkie oko mogło je dostrzec. Potencjalnie sygnały LiFi mogą osiągać stukrotnie większą przepustowość w porównaniu z WiFi, nawet do 224 GB/s.

Firmy pracujące nad tą technologią, w tym pureLiFi, Fraunhofer HHI i Philips, zintegrowały ją z systemami oświetleniowymi, dzięki czemu urządzenia mogą odbierać Internet za pośrednictwem lamp sufitowych w domach lub biurach. Fraunhofer HHI zaproponował wykorzystanie LiFi do usprawnienia transportu poprzez transmisję za pomocą lamp ulicznych, światel stopu i reflektorów pojazdów, potencjalnie umożliwiając komunikację między pojazdami. Jednak, aby sygnał LiFi był odbierany nadajnik i odbiornik muszą się „widzieć” w sensie dosłownym. Inną wadą jest wymaganie specjalnych, innych niż znane odbiorników. ■



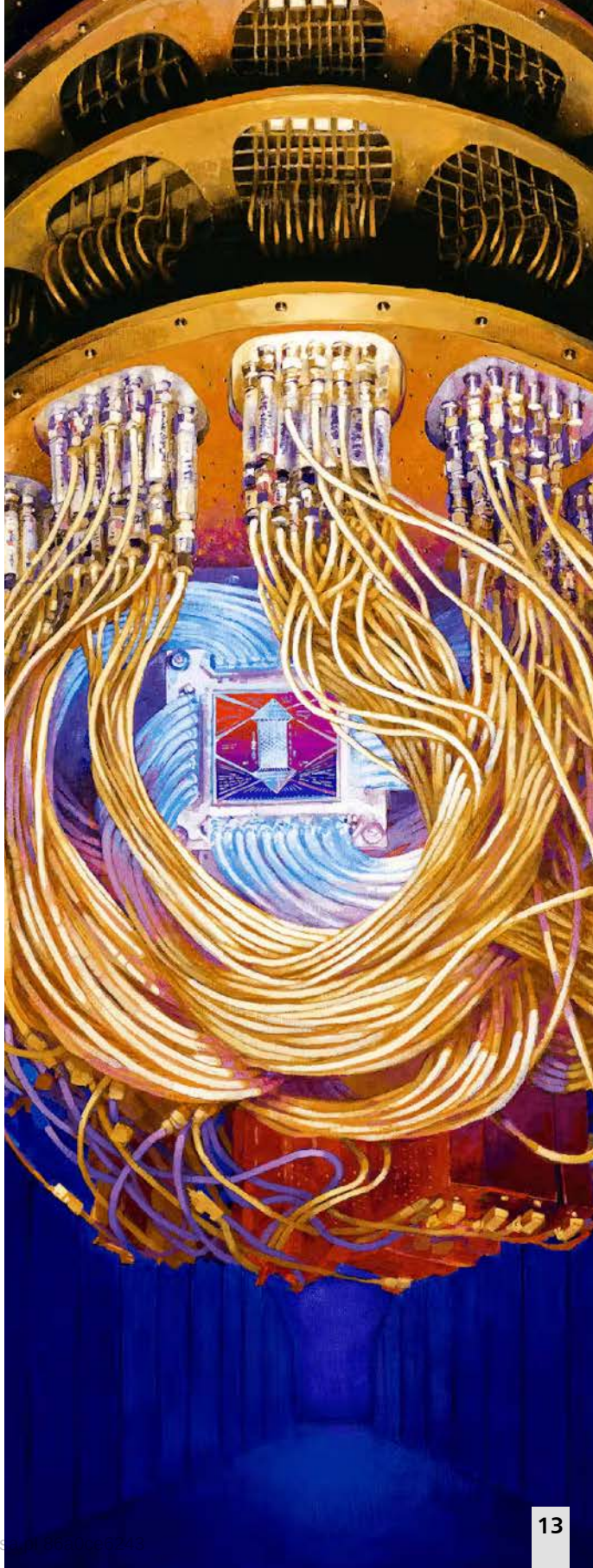
Wyjaśnienie sposobu działania techniki LiFi: <https://tiny.pl/cw235>

Google: nasz komputer kwantowy robi w chwilę to, co superkomputerom zajmuje pół wieku

Firma Google opracowała komputer kwantowy, który w mgnieniu oka wykonuje obliczenia, które zajęłyby najlepszym istniejącym superkomputerom 47 lat, co, zdaniem przedstawicieli Google wypowiadających się w publikacji, która ukazała się w internecie, stanowi przełom dowodzący ponad wszelką wątpliwość, że eksperymentalne maszyny kwantowe mogą przewyższać konwencjonalną konkurencję.

Przypomnijmy, że już cztery lata temu Google ogłosiło, że jest pierwszą firmą, która osiągnęła „kwantową supremację” – tak umownie nazywa się kamień milowy, w którym komputery kwantowe przewyższają istniejące maszyny. Tamten układ miał 53 kubity. Maszyna opisana obecnie w artykule pt. „Phase Transition in Random Circuit Sampling”, opublikowanym w serwisie ArXiv, ma mieć kubitów 70.

Dodanie większej liczby kubitów zwiększa moc komputera kwantowego wykładniczo, co oznacza, że nowa maszyna jest 241 milionów razy potężniejsza niż maszyna z 2019 roku. Badacze z Google podali, że Frontier, wiodący na świecie superkomputer, potrzebowałby 6,18 sekundy, aby dopasować obliczenia z 53-kubitowego komputera Google z 2019 roku. Dla porównania, dorównanie najnowszemu zajęłoby 47,2 roku. ■





TECHNIKI WIZUALNE

Sfera w Las Vegas – nie tylko gigantyczny wyświetlacz

111,5 metra wysokości i ponad 150 metrów szerokości, ponad dwadzieścia tysięcy metrów kwadratowych powierzchni, masa ok. 13 tysięcy ton, gigantyczny układ lamp LED, pozwalający wyświetlać 256 milionów kolorów – m.in. takie liczby podawane są w opisie MSG Sphere gigantycznego wyświetlacza, który stanął w amerykańskim Las Vegas.

Sferyczna konstrukcja nie służy jedynie do wyświetlania obrazów na zewnątrz, nie przede wszystkim. Wnętrze może pomieścić nawet 20 tysięcy widzów. Jego centralnym punktem jest ekran multimedialny o wysokiej rozdzielczości, który jest zakrzywiony tak jak to ma miejsce w planetarium z powierzchnią prawie 15 tysięcy m² pokrytą również diodami LED, kreującymi obrazy w rozdzielczości 16K na 16K.

Obiekt jest w swoim chyba najważniejszym przeznaczeniu jest salą koncertową. Za ogromną płaszczyzną MSG Sphere kryje się 168 tys. głośników, mających generować realistyczny i wciągający dźwięk przestrzenny. Zdaniem twórców, dzięki sferycznemu kształtowi przestrzeni, rozpraszanie dźwięku staje się mniejszym problemem. W pierwszym koncercie, jaki został tam zaplanowany, zagra grupa U2. ■



BEZZAŁOGOWCE

Francuski dron spędzi dobę w powietrzu, rozpozna, wysledzi i powalczy, gdy trzeba

24 godziny bez przerwy w powietrzu i 1,5 tony ładunku to tylko niektóre parametry nowego drona bojowego AAROK, skonstruowanego przez francuską firmę Turgis & Gaillard i zaprezentowanego Paris Air Show. Maszyna jest przeznaczona dla sił zbrojnych państw NATO, stanowiąc mniejszą i w założeniu bardziej ekonomiczną alternatywę dla amerykańskiego MK Reapera.

Wyposażony w napęd śmigłowy AAROK może na pokład zabierać sprzęt zwiadowczy, czujniki elektrooptyczne, radar wielomodowy, precyzyjną amunicję kierowaną i pociski powietrze-ziemia. Może więc oprócz misji rekonesansowych i zwiadowczych, uczestniczyć również w działaniach bojowych.

Francuski dron może także służyć jako powietrzny terminal komunikacyjny dla rejonów objętych działaniami bojowymi i walkami. Łącze satelitarne SATCOM pozwoli na operowanie maszyną z dowolnego miejsca na świecie. Jak zapewniają twórcy, będzie miał możliwość startu z lotnisk o niewysokim stopniu przygotowania. Jest także możliwość jego transportu na pokładzie taktycznego samolotu transportowego. ■



PODOBŃ KOSMOSU

Produkcja tlenu na Marsie ze zdwojoną wydajnością

Dwukrotnie większą niż w poprzednich próbach ilość tlenu z marsjańskiej atmosfery, a konkretne z zawartego w niej dwutlenku węgla, udało się wyprodukować na Marsie za pomocą instrumentu MOXIE (Mars Oxygen In-Situ Resource Utilization Experiment) zainstalowanym na łaziku marsjańskim Perseverance.

Instrument działa poprzez pompowanie marsjańskiego powietrza do zbiornika, a następnie wykorzystanie procesu elektrochemicznego do pozyskiwania atomu tlenu z każdej cząsteczki dwutlenku węgla. Proces ten nie jest pozbawiony ryzyka, ponieważ węglowy produkt uboczny w postaci stałej może gromadzić się wewnątrz urządzenia. Jednak

badacze odważyli się zaryzykować intensywniejszy proces. W ciągu testu trwającego 58 minut MOXIE osiągnął tempo produkcji około 12 gramów tlenu na godzinę, czyli około dwa razy wydajniej niż oczekiwano.

W 2021 roku MOXIE przeprowadził jednogodzinne eksperymenty siedem razy, nawet podczas surowej marsjańskiej zimy. Tym razem jednak jego wyniki były znacznie lepsze niż w poprzednich testach. Instrument jest demonstracyjną wersją, której finansowanie zakończy się w tym roku. Aby podtrzymać załogową misję i jej pobyt na Marsie potrzebna będzie produkcja tlenu rzędu nie gramów jest dziesiątek ton. ■

**BADANIA KOSMOSU**

Studia przeprowadzone przez międzynarodowy zespół uczonych sugerują, że w Obłoku Oorta na obrzeżach naszego Układu Słonecznego ukrywa się obca, przybyła z zewnątrz i przechwycona przez grawitację Słońca, egzoplaneta o rozmiarach Neptuna, a nie, jak czasem się podejrzewa, powstała w naszym układzie obiekt nazywany „Planetą X” lub Nemezis. ♦ Naukowcy bazujący na danych z badań przeprowadzonych przez lądownik NASA InSight, znaleźli dowody na anomalie rozkładu masy głęboko pod powierzchnią czerwonej planety, przy czym nowe wyniki sugerują niewielkie przyspieszenie rotacji Marsa i potwierdzają przypuszczenie, że nie ma on stałego jądra wewnętrznego. ♦

**ROBOTY**

♦ Firma o nazwie Throwflame sprzedaje pso-podobnego robota nazwanego „Theronator” z miotaczem ognia przypiętym do grzbietu, który może wystrzeliwać nawet trzydziestometrowe strumienie ognia, choć trzeba przyznać, że firma nie reklamuje urządzenia jako broni, lecz raczej jako urządzenie do wykorzystania w generowaniu pokazów i efektów specjalnych, w widowiskach a także w filmach. ♦ Siemeon Adebola z Uniwersytetu Kalifornijskiego w Berkeley i jego koledzy opracowali zrobotyzowany, zasilany systemami uczenia maszynowego system ogrodniczy o nazwie AlphaGarden, która zajmuje się uprawą mini-ogrodu z wyselekcjonowaną mieszanką gatunków, i przetestowali, czy może działać tak samo dobrze, jak praca zespołu ekspertów-ogrodników, dowodząc w rezultacie, że wzrost roślin był podobny, zaś roboty zużywały ok. 40 proc. mniej wody niż ogrodnicy-ludzie z wielo-

letnim stażem w pracy przy pielęgnacji roślin. ♦

AERONAUTYKA

♦ Latający samochód „Model A” Alef Aeronautics o zasięgu lotu ok. 170 km otrzymała specjalny certyfikat zdolności do lotów od amerykańskiej Federalnej Administracji Lotnictwa (FAA), co oznacza, że producent będzie mógł przeprowadzić testy drogowe/lotnicze pojazdu, poinformowała firma w komunikacie prasowym, dodając zarazem, że zaczyna zbierać zamówienia od chętnych klientów. ♦ Inżynierowie z Massachusetts Institute of Technology z powodzeniem zaprojektowali i przetestowali główne komponenty nowego typu silnika elektrycznego dla samolotów oraz przeprowadzili analizę obliczeniową, dowodzącą, że komponenty są w stanie współpracować w celu wygenerowania jednego megawata mocy, co pozwoliłoby efektywnie napędzić elektrycznie nawet duże jednostki latające. ♦

**ENERGIA**

♦ Firma Zero Mass Water zaprezentowała nowy typ hydropaneli słonecznych nazwanych SOURCE, które dzięki energii na nie padającej wytwarzają wodę „z powietrza”, czyli wydobywają wilgoć atmosferyczną i skraplają ją przetwarzając na wodę użytkową z wydajnością nawet do dziesięciu litrów dziennie. ♦ Fiński startup o nazwie Warming Surfaces opracował cienką jak papier folię ocieplającą Halia, z wbudowanym płaskim, niskonapięciowym elementem ogrzewającym, którą można pokrywać ściany i meble, co, zdaniem twórców rozwiązania, pozwala w energooszczędny sposób ocieplać i ogrzewać pomieszczenia. ■

M.U.



1. Dron Wing przenoszący paczkę

Kiedy wreszcie nadejdą dostawy dronami?

Marzenia sprowadzone na ziemię

Jeff Bezos, założyciel Amazona, w jednym z wywiadów dla amerykańskiej telewizji w 2013 roku zapowiedział wypełnienie nieba flotą dronów dostawczych, które miały dostarczać paczki do domów klientów w ciągu 30 minut od zamówienia. Miało to się stać w ciągu maksymalnie pięciu lat. Po dziesięciu latach wydaje się, że do realizacji nie jest bliżej niż w 2013.

Jak napisał niedawno po przeprowadzeniu dokładnych badań, śledztw i analiz, serwis Bloomberg, pomimo wydania ponad dwóch miliardów dolarów i zaangażowania specjalnego zespołu zajmującego się rozwojem techniki dostaw dronami latającym i liczącego ponad tysiąc osób na całym świecie, Amazon nie wydaje się obecnie bliski uruchomienia takich usług na większą skalę.

Poważna katastrofa drona Amazona wiosną 2022 r. skłoniła amerykańskie federalne organy regulacyjne do zakwestionowania zdatości tej maszyny do lotów. Zawiodło zbyt wiele funkcji bezpieczeństwa, a maszyna wymknęła się spod kontroli, powodując pożar zarośli, w które wpadła. Choć nikomu nic się nie stało, technika ta nie dowiodła, że jest bezpieczna. Wręcz przeciwnie. Nie był to zresztą jedyny taki wypadek a raczej kolejny w serii, jedynie bardziej spektakularny.

Według Bloomberg, Amazon nie zdołał przeprowadzić planowanych na 2021 roku 2500 lotów testowych w zeszłym roku, wyznaczając sobie na 2022 ambitniejszy cel – 12 tysięcy. Część z nich miała odbywać się w trybie autonomicznym. Firma planowała dodać nowe lokalizacje testowe w College Station w Teksasie oraz Lockeford w Kalifornii, niedaleko Stockton. Maszyny Amazona przeszły poważne zmiany, ewoluując od quadkoptera do jednostki hybrydowej, która może startować pionowo i latać jak samolot, zanim ponownie opadnie w miejscu docelowym. Wciąż jest w fazie testów, w tym na poligonie dronów we wschodnim Oregonie.

Firma Bezosa jest też pod rosnącą presją konkurencyjną. W zeszłym tygodniu należąca do Alphabet Inc. firma Google Wing (1) przyspieszyła swój własny program testowania dronów, rozpoczynając przewożenie



paczek do kupujących z Walgreens na obszarze podmiejskim o powierzchni 90 mil kwadratowych na północ od Dallas. Sieć supermarketów Walmart Inc. i firma kurierska United Parcel Service mają własne programy dronów na różnych etapach rozwoju. Niektóre z tych podmiotów wydają się być nieco bardziej zaawansowane na drodze do uruchomienia dostaw dronowych, choć zaczęły później niż Amazon, np. UPS w 2019 r. stał się pierwszą firmą, która otrzymała zgodę FAA na prowadzenie dronowej linii lotniczej, wykorzystując drony M2 firmy Matternet do dostarczania przesyłek do wybranych kampusów szpitalnych, uniwersyteckich i społeczności ośrodków dla emerytów.

Jednak w obecnej sytuacji, jak sądzą eksperci, miną lata, zanim Federalna Administracja Lotnictwa (FAA) USA zatwierdzi regularne komercyjne dostawy dronami. Z drugiej strony agencja nie ma nic przeciwko testom i zezwala na programy badawcze, pod warunkiem zachowania zasad bezpieczeństwa „bez stwarzania zagrożenia dla ludzi, mienia lub innych statków powietrznych”. „Testy w locie są krytyczną częścią wszystkich projektów certyfikacji samolotów”, głosi oświadczenie agencji FAA.

Długa historia zmagania Amazona z dronami

Jeszcze dekadę temu, w tym samym mniej więcej czasie, gdy Jeff Bezos wygłaszał w telewizji szumne zapowiedzi, Amazon zatrudnił inżyniera lotnictwa i oprogramowania Gura Kimchi do prowadzenia rodzącego się programu dronów, znanego obecnie jako Prime Air (2). Projektowanie dronów dostawczych zapowiadało się na ciężką pracę, a Amazon uczynił to wyzwaniem jeszcze trudniejszym, decydując się na samodzielne stworzenie zupełnie nowej maszyny, zamiast zlecać projektowanie i budowę prototypów innym firmom.

2. Dron Prime Air firmy Amazon



Kimchi preferował podejście „zrób to sam” ponieważ dawało to zespołowi kontrolę nad ostatecznym projektem, ale byli i obecni pracownicy twierdząc, że decyzja ta spowolniła rozwój. Na przykład, personel samodzielnie nawijał drut miedziany wokół magnesów silnika elektrycznego, gdy zewnętrzny dostawca mógł to zrobić szybciej. Nawet prototypy były budowane ręcznie.

Pierwsze maszyny ujawnione przez Bezosa nie były w stanie sprostać zadaniu dostaw paczek z zamówieniami. Ledwo mogły przelecieć półtora kilometra i były bardzo wrażliwe na podmuchy wiatru. Amazon potrzebował drona, który łączyłby zdolność samolotu do latania na duże odległości z manewrowością helikoptera, który może szybko zmieniać kierunek, aby omijać drzewa i linie energetyczne oraz unosić się szybko, także podczas niesprzyjającej pogody. Drony musiały również latać i znajdować miejsce docelowe bez interwencji człowieka, czyli w trybie autonomicznym, co stanowiło wymaganie podnoszące poziom trudności o kolejne kilka poprzeczek.

Według publikowanych w mediach informacji, przez pierwsze lata prac zespołu nie było silnej presji czasowej a troska o względy bezpieczeństwa zawsze wygrywała presją czasu. To się jednak w okolicach 2029 roku zaczęło zmieniać, gdy w kierownictwie Amazona zaczęło narastać zniecierpliwienie powolnym tempem prac. Wtedy odszedł Kimchi, który, zdaniem anonimowo wypowiadających się w mediach pracowników, nie wywiązał się ze złożonych obietnic.

W marcu 2020 roku Amazon zatrudnił do prowadzenia programu dronów Davida Carbona, doświadczonego specjalistę z Boeinga. Według informacji „New York Timesa”, fabryka Boeinga 787 w Południowej Karolinie, którą Carbon kierował, miała tendencję do przedkładania szybkości produkcji nad względę bezpieczeństwa. Choć doceniano jego wielkie doświadczenie, szybko pojawiły się przecieki na temat presji na tempo prac i lekceważenie procedur bezpieczeństwa.

W ciągu czterech miesięcy 2022 roku doszło do pięciu wypadków na terenie testowym w Pendleton w stanie Oregon (3). Wypadki są nieuniknione w programie testów lotniczych, w którym sprzęt jest celowo wykorzystywany na maksimum, aby określić punkty przerywania i ulepszyć konstrukcję pojazdu. Były to jednak pojazdy, które Amazon miał nadzieję wdrożyć do testów publicznych.

W maju 2022 śmigło drona odczepiło się, powodując upadek pojazdu i zderzenie z ziemią, gdy inne silniki nadal działały. Maszyna doznała znacznych uszkodzeń. Pracownicy Amazona usunęli wrak przed powiadomieniem urzędników federalnych, więc nie



3. Ośrodek testowania dronów w Pendleton mający powielać warunki typowego podmiejskiego lub wiejskiego domu

przeprowadzono inspekcji. FAA zaleciła firmie, aby w przyszłości nie zakłócała miejsc katastrof.

W kolejnym miesiącu silnik drona wysiadł w trakcie testowego lotu, gdy pojazd przechodził z pionowego wznoszenia do ruchu do przodu. Automatyczna funkcja bezpieczeństwa zaprojektowana do lądowania maszyny w takich przypadkach nie zadziałała. Samolot obrócił się do góry nogami, a stabilizująca funkcja bezpieczeństwa również zawiodła. „Zamiast kontrolowanego zejścia do bezpiecznego lądowania, dron spadał z wysokości ponad 50 metrów w niekontrolowanym pionowym upadku i stanął w ogniu”, napisała FAA w raporcie dotyczącym incydentu. Powstały pożar objął 25 akrów powierzchni zarośli i został ugaszony dopiero przez straż pożarną.

„Po tych wszystkich latach i wszystkich zainwestowanych pieniądzach można by oczekiwać czegoś lepszego”, powiedział cytowany przez „Wired” Antoine

Deux, który był starszym inżynierem w programie dronów przez cztery lata przed odejściem w 2018 roku. Jak dodał, dron Amazona jest zbyt ciężki w porównaniu z maszyną Google, która waży około ok. 5 kg. Co nie znaczy, że dron Google da radę z maksymalnie 2,5-kilogramowymi paczkami, które planuje dostarczać Amazon a które stanowią ok. 80 proc. wszystkich przesyłek w dostawach firmy. „Za każdym razem, gdy zwiększasz wagę ładunku, dron staje się cięższy i potrzebuje więcej mocy z baterii, mówi Deux. „To błędne koło”.

Choć media ostatnio skupiają się na problemach programu dronowego Amazona, nie oznacza to, że inne firmy poradziły sobie z fundamentalnymi problemami, ograniczeniami fizycznymi, gęstością mocy w akumulatorach, bezpieczeństwem, i wieloma innymi problemami z jakimi zderza się ta dziedzina, na którą dekadę temu z optymizmem patrzył nie tylko Jeff Bezos. ■

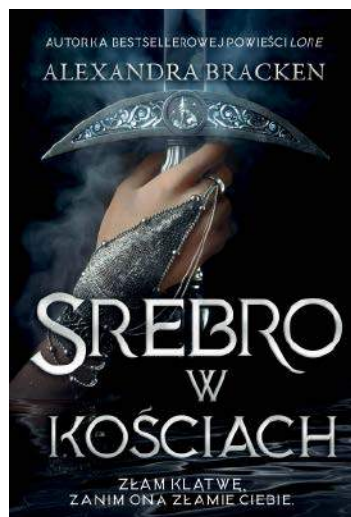
Mirosław Usidus

Srebro w kościach (tom 1)

Alexandra Bracken

Wydawnictwo Jaguar, liczba stron: 576, cena: 59,90 zł

Pierwszy tom nowej serii fantasy Alexandry Bracken, autorki cyklu Mroczne umysły oraz bestsellerowej powieści „Lore”. Historia inspirowana legendą arturiańską, pełna miłości, pragnienia zemsty i czystej adrenaliny! Czarna magia, ryzykowne wyprawy do starożytnych krypt w poszukiwaniu skarbów, tajemnicze zniknięcia i śmiertelny sekret, zdolny przebudzić duchy przeszłości – czy Tamsin podejmie wyzwanie i postawi na szali własne życie, by uratować brata? Alexandra Bracken jest jedną z czołowych amerykańskich autorek fantasy, jej książki stale goszczą na listach bestsellerów New York Timesa, a na ich podstawie powstają popularne produkcje filmowe. W przygotowaniu ekranizacja książki „Lore”. Adaptacji ma dokonać studio Universal. Projekt, opisywany jako połączenie „Igrzysk śmierci” z grecką mitologią, będzie filmową petardą! Tamsin Lark nie prosiła, żeby zostać Zgłębiającym. Jako śmiertelniczka bez talentu magicznego, miała nigdy nie włączyć się do starożytnych krypt ani też nigdy nie rywalizować z czarodziejkami i Oświeconymi o znajdujące się w środku skarby. Ale po tym, jak jej przybrany ojciec zniknął bez pożegnania, tylko w ten sposób mogła utrzymać przy życiu siebie i swojego brata, Cabella.





1. Dwie fiołki z fotoluminescencyjnymi kropkami kwantowymi obok niebieskiego elektroluminescencyjnego emitery QD

Czy elektroluminescencyjne kropki kwantowe zdobędą ekrany i wyświetlacze?

Kwantowe maleństwa zamiast LED i LCD

Kropki kwantowe, o czym może nie każdy wie, można już znaleźć w dostępnych na rynku telewizorach. Jednak nowa ich generacja nazwana kropkami elektroluminescencyjnymi, zademonstrowano po raz pierwszy na początku 2023 roku na targach technologicznych CES w USA mogłaby, zdaniem ekspertów, zastąpić LCD i OLED we wszystkich urządzeniach z wyświetlaczem.

Mogą poprawić jakość obrazu, dać oszczędności energii i wydajności produkcji. Prostsza struktura ma sprawiać, że wyświetlacze z kropek kwantowych są teoretycznie znacznie tańsze i łatwiejsze w produkcji.

Firma Nanosys zaprezentowała prototyp nowego wyświetlacza jedynie wybranym, specjalnie zaproszonym dziennikarzom. Ze skąpych informacji wynika, że typowy dla wyświetlaczy LED i LCD złożony układ warstw ma zostać znacznie uproszczony z zastosowaniem QD (kropek kwantowych) nakładanych za pomocą nowych metod. Konstrukcja ma być

znaczco uproszczona, co wpłynie na koszty i tańszą produkcję. Jest to jednak dopiero prototyp.

Ta sprawność

Definicyjnie kropka kwantowa to niewielki obszar przestrzeni ograniczony w trzech wymiarach barierami potencjału, nazywany tak, gdy wewnątrz uwięziona jest cząstka o długości fali porównywalnej z rozmiarami kropki. Oznacza to, że opis zachowania cząstki musi być przeprowadzony z użyciem mechaniki kwantowej. Ograniczenie ruchu cząstki

w trzech wymiarach oznacza kwantyzację w każdym z poszczególnych kierunków. Prowadzi to do sytuacji, gdy cząstka może znajdować się jedynie w pewnych stanach kwantowych, określonych równaniem Schrödingera. Tylko dobrze określone, dyskretne poziomy energetyczne mogą być zajęte przez cząstkę. Z tego powodu kropki kwantowe nazywa się czasem sztucznymi atomami.

Z punktu widzenia, który nas interesuje kropki kwantowe po dostarczeniu energii emitują światło o określonej długości fali, przy czym kropki o różnych rozmiarach emitują różne długości fal. Innymi słowy, niektóre kropki emitują światło czerwone, inne zielone, a jeszcze inne niebieskie. Możliwością jest więcej, ale w przypadku technologii wyświetlania, RGB to wszystko, czego potrzeba. Warto dodać, iż cechuje je wybitna sprawność. Emitują niemal idealnie taką samą ilość energii, jaką pochłonęły.

W ciągu ostatnich kilku lat kropki kwantowe były wykorzystywane przez producentów telewizorów do zwiększania jasności i kolorów telewizorów LCD. Litera „Q” w akronimie QLED TV oznacza właśnie technikę „kwantową”. Początkowo spotykane były tylko w telewizorach z wyższej półki. Obecnie kropki kwantowe można znaleźć w telewizorach średniej i niższej klasy takich marek jak Samsung, TCL, Hisense, LG i Vizio. Ich celem jest poprawa wyświetlania kolorów, wyższa jasność HDR i nie tylko.

W wyświetlaczach QD-OLED dodawane są do panelu wśród pikseli OLED. W technice microLED wbudowane są diody. Są dodatkiem poprawiającym parametry a nie samodzielną, odrębną techniką wyświetlania.

Kropki kwantowe stosowane do tej pory w technologii wyświetlaczy określa się jako „fotoluminescencyjne”. Pochłaniają one światło, a następnie je emitują. W przypadku telewizorów LED LCD zwykle oznacza, iż diody LED emitują niebieskie światło, widoczne na ekranie telewizora, ale zarazem używane do wywołania czerwonych i zielonych kropek kwantowych (1), które emitują własne kolorowe światło. To, co wiadać na ekranie, to niebieskie światło z diod LED oraz czerwone i zielone światło z kropek kwantowych, które łączą się, współtworząc obraz. Istnieje wiele sposobów na wdrożenie tego procesu, ale zasada jest mniej więcej taka.

Trzeba będzie jeszcze trochę zczekać

Rozwiązanie opisywane przez dziennikarzy, którzy mieli dostęp do nowej techniki elektroluminescencyjnych kropek kwantowych Nanosys, nie stosuje tradycyjnie rozumianych diod LED lub OLED. Zamiast

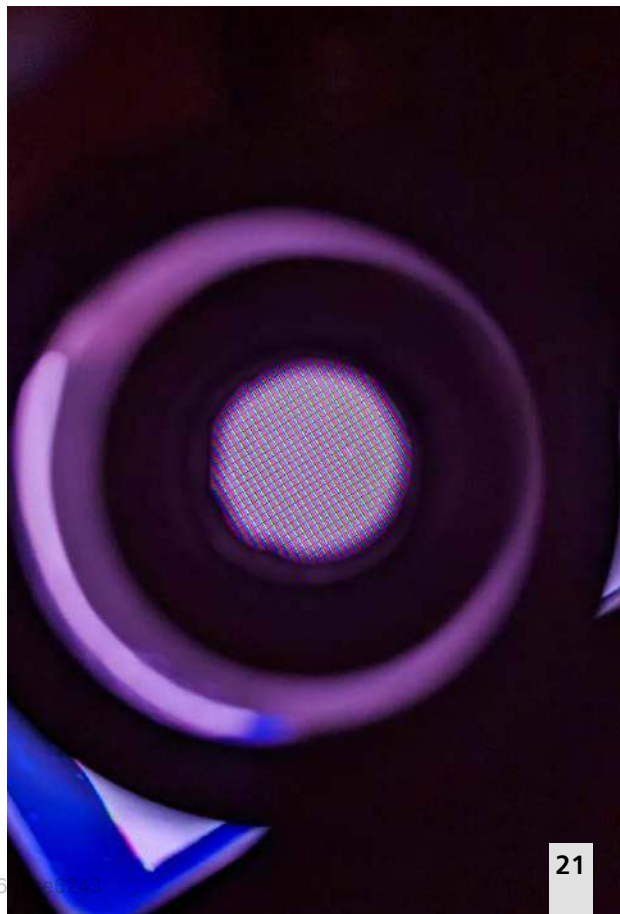
wykorzystywać światło do wzbudzania kropek kwantowych do emitowania światła, wykorzystuje bezpośrednio energię elektryczną. A kropki są bezpośrednim i jedynym źródłem wyświetlanego obrazu (2).

Usunięcie typowych dla ekranów warstw ma prowadzić do budowy cieńszych, bardziej energooszczędnych wyświetlaczy. Mają być tańsze w produkcji, czyli tańsze dla odbiorcy końcowego. Ocenia się, że poziom jakości obrazu jest nowej technice co najmniej tak dobry jak QD-OLED, jeśli nie lepszy. Ponadto ma nadawać się doskonale do małych, lekkich wyświetlaczy o wysokiej jasności w zestawach VR nowej generacji, wydajnych ekranów telefonów, nie mówiąc już o udoskonalonych telewizorach z płaskim ekranem.

Nanosys nazywa tę technikę elektroluminescencyjnych kropkami kwantowymi – „nanoLED”, co wydaje się nieco mylące, gdyż nie o rozwiązania LED-owe chodzi. Warto jednak nie zapominać, że to dopiero prototyp. Nawet Nanosys przyznaje, że wyświetlacze z bezpośrednim widokiem kropek kwantowych dzieli jeszcze co najmniej kilka lat od ewentualnej masowej produkcji. ■

Mirosław Usidus

2. Zbliżenie pikseli generowanych przez kropki kwantowe





AI uczona na danych wygenerowanych przez AI

Wąż zjadający własny ogon

Wielu ekspertów stawia sprawę jasno – sztuczna inteligencja, która szkoli się na danych wygenerowanych przez sztuczną inteligencję to początek końca AI. Ten sposób zaspokajania nigdy nie gasnącego apetytu algorytmów na dane to, w ich ocenie, droga donikąd. Modele sztucznej inteligencji potrzebują rzeczywistych danych, aby dokładnie zrozumieć i symulować nasz świat.

Problemem, który często kojarzy się z symboliką węża zjadającego własny ogon (1) zajęła się metodycznie grupa badaczy z Wielkiej Brytanii i Kanady. Wydała kilka miesięcy temu ostrzeżenie przed „załamaniem modelu” w nauce maszynowej, degenerującym procesem, który może całkowicie odseparować sztuczną inteligencję od rzeczywistości. W ich artykule, zatytułowanym „The Curse of Recursion: Training on Generated Data Makes Models Forget”, naukowcy z uniwersytetów w Cambridge i Oksfordzie, Uniwersytetu w Toronto i Imperial College w Londynie wyjaśniają, że załamanie modelu zachodzi, gdy „generowane dane zanieczyszczają zestaw treningowy następnej generacji modeli”. Jak dodają, „jako szkolone na zanieczyszczonych danych, błędnie postrzegają rzeczywistość”.

Innymi słowy, szeroko rozpowszechnione treści generowane przez sztuczną inteligencję, publikowane obecnie w coraz większych masach w Internecie mogą być zasysane niejako „z powrotem” do systemów sztucznej inteligencji, prowadząc do zniekształceń i nieścisłości.

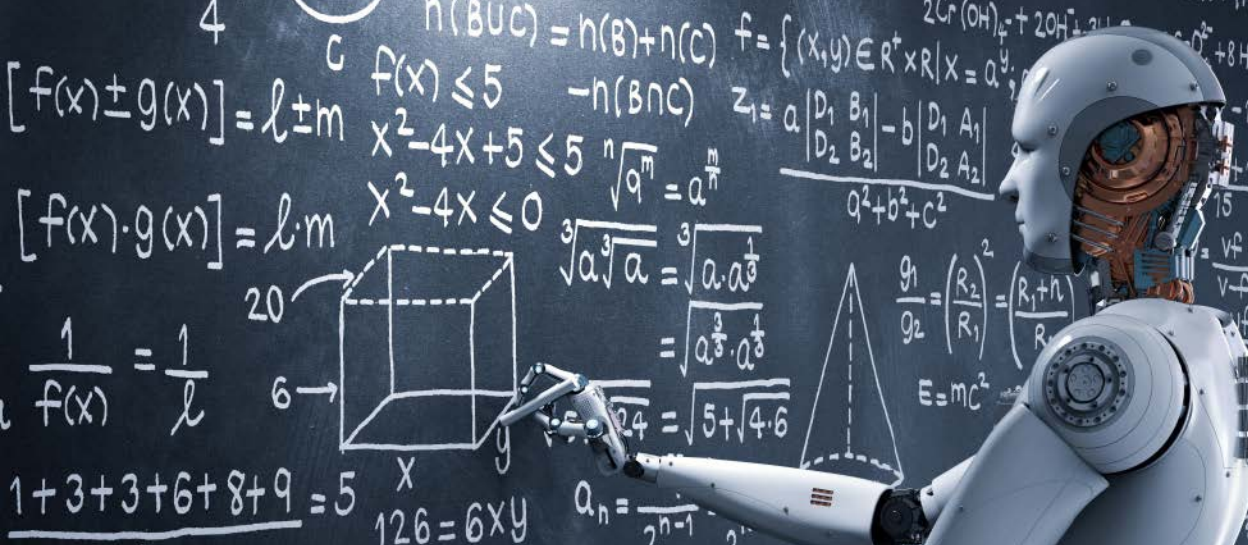
Problem ten został wykryty w szeregu szkolonych modeli generatywnych i podobnych systemach, m.in. w tym w dużych modelach językowych (LLM), autoenkoderach wariacyjnych i gaussowskich modelach mieszanych. Z biegiem czasu modele AI „żywione” danymi wygenerowanymi przez modele AI zaczynają „zapominać o prawdziwym rozkładzie danych”, co prowadzi do niedokładnych reprezentacji rzeczywistości, ponieważ oryginalne informacje stają się tak zniekształcone, że przestają przypominać rzeczywiste dane.



1. Uroboros – mityczny wąż pożerający własny ogon

Istnieją już przypadki, w których modele uczenia maszynowego są całkowicie świadomie szkolone na danych generowanych przez sztuczną inteligencję. O wykorzystywaniu takich „danych syntetycznych” pisaliśmy kilka miesięcy temu w MT. Na przykład modele uczenia się języka (LLM) są celowo szkolone na danych wyjściowych z GPT-4. DeviantArt, platforma internetowa dla artystów, umożliwia publikowanie dzieł sztuki stworzonych przez sztuczną inteligencję i wykorzystywanie ich jako danych treningowych dla nowszych modeli sztucznej inteligencji.

Podobnie jak próba kopiowania lub klonowania czegoś w nieskończoność, praktyki te, zdaniem



naukowców, mogą prowadzić do większej liczby przypadków załamania modelu. Artykuł, o którym mowa opisuje dwie główne przyczyny załamania modelu. Podstawową z nich jest „statystyczny błąd aproksymacji”, który wiąże się ze skończoną liczbą próbek danych. Druga możliwa przyczyna to „funkcjonalny błąd aproksymacji”, wynikająca z niewłaściwie skonfigurowanego marginesu błędu wykorzystywanego podczas uczenia maszynowego. Błędy te mogą narastać z generacji na generację, powodując kaskadowy efekt pogarszania się jakości danych i wzrostu nieścisłości.

Artykuł przedstawia też tzw. „przewagę pierwszego gracza” w szkoleniu modeli sztucznej inteligencji. Jeśli uda nam się zachować dostęp do oryginalnego źródła danych wygenerowanego przez człowieka, możemy

zapobiec szkodliwej zmianie rozkładu, a tym samym załamaniu modelu.

Rozróżnianie treści generowanych przez sztuczną inteligencję na dużą skalę jest jednak trudnym wyzwaniem, które może wymagać koordynacji w całym środowisku badaczy AI.

To ostatnie wydaje się niewykonalne, choćby dlatego, że duża część projektów AI jest prowadzona w krajach, które nie we wszystkich są skłonne do współpracy z zachodnimi specjalistami, np. w Chinach. Z drugiej strony – jeśli te kraje będą rozwijać swoje systemy na danych nie mających związku z rzeczywistością, to na dłuższą metę wyszkolone modele AI też będą mało przydatne, gdyż przestaną rozwiązywać problemy mające związek z realnym światem. ■

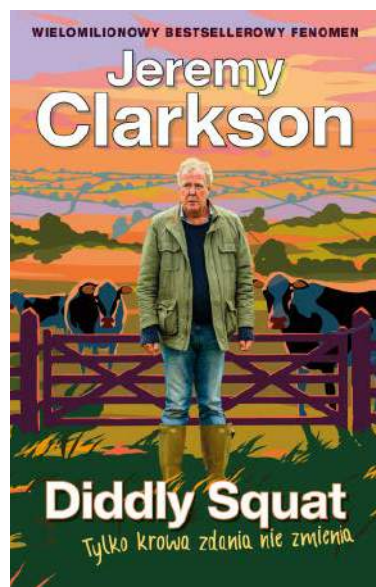
Mirosław Usidus

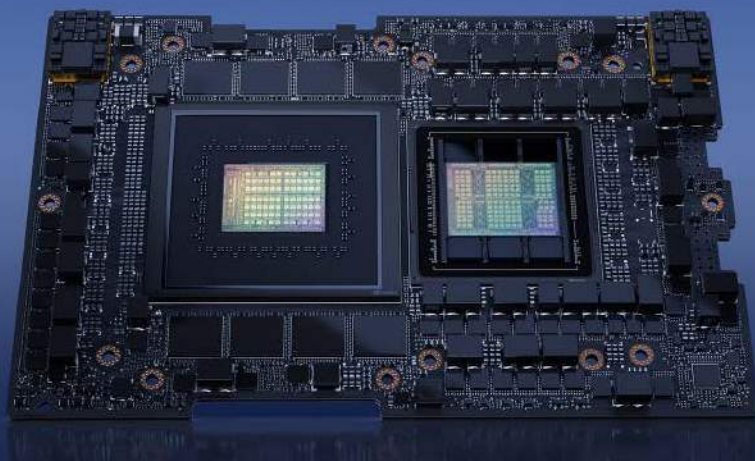
Diddly Squat. Tylko krowa zdania nie zmienia

Jeremy Clarkson

Wydawnictwo Insignis, liczba stron: 264, cena: 39,99 zł

Witajcie ponownie na farmie Clarksona! Pierwszy rok spędzony za kierownicą traktora na gospodarstwie Diddly Squat przyniósł Jeremy'emu astronomiczny zysk: aż 144 funtów! I choć sam postrzega się jako „rolnika praktykanta” (na wyrost, jak twierdzi Kaleb), musi coś zmienić w swoim podejściu, by zacząć w końcu wiązać koniec z końcem. Rzecz w tym, że Clarkson na razie opanował do perfekcji wyłącznie umiejętność biadolenia na prawie każdy temat, a wiele innych potrzebnych rolnikom kompetencji wciąż stanowi dla niego nie lada wyzwanie. Któż mógł przewidzieć, na przykład... że ładowanie ziarna na przyczepę wymaga lepszej koordynacji manualnej niż pilotowanie śmigłowca szturmowego Apache? Że krowy są większym zagrożeniem dla życia niż udział w wyścigach samochodowych? Że łatwiej byłoby uzyskać pozwolenie na budowę elektrowni atomowej niż na otwarcie restauracji w dawnej owczarni? Na szczęście na przekór gminnemu działowi planowania przestrzennego i mieszkających w okolicy nadętych bogaczy, którzy starają się robić wszystko, żeby pokrzyżować mu szyki, Jeremy może liczyć na niezawodne wsparcie ze strony Lisy, Kaleba, Radosnego Charliego i Geralda, specja od bezzaprawowych murków i szefa ochrony gospodarstwa. Życie na farmie Clarksona nie zawsze płynie zgodnie z planem. Ba, czasami w ogóle nie można go zaplanować. Ale za to nie ma takiego dnia, żeby Jeremy nie mógł szczerze powiedzieć: „Ha, zrobiłem... cokolwiek”.





1. Chipset Nvidia

Rynek AI jest już w dużej mierze zajęty i poukładany. Nie oznacza to, że wszyscy są z tego zadowoleni i nie będzie prób wprowadzenia nowych (najchętniej tańszych) rozwiązań i zmiany układu sił.

Czy AI jest skazane na karty graficzne?

NVIDIA I DŁUGO NIC... ALE NADZIEJA W RYWALACH NIE GAŚNIE

Nvidia, która posiada około 95 proc. rynku chipów AI (1), przewiduje, że w ciągu dekady lat modele sztucznej inteligencji będą milion razy bardziej potężne niż ChatGPT. Taką opinię wygłosił szef firmy Jen-Hsun Huang, podczas prezentacji na GTC, dodając, że prawdziwe przełomy dopiero przed nami, a modele sztucznej inteligencji będą w stanie wykonywać jeszcze bardziej skomplikowane zadania, w tym przewidywanie zachowania ludzi i zwierząt, projektowanie nowych leków, materiałów i inne.

W tej chwili świat jest pod wielkim wrażeniem sukcesu i dominującej pozycji Nvidii na rynku uczenia maszynowego. Nvidia jest wyraźnym liderem na rynku chipów używanych do szkolenia systemów AI, ale, co warto pamiętać, stanowi to tylko około 10 do 20 proc. popytu na chipy obsługujące sztuczną

inteligencję. Obok kart graficznych (GPU) funkcjonuje większy rynek układów służących do obsługi procesu wnioskowania, które interpretują wyszkolone modele i odpowiadają na zapytania użytkowników. Żaden pojedynczy podmiot, w tym Nvidia, nie może mówić, że na nim dominuje. Wiele firm z sektora Big Tech korzysta w przetwarzaniu AI ze swoich indywidualnych i zastrzeżonych rozwiązań. AWS z urządzeniami znanych pod nazwą Inferentia (2) a Google z opracowanych przez siebie jednostek TPU. Znaczna część obsługi systemów sztucznej inteligencji jest wciąż wykonywana na „zwykłych” procesorach CPU, zwłaszcza w sytuacji niedoboru układów GPU wysokiej klasy. Dominującym obecnie pytaniem w dziedzinie sztucznej inteligencji jest ilość mocy obliczeniowej potrzebnej do uruchomienia dużych modeli językowych (LLM), takich jak GPT, Bard lub Midjourney i Stable Diffusion. Poszukiwane są rozwiązania pozwalające na uruchomienie takiego modelu na jak najmniejszej mocy obliczeniowej. W tej mierze duże postępy osiągnęła społeczność open-source.

Obecnie przybliżone szacunki sugerują, że rynek „krzemu AI” rozkłada się tak: około 15 proc. mocy na szkolenie modeli, 45 proc. na wnioskowanie w centrach danych i 40 proc. na wnioskowanie brzegowe (edge), czyli przetwarzanie, analizowanie i przechowywanie danych bliżej miejsca, gdzie są generowane, co umożliwia szybką analizę i reakcję, niemalże w czasie rzeczywistym. Uważa się, że w bliższej perspektywie Nvidia utrzyma kontrolę nad rynkiem

AWS Inferentia

High performance machine learning inference chip,
custom designed by AWS

AVAILABLE LATE 2019



2. Prezentacja chipa AWS Inferentia

szkoliowym. Rynek wnioskowania w centrach danych to już nie tylko Nvidia, ale także firmy AMD i Intel, a także rozwiązania specjalne, indywidualnej i specjalistyczne.

Nvidia nie jest więc osamotniona na tym rynku, ale ma dużą przewagę. Nvidia jest postrzegana jako dominująca siła na rynku sztucznej inteligencji, napędzana przez technologię GPU i oprogramowanie CUDA. Oczekuje się, że duże inwestycje Apple i Tesli w sieci neuronowe będą nadal wpływać na projektowanie sprzętu AI.

Przejęcie Annapurna Labs przez AWS pokazało, że potentaci mogą obniżyć koszty i poprawić wydajność, przenosząc projektowanie do siebie. Potencjał AWS lub innych gigantów do przejmowania start-upów lub wprowadzania innowacji we własnym zakresie w celu konkurowania z Nvidią jest możliwy, ale w dającej się przewidzieć przyszłości będą one zależne od Nvidii.

Alternatywy dla drogiej kart

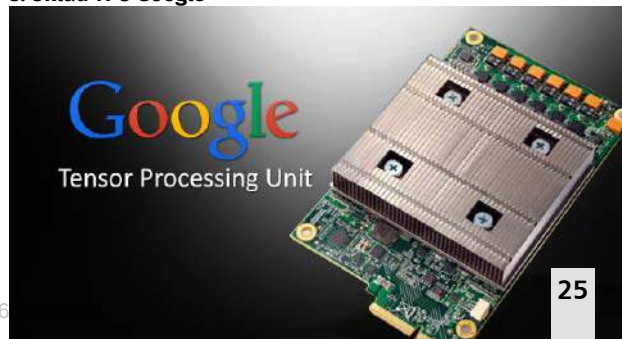
O wiele mniej rozpowszechniony niż GPU Nvidii, ale znany w świecie AI, jest procesor Tensor Processing Unit (TPU). To układ scalony opracowany przez Google do uczenia maszynowego sieci neuronowych, wykorzystujący autorskie oprogramowanie Google TensorFlow. Google zaczął używać TPU (3) wewnętrznie w 2015 roku, a w 2018 roku udostępnił je do użytku innym, zarówno jako część swojej infrastruktury chmury, jak też oferując mniejszą wersję układu na sprzedaż. W porównaniu z procesorem graficznym, TPU są przeznaczone do dużych ilościowo obliczeń o niskiej precyzji (np. zaledwie 8-bitowej) z większą liczbą operacji wejścia/wyjścia na dżul. TPU są dobrze przystosowane do neuronowych sieci konwulucyjnych (CNN). Między innymi miały być wykorzystane w szkoleniu algorytmu AlphaGo przed rozgrywką

z arcymistrzem Lee Sedolem, a także w systemie AlphaZero, który produkował programy do gry w szachy. Google wykorzystał również TPU do przetwarzania tekstu w Google Street View. W usłudze Zdjęć Google pojedyncza TPU może przetwarzać ponad sto milionów zdjęć dziennie. Jest ona również wykorzystywana w algorytmie RankBrain, którego Google używa do dostarczania wyników wyszukiwania.

Konkurująca przez lata z sukcesem na rynku chipów zwłaszcza z Intelem, firma AMD, podejmuje od niedawna próby wejścia na większą skalę do świata AI. Zaprezentowała w styczniu 2023 r. procesory AMD Ryzen z serii 7040, które są pierwszymi jednostkami x86 z wbudowanym akceleratorem sztucznej inteligencji. Układy są produkowane w 4-nanometrowym procesie technologicznym, oferując najmocniejszy na świecie zintegrowany układ graficzny, bazujący na architekturze oznaczonej RDNA 3.

AMD w czerwcu 2023 ujawnił swój nowy układ AI – MI300X. MI300X to układ skrojony na miarę potrzeb AI. Zaprojektowany z myślą o efektywnej pracy z ogromnymi modelami jest wręcz stworzony do wykonywania zadań wymagających olbrzymiej mocy obliczeniowej. Pamięć to jedna z najważniejszych kart w rękawie AMD – aż do 192 GB. To wyjątkowo dużo, nawet jak na standardy branży, i to dzięki tej konfiguracji MI300X może poradzić sobie z modelami, które

3. Układ TPU Google



przekraczają możliwości innych układów. Jednak pamięć to nie wszystko – architektura MI300X została stworzona do płynnej obsługi generatywnych zadań AI. Lisa Su, prezes AMD, podkreśla, że najnowsze modele mogą bez problemu „zamieszkać” w 192 GB pamięci HBM3 (high-bandwidth memory) MI300X. „Dzięki tej dodatkowej pojemności pamięci mamy przewagę przy dużych modelach językowych”, uważa Su, dodając, że użytkownicy potrzebować będą mniej układów GPU do wykonania tych samych zadań w krótszym czasie.

Wartość rynkowa NVIDIA przekroczyła bilion dolarów. AMD, z „zaledwie” 207 miliardami, ma jeszcze długą drogę do przebycia. AMD uruchomiło model Falcon 40B LLM na MI300X podczas prezentacji nowego układu. Według Su, to „pierwszy raz, kiedy model LLM tej wielkości może być uruchomiony całkowicie w pamięci”. Firma nie zamierza jednak zatrzymać się na hardware. AMD jednocześnie zaprezentowało aktualizację do RocM, czyli swojego oprogramowania do programowania GPU, które stanowi konkurencję dla języka CUDA od NVIDIA. Dzięki dużej przepustowości pamięci dostępnej z MI300X, firmy mogą decydować się na zakup mniejszej liczby układów GPU. To czyni AMD atrakcyjnym rozwiązaniem, szczególnie dla mniejszych przedsiębiorstw o średnich obciążeniach AI.

Idąc mniej więcej w tym samym kierunku, Apple, Nvidia i Intel także zapowiedziały budowę procesorów specjalnie dla AI tworzonymi rozwiązaniami wbudowanymi w tradycyjne procesory. Amazon dla swoich usług AWS tworzy nowy procesor Inferentia2. Te rozwiązania wciąż wykorzystują tradycyjną architekturę von Neumanna, zintegrowanej pamięci SRAM i zewnętrznej pamięci DRAM, wszystkie wymagają energii elektrycznej do przenoszenia danych do i z pamięci.

Opóźniony Intel

Intel podał w maju 2023 r. kilka nowych szczegółów na temat chipu do obliczeń sztucznej inteligencji (AI), który planuje wprowadzić w 2025 r., zmieniając strategię konkurowania z Nvidia i Advanced Micro Devices Inc (AMD). Na konferencji superkomputerowej w Niemczech w poniedziałek Intel powiedział, że jego nadchodzący chip „Falcon Shores” będzie miał 288 gigabajtów pamięci i będzie obsługiwał 8-bitowe obliczenia zmiennoprzecinkowe. Te specyfikacje techniczne są ważne, ponieważ modele sztucznej inteligencji podobne do usług takich jak ChatGPT eksplodowały, a firmy szukają mocniejszych chipów do ich obsługi.

Intel nie ma praktycznie żadnego udziału w rynku AI po tym, jak jego niedoszły konkurent Nvidii, chip o nazwie Ponte Vecchio, doznał wieloletnich opóźnień. Intel poinformował niedawno, że prawie zakończył dostawy superkomputera Aurora dla Argonne National Lab opartego na Ponte Vecchio, który według Intela ma lepszą wydajność niż najnowszy układ AI Nvidii, H100. Kolejny chip Intela, Falcon Shores, trafi na rynek dopiero w 2025 roku, kiedy to Nvidia prawdopodobnie będzie miała już inny własny układ. Jeff McVeigh, wiceprezes korporacyjny grupy superkomputerów Intela, powiedział, że firma potrzebuje czasu na przerobienie chipu po rezygnacji z wcześniejszej strategii łączenia procesorów graficznych (GPU) z jednostkami centralnymi (CPU). Firma na targach Computex 2023 snuła wizję wykorzystania swojego „silnika AI”, Vision Processing Unit (VPU), zapowiadając, że procesory te będą dostępne w każdym układzie Meteor Lake, którego premiera zapowiadana jest też na ten rok. Komponent ten ma obsługiwać zadania AI po stronie klienta poprzez znaczne zmniejszenie wymagań obliczeniowych związanych z wnioskowaniem AI.

Pomysł Intela jest taki – przenieść przetwarzanie sztucznej inteligencji, które już odbywa się na CPU i GPU, do specjalnego procesora tylko do AI. Według Intela ponad sto aplikacji wykorzystuje już sztuczną inteligencję na CPU lub GPU. To m.in. pakiet Adobe, Microsoft Teams, Avid Pro Tools, xSplit, Zoom i Unreal Engine. Jednostki VPU są nie tylko bardziej energooszczędne, ale także pozwalają na uruchamianie znacznie bardziej złożonych modeli AI. Jednocześnie Intel nie chce przenosić całej pracy na VPU. CPU i GPU nadal mają swoje miejsce. Według Intela, GPU jest nadal idealną opcją do zadań związanych z tworzeniem multimediów z wykorzystaniem sztucznej inteligencji, podczas gdy CPU może obsługiwać prostsze zadania AI.

Zazwyczaj większość obciążeń związanych z przetwarzaniem i obliczaniem na potrzeby sztucznej inteligencji jest obsługiwana w chmurze, co może skutkować zwiększonymi kosztami dla dostawców oprogramowania, ponieważ narzędzia te stają się bardziej zaawansowane i wymagające pod względem zasobów obliczeniowych. I tu swoją rolę do odegrania ma VPU, jednostka przetwarzania wizji na platformie Meteor Lake. Ten wyspecjalizowany procesor został specjalnie zaprojektowany do obsługi przetwarzania dla AI z większą wydajnością niż procesor ogólnego przeznaczenia, a nawet wydajny układ GPU. Jednostki VPU mają umożliwiać przetwarzanie tych obciążeń lokalnie na komputerze stacjonarnym lub laptopie.

Firma Intel przedstawiła przykład w postaci generacji obrazu przez Stable Diffusion w programie GIMP na platformie Meteor Lake. W tym scenariuszu system wygenerował złożony obraz na podstawie zapytania tekstowego i wykorzystał łącznie CPU, GPU i VPU. Proces ten trwał około 20 sekund. Dodatkowo, narzędzie Super Resolution, które może być zastosowane wyłącznie na VPU, dostarcza wersję obrazu o wyższej rozdzielczości w zaledwie kilka sekund.

Mobilne procesory AMD Ryzen 7000, które również posiadają procesor AI, nazwany Ryzen AI, powinny mieć podobne zastosowania jak VPU Intela, choć AMD nie wyszczególniła jeszcze jego możliwości.

Szukając własnej ścieżki

Meta buduje swój pierwszy niestandardowy chip specjalnie do uruchamiania modeli sztucznej inteligencji, ogłosiła firma w maju 2023. Nowy układ MTIA firmy Meta, będący skrótem od Meta Training and Inference Accelerator (4), jest „wewnętrzna, niestandardowa rodzina układów akceleratorów ukierunkowanych na obciążenia związane z wnioskowaniem” napisał w poście na blogu wiceprezes firmy Meta i szef działu infrastruktury, Santosh Janardhan. Układ MTIA ma się pojawić dopiero w 2025 roku, donosił TechCrunch.

Oprócz MTIA, Meta wprowadza również nowy układ ASIC, który ma pomóc w transkodowaniu wideo, który nazywa „MSVP” lub Meta Scalable Video Processor. Został on zaprojektowany do obsługi zarówno „wysokiej jakości transkodowania potrzebnego do VOD, jak i niskiego opóźnienia i krótszych czasów przetwarzania, których wymaga transmisja na żywo”, informowała Meta w osobnym poście na blogu, a „w przyszłości” pomoże wprowadzić takie rzeczy, jak treści stworzone przez sztuczną inteligencję oraz treści specyficzne dla AR i VR do aplikacji Meta.

Także IBM zaprezentował nowy układ AI, który według niego może uruchamiać i trenować modele głębokiego uczenia się szybciej niż procesor ogólnego przeznaczenia. Prototypowy układ, IBM Research Artificial Intelligence Unit lub AIU, jest pierwszym kompletnym systemem na chipie zbudowanym specjalnie do głębokiego uczenia. Jest to półprzewodnik zaprojektowany w technologii nazywanej ASIC, który można zaprogramować do wykonywania dowolnego rodzaju zadań głębokiego uczenia, w tym przetwarzania języka mówionego lub słów i obrazów na ekranie. Prototyp AIU ma 32 rdzenie przetwarzające i 23 miliardy tranzystorów. IBM AIU został również zaprojektowany tak, aby był tak prosty jak

MTIA Meta Training Inference Accelerator

Niestandardowy chip AI



4. Meta Training and Inference Accelerator

karta graficzna, którą można podłączyć do dowolnego komputera lub serwera ze slotem PCIe. Układ wykorzystuje technikę obliczeń przybliżonych IBM, aby zmniejszyć ilość obliczeń potrzebnych do wytrenowania i uruchomienia modelu sztucznej inteligencji, bez poświęcania dokładności. AIU wykorzystuje mniejsze formaty bitów do uruchamiania modelu AI w tempie wymagającym znacznie mniej pamięci. Jest to rozwinięta wersja już sprawdzonego akceleratora AI wbudowanego w układ Telum. Telum wykorzystuje tranzystory o rozmiarze 7 nm, podczas gdy AIU będzie wyposażony w szybsze, jeszcze mniejsze tranzystory 5 nm.

Dzięki technice zapoczątkowanej przez IBM, zwanej obliczeniami przybliżonymi, możemy zrezygnować z 32-bitowej arytmetyki zmiennoprzecinkowej na rzecz formatów bitowych przechowujących jedną czwartą informacji. Ten uproszczony format radykalnie zmniejsza ilość obliczeń potrzebnych do wytrenowania i uruchomienia modelu sztucznej inteligencji, bez poświęcania dokładności. Mniejsze formaty bitowe zmniejszają również inny czynnik wpływający na szybkość: przenoszenie danych do i z pamięci. AIU wykorzystuje szereg mniejszych formatów bitowych, w tym zarówno reprezentacje zmiennoprzecinkowe, jak i całkowite, dzięki czemu uruchamianie modelu AI jest znacznie mniej „pamięciochłonne”.

W powyższym przeglądzie mowa głównie o projektach kolosów rynkowych. Nie brakuje też start-upów próbujących wykorzystać wielkie poruszenie związane ze sztuczną inteligencją by zaproponować własne procesory. To opisywane w MT kilka miesięcy temu Cerebras Systems, Nuvia, czy izraelski Hailo. Warto wspomnieć także o takich firmach jak SambaNova, Graphcore czy Wave Computing. Stworzyć sprzęt, który przebiję się przez szklany sufit, nad którym operują giganci, jest trudno, ale miejmy nadzieję, że nie jest to niemożliwe. ■

Mirosław Usidus

TSMC przygotowuje się do próbnej produkcji układów 2 nm (o wielkości dwóch nanometrów) z wykorzystaniem procesu wspomaganej sztuczną inteligencją. W miarę jak masowa produkcja półprzewodników zaczyna wkraczać w erę 3 nm, wszyscy główni konkurenci intensyfikują wyścig w kierunku 2 nm. Źródła poinformowały tajwański dziennik Economic Daily, że TSMC rozpoczęło prace przedprodukcyjne nad półprzewodnikami 2 nm. Obecny harmonogram firmy może pozwolić jej wyprzedzić takich konkurentów jak Intel i Samsung. Według doniesień, TSMC zmobilizowało inżynierów i naukowców, stawiając cel wyprodukowania tysiąca waflów w tym procesie jeszcze w tym roku. Wcześniej potwierdzone mapy drogowe wskazują na plany próbnej produkcji w procesie 2 nm w 2024 roku, i masowej produkcji w 2025 roku.

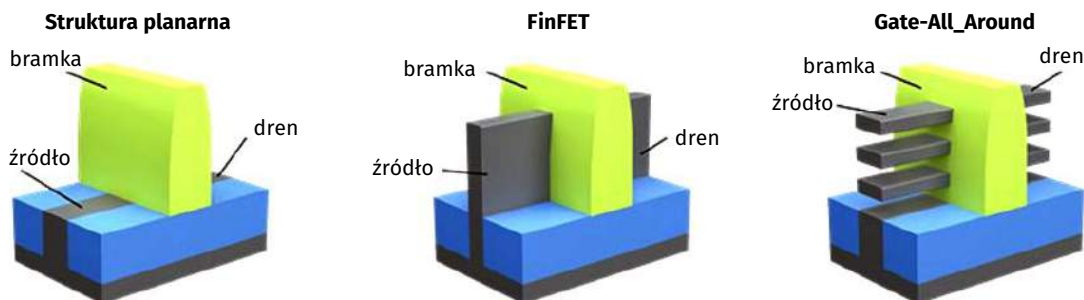
Dokąd zmierzamy idąc krzemową doliną i poza nią?

EGZOTYKA NIE- OGRANICZONYCH MOŻLIWOŚCI

Nieoficjalne źródła tych doniesień podają, że sięgający nowych rubieży miniaturyzacji proces produkcyjny TSMC wykorzystuje metodę wspomaganą sztuczną inteligencją o nazwie AutoDMP, która jest zasilana przez chipy DGX H100 firmy Nvidia. Sztuczna inteligencja ma optymalizować projekty chipów trzydzieści razy szybciej niż inne techniki.

W ramach procesu 2 nm TSMC przejdzie na tranzystory GAAFET (gate-all-around), które mają zwiększyć

gęstość tranzystorów 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. Tranzystor FET typu gate-all-around, w skrócie GAAFET, znany również jako tranzystor z bramką otaczającą (SGT), jest podobny w koncepcji do MOSFET – FinFET (o strukturze 3D zamiast planarnej), z wyjątkiem tego, że materiał bramki otacza obszar kanału ze wszystkich stron (1). W zależności od konstrukcji, tranzystory FET typu gate-all-around mogą mieć dwie lub cztery efektywne bramki. FET typu gate-all-around zostały z powodzeniem wytrawione na nanodrutach InGaAs, które mają 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 Unisantis Electronics w 2004 roku, aby wraz z Uniwersytetem Tohoku badać technologię otaczającej bramki. W 2006 roku zespół koreańskich naukowców z Korea Advanced



1. Porównanie struktur typów tranzystorów

Institute of Science and Technology (KAIST) i National Nano Fab Center opracował tranzystor 3 nm, najmniejsze na świecie urządzenie nanoelektroniczne, oparte na technologii GAA FinFET. GAAFET zostały wykorzystane przez IBM do zademonstrowania technologii procesowej 5 nm.

Od 2020 r. Samsung i Intel ogłosiły plany masowej produkcji tranzystorów GAAFET (w szczególności tranzystorów MBCFET), zaś TSMC podało, że będzie nadal używać FinFET w swoim węźle 3 nm, choć pracowało już nad tranzystorami GAAFET. Intel chce wypuścić swój 2-nanometrowy węzeł w 2024 roku, a Samsung planuje rozpocząć masową produkcję 2 nm w 2025 roku. Oba układy będą wykorzystywać GAAFET i tylne zasilanie w celu zwiększenia gęstości logicznej i zmniejszenia wycieków mocy. Również japońska firma Rapidus zamierza rozpocząć masową produkcję półprzewodników 2 nm do 2027 roku. Są plany wybiegające dalej w przyszłość, Samsung zamierza osiągnąć 1,4 nm do 2027 roku, podczas gdy TSMC planuje w tej perspektywie 1 nm.

Na razie Taiwan Semiconductor Manufacturing Co. (TSMC) stara się zaspokoić popyt ze strony głównego klienta Apple na chipy 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, aby wieść prym w produkcji 3 nm dla klientów takich jak Apple i Nvidia w jej wysokowydajnych komputerach (HPC) i w smartfonach. TSMC ogłosiło niedawno, że jest liderem w produkcji 3 nm.

Ostateczny lider zdobędzie lwią część zysków w branży. Według Mehdiego Hosseiniego, starszego analityka ds. badań kapitałowych w Susquehanna International Group, na razie TSMC utrzymuje pierwsze miejsce. „TSMC pozostaje preferowanym wyborem odlewni dla wiodących węzłów, ponieważ Samsung Foundry musi jeszcze zademonstrować stabilną, wiodącą technologię procesową, podczas gdy IFS [Intel Foundry Services] jest o lata świetlne od zaoferowania konkurencyjnego rozwiązania”, pisze Hosseini w „EE Times”. Według niego, w drugiej połowie 2023 roku TSMC będzie produkować procesory Apple A17 i M3 w węźle N3 (trzy nanometry), a także procesory serwerowe oparte na ASIC w węzłach N4 i N3. TSMC będzie również produkować chipy graficzne Intel Meteor Lake w N5, procesory AMD Genoa i Nvidia Grace w N5 i N4, a także GPU Nvidia H100 w N5. Brett Simpson, analityk w Arete Research, w raporcie „EE Times” dodaje: „Obecnie uważamy, że wydajność N3 w TSMC dla procesorów A17 i M3

wynosi około 55 proc. a TSMC wydaje się zgodnie z harmonogramem zwiększać wydajność o około 5+ punktów co kwartał”.

Przezwyciężyć CMOS

Technika obliczeniowa oparta na półprzewodnikach mają historię sięgającą lat 50. ubiegłego wieku, kiedy tranzystory zaczęły zastępować lampy próżniowe w układach elektronicznych. Generacje nowych urządzeń półprzewodnikowych pojawiały się i znikaly. Tranzystory germanowe zostały zastąpione krzemowymi, a następnie układami scalonymi, a potem coraz bardziej złożonymi chipami wypełnionymi coraz większą liczbą mniejszych tranzystorów. Od lat sześćdziesiątych branża mierzy swoje postępy prawem Moore’a, czyli prognozami Gordona Moore’a, współzałożyciela Intela, że coraz mniejsze urządzenia będą prowadzić do stałej poprawy wydajności obliczeniowej i efektywności energetycznej.

Postępy w nanotechnologii pozwoliły na zmniejszenie najmniejszych elementów najbardziej zaawansowanych obecnie układów scalonych do skali atomowej. Kolejny ważny krok w dziedzinie obliczeń wymaga nie tylko nowych nanomateriałów, ale także nowej architektury. Tranzystory CMOS (complementary metal-oxide-semiconductor) były standardowym budulcem układów scalonych od lat 80. ubiegłego wieku. Opierają się na architekturze, którą John von Neumann wybrał w połowie XX wieku. Ta została zaprojektowana tak, by oddzielać elektronikę przechowującą dane w komputerach od elektroniki przetwarzającej informacje cyfrowe. Oddzielenie przechowywanej pamięci od procesora zapobiegało wzajemnemu zakłócaniu się sygnałów i zachowywała dokładność potrzebną do obliczeń cyfrowych. Jednak przenoszenie danych z pamięci do procesorów stało się wąskim gardłem. Deweloperzy poszukują obecnie rozwiązań alternatywnych dla architektury von Neumanna. Dążą do wykonywania obliczeń „w pamięci”, aby uniknąć marnowania czasu i energii na przenoszenie danych.

Pierwsze oznaki poważnych zmian w pojawiły się około 2012 roku, gdy postęp wynikający z prawa Moore’a zaczął słabnąć, a twórcy deep learning zdali sobie sprawę, że jednostki centralne (CPU) ogólnego przeznaczenia stosowane w konwencjonalnych komputerach nie są w stanie zaspokoić ich potrzeb. Siłą procesorów była ich wszechstronność, ale z punktu widzenia głębokiej nauki maszynowej to niekoniecznie atut. Poszukując lepszych procesorów do głębokiego uczenia, programiści IBM zwrócili się w stronę procesorów graficznych (GPU), zaprojektowanych do wykonywania zaawansowanej

matematyki wykorzystywanej do szybkiego, trójwymiarowego obrazowania w grach komputerowych. Odkryto, że procesory graficzne mogą uruchamiać algorytmy głębokiego uczenia się znacznie wydajniej niż procesory centralne. W maszynach opartych na przepływie danych pewne instrukcje są wbudowane w procesor, więc nie trzeba ich ładować. W przypadku algorytmu głębokiej nauki maszynowej, około 80 proc. operacji wykorzystuje te same zaawansowane obliczenia, co przetwarzanie obrazu.

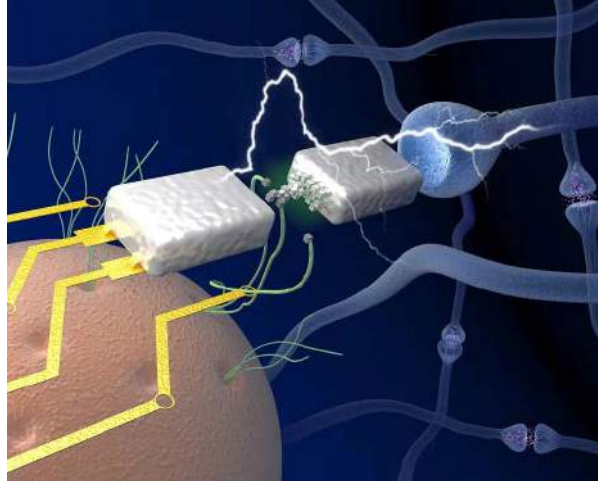
Procesory graficzne, które „już tu były” przydały się bardzo. Nie oznacza to, że wszyscy traktują je jako rozwiązanie ostateczne. Istnieje wiele nowych pomysłów, nowych urządzeń i nowych nanostruktur, ale wciąż nie wydają się gotowe do zastąpienia CMOS, i kart graficznych, ma się rozumieć.

Próby odtwarzania neuronów mózgu

Jednym z kierunków poszukiwań są systemy neuromorficzne, które wykorzystują algorytmy i projekty sieci, naśladujące 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 (2). 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, oddzielną 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 z laboratoriów Hewletta-Packarda w Kalifornii, stworzył pierwszy cienkowarstwowy memrystor półprzewodnikowy, który można było wykorzystać w obwodzie.

Naukowcy poszukują nowych klas materiałów, aby zaspokoić potrzeby zaawansowanych obliczeń. Mark Hersam i 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 heterostrukтуры 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 rurkowate aksony,



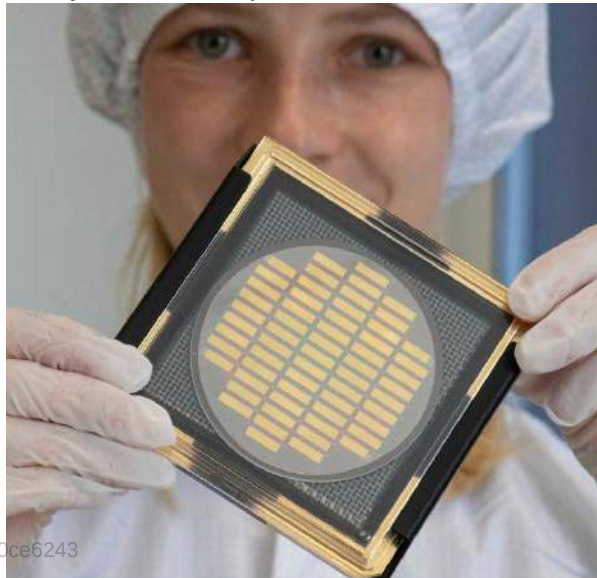
2. Wizualizacja neuromorficznych memrystorów

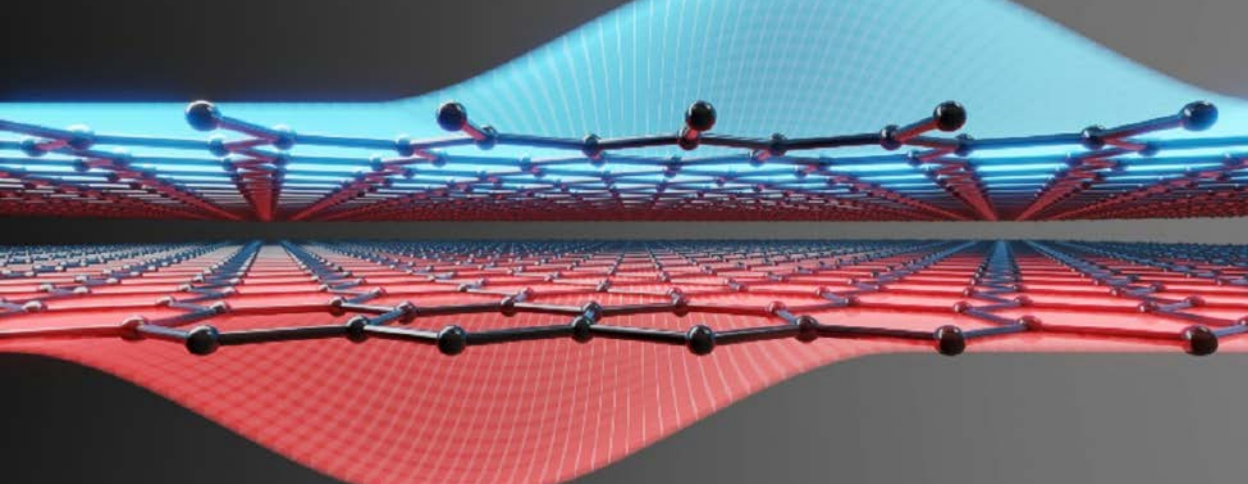
przez które komórki nerwowe przesyłają sygnały elektryczne w systemach biologicznych. Jednak, zdania co do tego, jak materiały te wpłyną na przyszłość informatyki, są podzielone.

Kwanty i egzotyka

Omawiając futurystyczne perspektywy elektroniki, nie sposób pominąć poszukiwań związanych z innym nurtem świata obliczeń przyszłości – komputerami kwantowymi (3). Tu też się poszukuje i to w nieco podobnych do systemów neuromorficznych obszarach, np. w świecie kropek kwantowych wykonanych z dwuwarstwowego grafenu. Badania nad nimi zostały przeprowadzone m.in. przez Christopa Stampfera z Uniwersytetu RWTH w Akwizgramie. Jego zespół zademonstrował, że struktura ta może pomieścić elektron w jednej warstwie i dziurę w drugiej. Co więcej, kwantowe stany spinowe tych dwóch jednostek są niemal idealnymi swoimi lustrzanymi odbiciami.

3. Płytką z fotonicznymi chipami do zastosowania w komputerach kwantowych





4. Wizualizacja warstw grafenowych

Kropka kwantowa to niewielki kawałek półprzewodnika o właściwościach elektronicznych, które bardziej przypominają atom niż materiał objętościowy. Na przykład, elektron w kropce kwantowej jest wzbudzany do skwantowanych poziomów energetycznych, podobnie jak w atomie. To przypominające atom zachowanie można precyzyjnie dostroić, dostosowując rozmiar i kształt kropki kwantowej. Kropka kwantowa może być wykonana z małych kawałków grafenu, arkusza węgla o grubości zaledwie jednego atomu. Takie kropki kwantowe mogą być wykonane tylko z jednego arkusza grafenu, dwóch arkuszy (grafen dwuwarstwowy) lub więcej (4).

Jednym z obiecujących zastosowań grafenowych kropek kwantowych jest tworzenie bitów kwantowych (kubitów), które przechowują informacje w stanach spinowych elektronów. Dziury, jak wiadomo z fizyki tranzystorów, powstają w półprzewodnikach, gdy elektron jest wzbudzony. Dwuwarstwowy grafen może uwięzić elektrony i dziury po przyłożeniu do nich zewnętrznego napięcia – tworząc strukturę bramki. W 2018 r. udało się po raz pierwszy osiągnąć

w takim układzie w dwuwarstwowym grafenie przerwę pasmową indukowaną przez pole elektryczne. Była to duża rzecz, gdyż wcześniej grafen zawodził swoich entuzjastów w dziedzinie elektroniki właśnie dlatego, że nie oferował pasma zabronionego, tak jak półprzewodniki. Osiągnięcie to pozwalało na transport przez tworzenie i anihilację pojedynczych par elektron-dziura o przeciwnych liczbach kwantowych, Zdaniem uczonych, powinno być możliwe łączenie kubitów elektronowo-spinowych na większe odległości, przy jednoczesnym bardziej niezawodnym odczytywaniu ich spinowo symetrycznych stanów. Dzięki temu komputery kwantowe mogłyby zyskać na skali i bezbłędności.

Poszukiwania nowej elektroniki kierują się coraz bardziej egzotyczne obszary fizyki – w tym wydaniu MT piszemy w Infozoomie o postępach firmy Microsoft, która buduje kwantowe układy obliczeniowe oparte na tzw. fermionach Majorany. A to nie koniec egzotyki w badaniach nad nową elektroniką. Według niedawnej publikacji w „Advanced Materials”, Korea Research Institute of Standards and Science (KRISS)

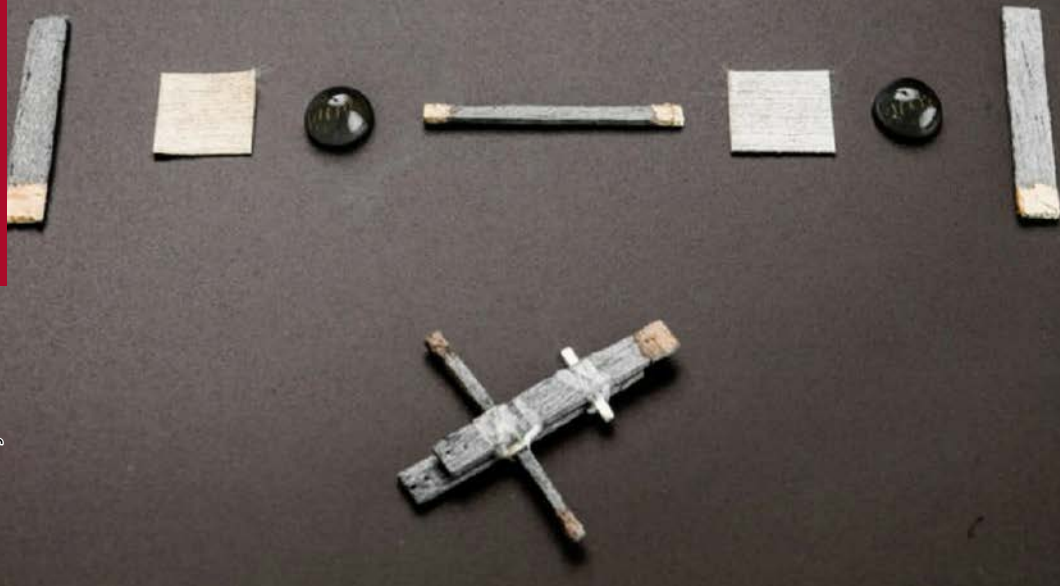
Na przekór. Superkot. Klub komiksowy. Tom 3

Dav Pilkey

Wydawnictwo Jaguar, liczba stron: 224, cena: 39,90 zł

Klub komiksowy rozwija się w różnych kierunkach! Naomi, Melvin i rodzeństwo próbują znaleźć swój cel. Naomi wpada na pomysł szybkiego wzbogacenia się, co wywołuje wiele zamieszania i emocji. Pomysł spotyka się z odrzuceniem, ale przyjaciele próbują pozostać wierni swojej wizji. Na dodatek do klasy przychodzi niespodziewany gość, aby namieszać jeszcze bardziej... Czy żądza pieniędzy i władzy przyćmi Naomi cel? Czy to oznacza koniec Klubu Komiksowego? Czy Klub będzie kiedyś taki jak dawniej? W tej wyjątkowej komiksowej serii wielokrotnie nagradzany autor i ilustrator Dav Pilkey używa różnych technik – w tym farb akrylowych, kolorowych ołówków, fotografii, kolażu, gwaszu, akwareli i wielu innych – aby zilustrować twórczy cel każdej żaby i zachęcić do pracy w zespole. Kalejdoskop stylów artystycznych w połączeniu ze znakami rozpoznawczymi Pilkeya i humorem sprzyja rozwijaniu kreatywności, uczenia współpracy, niezależności i empatii. Czytelnicy w każdym wieku będą cieszyć się tym zabawnym i ekscytującym komiksem!





5. Komponenty drewnianego tranzystora

opracował niedawno 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 spintroniki. Zespół KRISS opracował metodę jednolitej kontroli anizotropii magnetycznej przez wykorzystanie wodoru w izolatorach z tlenku glinu. Opracowanie tranzystora spintronicznego ma przyspieszyć rozwój m.in. urządzeń neuromorficznych i bramek logicznych, przy znacznych korzyściach w zużyciu energii, stabilności i szybkości w porównaniu z tradycyjnymi urządzeniami elektronicznymi.

Z drugiej strony dla elektroniki mniej wymagającej szuka się rozwiązań może nie z obszarów fizyki egzotycznych cząstek, ale wciąż wykorzystującej dość egzotyczne pomysły i materiały, takie jak... drewno. W ostatnim

czasie naukowcy z uniwersytetów szwedzkich Linköping i KTH stworzyli pierwszy na świecie drewniany tranzystor. Twarde drewno balsa pozbawili ligniny i wypełnili mieszaniną polimerów przewodzących o nazwie poli(3,4-etylenodioksytyofen)-polistyrenosulfonian lub PEDOT:PSS. Układając jednostki o grubości milimetra, które działały jako elektrody i kanały, zespół odkrył, że powstaje w ten sposób rodzaj tranzystora (5). Po podaniu napięcia 6 V, kanał wypełnia się elektronami, stopniowo zamykając się i przełączając przełącznik w pozycję „wyłączony”. Wyłączenie drewnianego urządzenia zajmuje około sekundy, podczas gdy przywrócenie go do pozycji włączonej zajmuje około pięciu sekund. Biodegradowalna elektronika mogłaby, zdaniem uczonych, być wykorzystywana w zdalnych czujnikach, które muszą się łatwo zepsuć lub niepozornych urządzeniach zasilanych w środowisku. ■

Mirosław Usidus

Pocztówka

Anne Berest

Wydawnictwo W.A.B., liczba stron: 496, cena: 54,99 zł

W skrzynce pocztowej matki Anny Berest znalazła się dziwna pocztówka. Nie była podpisana. Zdjęcie opery paryskiej Garniera, a na odwrocie imiona dziadków jej matki oraz jej ciotki i wujka, którzy zginęli w Auschwitz w 1942 roku. Autorka postanowiła dowiedzieć się, kto ją przysłał. Wynajęła prywatnego detektywa i kryminologa, przepytła mieszkańców wioski, w której zatrzymano jej przodków. Udało jej się odtworzyć epopeję rodziny Rabinowiczów: ucieczkę z Rosji, przeprowadzkę na Łotwę, a potem do Palestyny, osiedlenie w Paryżu i tragiczne wojenne losy. I udało się wreszcie odkryć, kto był nadawcą dziwnej pocztówki... „Pocztówka” to zarówno powieść detektywistyczna, saga rodzinna, jak i próba zrozumienia własnej tożsamości.



Maszyny cyfrowe zmieniły i ukształtowały bieg naszej cywilizacji. Coraz częściej jednak szukamy czegoś nowego, lub nie tak nowego, ale na pewno alternatywnego sposobu przetwarzania i obliczania, gdyż mamy wrażenie, że cyfrowa technika dociera do granic swoich możliwości.

Komputery jakie znamy mogą nie przetrwać rewolucji w technice i nauce

CZY TO JUŻ KRES ŚWIATA ZER I JEDYNEK?

Bernd Ulmann (1), matematyk znany jako „ewangelista komputerów analogowych”, opisuje działanie komputera cyfrowego w następujący sposób: dzielisz problem na wiele cienkich plasterków i rozwiązujesz go krok po kroku pod centralną kontrolą precyzyjnie zdefiniowanego zestawu instrukcji – algorytmu. Niezbędne instrukcje są przechowywane w pamięci komputera. Chociaż proces ten był historycznie nadzwyczaj wydajny, dociera dziś do granic efektywności. Jednym z głównych problemów jest ogromne zużycie energii na miliardy operacji arytmetycznych, które komputer musi wykonać w każdej sekundzie, przetwarzając złożone problemy. Nie można bez końca zwiększać liczby równoległych jednostek obliczeniowych. Ostatecznie niemożliwe jest uczynienie procesorów roboczych w tych maszynach mniejszymi niż atomy substancji używanych do budowy obwodów.

Dlatego Ulmann, i wielu innych, planuje powrót do idei komputera analogowego, który przeprowadza obliczenia w zupełnie inny sposób. Warto przypomnieć, że była to dominująca forma maszyny obliczeniowej od wczesnych lat czterdziestych do lat siedemdziesiątych XX wieku. Komputer cyfrowy ma tylko kilka jednostek arytmetycznych, które przetwarzają proste instrukcje niezwykle szybko. Komputer analogowy składa się z dużej liczby elementów obliczeniowych, w najdalej idących założeniach nawet milionów, które są ze sobą połączone. Zamiast postępować krok po kroku i transferować miliardy

bitów tam i z powrotem, przetwarzany problem jest w układzie analogowym w pewnym sensie odtwarzany elektronicznie w obwodach. Wynik odczytuje się przez pomiar.

Komputery analogowe dążą do odtworzenia sposobu działania biologicznych obwodów neuronowych w mózgu, które mogą przetwarzać ogromne ilości danych, nie zużywając przy tym więcej energii niż 30-watowa żarówka. Na razie z wydajnością tych systemów nie jest tak dobrze jak biologicznym pierwowzorze, ale to może się zmienić. Ulmann przewiduje więc, że głównym obszarem ich zastosowań będzie dziedzina sztucznej inteligencji, która stara się naśladować działanie sieci neuronowych. Sam pracuje nad analogowym chipem komputerowym, w którym poszczególne obwody można dowolnie konfigurować. Byłby on kontrolowany przez komputer cyfrowy, który wraz z analogowym tworzy system hybrydowy.



1. Bernd Ulmann

Obliczenia w hiperwymiarze

Wielu ekspertów narzeka, że znane i używane obecnie sieci neuronowe są bardzo energochłonne. Istotnym problemem jest także brak przejrzystości ich działania. Systemy te są bowiem tak skomplikowane, że nikt tak naprawdę nie rozumie, co robią i dlaczego działają tak dobrze. To problem znany i niejednokrotnie opisywany w MT, nazywany problemem „czarnej skrzynki”.

Rozważmy na przykład sztuczną sieć neuronową (ANN), która odróżnia koła od kwadratów. Można sobie wyobrazić funkcjonowanie w systemie dwóch neuronów w warstwie wyjściowej, jednego wskazującego koło i drugiego – kwadrat. Jeśli chcemy, by ANN rozpoznawała również kolor kształtu – niebieski lub czerwony, to potrzeba czterech neuronów wyjściowych, po jednym dla niebieskiego koła, niebieskiego kwadratu, czerwonego koła i czerwonego kwadratu. Obsługa większej liczby funkcji oznacza potrzebę korzystania z jeszcze większej liczby neuronów – jednostek obliczeniowych.

Badacze zwracają uwagę, że w mózgu jest inaczej – informacje są reprezentowane przez równoczesną aktywność wielkich liczb neuronów. np. postrzeganie fioletowego mercedesa nie jest zakodowane jako funkcja pojedynczego neuronu, ale jako aktywność tysięcy neuronów. Ten sam zestaw neuronów, opalając impulsy inaczej, może reprezentować zupełnie inną koncepcję (choćby różowego Cadillaca). Stanowi to punkt wyjścia dla nowego podejścia do obliczeń komputerowych, znanego jako hiperwymiarowe, w którym każda informacja, taka jak idea samochodu, jego marka, model lub kolor, lub wszystkie razem, jest reprezentowana jako jednostka nazywana hiperwymiarowym wektorem.

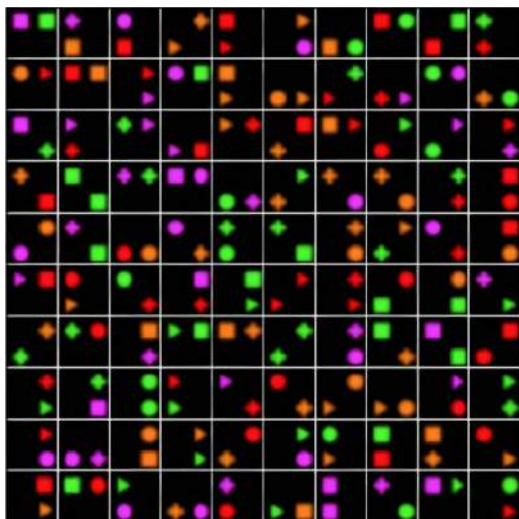
Wektor to po prostu uporządkowana kombinacja liczb. Na przykład wektor 3D opisany jest za pomocą trzech liczb: współrzędnych x , y i z punktu w przestrzeni 3D. Wektor hiperwymiarowy lub hiperwektor może być kombinacją 10 tys. liczb reprezentujących punkt w 10000-wymiarowej przestrzeni. Wróćmy do przykładu obrazów kolorowych kół i kwadratów. Najpierw potrzebujemy wektorów do reprezentowania zmiennych – KSZTAŁT i KOLOR. Następnie potrzebujemy również wektorów dla wartości, które można przypisać do zmiennych – KOŁO, KWADRAT, NIEBIESKI i CZERWONY. Wektory muszą być rozróżnialne. Odrębność ta może być określona ilościowo przez właściwość zwaną ortogonalnością, która oznacza prostopadłość. W przestrzeni 3D istnieją trzy wektory, które są względem siebie ortogonalne, osie x , y i z . W przestrzeni 10 tys. wymiarów istnieje 10 tys.

takich wzajemnie ortogonalnych wektorów. Jeśli jednak pozwolimy wektorom być prawie ortogonalnymi, liczba takich odrębnych wektorów w przestrzeni wielowymiarowej eksploduje. W przestrzeni 10 000 wymiarów istnieją miliony wektorów.

Tworzymy odrębne wektory reprezentujące KSZTAŁT, KOLOR, KOŁO, KWADRAT, NIEBIESKI i CZERWONY. Ponieważ istnieje tak wiele możliwych prawie ortogonalnych wektorów w wielowymiarowej przestrzeni, można po prostu przypisać sześć losowych wektorów do reprezentowania sześciu elementów. Prawie na pewno będą one prawie ortogonalne. Jak nimi manipulować za pomocą operacji matematycznych? Pierwszą operacją jest mnożenie. Jest to sposób łączenia idei. Na przykład pomnożenie wektora KSZTAŁT z wektorem KOŁO łączy te dwa wektory w reprezentację idei „KSZTAŁT jest KOŁEM”. Ten nowy „związany” wektor jest prawie ortogonalny zarówno do KSZTAŁT, jak i KOŁO. Druga operacja, dodawanie, tworzy nowy wektor, który reprezentuje tak zwaną superpozycję pojęć. Na przykład, można wziąć dwa powiązane wektory „KSZTAŁT jest KOŁEM” i „KOLOR jest CZERWONY” i dodać je do siebie, aby utworzyć wektor reprezentujący okrągły kształt w kolorze czerwonym. Trzecią operacją jest permutacja – polega ona na przedstawianiu poszczególnych elementów wektorów. Na przykład, jeśli trójwymiarowy wektor ma wartości oznaczone x , y i z , permutacja może przenieść wartość x na y , y na z , a z na x . Permutacja pozwala budować strukturę i radzić sobie z sekwencjami, zdarzeniami, które dzieją się jedno po drugim. Jeśli mamy dwa zdarzenia, reprezentowane przez hiperwektory A i B , to gdyby nałożyć je na jeden wektor, zniszczyłoby to informacje o kolejności zdarzeń. Połączenie dodawania z permutacją zachowuje kolejność – zdarzenia można odtworzyć w kolejności, odwracając operacje.

Razem te trzy operacje okazały się wystarczające do stworzenia formalnej algebry hiperwektorów, która pozwoliła na symboliczne rozumowanie. W 2015 roku Eric Weiss, zademonstrował, jak odwzorować złożony obraz w postaci pojedynczego hiperwymiarowego wektora, który zawiera informacje o wszystkich obiektach na obrazie, w tym o ich właściwościach, takich jak kolory, pozycje i rozmiary. Wkrótce kolejni badacze zaczęli opracowywać algorytmy hiperwymiarowe, aby odtworzyć proste zadania dla głębokich sieci neuronowych, takie jak np. klasyfikowanie obrazów (2).

To jednak dopiero początek. Mocną stroną obliczeń hiperwymiarowych jest zdolność do komponowania i dekomponowania hiperwektorów na potrzeby



2. Zestaw do treningu wizualnego komputera hiperwymiarowego

rozumowania. Najnowsza demonstracja tego zjawiska miała miejsce w marcu 2023 r., kiedy Abbas Rahimi i współpracownicy z IBM Research w Zurychu wykorzystali obliczenia hiperwymiarowe w sieciach neuronowych do rozwiązania klasycznego problemu w abstrakcyjnym rozumowaniu wizualnym, które jest sporym wyzwaniem dla typowych sieci neuronowych, a nawet niektórych ludzi. Chodzi test z obrazami obiektów geometrycznych na siatce, np. 3 na 3 pola. Jedna pozycja w siatce jest pusta. Zadanie polega na wyborze z zestawu kandydujących obrazów tego, który najlepiej pasuje do pustego miejsca. Znamy to m.in. z testów na inteligencję. Aby rozwiązać ten problem przy użyciu obliczeń hiperwymiarowych, naukowcy najpierw stworzyli słownik hiperwektorów do reprezentowania obiektów na każdym obrazie. Każdy hiperwektor w słowniku reprezentuje obiekt i pewną kombinację jego atrybutów. Następnie zespół przeszkolił sieć neuronową do badania obrazu i generowania dwubiegunowego hiperwektora – +1 lub -1 – który jest zbliżony do pewnej superpozycji hiperwektorów w słowniku. Wygenerowany hiperwektor zawiera zatem informacje o wszystkich obiektach i ich atrybutach na obrazie. Gdy sieć wygeneruje hiperwektory dla każdego z obrazów kontekstowych i dla każdego kandydata na puste miejsce, inny algorytm analizuje hiperwektory, aby utworzyć rozkłady prawdopodobieństwa dla liczby obiektów na każdym obrazie, ich wielkości i innych cech. Te można przekształcić w hiperwektory, umożliwiając wykorzystanie algebry do przewidywania najbardziej prawdopodobnego obrazu do wypełnienia wolnego miejsca. To podejście

dało 88-procentową dokładność w przypadku jednego zestawu problemów, zaś rozwiązania oparte wyłącznie na sieciach neuronowych wykazywały 61 proc. dokładności. Zespół wykazał również, że dla siatek 3 na 3, ich system był prawie 250 razy szybszy niż tradycyjna metoda, która wykorzystuje reguły logiki symbolicznej do rozumowania.

Wydajność dzisiejszych komputerów gwałtownie spada, jeśli błędy spowodowane np. losową zamianą bitów (0 staje się 1 lub odwrotnie) nie mogą zostać skorygowane przez wbudowane mechanizmy korekcji błędów. Obliczenia hiperwymiarowe lepiej tolerują błędy, ponieważ nawet jeśli hiperwektor narażony jest na znaczną liczbę losowych odwróceń bitów, nadal jest zbliżony do oryginalnego wektora. Wykazano, że systemy te są co najmniej dziesięć razy bardziej odporne na błędy sprzętowe niż tradycyjne sieci neuronowe, które same w sobie są o rzędy wielkości bardziej odporne niż klasyczne architektury obliczeniowe. Kolejną zaletą obliczeń hiperwymiarowych jest przejrzystość – metoda wyraźnie wskazuje, dlaczego system wybrał daną odpowiedź. Nie ma więc „czarnej skrzynki”.

Wszystkie te zalety w porównaniu z tradycyjnymi obliczeniami sugerują, że obliczenia hiperwymiarowe dobrze nadają się do nowej generacji niezwykle wytrzymałego sprzętu o niskim poborze mocy. Są one również kompatybilne z „systemami obliczeniowymi w pamięci”, które wykonują obliczenia na tym samym sprzęcie, który przechowuje dane (w przeciwieństwie do tradycyjnych cyfrowych komputerów opartych na architekturze von Neumanna, które nieefektywnie przenoszą dane między pamięcią a jednostką centralną). Urządzenia nowego typu mogą mieć konstrukcję analogową, działając przy bardzo niskim napięciu, co czyni je energooszczędnymi, ale także podatnymi na przypadkowe zakłócenia

Pomimo zalet, konstrukcje komputerów opartych na obliczeniach hiperwymiarowych są nadal w powijakach. Wciąż muszą zostać przetestowane pod kątem rzeczywistych problemów i w znacznie większej niż laboratoryjna skali, bliższej rozmiarom nowoczesnych sieci neuronowych. Wymagają niezwykle wydajnego sprzętu, co w tej chwili stawia pod znakiem zapytania ich opłacalność.

Wobec kwantów znane komputery mogą być bezsilne

Oczywiście obliczenia hiperwymiarowe to nie jedyny kierunek w poszukiwaniu alternatywnych wobec cyfrowych systemów opartych na architekturze von Neumanna. Piszemy o nich obszernie, gdyż



3. Komputer kwantowy IBM

są mniej znaną koncepcją. Mniej znaną, niż np. wielokrotnie opisywana w „Młodych Techniku” nadzieja na rewolucję komputerową – maszyny kwantowe (3), które budzą zarówno entuzjazm, trochę przedwcześnie, jak też wielkie obawy, związane przede wszystkim z bezpieczeństwem naszych systemów, kryptografii, transakcji, kont bankowych i internetu jako całości. Systemy kwantowe mogą sprawić, że komputery, jakie znamy, nie przetrwają, nie tylko dlatego, że nie sprostają konkurencji.

Cały nasz Internet opiera się na dużych liczbach pierwszych, które są trudne do rozpoznania. Nowoczesne komputery używają szyfrowania znanego jako RSA, które opiera się na asymetrycznych kluczach, więc masz jeden klucz do szyfrowania danych i inny klucz do ich odszyfrowania. Klucz blokujący dane jest publicznie znany, ale tylko ty masz klucz deszyfrujący, klucz prywatny. Teoretycznie można obliczyć czyjś klucz prywatny na podstawie jego klucza publicznego. Ale nawet przy użyciu najpotężniejszego na świecie konwencjonalnego komputera zajęłoby to lata. Wynika to z faktu, że podstawą kluczy są dwie ogromne liczby pierwsze mnożone przez siebie. Klucz publiczny zawiera wynik tego mnożenia, ale trzeba znać dwie liczby pierwsze, aby utworzyć klucz prywatny. Ponieważ problem odkrycia tych liczb jest w zasadzie nierozwiązywalny dla zwykłych komputerów, zanim ktoś je złamie, klucz nie jest już używany. Jednak potężny komputer kwantowy może wykonać te obliczenia w ciągu kilku sekund. Komputery kwantowe są obecnie stosunkowo słabe

i przede wszystkim niestabilne. Uważa się jednak, że kiedy uda się zbudować maszyny z trwałymi kubitami w dużej liczbie, każdy, kogo na to będzie stać, będzie mógł złamać szyfrowanie internetowe i ukraść dowolny fragment danych.

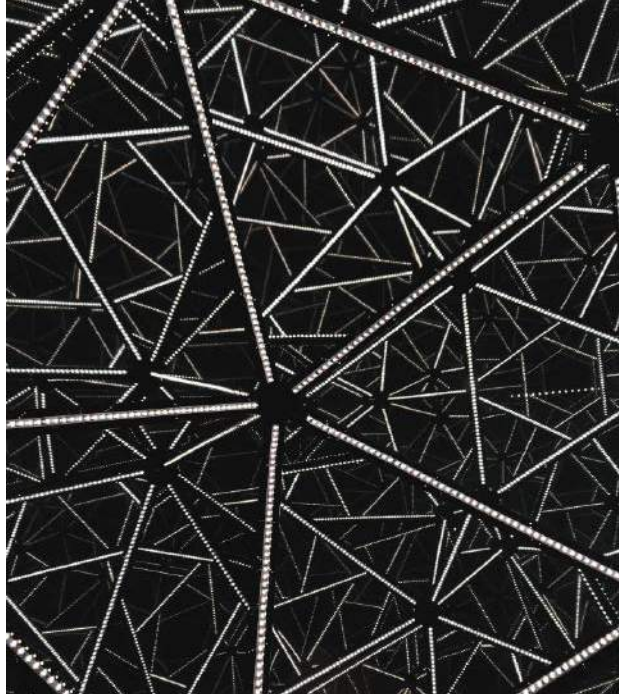
Świadomość zagrożenia sięga lat 90., kiedy to amerykański profesor Peter Shor opracował matematyczny opis sposobu, w jaki sposób komputer kwantowy mógłby złamać „kryptografię klucza publicznego”, na której nadal polegamy. Zespół chińskich naukowców ogłosił w styczniu tego roku, że znalazł praktyczny sposób realizacji tego złowrogiemu planu. Potem amerykańscy eksperci uspokajali, że Chińczycy opracowali wciąż jedynie teoretyczną koncepcję, a nie maszynę, która faktycznie jest do tego zdolna. Wkrótce potem, w lutym Google, ogłosiło przełom w dziedzinie obliczeń kwantowych (kolejny już – w 2019 Google ogłosiło, że osiągnęło supremację kwantową). Jednak nawet najbardziej optymistyczni eksperci podają rok 2027 jako najwcześniejszy termin, w którym powstanie w pełni sprawna maszyna kwantowa, choć to są optymiści – wielu innych przesuwają tę perspektywę o lata i dekady w przyszłość.

Mimo to liczne służby wywiadowcze już twierdzą, że zagrożenia są realne. Brytyjskie Narodowe Centrum Cyberbezpieczeństwa (NCSC), „uznaje za poważne zagrożenie, jakie komputery kwantowe stanowią dla długoterminowego bezpieczeństwa kryptograficznego”. Rząd USA twierdzi, że mogą one „poważnie zagrozić poufności i integralności komunikacji cyfrowej”, a amerykańska

Agencja Bezpieczeństwa Narodowego (NSA) stwierdza, że to teraz jest „czas na planowanie, przygotowanie i budżetowanie” nowych zabezpieczeń. W grudniu 2022 r. Joe Biden podpisał ustawę Quantum Computing Cybersecurity Preparedness Act. Zobowiązuje ona amerykańskie agencje rządowe do rozpoczęcia tworzenia listy wszystkich systemów, które mogą być podatne na działanie komputera kwantowego i opracowania sposobu ich aktualizacji do kryptografii post-quantowej, z wykorzystaniem zwycięzców konkursu NIST.

Gdyby ktoś o niedobrych intencjach zyskał dostęp do działającego, stabilnego i silnego systemu kwantowego, to specjaliści są dość zgodni, że zniszczyłoby to Internet, jaki znamy i zniknęłoby bezpieczeństwo, jakie znamy. Niemożliwe stałyby się prywatne e-maile, historia wyszukiwania i korzystania w sieci, bezpieczne płatności, szyfrowanie wiadomości i danych. Łatwy byłby dostęp do cudzego konta w mediach społecznościowych. A to tylko początek listy zagrożeń.

Co robić w tej sytuacji. Sytuacja nie jest beznadziejna nie tylko dlatego, że do operacyjnych komputerów kwantowych jeszcze nam (i cyberprzestępcom) sporo brakuje. Także dlatego, że już teraz pracuje się nad metodami zabezpieczeń w przyszłej erze kwantowej, np. nad kryptografią opartą na kratkach czy też siatkach (4). Ta metoda szyfrowania nie wykorzystuje złożonych algorytmów, jak RSA. Zamiast tego wykorzystuje geometrię w aż stu wymiarach. Szyfrowanie to działa poprzez przedstawienie siatki składającej się z wielu węzłów (punktów w tej teoretycznej przestrzeni). Kod bazowy do szyfrowania jest ukryty w tych dwóch węzłach, podobnie jak w przypadku dwóch liczb pierwszych w RSA. Aby złamać



4. Geometria matematyczna jako zbawienie kryptografii

kod, komputer musi znaleźć te dwa węzły. W trzech wymiarach jest to łatwe. Ale w stu wymiarach jest to zadanie trudno wykonalne, nawet dla komputerów kwantowych. Problem w tym, że sieci szyfrujące w ten sposób mają gigantyczne rozmiary. Klucz RSA składa się z 256 bitów nie zajmując wiele przepustowości łączy. Klucz kratowy jest o rzędy wielkości większy, dramatycznie spowalniając połączenie. Jest to duże wyzwanie, z którym niełatwo sobie poradzić. W efekcie, podobnie zresztą jak komputery kwantowe, technika ta jest daleka od użyteczności. ■

Mirosław Usidus

Umowa na miłość

Falon Ballard

Wydawnictwo MUZA SA, liczba stron: 384, cena: 46,90 zł

Kiedy awans kolejny raz przechodzi jej koło nosa, Sadie nie potrafi ukryć emocji i w burzliwej atmosferze rozstaje się z branżą finansową. Poświęca pracę tak dużo czasu, że nie miała nawet chwili na życie osobiste. Rozżalona i zawiedziona marzy tylko o tym, by odreagować. Loguje się na portalu randkowym i następnego dnia umawia ze swoim wybrankiem. Jack na pierwszy rzut oka wydaje się nudnym i nieciekawym facetem. Sadie jest rozczarowana. W dodatku okazuje się, że omyłkowo zalogowała się na portalu z nieruchomościami i młody okularnik nie jest wcale kandydatem na randkę, tylko właścicielem kamienicy na Brooklynie. I proponuje Sadie wynajem pokoju po korzystnej cenie. Atrakcyjne warunki najmu i mało atrakcyjny właściciel. Sadie zależy, by zacząć życie od nowa i spełnić marzenia o własnej firmie florystycznej, chętnie więc przyjmuje tę propozycję. Wkrótce dzięki jej zdolnościom plastycznym kamienica zyskuje niepowtarzalny klimat. Stopniowo ociepla się także relacja między Jackiem a lokatorką.





1. Chipowe starcie USA – Chiny

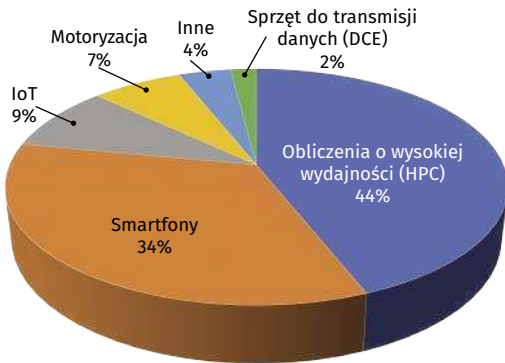
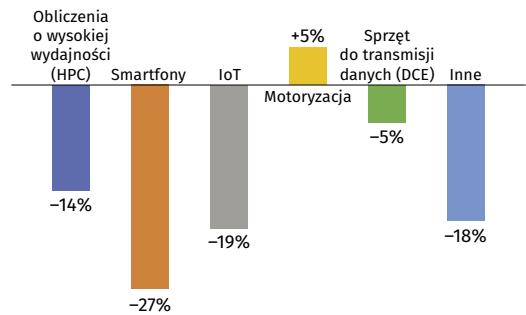
Semiconductor Manufacturing International Corp. (SMIC), największy chiński producent chipów, bez rozgłosu, w maju 2023 r. usunął proces produkcji 14 nm z listy swoich usług na swojej stronie internetowej. Ta niepozorna (z pozoru) informacja to znak obecnego czasu, czasu bezwzględnej wojny na światowym rynku półprzewodnikowym.

W świecie półprzewodników trwa bezwzględna rywalizacja mocarstw i firm

WOJNA O CHIPY

Wprowadzane od 2022 roku nowe zasady kontroli eksportu rządu USA zabraniają chińskim producentom (1) nabywania instrumentów i technologii wymaganych do produkcji nieplanarnych

tranzystorowych układów logicznych o wymiarach 14nm/16nm lub mniejszych, układów 3D NAND ze 128 lub więcej aktywnymi warstwami oraz układów scalonych DRAM, również o określonym poziomie miniaturyzacji – w dół. Wraz z podobnymi ograniczeniami wprowadzonymi przez Holandię, Japonię i Tajwan, które mają wejść lub już weszły w życie w 2023 roku, chińskie firmy takie jak SMIC i inny czołowy na tamtym rynku producent komponentów, YMTC, zostają odcięte od możliwości pozyskania sprzętu potrzebnego do produkcji chipów w najnowszych węzłach produkcyjnych. Bez dostępu do zaawansowanego sprzętu i części zamiennych od dostawców takich jak holenderski ASML, Applied Materials, KLA i Lam Research, SMIC raczej nie jest na razie w stanie budować chipów przy użyciu najnowszych technologii produkcji, stąd usunięcie platformy 14 nm z listy swoich technologii.

Przychody według platform w 1 kw. 2023 r.**Tempo wzrostu według platformy (kw/kw)****2. Udział platform w dochodach ze sprzedaży procesorów****Utrudnić Chińczykom i ich spowolnić**

Chipy półprzewodnikowe stały się najgorętszym towarem w globalnym wyścigu o technologiczną dominację. Stanowią podstawowy składnik urządzeń elektronicznych, krytyczny w wielu branżach, od telekomunikacji, smartfonów, nowoczesnych samochodów (2), sztucznej inteligencji i informatykę, po opiekę zdrowotną, czystą energię i zastosowania wojskowe. Oznaki eskalacji wojny handlowej na rynku chipów były widoczne już w październiku 2021 roku. Wprowadzone przez USA a potem kolejne kraje sankcje zaogroiły konflikt.

Biorąc pod uwagę, że 95 proc. wszystkich chipów AI używanych obecnie w Chinach (tak jak na całym świecie) to amerykańskie procesory graficzne Nvidia, a większość pozostałych to również amerykańskie chipy AMD, zakaz ten, według przewidywań, ma mieć katastrofalne skutki dla chińskiej branży AI. Ale rząd USA nie poprzestał na jedynie na spręcie. Zidentyfikował szereg strategicznych „wąskich gardeł”, bez których nie można utrzymać produkcji chipów AI, i odciął Chinom dostęp również do nich. Dotyczy to oprogramowania potrzebnego do projektowania układów chipów, przede

wszystkim rozwiązań automatyzacji projektowania (EDA), sprzętu produkcyjnego potrzebnego do budowy chipów, a nawet komponentów, które trafiają do tego sprzętu produkcyjnego. Wszystkie trzy wiodące na świecie firmy EDA – Mentor Graphics, Cadence Design Systems i Synopsys – to firmy amerykańskie. Najważniejszy sprzęt do produkcji półprzewodników pochodzi ze Stanów Zjednoczonych lub ich sojuszników. Każdy podmiot działający w Chinach ma obecnie zakaz dostępu do tych produktów.

Dlaczego Amerykanie to robią. Najprostsza i najbardziej odpowiadająca brzmie – by utrudnić i spowolnić chińskie przygotowania do starcia z USA, na każdym polu, także militarnym.

Europa chce sama produkować a Chiny kontratakują

Inne kraje, w tym sojusznicy USA, choć po części współpracują z USA, to jednak mają nieco inne podejście. Europa np. chce odtworzyć, kiedyś silny na naszym kontynencie własny przemysł półprzewodnikowy. Parlament Europejski i Rada osiągnęły wiosną 2023 wstępne porozumienie w sprawie rozporządzenia

3. Wizualizacja powstającego zakładu wytwórczego komponentów elektronicznych Infineon pod Dreznem w Niemczech

mającego na celu zwiększenie produkcji półprzewodników w regionie (3). Ostatecznym celem jest zwiększenie udziału UE w rynku globalnym do co najmniej 20 procent do 2030 roku. Kiedy początkowo zaproponowano ten ruch, udział regionu wynosił 10 procent. Ustawa obejmuje inicjatywę „Chipy dla Europy”, która ma przyciągnąć 43 mld euro z funduszy publicznych i inwestycji prywatnych, w tym 3,3 mld euro z budżetu UE. „Mikroprocesory są podstawą innowacji i konkurencyjności przemysłowej Europy w cyfrowym świecie”, mówiła cytowana w mediach Margrethe Vestager z Komisji Europejskiej. Warto tu dodać, że Stany Zjednoczone przeznaczają 52 mld USD na własny „Chips and Science Act”, aby wzmocnić produkcję i łańcuchy dostaw oraz „przeciwwstawić się Chinom”.

UE przyjmuje bardziej zniuansowane podejście wobec Chin niż USA. „Uważam, że oddzielenie się od Chin nie jest ani opłacalne, ani nie leży w interesie Europy”, powiedziała w marcu przewodnicząca Komisji Europejskiej Von der Leyen. „Nasze relacje nie są czarne lub białe – i nasza reakcja też nie może taka być”. Wcześniej jednak Thierry Breton z Komisji Europejskiej podkreślił zaangażowanie Unii we wspieranie amerykańskiego celu, jakim jest ograniczenie chińskiego przemysłu półprzewodników. „W pełni zgadzamy się z celem pozbawienia Chin najbardziej zaawansowanych chipów”, powiedział podczas przemówienia w Centrum Studiów Strategicznych i Międzynarodowych. „Nie możemy pozwolić Chinom na dostęp do najbardziej zaawansowanych technologii w zakresie półprzewodników, kwantowych, chmurowych, brzegowych, sztucznej inteligencji, łączności i tak dalej”.

W czerwcu 2023 r., po miesiącach nacisków ze strony USA, Holandia oficjalnie wprowadziła bardziej rygorystyczne kontrole eksportu zaawansowanych maszyn do produkcji chipów, powołując się na „interesy bezpieczeństwa narodowego”. Począwszy od września, holenderskie firmy z branży będą musiały najpierw uzyskać pozwolenie przed wysyłką do krajów, które mogą wykorzystywać technologię do zastosowań wojskowych, takich jak Chiny. Posunięcie to stanowi poważny cios, ponieważ Holandia to ASML, wiodący na świecie producent wysokiej klasy sprzętu do produkcji chipów. ASML miał już od 2019 r. zakaz sprzedaży swoich najbardziej zaawansowanych maszyn do Chin. Według firmy, nowe sankcje obejmą również jej najbardziej zaawansowane systemy litografii zanurzeniowej DUV.

W reakcji na decyzje Holandii, chiński ambasador Tan Jian zauważył, że takie praktyki, stosowane przez USA i ich sojuszników, „naruszają międzynarodowe zasady handlu” i „podważają globalne łańcuchy dostaw”. Dodał, że Chiny nigdy nie zaszkodziły

europejskim interesom, ale ostrzegł, że „będą konsekwencje”, jeśli holenderski rząd wprowadzi ograniczenia. I wkrótce po ogłoszeniu przez Holandię, chińskie Ministerstwo Handlu ogłosiło kontrolę eksportu galu i germanu, dwóch metali o kluczowym znaczeniu dla produkcji elektroniki i półprzewodników. Począwszy od sierpnia 2023 r., firmy, które chcą dostarczać metale za granicę, będą musiały ubiegać się o licencję i zgłaszać szczegóły dotyczące nabywców i ich wniosków. Jako powód tego posunięcia po raz kolejny podano „obawy o bezpieczeństwo narodowe” Chin. Według europejskiego stowarzyszenia branżowego Critical Raw Materials Alliance (CRMA), Chiny produkują 80 proc. światowego galu i 60 proc. germanu. Oba metale znajdują się na unijnej liście surowców krytycznych, uznawanych za „kluczowe dla europejskiej gospodarki”. Metale te nie są szczególnie rzadkie lub trudne do znalezienia, ale dostarcza je niewielka grupa krajów. Jeśli chodzi o gal to, oprócz Chin także Japonia, Korea Południowa, Rosja i Ukraina – w przypadku germanu są to Belgia, Kanada i Rosja. Chiny dominują w produkcji tych dwóch metali nie dlatego, że są one rzadkie, ale dlatego, że dostarczają je relatywnie tanio. Analitycy z Eurasia Group określają działania Pekinu jako „strzał ostrzegawczy, a nie śmiertelny cios, mający przypomnieć Zachodowi, że Chiny mają opcje odwetowe”.

Zauważa się też, że Chiny mogą pokrzyżować szumne plany zielonej transformacji w UE. UE jest obecnie w 98 proc. zależna od Pekinu w zakresie dostępu do metali ziem rzadkich, które mają kluczowe znaczenie m.in. w takich dziedzinach jak magazynowanie wodoru, energia słoneczna i wiatrowa. Jest również całkowicie zależna od importu rafinowanego litu (głównego składnika baterii EV), który jest w około 80 proc. zmonopolizowany przez Chiny. Robert Habeck, niemiecki federalny minister gospodarki i działań na rzecz klimatu, podczas wydarzenia branżowego: „Jeśli zostanie to rozszerzone na przykład na lit, będziemy mieli zupełnie inny problem”.

Chiński organ nadzorujący cyberbezpieczeństwo zdecydował też niedawno, że amerykańska firma Micron, jeden z największych na świecie producentów chipów, nie przeszła przeglądu bezpieczeństwa narodowego. „Operatorzy krytycznej infrastruktury informatycznej w Chinach powinni zaprzestać kupowania produktów Micron”, powiedział organ nadzorczy.

Wydaje się, że ostrzeżenia Pekinu spełniło swój cel, a Stany Zjednoczone dążą do ożywienia stosunków z Chinami. Podczas lipcowej wizyty w chińskiej stolicy, sekretarz skarbu USA Janet Yellen powiedziała chińskiemu premierowi Li Qiangowi: „Dążymy do zdrowej konkurencji gospodarczej, która nie będzie polegała

na zasadzie zwycięzca bierze wszystko, ale która, przy sprawiedliwym zestawie zasad, może z czasem przynieść korzyści obu krajom”. Tonując wcześniejszą retorykę Waszymingtonu, Yellen podkreśliła, że Stany Zjednoczone są zainteresowane zmniejszeniem ryzyka handlowego, a nie oddzieleniem obu gospodarek.

Tajwańskie serce krzemowego globu

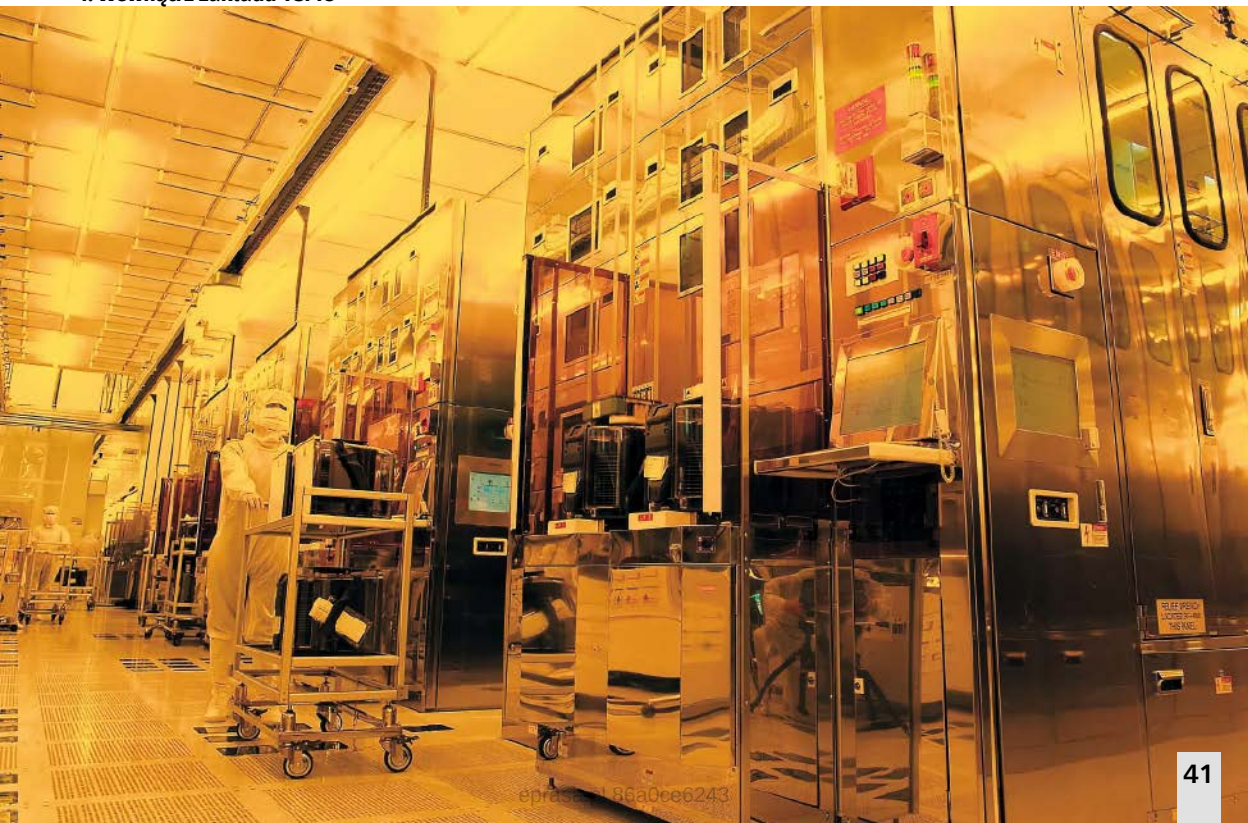
Tyle geopolityka. Przyjdźmy do przemysłu, w którym jest równie ciekawie. Obecnie jest tak, że właściwie wszystkie zaawansowane chipy sztucznej inteligencji na świecie wytwarza Taiwan Semiconductor Manufacturing Company (TSMC). Jak wszystkie to wszystkie, czyli nie tylko procesory graficzne Nvidii ale także chipy AI od Google, AMD, Amazon, Microsoft, Cerebras, SambaNova, Untether i każdego innego podmiotu, który stanął w szranki rywalizacji zbudowanie procesorów, które obsługują AI. Nic więc dziwnego, że magazyn „Time” opisał TSMC jako „najważniejszą firmę na świecie”. Dyrektor generalny Nvidii, Jensen Huang, ujął to bardziej metaforycznie: „Zasadniczo istnieje powietrze i TSMC”.

Zakłady produkcji chipów TSMC, budynki, w których fizycznie budowane są chipy – znajdują się na zachodnim wybrzeżu Tajwanu, zaledwie 180 km od Chin kontynentalnych. Obecnie Tajwan i Chiny znajdują się bliżej wojny niż od dziesięcioleci. Wraz

z eskalacją napięć, Chiny przeprowadzają ćwiczenia wojskowe wokół Tajwanu o niespotykanej dotąd skali i intensywności. Wielu decydentów politycznych w Waszyngtonie przewiduje, że Chiny dokonają inwazji na Tajwan do 2027 lub nawet 2025 roku. Konflikt między Chinami a Tajwanem byłby niszczycielski z wielu powodów. Ale dla przemysłu półprzewodnikowego w obecnym kształcie byłby kataklizmem.

Szybka myśl, by w te pędy wyprowadzić z tego niepewnego miejsca tak strategicznie ważną produkcję, równie szybko rozbija się o komplikacje i problemy. Produkcja półprzewodnikowa wymaga najczystszych metali na świecie, najdroższych maszyn na świecie (zdolnych do budowy tranzystorów o rozmiarze mniejszym niż 100 atomów), legionów wysoce wyspecjalizowanych inżynierów i niewiarygodnej precyzji, gdyż jedna drobina kurzu może zrujnować całą serię produkcyjną chipów, marnując miliony dolarów. Łańcuch dostaw półprzewodników jest skomplikowany i zglobalizowany. Jest również szokująco skoncentrowany w niektórych obszarach. Dla przykładu – 100 proc. światowej podaży maszyn do litografii w ultrafiolecie, złożonego sprzętu wymaganego do budowy zaawansowanych chipów, pochodzi od jednej firmy z Holandii – wspomnianej ASML. Jak ujął to Chris Miller w książce „Chip War”, „Żaden inny aspekt gospodarki nie jest tak zależny od tak niewielu firm”.

4. Wewnątrz zakładu TSMC



W początkach przemysłu półprzewodnikowego, od lat 50. do 70. ubiegłego wieku, wszystkie firmy produkujące chipy były zintegrowane pionowo. Producenci chipów, tacy jak Fairchild Semiconductor i Texas Instruments, realizowali każdy etap procesu produkcji półprzewodników we własnym zakresie. Sami projektowali, produkowali i sprzedawali swoje chipy. Każda firma produkująca chipy posiadała i zarządzała własnymi zakładami produkcyjnymi. Jednak poczynając od lat 80. XX wieku rozwinął się nowy model, oparty na specjalizacji. Pojawiły się dwa różne rodzaje firm produkujących chipy „producenci” chipów, którzy nie mają fabryk, którzy projektują, ale fizycznie nie produkują procesorów, oraz „odlewnie”, które produkują chipy zaprojektowane przez inne firmy.

Obecnie większość znanych firm na rynku układów scalonych, w tym Nvidia, Qualcomm, AMD, Broadcom, nie ma własnych fabryk. Nie produkują one własnych chipów. Zamiast tego polegają na zakładach takich jak TSMC, „odlewniach” (4). Tranzystory stają się coraz mniejsze, a układy scalone stają się coraz bardziej zaawansowane z każdym rokiem (nawet jeśli Prawo Moore’a zwalnia). Proces produkcji najnowocześniejszych chipów staje się zatem coraz bardziej skomplikowany, co dodatkowo wzmacnia logikę i potrzebę istnienia firm specjalizujących się w sztuce produkcji chipów.

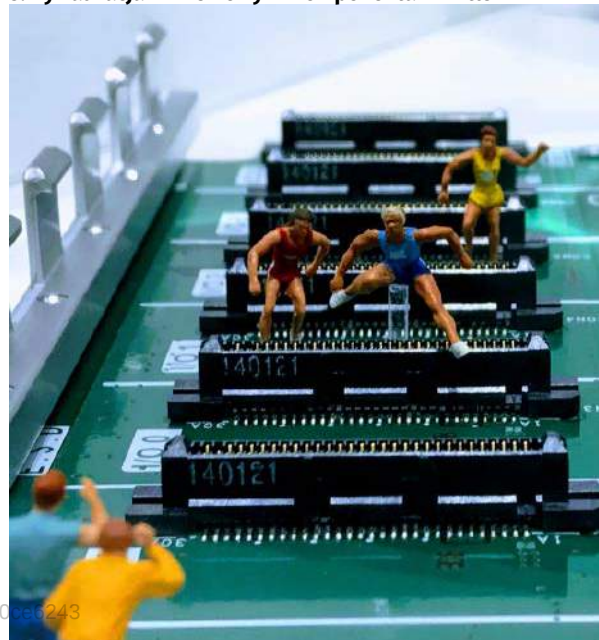
W 1970 roku najmniejsze tranzystory półprzewodnikowe miały około 12 tys. nanometrów szerokości. Dziś najważniejszy i najszerzej stosowany obecnie układ AI na świecie, procesor graficzny A100 firmy Nvidia, ma tranzystory o rozmiarze 7 nanometrów. Najnowsza jednostka przetwarzania tensorowego (TPU) firmy Google, najbardziej znana alternatywa dla procesorów graficznych Nvidii, podobnie wykorzystuje technologię 7-nanometrową. Niecierpliwie oczekiwany układ AI nowej generacji firmy Nvidia, H100, ma tranzystory 4-nanometrowe. Istnieją tylko trzy firmy na świecie, które potrafią wytwarzać zaawansowane chipy w tych zakresach miniaturyzacji i komplikacji – TSMC, Samsung i Intel. Jednak te najbardziej zaawansowane chipy, takie jak procesory graficzne H100 firmy Nvidia potrafi na razie produkować jedynie TSMC. W jaki sposób TSMC stało się tak dominującą siłą? I dlaczego tak trudno jest jakiegokolwiek innej firmie na świecie powtórzyć to, co robi TSMC?

W świecie produkcji chipów rządzi niepodzielnie ekonomia skali, co nieuchronnie prowadzi do reguły „zwycięzca bierze wszystko” (5). Cytowany Miller pisze w swojej książce: „Ekonomia produkcji chipów wymaga nieustannej konsolidacji. Niezależnie od tego, która firma produkuje najwięcej chipów, ma wbudowaną przewagę, poprawiając wydajność i rozkładając koszty inwestycji kapitałowych na większą liczbę

klientów”. Produkcja chipów wymaga ogromnych początkowych nakładów. W 2021 roku TSMC ogłosiło, że w ciągu trzech lat wyda aż 100 miliardów dolarów na rozszerzenie swoich możliwości produkcyjnych. Oprócz nakładów inwestycyjnych, zaawansowana produkcja chipów wymaga wielomiliardowych inwestycji w badania i rozwój. W przeciwieństwie do jakiegokolwiek innej firmy na świecie, TSMC może uzasadnić te ogromne inwestycje ze względu na samą ilość produkowanych chipów, znacznie większą niż jakakolwiek inna firma na świecie. Tworzy to dodatnie sprzężenie zwrotne dla TSMC. Firmy poszukujące chipów, od Apple przez Teslę po Nvidię, wybierają TSMC, ponieważ oferuje najbardziej zaawansowane możliwości produkcji chipów. Daje to TSMC uzasadnienie dla ciągłych inwestycji w celu utrzymania pozycji lidera, a to dodatkowo zwiększa przewagę firmy nad rywalami, czyniąc ją najlepszym wyborem dla kolejnych klientów.

W 2010 roku firma o nazwie GlobalFoundries stała się rzucić wyzwanie TSMC w produkcji chipów. GlobalFoundries powstała w 2009 roku. Nowy podmiot był finansowany przez Mubadala, państwowy fundusz ze Zjednoczonych Emiratów Arabskich. Firma zainwestowała miliardy dolarów w celu opracowania najnowocześniejszej technologii węzłowej i zbudowania najbardziej zaawansowanych chipów na świecie. Jednak w 2018 roku, po niecałej dekadzie, kierownictwo GlobalFoundries doszło do wniosku, że biorąc pod uwagę jego skalę, nigdy nie będzie miało sensu finansowego dokonywanie wielomiliardowych inwestycji potrzebnych rok po roku i pozostać w czołówce produkcji chipów. GlobalFoundries zrezygnowało z rozwijania wiodących rozwiązań, obniżyło koszty badań i rozwoju i przestało

5. Rywalizacja z krzemowymi komponentami w tle



konkurować z TSMC o budowę najbardziej zaawansowanych chipów. Firma koncentruje się teraz na produkcji mniej zaawansowanych komponentów.

Jednak oprócz biznesowych są również polityczne powody, aby zmniejszać uzależnienie od TSMC, a właściwie niekoniecznie od samej firmy TSMC ile od produkcji w tak zapalnym miejscu. Pod koniec 2022 roku TSMC ogłosiło, że zainwestuje 40 miliardów dolarów w budowę dwóch najnowocześniejszych fabryk w Stanach Zjednoczonych, w stanie Arizona. Pierwsza z tych dwóch fabryk rozpocznie produkcję w przyszłym roku i będzie wyposażona w sprzęt do produkcji chipów 4-nanometrowych. Druga fabryka w Arizonie ma zostać uruchomiona w 2026 r.; będzie zdolna do produkcji chipów 3-nanometrowych, nowej generacji wiodącej technologii półprzewodnikowej. Decyzja TSMC o budowie tak zaawansowanych fabryk w Stanach Zjednoczonych, innymi słowy, o podzieleniu się swoimi klejnotami koronnymi z USA, była wynikiem silnej presji i hojnych dotacji ze strony amerykańskiego rządu. Ostatecznie, tak bardzo jak Stany Zjednoczone potrzebują Tajwanu, Tajwan potrzebuje Stanów Zjednoczonych. TSMC nie miało innego wyboru.

Przeniesienie zaawansowanej produkcji chipów do Stanów Zjednoczonych pomoże złagodzić absolutną zależność branży sztucznej inteligencji od fabryk na Tajwanie. Fabryki w Arizonie nie są jednak panaceum. Po pierwsze, ich moce produkcyjne będą stosunkowo skromne – w sumie oczekuje się, że te dwie fabryki będą produkować 600 tys. wafli krzemowych rocznie. Dla porównania, TSMC produkuje ponad 13 milionów wafli rocznie, co oznacza, że amerykańskie fabryki będą stanowić mniej niż 5 proc. całkowitej produkcji. Co więcej, najbardziej zaawansowane węzły produkcji półprzewodników pozostaną na Tajwanie. Do czasu, gdy 4-nanometrowa fabryka w Arizonie rozpocznie produkcję w 2024 roku, a następnie 3-nanometrowa fabryka w Arizonie rozpocznie produkcję w 2026 roku, zakłady te będą o jedną generację za wiodącymi węzłami, które będą mogły produkować tylko fabryki z Tajwanu.

Świat ma opcje na wypadek wojny, ale postęp zwolni

Chiny w tak dużym stopniu polegają na Tajwanie, jeśli chodzi o dostawy chipów, których potrzebują do napędzania swojej gospodarki (70 proc. wszystkich chipów w Chinach jest wytwarzanych przez TSMC). Ewentualna inwazja tego kraju na Tajwan w tej sytuacji mogłaby być potężnym strzałem w stopę, bo trudno uwierzyć, by udało się przejąć te wytwórnie tak po prostu. Nawet jeśli fizyczne budynki

pozostałyby nienaruszone po chińskiej inwazji, nie-realne jest, aby KPCh była w stanie kontynuować ich eksploatację w celu produkcji najnowocześniejszych chipów. Utrzymanie najnowocześniejszych fabryk wymaga współpracy z partnerami z całego globalnego ekosystemu półprzewodników oraz stałego napływu materiałów, sprzętu i usług od dostawców, których odmówiono by najeźdźcy. Krzemowa tarcza to jednak tylko teoria, a nie gwarancja. Chińska kalkulacja dotycząca Tajwanu nie zależy wyłącznie od półprzewodników.

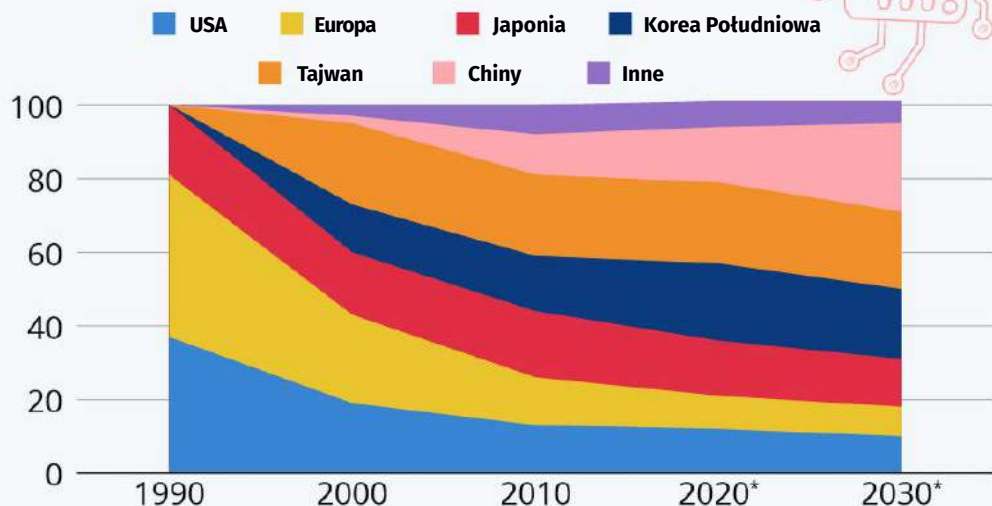
Jednak Stany Zjednoczone, jak się wydaje, będą bronić wyspy i chronić jej suwerenność, nie tylko ze względu na TSMC, od którego są tak jak cały świat zależne. Amerykańska Rada Bezpieczeństwa Narodowego oszacowała niedawno, że konflikt zbrojny między Chinami a Tajwanem mógłby kosztować światową gospodarkę ponad bilion dolarów rocznie z powodu zakłóceń w produkcji półprzewodników.

Gdyby stało się to najgorsze, to TSMC, firmą najlepiej przygotowaną do produkcji najnowocześniejszych chipów AI jest Samsung. Samsung jest obecnie jedyną firmą na świecie, poza TSMC, która może produkować chipy w procesie 3 nanometrów, obecnie najnowocześniejszym z dostępnych. Niestety, jak wynika z danych, jakość produkcji Samsunga uchodzi za znacznie gorszą niż TSMC. „Wydajność” to ważny wskaźnik branżowy, który wskazuje procent wafli krzemowych wprowadzonych do procesu produkcyjnego, które kończą jako działające chipy. Wydajność TSMC w przypadku 3-nanometrowych chipów szacuje się na 80 proc. Tymczasem wydajność Samsunga wynosiła od 10 do 20 proc., gdy w zeszłym roku rozpoczął produkcję 3-nanometrowych układów (choć najnowsze doniesienia sugerują, że może się ona poprawiać). Niska jakość produkcji Samsunga skłoniła niedawno Nvidię do przeniesienia produkcji wszystkich swoich procesorów graficznych, nie tylko wysokiej klasy chipów AI, z Samsunga do TSMC. W najlepszym scenariuszu, Samsungowi zajęłoby lata, aby skalować się do obecnego poziomu produkcji chipów AI i wydajności TSMC. W dodatku, z perspektywy Zachodu, fabryki Samsunga same w sobie znajdują się we wrażliwym miejscu.

To prowadzi wywód do dawnego amerykańskiego championa chipów, firmy Intel. W ostatnich latach, w wyniku strategicznych błędów i niepowodzeń produkcyjnych, firma pozostała w tyle. Intel zmagał się z przejściem zarówno na proces 10-nanometrowy, jak i 7-nanometrowy, a jego 7-nanometrowe chipy weszły do pełnej produkcji dopiero w tym roku. Firma uciekła się nawet do outsourcingu niektórych części swojego 7-nanometrowego procesu

Produkcja chipów oddala się od tradycyjnych ośrodków

Globalna produkcja komponentów półprzewodnikowych według lokalizacji (w proc.)



*prognoza

Źródło: Boston Consulting Group, Semiconductor Industry Association

6. Udział w produkcji chipów

produkcyjnego do konkurencyjnego TSMC, co jest upokarzającym posunięciem dla dumnego pioniera branży chipów. Pod wodzą nowego CEO, Pata Gelsingera, Intel aspiruje do odzyskania prymatu w produkcji chipów. Firma wyznaczyła sobie ambitne cele, aby rozpocząć produkcję 2-nanometrowych chipów do 2024 roku i dostarczyć pięć nowych nanometrowych węzłów w ciągu najbliższych czterech lat, przeskakując TSMC.

Niektórzy obserwatorzy uważają, że gdyby fabryki TSMC przestały działać, rząd USA podjąłby zdecydowane kroki w celu pośredniczenia w partnerstwie

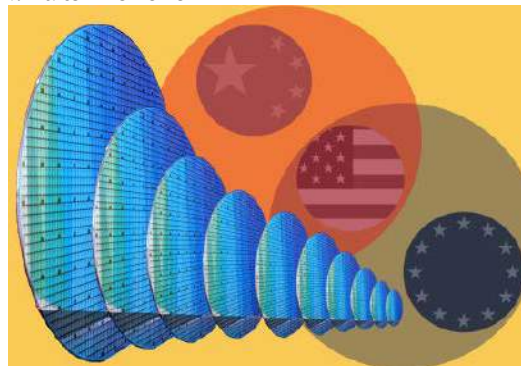
między amerykańskimi konkurentami Intel i Nvidia. W ogólnym zarysie, pomysł polegałby na tym, że Intel agresywnie zwiększyłby swoje możliwości produkcyjne wszelkimi niezbędnymi środkami, aby jak najszybciej wesprzeć produkcję procesorów graficznych Nvidii. Nie jest jednak jasne, czy taka „wroga” współpraca jest możliwa, a w szczególności, czy Intel ma możliwości produkcyjne, aby zrealizować to w jakimkolwiek rozsądnym terminie.

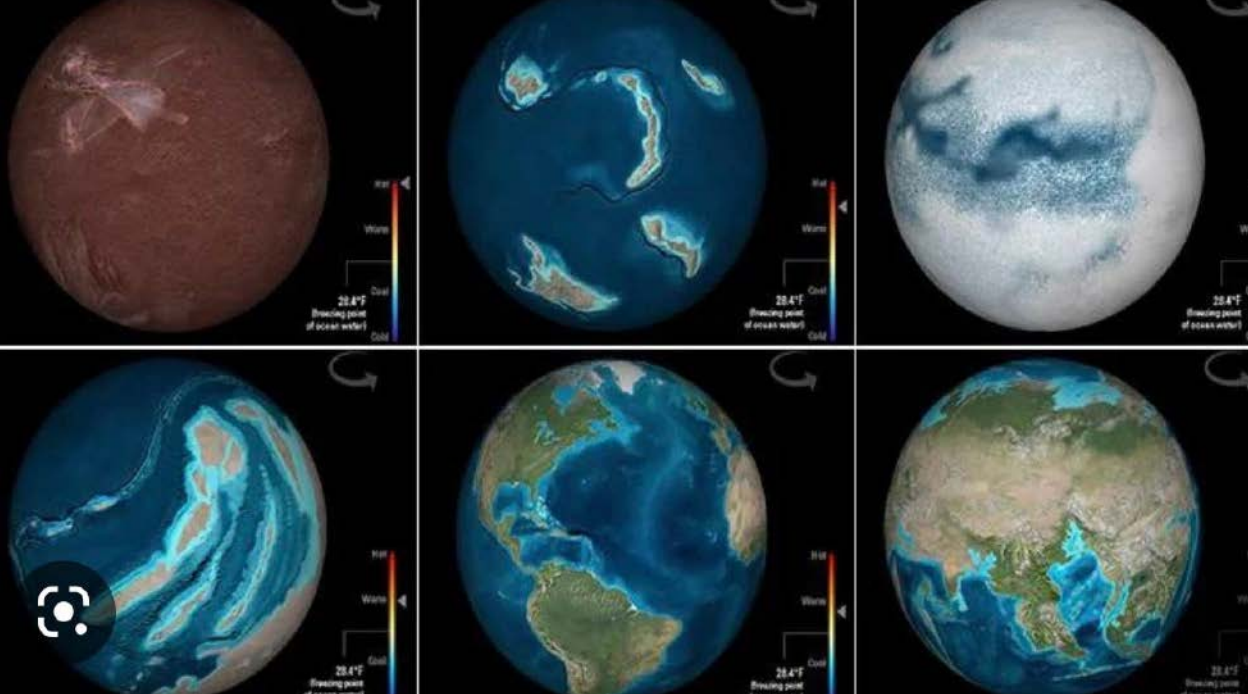
Choć najbardziej zaawansowane chipy, takie jak H100 firmy Nvidia i TPU firmy Google, mogą być produkowane tylko na Tajwanie, istnieje wiele fabryk na całym świecie (6), od Stanów Zjednoczonych po Europę i Izrael, zdolnych do produkcji mniej zminiaturyzowanych układów logicznych na dużą skalę. Choć opóźnia to postęp w dziedzinie obliczeń, nie tylko AI, jednak jako awaryjne rozwiązanie w złym scenariuszu, też jest brane pod uwagę.

Przede wszystkim trwa wszędzie budowa własnych krzemowych i swoistych tarcz „obronnych” z krzemowych wafli (7), co praktycznie oznacza rozwijanie u siebie zdolności produkcji elektronicznych komponentów. Że nie jest to łatwe i tanie, to oczywiste, ale może nie być wyjątkiem. ■

Mirosław Usidus

7. Wafle krzemowe





1. Zrzut ekranowy z interaktywnej prezentacji Instytutu Smithsonian obrazujący zmiany wyglądu planety Ziemia w jej historii

Luki z historii Ziemi

RAPORT

Planeta dotknięta amnezją

Być może kiedyś wymyślimy taką maszynę, która pozwoli nam się cofnąć w czasie i dokładnie przyrzeć się temu, co się działo miliony i miliardy lata temu. Na razie, jeśli chodzi o badania historii naszej planety (1), mamy głównie geologię i archeologię, choć eksploracja najbliższej kosmicznej okolicy, też może pomóc.

W trakcie ewolucji Ziemi warstwy osadów skalnych tworzyły się jedna na drugiej. Każda warstwa reprezentuje osobny okres w historii Ziemi. Dość dawno zauważono jednak, iż w tym zapisie brakuje w wielu miejscach warstw odpowiadających setkom milionów lat geologicznej historii planety. Geologiczne luki w historii Ziemi nazywa się „niezgodnościami”, a największa i najbardziej znana z nich, określana jako Wielka Niezgodność (2), która „kończy się” około 550 milionów lat temu może mieć początek sięgający nawet pół miliarda lat wstecz.

Zjawisko „niezgodności” polega mówiąc w uproszczeniu na tym, że pomiędzy skałami ułożonymi powyżej, młodszymi niż skały poniżej, nie zachowują się żadne osady, które powinny powstać w okresie dzielącym dwie odległe od siebie w czasie warstwy.

Brakuje lokalnego zapisu dla określonego przedziału czasowego. Na przykład w Wielkim Kanionie warstwa paleoproterozoicznych tzw. łupków Wisznu, znajduje się bezpośrednio pod kambryjskim piaskowcem Tapeats, mimo że te dwa okresy dzieli setki milionów lat. Nigdzie nie można znaleźć warstw skalnych datowanych na tę lukę. Właśnie to dziwne zjawisko nazywane jest Wielką Niezgodnością i od ponad wieku stanowiło jedną z największych tajemnic w geologii. Luki analogiczne Wielkiej Niezgodności, znajdowano nie tylko w tej słynnej szczelinie, ale w miejscach na całym świecie. Wzór jest taki, że w jednej, młodszej, warstwie mamy okres kambryjski, który rozpoczął się około 540 milionów lat temu i pozostawił po sobie skały osadowe wypełnione skamieniałościami złożonego, wielokomórkowego życia. Bezpośrednio



2. Warstwy skał luki w zapisie geologicznym w miejscu zwanym Wielką Niezgodnością

poniżej znajduje się wolna od skamieniałości krystaliczna skała podstawowa, która uformowała się około miliarda lub więcej lat temu. Gdzie więc podzieliła się cała skała, która znajduje się pomiędzy tymi okresami?

Geologiczne niezgodności, odzwierciedlające długoterminowe zmiany we wzorze akumulacji warstw osadowych lub magmowych, dotyczą nie tylko jednego fragmentu dziejów Ziemi. Na naszej planecie znaleźć ich znacznie więcej, często basenach oceanicznych, takich jak Zatoka Meksykańska lub Morze Północne, ale także w Bangladeszu i w Brazylii, w Górach Skalistych lub Appalachach. W szkockich górach w Siccar Point, w hrabstwie Berwickshire występuje tzw. niezgodność Huttona, wyraźne ścięcie niemal pionowych sylurskich warstw osadowych przez dobrze osadzone zlepnie i piaskowce należące do górnego starego czerwonego piaskowca pozwoliło odkrywczy na zidentyfikowanie jednej z pierwszych znanych luk w historii geologicznej Ziemi.

Erozja podlodowcowa wymyła osady

Od chwili publikacji w „Proceedings of the National Academy of Sciences” w 2020 r. nowych wyników badań tej ogromnej luki, być może mamy wyjaśnienie lepsze niż poprzednie, ale jak wszystko w tej

dziedzinie, zawsze jest to opatrzone mniejszym lub większym znakiem zapytania. Naukowcy powszechnie stawiali wcześniej hipotezę, że Wielka Niezgodność została spowodowana globalnym zdarzeniem erozyjnym podczas fazy ewolucji Ziemi, w której Ziemia była niemal całkowicie zamrożona, pokryta warstwą lodu i śniegu, (Ziemia-śnieżka, ang. „Snowball Earth”), a która miała miejsce dwukrotnie w przedziale czasowym od 715 a 640 mln lat temu (3).

Według tej starszej teorii, w odstępach około miliarda lat, nawet jedna trzecia powierzchni skorupy ziemskiej została zamknięta przez wędrujące lodowce Ziemi-śnieżki, które oddziaływały na powierzchnię planety erozyjnie. Podobnie do tego, co widzimy dziś na Antarktydzie, wiele lodowców Snowball Earth było potężnymi czynnikami erozyjnymi: Nacisk lodu na powierzchnię tworzy mokre podłoże, które może przemieszczać osady, pomimo ekstremalnie niskich temperatur na powierzchni. Powstały osad został wyrzucony do pokrytych błotem oceanów, gdzie został następnie wessany do płaszcza przez napierające na siebie płyty tektoniczne. W efekcie, jak twierdzą naukowcy w „Proceedings of the National Academy of Sciences”, w wielu miejscach Ziemia zakopała do wody około jednej piątej swojej historii geologicznej.

To były jednak jedynie hipotezy, gdyż sam pomysł, że Ziemia była w pewnym okresie swojej historii mroźną „kulą śnieżną” jest wciąż hipotezą, choć trzeba przyznać, że jest coraz częściej akceptowany przez społeczność naukową.

Teorię o gigantycznej erozji potwierdzały znajdowane przez geologów w starych warstwach cyrkonie, charakteryzując się tym, że utrwalają na miliardy lat stan rzeczy z zamierzchłej przeszłości Ziemi, potwierdzając te hipotezy. Cyrkonie zawierają różne radioaktywne izotopy, pozwalające jak w przypadku uranu, ustalić wiek powstania kryształów z niezwykle precyzją, inne zaś, takie jak izotopy hafnu, ujawniają, co działo się ze skorupą i płaszczem, ponieważ niektóre izotopy wolą jedno ustawienie geologiczne od drugiego. Dzięki cyrkoniom wiemy, że w okolicach domniemanego początku Ziemi-śnieżki, nastąpiła kolosalna zmiana geochemiczna na całej planecie, co można wytłumaczyć tym, że duża część skorupy ziemskiej była poddawana recyklingowi do nowych zbiorników magmy. Izotopy tlenu w znalezionych cyrkoniach wskazywały, że skorupa ziemska przeszła niskotemperaturowe zmiany hydrotermalne. Oznaczało to, że górna część skorupy, w kontakcie z wodą i lodem została odcięta i zapadła się. W sumie dowody te sugerują, że na powierzchni miało miejsce gigantyczne wydarzenie erozyjne i warstwa osadów została zmieciona.

Hipotezę erozji pod wpływem mas lodu wspiera fakt, że na powierzchni Ziemi zanikły kratery, który pełno było we wcześniejszym okresie, jako pozostałość po pradziejowych bombardowaniach meteorytowych i kometarnych.

4. Wizualizacja kambryjskiej eksplozji życia



3. Ziemia w fazie kuli śnieżnej – wizualizacja

Te kolosalne „przeoranie” skorupy ziemskiej mogło, zdaniem proponujących tę wersję, zbiec się w czasie z poważnymi zmianami geochemicznymi i środowiskowymi, które na kolejnym etapie były potencjalnie korzystne dla ewolucji biologicznej. Wraz z ociepleniem planety, po lodowcach pozostały płytkie zbiorniki wodne, doskonale nadające się do rozwoju i różnicowania się życia. W kolejnej fazie nadeszła tzw. kambryjska eksplozja życia (4).

Wszystko składałoby się w świetnie pasujące wyjaśnienie. Jednak nie wszyscy byli zadowoleni z tej wizji. Wspomniane badania z 2020 r. wykazują, że za brakujące warstwy w Wielkiej Niezgodności i w innych podobnych skalnych pozostałościach na całym świecie odpowiadają ruchy tektoniczne. Naukowcy zbadałi próbki minerałów i kryształów ze skał w celu





określenia historii termicznej warstw skalnych. Ich analiza wykazała, że starsza warstwa skalna uległa erozji jeszcze przed pierwszą fazą Ziemi-śnieżki, co sugeruje, że erozja lodowcowa nie może być odpowiedzialna za Wielką Niezgodność, przynajmniej w tym zapisie warstw skalnych na terenie Kolorado, który badali (Pikes Peak). Badacze ci uważają, że warstwy osadów z ziemskiego zapisu geologicznego zostały wymazane przez procesy tektoniczne związane z formowaniem się i rozpadem Rodinii (5), superkontynentu istniejącego na Ziemi pod koniec eonu proterozoicznego, ok. miliard lat temu. Rodinia otoczona była wszechoceanem Mirowia. Badania te obalają teorię, że erozja na dużą skalę sprzed około 540 milionów lat mogła zasiać na Ziemi składniki odżywcze i doprowadzić do powstania złożonego życia.

„Naukowcy od dawna postrzegają to jako fundamentalną granicę w historii geologicznej”, mówiła Rebecca Flowers, profesor nadzwyczajny na Wydziale Nauk Geologicznych, w wywiadzie dla serwisu „Science Daily”. „W zapisie geologicznym brakuje wielu elementów, ale to, że ich brakuje, nie oznacza, że ta historia jest prosta”.

Tajemnice od bombardowania po szlak Homo sapiens

Dotyczy to oczywiście nie tylko tego fragmentu dziejów planety. Gdy patrzymy głębiej w jej historię, luki w naszej wiedzy dotyczą przede wszystkim dokładnego mechanizmu i przebiegu procesu powstania Ziemi i Układu Słonecznego. Wciąż nie do końca rozumiemy, w jaki sposób powstały planety i gwiazdy. Dalej nie jest znana precyzyjna chronologia i przebieg wydarzeń z okresu tzw. bombardowania późnego, ok. 4 mld lat temu (6). Wiemy, że Ziemia była intensywnie bombardowana przez asteroidy i komety, ale szczegóły tego procesu nie są dobrze znane. Nie mamy też pełnej jasności, jedynie przypuszczenia, jakie było źródło i przyczyna bombardowania.

Wielkimi lukami w naszej wiedzy o historii Ziemi są przyczyny masowych wymierań w historii Ziemi. Nie znamy dokładnych przyczyn



5. Wyobrażenie superkontynentu Rodinia

wymarcia dinozaurów czy innych grup organizmów, choć często wiążemy je z jakimiś wydarzeniami, np. uderzeniem wielkiego obiektu w naszą planetę lub domniemanym wybuchem supernowej w kosmicznej okolicy. Dokładny przebieg i spłot przy czynowo-skutkowy nie jest jednak wcale dobrze poznany.

Wcale nie wiemy też gdzie na osi czasu mamy umieścić powstanie pierwszego życia na Ziemi.

Podajemy w tym raporcie pewne skrajne szacunki (3,95 mld lat temu).

Nie są to jednak w dobrze udowodnione fakty a raczej hipotezy, wciąż



zresztą przesuwające powstanie życia w przeszłość, aż do epoki, w której trudno sobie nam wyobrazić istnienie żywych organizmów. Nie znamy też mechanizmu wydarzenia (wydarzeń?) pojawienia się życia. Nie wiemy, kiedy i w jaki sposób powstały pierwsze organizmy żywe.

Pomimo pewnych poszlak wcale nie jest znany przebieg kluczowych wydarzeń w okresie wczesnej historii gatunku ludzkiego, np. wyjścia z Afryki, rozprzestrzenienia się po Eurazji, kolonizowania obu Ameryk. Wiele szczegółów ewolucji i migracji pierwszych ludzi jest nieznanych.

Śnieżny glob we mgle

Wielką luką jest wspominany mechanizm powstania domniemanego globalnego zlodowacenia, które doprowadzić miało do fazy Ziemi-śnieżki, o której była mowa przy okazji omawiania luk w warstwach geologicznych. Pierwszym, który zidentyfikował i opisał zjawisko zlodowacenia całego globu, używając terminu „Snowball Earth”, był John Kirschvink. Wykazał

on na podstawie pomiarów, że kluczowe dla tej idei osady lodowcowe ery neoproterozoicznej tworzyły się w obszarach okołorównikowych (podzwrotnikowych). Jednak do popularyzacji tego poglądu doprowadził dopiero zespół autorski na czele z Paulem Felixem Hoffmanem w 1998 r. Stworzył on model globalnego zlodowacenia tłumaczący szereg czynników, które miałyby być odpowiedzialne za to niezwykle zjawisko geologiczne. Zgodnie z ich hipotezą, w prekambrze pojawiła się epoka lodowcowa tak silna, że wszystkie oceany Ziemi zamarły. Cała planeta została pokryta ponad kilometrową czapą lodu. Jedynie głębie oceaniczne podgrzewane wewnętrznym ciepłem Ziemi pozostały niezamrożone. Nie było deszczu, śniegu, ani chmur.

Jeżeli równik został pokryty lodem, oznaczało to pokrycie białą czapą całej planety. Podstawą tego rozumowania był fakt, że wzrost ilości lodu podwyższa albedo powierzchni planety. Biały lód odbija znacznie więcej ogrzewającego planetę światła słonecznego niż woda lub ziemia. Jeżeli lodu

6. Jedna z wizji późnego bombardowania Ziemi





zaczyna przybywać, to Ziemia odbija coraz więcej ciepła w kosmos i lodu przybywa jeszcze więcej – pojawia się dodatnie sprzężenie zwrotne. Wyliczono, że przekroczenie przez lodowce równoleżnika 30 stopni na obu półkulach lub 20 na jednej półkuli oznacza nieodwracalnie zamrożenie całej planety. W takiej sytuacji bowiem ilość energii absorbowanej jest już mniejsza od odbijanej przez czapę lodową. Pomiędzy okresami całkowitego zamarznięcia powierzchni Ziemi pojawiały się ograniczone w zasięgu, krótkotrwałe okresy ocieplenia (interglacjalny), ale nie doprowadzały one do wycofania się lodowców poza granicę 30. stopnia szerokości geograficznej.

Hipoteza Ziemi-śnieżki wywołała ogromne kontrowersje. Niewyobrażalne wydawało się zamarznięcie całej planety. Na dodatek wiadomo, że przed epoką lodu i po niej istniało życie. Co więcej, niedługo po stopieniu lodu, na Ziemi miała miejsce tzw. eksplozja kambryjska. Planetę zasiedliła niezliczona liczba nowych organizmów wielokomórkowych, które pozostawiły ogromną ilość skamieniałości. Wszystkie te argumenty przemawiały przeciw hipotezie Ziemi śnieżki.

Naukowcy powszechnie zgadzają się, że ważnym czynnikiem określającym temperaturę na Ziemi jest ilość dwutlenku węgla w atmosferze. Wzrost poziomu CO_2 powoduje wzrost średnich temperatur na całej planecie. Zjawisko to nosi nazwę efektu cieplarnianego. W sytuacji ubytku gazów cieplarnianych cała planeta może się znacznie ochłodzić. Być może erozja krzemianów może doprowadzić do szybszego niż zwykle wiązania dwutlenku węgla w glebie. Dodatkowo 750 mln lat temu na Ziemi istniały już zdolne do fotosyntezy bakterie. One również są w stanie wiązać CO_2 , produkując związki organiczne. Nagły wzrost aktywności takich organizmów oraz pochłaniania gazów cieplarnianych przez glebę mógł spowodować tak silny spadek sprawności efektu cieplarnianego, że średnie temperatury na całej planecie znacznie się obniżyły. Powstający lód spowodował jeszcze większe obniżenie temperatury i w końcu cała Ziemia zamarzła. Podczas gdy w czasie ostatnich epok lodowych temp. obniżała się o ok. 3...12°C, to po zamarznięciu całej planety jej średnia temperatura mogła wynosić nawet -50°C.

Dokładniejsze analizy wykazały, że zamrożenie całej planety ma pewne dodatkowe skutki. Wytwarzany dziś dwutlenek węgla rozpuszcza się w ogromnych ilościach w oceanach, gdzie pochłaniają go glony. Jednak pokrycie oceanów lodem zablokowało kontakt atmosfery z hydrosferą. Każda ilość CO_2 wyrzucana do atmosfery musiała tam pozostać. Na lądach lód zalegał w jeszcze większych ilościach niż w wodzie

i prawdopodobnie całkowicie zasłonił cały grunt. Reakcja skał z dwutlenkiem węgla stała się całkowicie niemożliwa. Przez wiele tysięcy lat cały dwutlenek węgla wyrzucany co jakiś czas przez wulkany pozostawał w atmosferze. Efekt cieplarniany stał się tak silny, że okowy lodu puścili. Jeżeli niewielki fragment oceanu uwolnił się od białej pokrywy, to zaczął przyjmować coraz więcej słonecznego ciepła. Pojawiło się dodatnie sprzężenie zwrotne, które w bardzo krótkim (w skali geologicznej) czasie – rzędu jednego tysiąca lat – mogło uwolnić całą planetę od lodu. Rozumowanie to jednak wcale nie wyjaśnia długości trwania epoki lodowej oraz występowania okresów ocieplenia – interglacjalów.

Jeżeli lód pokrył całą planetę, a na powierzchni zapanowały niespotykane mrozy, uniemożliwiłoby to trwanie życia. Niemniej według hipotezy Ziemi śnieżki nie cała biosfera uległa zamrożeniu. Bez przeszkód mogły funkcjonować ekosystemy kominów geotermalnych na dnie oceanów. Badania podmorskich grzbietów ujawniły istnienie tam potężnych źródeł gorącej i bardzo zanieczyszczonej wody. Jej źródłem są procesy wulkaniczne. Woda jest gorąca i zawiera wiele związków siarki oraz żelaza. W okolicach kominów geotermalnych rozwijają się bogate ekosystemy, które są oazami tętniącymi życiem. Związki siarki, żelaza i inne substancje mineralne są wykorzystywane przez liczne organizmy prokariotyczne – bakterie i archeony – zdolne do czerpania energii potrzebnej do procesów metabolicznych podczas ich przetwarzania. Prokarioty są zjadane przez inne organizmy i tak powstaje cały łańcuch pokarmowy, na którego szczycie obecnie znajdują zwierzęta, w tym krewetki i ryby. Szczególną cechą tych ekosystemów jest duża niezależność od światła Słońca. Tylko wyższe poziomy potrzebują tlenu do oddychania, a w czasach przed wyewoluowaniem skorupiaków i kręgowców na szczycie piramid troficznych mogły znajdować się organizmy beztlenowe. Pokrycie oceanów lodem dla organizmów żyjących wokół kominów hydrotermalnych byłoby niezauważalne, więc jest to miejsce, w którym na Ziemi śnieżce mogłoby się rozwijać życie. Problem w tym, że kominy hydrotermalne związane są z grzbietami śródoceanicznymi, a nie wiadomo, czy te struktury istniały wtedy na Ziemi.

Kominy geotermalne nie są środowiskiem, w którym mogłyby przetrwać organizmy zdolne do fotosyntezy. Zapisy kopalne wyraźnie pokazują, że przed okresem Ziemi śnieżki oraz potem istniały formy bakterii zdolnych do fotosyntezy. Oznacza to, że jakoś musiały przetrwać okres zamrożenia oceanów, w których żyją. Odpowiedzią okazały się badania współczesnych

lodowców szelfowych. Odnaleziono glony zdolne do przetrwania pod bardzo grubą taflą lodu. Lód jest bardzo przenikliwy dla światła. Z drugiej strony jego grubość na równiku nie była większa niż kilometr. Obserwacje współczesnych glonów wykazały, że byłyby one w stanie przeżyć w takich warunkach.

Zwolennicy teorii Ziemi śnieżki zauważyli, że zakończenie okresu zlodowacenia miało miejsce kilka milionów lat przed kambryjską eksplozją życia. Sugerują, że zamrożenie planety mogło być czynnikiem zwiększającym presję ewolucyjną, co zwiększyło tempo rozwoju prymitywnych organizmów. Z drugiej strony po zakończeniu epoki lodu warunki stały się bardzo ciepłarniane, co pozwoliło na rozwój ogromnej różnorodności żywych organizmów obserwowanej kilka milionów lat później w osadach.

Istnieją hipotezy, które sugerują, że już wcześniej, 2,3 mld lat temu, nasza planeta przeszła pierwszy okres Ziemi-śnieżki. W tym okresie powstały pierwsze cyjanobakterie zdolne do fotosyntezy. Wytworzony przez nie tlen spowodował rozkład zawartego w atmosferze metanu. Gaz ten jest dużo lepszym czynnikiem powodującym efekt cieplarniany niż dwutlenek węgla. Jednocześnie astrofizycy uważają, że 2,3 mld lat temu Słońce było trochę zimniejsze i dostarczało Ziemi mniej energii. Niestety, z tego okresu dotrwało do naszych czasów niewiele osadów i trudno jest jednoznacznie potwierdzić lub zaprzeczyć istnieniu Ziemi śnieżki 2,3 mld lat temu.

Naukowcy próbowali wyjaśnić obecność lodu na równiku na kilka innych sposobów. np. tak, że w tym okresie kąt pomiędzy płaszczyzną równika oraz płaszczyzną orbity był większy (być może wynosił nawet 54°), co powodowało nierównomierne nagrzewanie planety. Inna hipoteza głosi, że biegun geomagnetyczny Ziemi, znajdował się w większej odległości od osi obrotu planety, niż dzisiaj. Położenie lądów z tego okresu wyznacza się na podstawie zapisanych w skałach linii pola magnetycznego. Badania geologiczne potwierdzają, że osady formowały się daleko od ówczesnego bieguna geomagnetycznego, ale możliwe, że znajdowały się blisko prawdziwego bieguna geograficznego Ziemi.

Jak powstał nasz dom?

Ostatecznie nie wiemy dokładnie, w jaki sposób uformowała się Ziemia. Naukowcy wciąż nie mają dokładnej teorii ani modelu formowania się naszej planety i innych planet (czy w ogóle jest tylko jeden model). Obecnie walczą (a może się uzupełniają?) dwa modele. Pierwszy i najszerzej akceptowany, mówiący o akrecji jądra, sprawdza się w przypadku formowania się planet skalistych takich jak Ziemia, ale ma kłopoty z planetami dużymi. Drugi, oparty na niestabilności dysku, może odpowiadać za powstanie największych planet. Dochodzą teorie „komplikujące”, jak np. opublikowane kilka lat temu w czasopiśmie „Science Advances American Association for

7. Jedna z wizualizacji zderzenia Ziemi z Theią





the Advancement of Science”, badania sugerujące, że wiele oryginalnych planetarnych elementów naszego Układu Słonecznego mogło w rzeczywistości zacząć funkcjonować nie jako obiekty skaliste, ale jako gigantyczne kule ciepłego błota.

Chociaż populacja komet i asteroid przechodzących przez wewnętrzny Układ Słoneczny jest dziś niewielka, były one bardziej obfite, gdy planety i Słońce były młode. Zderzenia z tymi lodowymi ciałami prawdopodobnie zdeponowały na powierzchni Ziemi znaczną część jej wody. Planeta nasza znajduje się w strefie Złotowłosej, czyli w regionie, w którym woda może pozostać w stanie ciekłym, co zdaniem wielu naukowców odgrywa kluczową rolę w rozwoju życia.

Z jednej strony obowiązuje coś w rodzaju głównego nurtu przekonań naukowych na temat powstania i wczesnego okresu w historii Ziemi. Z drugiej mnóstwo jest zagadek, konkurujących hipotez i nowych koncepcji, dawniej nie branych pod uwagę. Należy do nich popularna dziś hipoteza zakładająca, że miliardy lat temu w wewnętrznym Układzie Słonecznym istniał obiekt wielkości Marsa, nazywany Theia, umiejscowiony tuż przy orbitalnym szlaku protoplanety, z której ukształtowała się ostatecznie Ziemia. Gdy doszło do kolizji Thei z obiektem, którego być może nie należy jeszcze nazywać Ziemią (7), w przestrzeń kosmiczną wyrzucone zostało mnóstwo materiału, z którego później powstała Ziemia i Księżyc. Co zaś stało się z Theią, tego nie określa się dokładnie.

Obecnie astronomowie z Harvardu przypuszczają, że takie kosmiczne katastrofy skutkują bardziej widowiskowymi efektami. Naukowcy stworzyli model komputerowy obrazujący zderzenie dwóch dużych i gorących obiektów o podobnej masie obrotowej. Stwierdzili, że najbardziej gwałtowne z takich kolizji wytworzyłyby synestię – osoblive halo gazowe powstałe z rozpuszczonych skał, wirujących wokół płynnego rdzenia. Nowo powstały obiekt, żyjący jedynie przez paręset lat, miałby znacznie większą objętość niż obiekty, które się zderzyły. Nie byłoby na nim ani cieczy, ani obiektów stałych. Zespół uważa również, że po około stuleciu trwania w tej formie synestia traci tyle ciepła, że z powrotem zastyga do skalistej postaci. Część skał synestii znalazłaby się na orbicie zestalonego ponownie obiektu, dając początek Księżycowi. Autorzy przypuszczają, że większość obecnych planet była początkowo synestiami. Model ten wywołał mieszane reakcje w szeregach planetologów, z których wielu zachowuje sceptycyzm wobec tej teorii.

Innym, bardziej akceptowanym scenariuszem, który miał i ma nadal duże znaczenie dla losów Ziemi jest teoria migracji Jowisza. Według

wyliczeń Francesco DeMeo, naukowiec z amerykańskiego Harvard-Smithsonian Center for Astrophysics, Jowisz w przeszłości znajdował się tak blisko Słońca jak obecnie Mars. Następnie, podczas migracji ku swojej obecnej orbicie, Jowisz zniszczył prawie w całości Pas Planetoid – pozostało w nim do dziś jedynie 0,1 proc. populacji asteroid. Z drugiej strony, ta migracja wyekspediowała mniejsze obiekty do krańców Układu Słonecznego, co miało potem duże znaczenie dla „spokoju” naszej planety, jednak na wczesnym jednak etapie mogło to doprowadzić do silnego bombardowania dostarczającego wczesnej Ziemi dużo ważnych dla dalszego rozwoju wypadków substancji, przede wszystkim wodę.

Do roli Jowisza, ale nie tylko jego nawiązuje zdobywający w ostatnich latach popularność dwuetapowy model formowania się Ziemi nazywany modelem pojawienia się biopierwiastków (ABEL), który przedstawia nieco inną opowieść o wczesnej Ziemi niż ta, do której przez dekady nas przyzwyczajano. Co nie znaczy, że różni się od innych teorii radykalnie, bo nie tyle wprowadza nowe elementy co reinterpretuje istniejące. Dwie wcześniejsze teorie powstania Ziemi naukowcy zakładały, że oceany powstały w wyniku ucieczki pary wodnej i innych gazów z atmosfery. Kiedy powierzchnia Ziemi ochłodziła się do temperatury poniżej punktu wrzenia wody, w wyniku opadów deszczu woda spływała do wielkich zagłębień w powierzchni Ziemi, tworzyły się oceany. Grawitacja uniemożliwiła wodzie opuszczenie Ziemi. Model ABEL zakłada, że Ziemia narodziła się jako sucha planeta bez atmosfery i składników oceanicznych 4,56 mld lata temu, zaś potem dopiero nastąpiła wtórna akrecja biopierwiastków, takich jak węgiel (C), wodór (H), tlen (O) i azot (N), która osiągnęła szczyt w 4,37...4,20 mld lat temu. Przyjmuje się zarazem założenie, że ziemska woda pochodzi głównie z węglowego materiału chondrytowego, co wynika z proporcji izotopów wodoru.

Gdyby bombardowanie ABEL nie miało miejsca, życie nigdy nie powstałoby na Ziemi, twierdzą zwolennicy tej hipotezy. Ponadto z powodu wstrzykiwania lotnych substancji do początkowej suchej Ziemi wyzwoliło przejście Ziemi od litej pokrywy do tektoniki płyt co również miało przełomowe znaczenie. Dlatego też jest jednym z najważniejszych dla naszej planety wydarzeń. Uważa się, że bombardowanie ABEL miało miejsce w okresie 4,37...4,20 mld lat temu.

Przez długi czas uważano, że formowanie się oceanu i atmosfery jest równoczesne z pierwotną akrecją Ziemi. Jednak rachunek ilości wody, także przy założeniu zderzenia z Theią i formowaniem Księżyca, nie bardzo się zgadzał. Model ABEL może przezwyciężyć

te trudności, starając się wyjaśnić dostawę wody i lotnych substancji potrzebnych do zapoczątkowania prebiotycznej ewolucji chemicznej, która w dalszej konsekwencji doprowadziła do powstania życia na Ziemi, dzięki mieszanemu materiałowi redukcyjnego i utleniającego podczas teoretycznego wczesnego bombardowania. Nastąpić miało ponad 200 milionów lat później niż etap T-tauri (dopiero co powstałego Słońca), dlatego uważa się, że to wydarzenie nie jest związane z główną częścią procesu formowania się planet. Jeśli tak, to powinien istnieć inny czynnik wyzwalający bombardowanie ABEL. Jeden z możliwych scenariuszy jest następujący: trzy gazowe giganty – Jowisz, Saturn i hipotetyczna planeta „czarna owca” zostały uformowane jako pierwsze. Grawitacyjne interakcje wyrzuciły „czarną owcę”, która być może znajduje się obecnie w Pasie Kuipera a może została całkiem wypchnięta z Układu Słonecznego.

To dramatyczne wydarzenie doprowadzić mogło do ciężkiego bombardowania suchej i pozbawionej atmosfery Ziemi, która właśnie nieco ostygła po pierwszym etapie formowania. Lodowe asteroidy złożone z węglowych chondrytów uderzały w Ziemię, kierowane z zewnętrznej części pasa asteroid w wyniku grawitacyjnego oddziaływania Jowisza, Saturna i zaginionej dużej planety. To one dostarczyły budulec dla atmosfery i oceanu na suchą Ziemię. Mieszanie się materiału redukcyjnego z utleniającym zapoczątkowało reakcje chemiczne, która stały pierwszym bodźcem do powstania życia na Ziemi. Dodatkowo, bombardowanie ABEL, w ramach której uderzać miały w pierwotną Ziemię obiekty o średnicy nawet tysiąca kilometrów, umożliwiło przejście od zastygłej tektoniki pokrywowej do tektoniki płyt.

Woda z importu czy własna?

W tym czy w innych modelach zakłada się, że pierwotna dostawa wody na Ziemię wiązała się z uderzeniami takich obiektów jak np. komety. Mamy jednak dane, które zdają się przeczyć teorii, że dostawa wody na Ziemię pochodzi właśnie z nich. Według wyników badań przeprowadzonych za pomocą sondy Rosetta, woda na komecie 67P/Czuriumow-Gierasimienko (8) ma inny skład izotopowy niż ta, która znajduje się w ziemskich zbiornikach wodnych. Oznaczać to może w praktyce, że ziemska woda nie pochodzi komet, a przynajmniej nie z komet takiego typu jak 67P.

Są inne wyniki, wykazujące, że przynajmniej połowa wody, która jest obecnie na planecie Ziemia powstała zanim powstała nasza gwiazda. Dla naukowców, którzy napisali o tym w magazynie „Science”, nie jest jasne

jednak, czy woda pochodzi z chmury będącej fazą złączkową Układu Słonecznego, czy w ogóle przybyła z zewnątrz. Skąd uczeni wiedzą, jak stara jest woda? Zbadano jej wiek za pomocą analizy stosunku izotopu deuteru w atomach wodoru. Woda powstająca w określonych warunkach różni się jego zawartością od wody formującej się w innych. np. zawartość deuteru w wodzie powstałej w warunkach międzygwiazdowych jest relatywnie wysoka.

Geolodzy z University of Saskatchewan w Kanadzie posłużyli się zaawansowaną symulacją, za pomocą której obliczyli, że wielkie ilości wody mogą powstawać nie gdzieś w komosie ale we wnętrzu Ziemi dzięki reakcjom chemicznym zachodzącym na głębokości nawet tysiąca kilometrów. Woda powstaje tam w temperaturze około 1400°C i pod ciśnieniem 20 tysięcy razy większym od atmosferycznego.

Trzy atmosfery Ziemi

Pochodzenie wody na Ziemi nie jest więc wcale jasne. A co z atmosferą? Uważa się, że nieostygła jeszcze kula magmy wydzielala gazy takie jak azot, wodór, węgiel i tlen, tworząc najstarszą znaną wersję atmosfery. „Stwierdziliśmy, że atmosfera, którą według naszych obliczeń była obecna na Ziemi miliardy lat temu, miała podobny skład do tego, co znajdujemy dziś na Wenus i Marsie”, pisał Paolo Sossi z ETH w Zurichu, współautor pracy na ten temat opublikowanej w czasopiśmie „Science Advances” w 2021 r.

Mówi się, że Ziemia miała w swojej historii trzy atmosfery. Pierwsza atmosfera, przechwycona być może z mgławicy słonecznej, składała się z lekkich pierwiastków z mgławicy słonecznej, głównie wodoru i helu. Działanie wiatru słonecznego i ciepła

8. Komet 67P Czuriumow-Gierasimienko



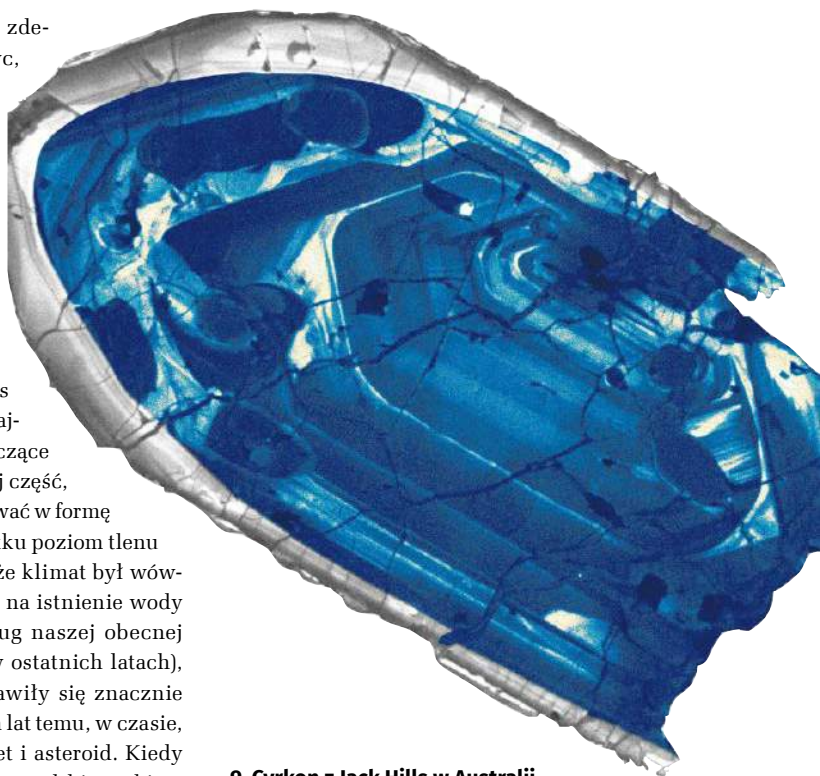


Ziemi wyparło tę atmosferę. Po zderzeniu, które stworzyło Księżyc, stopiona Ziemia uwolniła lotne gazy; a później więcej gazów zostało uwolnionych przez wulkany, uzupełniając drugą atmosferę bogatą w gazy cieplarniane, ale ubogą w tlen. Wreszcie trzecia atmosfera, bogata w tlen, pojawiła się, gdy bakterie zaczęły produkować tlen około 2,8 mld lat temu.

W rejonie wzgórz Jack Hills w Australii już wiele lat temu znajdowano ślady cyrkonu (9) świadczące o tym, że Ziemia, przynajmniej jej część, musiała się schłodzić i skryształizować w formę stałą już 4,4 mld lat temu. W dodatku poziom tlenu w tamtejszych skałach sugeruje, że klimat był wówczas na tyle łagodny, że pozwalał na istnienie wody w stanie ciekłym. Jednak według naszej obecnej wiedzy (publikacje w „Nature” w ostatnich latach), najwcześniejsze formy życia pojawiły się znacznie później, co najmniej 3,95 miliarda lat temu, w czasie, gdy trwał intensywny nalot komet i asteroid. Kiedy japońscy badacze w tym badaniu, pod kierunkiem geologa Tsuyoshi Komiya z Uniwersytetu w Tokio, zbadał próbki skał z Kanady, zauważyli ślady życia w postaci biogenicznie modyfikowanej skały grafitowej. Te wyniki zresztą są dla wielu zbyt kontrowersyjne – inni eksperci wątpią w ich datowanie. Ogólnie jednak w ostatnich latach rośnie grono badaczy spychających początki życia do czasów kiedy naukowcy kiedyś uważali Ziemię za niezdatną do zamieszkania.

Z czasów, gdy nasza planeta powstała około 4,5 miliarda lat temu na powierzchni praktycznie nie ma żadnych śladów skalnych. Być może są ukryte we wnętrzu planety, podobnie jak fragmenty tajemniczej Thei, w pewnym sensie naszej współmatki. To co udaje się znaleźć z tych najstarszych pozostałości prawie zawsze prowadzi do zamieszania w ustalonej wiedzy. np. wspomniane mikroskopijne minerały wydobyte w 2020 r. ze starożytnej odkrywki Jack Hills, w Zachodniej Australii, wydają się nosić ślady ziemskiego pola magnetycznego sięgającego aż 4,2 miliarda lat wstecz, czyli miliard lat wcześniej, niż zakładano. Może to sugerować, że pole magnetyczne ziemi istnieje praktycznie od samego początku istnienia naszej planety.

Pole magnetyczne odgrywa ważną rolę w przydatności naszej planety do życia. Jest kolejnym, po ciekłej wodzie, atmosferze, kluczowym warunkiem istnienia form biologicznych, działa bowiem jako swego rodzaju tarcza,



9. Cyrkon z Jack Hills w Australii

odchylająca zabójczy, wiatr słoneczny, który mógłby zniszczyć atmosferę. Dawniej uważano, że pole magnetyczne Ziemi istniało co najmniej 3,5 miliarda lat temu. Jednocześnie uważa się, że jądro planety zaczęło krzepnąć zaledwie miliard lat temu, co oznacza, że pole magnetyczne musiało być napędzane wcześniej przez jakiś inny mechanizm. Nie wiadomo jednak dokładnie jaki.

Zakładając więc, że wszystkie te elementy, energię, wodę w stanie ciekłym, atmosferę (beztlenową, ale dziś wiemy, że nie każdemu życiu tlen jest potrzebny) i w końcu ochronne pole magnetyczne, istniały już bardzo wcześnie w dziejach Ziemi, przestajemy się dziwić, że naukowcy znajdują coraz starsze ślady wskazujące na istnienie życia na naszej planecie.

A co jeśli życie nie powstało na Ziemi?

Od dawna podejrzewano, że przynajmniej część elementów, które przyczyniły się do powstania życia na Ziemi, mogła dotrzeć na naszą planetę wraz z odłatkami kosmicznych skał. Od niedawna mamy kolejne dane sprzyjające tej teorii. Europejska Agencja Kosmiczna ogłosiła w piątek, że sonda Rosetta odkryła w komie (ogonie) komety 67P – Czuriumow-Gierasimienko kilka z tak zwanych „budulców życia”, a w tym fosfor oraz jeden

z aminokwasów – glicynę. Oznacza to, że kolizje tego typu obiektów z Ziemią mogły doprowadzić do zwiększenia stężenia substancji, dzięki którym powstały pierwsze organizmy.

Wcześniej, jeszcze w 2006 roku NASA wykryła glicynę w próbkach pochodzących z ogona komety Wild 2, jednak ich analiza była utrudniona, gdyż podejrzewano, że zostały zanieczyszczone na Ziemi. Naukowcy uważają ostatnie wieści za niezwykle ważne, ponieważ tym przypadku możemy być pewni, że mamy do czynienia z czystymi próbkami.

Glicyna to powszechny składnik białek, natomiast fosfor jest kluczowym elementem budującym DNA. Co ciekawe, nie są to jedyne substancje potrzebne do powstania życia, które odnaleziono na kometach – pozostałe to cyjanowodor i siarkowodor. Według opublikowanego w magazynie „Science” raportu zespołu badawczego pracującego przy misji sondy Rosetta, instrument do badania składu chemicznego COSAC, wykrył w pyłe unoszącym się nad powierzchnią związki alkoholowe, karbonylowe, amidowe, nitrylowe oraz izocyjaniany. Niektóre z nich, jak np. aceton, acetaamid, izocyjanian metylowy, nigdy wcześniej nie zostały zaobserwowane na kometach.

Wymienione związki organiczne mogą być podstawą do powstawania bardziej złożonych struktur, takich jak cukry, aminokwasy i nawet podstawy DNA, ale żadnych tego rodzaju cząsteczek nie wykryto w sposób jednoznaczny.

Specjaliści z Uniwersytetu Edynburskiego przypuszczają, że żywe organizmy mogły zostać przeniesione

na Ziemię wraz ze strumieniami pyłu kosmicznego. Informację o tej hipotezie podała strona phys.org. Naukowcy uważają, że mocne strumienie pyłu kosmicznego przemieszczające się z prędkością do 70 km/s są w stanie przenosić najmniejsze żywe organizmy. Ekspertci między innymi twierdzą, że pył kosmiczny ciągle bombardujący atmosferę Ziemi może wyprzeć poza ziemskie pole grawitacyjne drobne cząsteczki znajdujące się na wysokości 150 km lub więcej nad powierzchnią i w końcu osiągnąć innych planet.

Badacze twierdzą, że podobna sytuacja może mieć miejsce na innych planetach, a więc żywe organizmy mogą trafić z jednego ciała niebieskiego na drugie. Jednocześnie specjaliści podkreślają, że niektóre żywe istoty, w tym bakterie, rośliny, mikroskopijne bezkręgowce niesporczaki (10) są w stanie przetrwać tylko w warunkach otwartej przestrzeni kosmicznej, co potwierdzałyby hipotezy ekspertów.

Teoria panspermii, siania życia na Ziemi z kosmosu, jest już stosunkowo stara. Obecnie nie brzmi już tak sensacyjnie jak sto lat temu. Współczesna zmiana w myśleniu o niej polega głównie na tym, że powiększyła się znacznie nasza wiedza o ciałach Układu Słonecznego. Dopuszczamy obecnie możliwość istnienia życia w wielu miejscach, które kiedyś uważaliśmy za martwe i wręcz wrogie życiu. Nie rozwiązujemy w tym ujęciu kwestii – jak powstało życie – lecz raczej przenosimy ją na inny szerszy plan kosmiczny. Nie oznacza to jednak, że przestajemy mieć tu lukę w wiedzy. ■

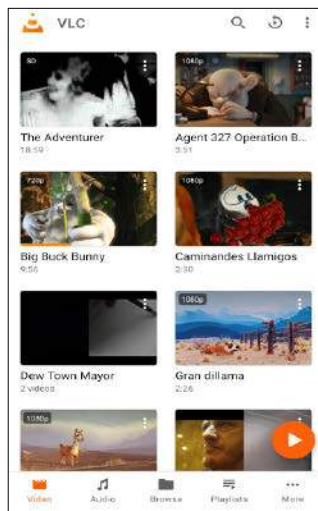
Mirosław Usidus

10. Niesporczaki w przestrzeni kosmicznej





Programy do obsługi wideo i multimedialnych



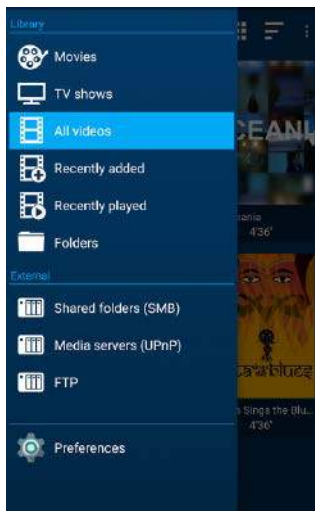
VLC

VLC to popularny i darmowy odtwarzacz multimedialny. Dostępna jest również mobilna aplikacja, dostępna na urządzenia z systemem Android i iOS. VLC jest darmowym i otwartoźródłowym oprogramowaniem. Jednak dostępne są również wersje płatne aplikacji z usługami bez reklam oraz dodatkowymi funkcjami.

VLC obsługuje większość popularnych formatów plików multimedialnych. Umożliwia odtwarzanie plików wideo i audio m.in. MP4, MKV, AVI, FLAC itp. Obsługuje strumieniowanie multimedialnych z urządzenia przez sieć Wi-Fi z podłączonego dysku. Umożliwia również odtwarzanie plików zapisanych na karcie SD telefonu lub bezpośrednio z pamięci wewnętrznej urządzenia.

VLC umożliwia przeglądanie i odtwarzanie wszystkich plików wideo, muzyki i zdjęć znajdujących się w pamięci urządzenia. Zapewnia też obsługę napisów, wyświetlając je do odtwarzanych filmów i seriali. Wady programu to m.in. brak wsparcia dla usługi Chromecast, czyli nie można np. streamować materiału na inne urządzenia, np. telewizory. Niektóre opcje i ustawienia VLC mogą być nieco skomplikowane dla przeciętnego użytkownika.

VLC Media Player	
Producent	VideoLAN Organization
Platforma	Android, iOS, macOS, Windows, Linux, Amazon Apps, Huawei App Gallery, F-Droid
Oceny	Możliwości
	Łatwość obsługi
	Ocena ogólna
	9/10
	7/10
	8/10



Archos Video Player

Archos Video Player jest mobilną aplikacją do odtwarzania multimedialnych stworzoną przez francuską firmę Archos. Służy do odtwarzania plików wideo (MP4, MKV, AVI itp.) oraz muzyki (MP3, FLAC, WAV, AAC itd.) na urządzeniach z systemem Android. Aplikacja obsługuje pliki wideo HD, wyświetlane w wysokiej jakości na ekranie smartfona. Zapewnia wbudowaną obsługę wielu kodeków wideo i audio bez konieczności ich instalowania.

Oferuje takie funkcje jak wstrzymanie i wznowienie odtwarzania, przewijanie do przodu i do tyłu, ustawianie głośności itp. Możliwe jest odtwarzanie filmu w tle, co pozwala korzystać z innych aplikacji podczas oglądania. Użytkownicy mogą korzystać z funkcji śledzenia postępu odtwarzania, usuwania zakładek oraz zmiany jasności, kontrastu i nasycenia obrazu.

Archos Video Player oferuje możliwość odtwarzania wideo na tabletach, telefonach i na urządzeniach AndroidTV i kilku innych, mniej znanych w Polsce systemach smart TV (Nexus Player, Nvidia SHIELD TV, Rockchip i AmLogic). Aby uzyskać pełną wersję aplikacji (wszystkie funkcje, brak reklam) można kupić płatną wersję aplikacji. Archos Video Player jest kompatybilny z systemem Android 4.0 i powyżej. Jego wadą, obniżającą ocenę, jest dostępność jedynie na Androida.

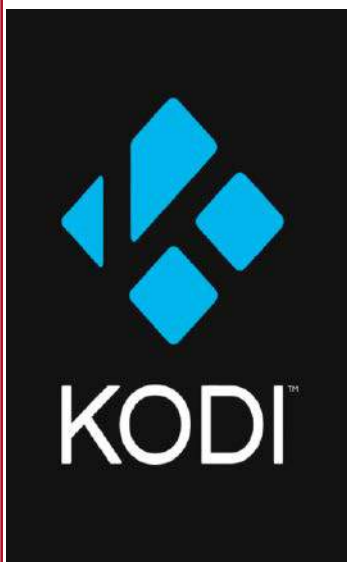
Archos Video Player	
Producent	Archos SA
Platforma	Android
Oceny	Możliwości
	Łatwość obsługi
	Ocena ogólna
	7/10
	9/10
	8/10

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 oczywiste, 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 produkowanych przez nią smartfonach z charakterystyczną, pełną klawiaturą QWERTY. Także w niektórych telefonach innych firm (HTC, Motorola, Nokia, Samsung, Sony Ericsson).



Kodi

Kodi (wcześniej XBMC) jest otwartą platformą multimedialną. Pozwala na odtwarzanie plików audio i wideo, oglądanie telewizji, słuchanie radia oraz korzystanie z wielu dostępnych dodatków. Kodi jest dostępny na wiele platform sprzętowych, w tym Android, iOS, Linux, macOS, Windows i innych, z tym, że nie zawsze są to proste metody instalacji, gdyż np. w urządzeniach wymagana jest wersja tzw. „jailbroken”.

Odtwarza niemal wszystkie popularne formaty plików wideo i audio. Obsługuje odtwarzanie strumieniowe z komputera, usług takich jak Netflix czy Amazon Prime Video lub z NAS i innych urządzeń sieciowych. Możliwe jest podłączenie do Kodi fizycznych tunerów telewizyjnych lub radiowych. Dostępny jest szeroki wybór dodatków (plug-inów), dzięki którym Kodi może świadczyć wiele dodatkowych funkcji.

Wydawcy podkreślają, że Kodi jest jedynie usługą pośredniczącą a aplikacja sama w sobie nie dostarcza i nie zawiera żadnych materiałów lub treści. Użytkownicy sami dostarczają własne treści lub instalują odpowiednie wtyczki stron trzecich. Program Kodi domyślnie uruchamia się w angielskiej wersji językowej. Aby korzystać z niego z polskiej wersji, należy zmienić opcję w ustawieniach programu.

Kodi	
Producent	Kodi Foundation
Platforma	Android, iOS, macOS, Windows, tvOS, Linux, Raspberry Pi, Huawei App Gallery
Oceny	
Możliwości	9,5/10
Łatwość obsługi	8,5/10
Ocena ogólna	9/10



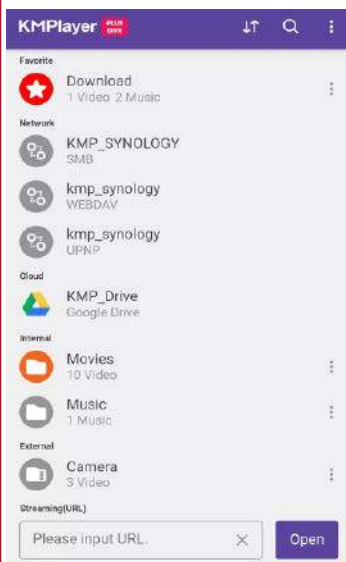
ALLPlayer

To najczęściej pobierany, darmowy i polski program do oglądania filmów oraz słuchania muzyki. Oferuje odtwarzanie filmów CD/DVD, multimediiów z plików .torrent oraz filmów spakowanych RARem, obsługę do 4 monitorów lub telewizorów, automatyczne odtwarzanie kolejnych części podzielonych filmów lub seriali, wsparcie Dolby Surround, DTS, 3D Audio, SPDIF i innych, funkcje equalizera i audiowizualizacji, obsługę strumieni audio-video, łącznie z filmami z YouTube.

Użytkownik otrzymuje w aplikacji szereg funkcji i możliwości, np. dowolne obracanie obrazu, korekty kolorów z funkcją poprawy jakości, wyłączenie komputera/monitora po zakończonym seansie, autoresume – wznowianie oglądania od miejsca zakończenia, a także funkcję rodzicielską – zabezpieczanie pliku hasłem.

Łącząc się ze znanym serwisem z napisami Napisy24.pl oraz takimi bazami jak Napirojekt.pl i Opensubtitles.org, ALLPlayer automatycznie pobiera dopasowane napisy w wybranym języku. Nie można też zapomnieć o wbudowanym edytorze napisów. Producent chwali się wynikami badań Mobience, prowadzonych przez Spicy Mobile, według których odtwarzacz ALLPlayer jest w pierwszej trójce najpopularniejszych odtwarzaczy wideo używanych na telefonach komórkowych z systemem Android, w kategorii Media i Video.

ALLPlayer	
Producent	ALLPlayer Group
Platforma	Android, iOS, Windows
Oceny	
Możliwości	8,5/10
Łatwość obsługi	9,5/10
Ocena ogólna	9/10



KMPlayer

Obsługuje wiele formatów, w tym VCD, DVD, AC3, DTS, AVI, MPEG-1/2/4, WMV, WMA, OGG, RealMedia, QuickTime i wiele innych oraz całą gamę formatów muzycznych i formatów napisów. Potrafi odtwarzać wideo o jakości do 4K, 8K UHD. Zawiera także m.in. klienta usługi Torrent w tzw. wersji VIP.

Przez The KMPlayer możemy ustawić specjalne efekty dla ścieżki audio i wideo, spowalniać lub przyspieszać prędkość odtwarzanego filmu, wskazywać ulubione fragmenty filmów i znacznie więcej. KMPlayer wyposażony został również w funkcje, umożliwiające odtwarzanie materiałów wideo w różnych trybach (np. klatka po klatce), wyświetlanie napisów przy użyciu kilku metod i mechanizmów, zarządzanie kanałami dźwiękowymi, tworzenie galerii miniatur z filmu, modyfikowanie rozmiaru i proporcji wyświetlanego filmu oraz obsługę napisów 3D.

W wersji płatnej (VIP), dostępne są funkcje edycji i personalizacji plików multimedialnych z konwersjami, wyodrębnianiem dźwięku itp. Appka wspiera Chromecasta, czyli możliwa jest współpraca np. z smart TV. Bezpłatna wersja obejmuje jedynie jedno konto Google Play.

Plex	
Producent	Plex Inc.
Platforma	Android, iOS, Windows
Oceny	
Możliwości	8/10
Łatwość obsługi	8/10
Ocena ogólna	8/10

Wcale nie takie rzadkie (1)

Nauczyciel z rzadka wspomni o nich na lekcjach chemii, wiele nie znajdziesz również w swoim podręczniku. Nominalnie należąc do grupy trzeciej, umieszczane są zazwyczaj na uboczu, pod całą tablicą układu okresowego. Zalicza się je nawet do metali rzadkich, ale po przeczytaniu artykułu dojdiesz do wniosku, że nie zasłużyły na takie miano.

O tym, że wcale nie są takie rzadkie przeczytasz dalej. Bohaterowie artykułu to metale, bez których trudno byłoby osiągnąć współczesny poziom technologii, przykładem choćby wszechobecne smartfony. Opowieść zacznie się jednak od pasjonującej historii ich odkrywania.

Kalendarium część 1

W wieku XIX **lantanowce** miały wielkie trudności z przyjściem na świat: często, już odkryte, znikwały, a zazwyczaj „pączkowały” wydając kolejne metale znajdujące się w ich związkach. Wielu odkrywców musiało więc pogodzić się z faktem, że pierwiastek okazał się mieszaniną kilku innych. Długo nie mogli również znaleźć miejsca w układzie okresowym, a kłopoty z ich zameldowaniem trwają właściwie do dziś. Lantanowce nadal są wyjątkowo trudne do identyfikacji. Nie jest to jednak wina niedbałości chemików, ale rezultat niespotykanego w innych grupach chemicznego podobieństwa. Dzieje odkryć lantanowców i nierozdzielnie związanych z nimi skandu i itru (bardzo podobnych ze względu na właściwości i występowanie), poznasz w formie kalendarium.

1794. Rozpoczyna się historia **metali ziem rzadkich** (taka jest wspólna nazwa dla skandu, itru i lantanowców). Fiński chemik **Johan Gadolin**, pracujący w Szwecji, wykonał analizę minerału o nazwie **iterbit**, znalezionej kilka lat wcześniej w kopalni **Ytterby** położonej na wyspie w pobliżu Sztokholmu. Gadolin stwierdził w minerale obecność nieznaną ziemi (jak wtedy nazywano tlenki metali). Wykazał również, że znajduje się w niej pierwiastek o właściwościach odrębnych od ówczesnie znanych. Odkrywcą zauważył (a ty dobrze zapamiętaj ten cytat), że: *niektóre własności nowej ziemi są zbliżone do ziemi alunowej* (tlenek glinu Al_2O_3), a inne do ziemi wapniowej (tlenek wapnia CaO). W następnych latach wyniki analizy potwierdzili inni chemicy. Metal otrzymał nazwę **itr** (skrót od Ytterby), tlenek – **itria** (ziemia itrowa), a nazwę minerału zmieniono na **gadolinit** (1).



1. Od niego zaczęła się historia metali ziem rzadkich – Johan Gadolin (1760-1852).

1803. Historia lubi się powtarzać. Niemiecki chemik **Martin Klaproth** wykonał analizę odkrytego przed kilkudziesięciu laty **bastnazytu**, minerału pochodzącego ze szwedzkiej kopalni w Bastnäs. I on znalazł w nim nieznaną wcześniej ziemię. W tym samym czasie również **Jöns Jacob Berzelius** wraz ze swym uczniem **Wilhelmem Hisingerem** wydzielili z bastnazytu nowy tlenek (zgodnie z ówczesnym nazewnictwem była to ziemia cerytowa lub ceria), a obecnemu w niej metalowi nadali nazwę **cer** (od odkrytej dwa lata wcześniej planetoidy Ceres). Metaliczny cer otrzymano dopiero po ponad 30 latach, a pierwiastek o wysokim stopniu czystości – w ubiegłym stuleciu, podobnie zresztą jak większość pozostałych lantanowców.

1839. **Carl Mosander**, uczeń Berzeliusa, w latach 20. XIX wieku zaczął podejrzewać, że ziemia cerytowa nie jest substancją jednorodną. Prace analityczne trwały długo – dopiero po kilkunastu latach

udało mu się rozdzielić tworzące ją tlenki. Metal znajdujący się w nowej ziemi otrzymał nazwę **lantan** (gr. *lanthanein* = ukryty) (2).

1842. Mosander nie spoczął na laurach i dalej badał otrzymaną uprzednio ziemię. Doszedł do wniosku, że i lantan ma „bliźniaka” skrywającego się w jego tlenku. Kilka lat pracy zaowocowało wydzielaniem jeszcze jednej ziemi, a obecny w tlenku metal otrzymał nazwę **didym** (gr. *didymoi* = bliźnięta) i symbol Di. Na próżno jednak szukać tego pierwiastka we współczesnym układzie okresowym, „żył” tylko 43 lata (o jego dalszych losach dowiesz się za miesiąc).

1843. Mosander, niekwestionowany lider wśród badaczy pierwiastków ziem rzadkich I połowy XIX wieku, przeanalizował dokładnie również ziemię itrową i rozdzielił ją aż na trzy nowe tlenki. Były to tlenek znanego już itru oraz dwie nowe ziemie – erbinowa i terbinowa. Metale zawarte w tlenkach otrzymały nazwy **erb** i **terb** (podobnie jak w przypadku itru, ich imiona powstały przez skrócenie nazwy kopalni Ytterby). To jednak nie koniec historii tych dwóch pierwiastków. Ponad 20 lat później szwajcarski chemik **Marc Delafontaine** ponownie zbadał ziemię itrową. Potwierdził wyniki uzyskane przez Mosandera, ale – nie wiadomo dlaczego – zmienił nazwy otrzymanych z niej tlenków. I tak brunatno-żółtej barwy erbia (ziemia erbinowa Mosandera) stała się tlenkiem terbu, a różowa terbia – tlenkiem erbu. Tym samym i pierwiastki również zamieniły się nazwami (3). I tak już zostało.

W tym miejscu następuje przerwa w opowieści o historii odkryć metali ziem rzadkich, zwłaszcza, że przez



2. Carl Gustav Mosander (1797-1858)
– odkrywca lantanu, didymu, erbu i terbu.

kolejne 35 lat w ich chemii właściwie nic nowego się nie działo. Za miesiąc część druga kalendarium.

Metale nowych technologii

Historycznie pierwszym zastosowaniem metali ziem rzadkich była technika oświetleniowa. **Auer von Welsbach**, odkrywca kilku z nich, zauważył, że niektóre tlenki emitują silne światło

3. Tlenki wybranych lantanowców (od góry zgodnie z kierunkiem ruchu wskazówek zegara): prazeodymu (ciemny), ceru, lantanu, neodymu, samaru i gadolinu.



po ogrzaniu w płomieniu palnika. Auer nasycił bawełnianą siatkę mieszaniną związków toru i lantanowców. Po umieszczeniu siatki w płomieniu lampy gazowej, włókna bawełniane spalały się i pozostawał szkielet z trudnotopliwych tlenków. W latach 80. XIX wieku w domach powszechnie stosowano oświetlenie gazowe, siatki auerowskie okazały się wielkim sukcesem rynkowym oraz ratunkiem dla gazowni produkujących gaz z węgla, zagrożonych przez wchodzące do użycia oświetlenie elektryczne. Być może widziałeś taką siatkę, ponieważ do dzisiaj używa się ich w turystycznych lampach montowanych na butlach gazowych (4).



4. Koszulki auerowskie stosuje się także w lampach naftowych.



5. Piroforyczne właściwości lantanowców wykorzystano do produkcji kamieni do zapalniczek (i-eagle123, Pixabay.com)

Konsekwencją niespotykanego w innych grupach chemicznego podobieństwa jest możliwości zastosowania lantanowców w postaci nierozdzielonej mieszaniny związków lub stopu metali, co znacznie obniża koszty otrzymywania. **Metal mieszany** (z niem. *Mischmetall*) to stop lantanowców (głównie Ce, La, Nd, Pr; skład zależy od źródła pochodzenia, czyli minerału zawierającego lantanowce) służący do odwęglania, odtleniania i odsiarczania stali oraz jako dodatek stopowy zwiększający odporność na korozję i wytrzymałość mechaniczną, np. aluminium. Metal mieszany w postaci rozdrobionej ma właściwości piroforyczne i stosowany jest jako kamień do zapalniczek (5). Mieszanina tlenków lantanowców to z kolei dobry proszek polerski, niektóre z nich używa się jako środki do barwienia i odbarwianie szkła.

W luminoforach kineskopów telewizji kolorowej i rtęciowych świetlówek stosowane były związki lantanowców, ale to już historia. Natomiast współczesnością jest coraz szersze ich użycie w wyświetlaczach nowej generacji: ekrany i diody LED oparte na związkach metali ziem rzadkich to nieodłączny element smartfonów, laptopów i innych urządzeń (6).

Technika kosmiczna i lotnicza również wiele zawdzięcza lantanowcom. Oczywiście wytrzymałe stopy metali mają znaczenie kluczowe, ale równie ważne są akumulatory, nadprzewodniki wysokotemperaturowe (do ich produkcji stosuje się głównie związki

lantanu) i silne magnesy neodymowe i samarowo-kobaltowe. Technika laserowa to z kolei urządzenia oparte na itrze, neodymie, erbie, tulu i holmie.

Samar, europ i gadolin, ze względu na dużą zdolność do pochłaniania neutronów, wchodzi w skład prętów kontrolnych oraz szkła okien w obudowie reaktorów jądrowych. Lantanowce i ich związki to również katalizatory krakingu, procesu w którym z ciężkich frakcji ropy naftowej

6. Wyświetlacze smartfonów to współcześnie ważne zastosowanie lantanowców.





7. Z lewej gadolinit z kopalni Ytterby (autor: Rob Lavinsky, iRocks.com), z prawej bastnazyt z Koronge w Burundi.

uzyskuje się przeważającą część produkcji benzyn. Petrochemia wraz z metalurgią pochłaniają zresztą zdecydowaną większość światowej produkcji lantanowców.

Jak więc widzisz, duży fragment naszej cywilizacji nie istniałaby bez lantanowców (w każdym razie w znanej obecnie formie).

Nie takie rzadkie

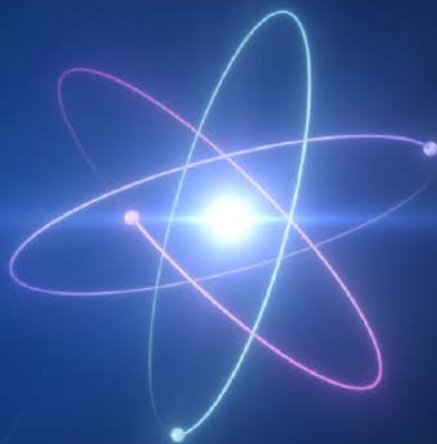
Skandowce i lantanowce tworzą tylko nieliczne własne minerały, a i te znajdują się w niewielu miejscach świata. To powód, dla którego w I połowie XIX wieku otrzymały miano metali ziem rzadkich. Najczęściej spotykany z nich cer zajmuje 27 miejsce na liście rozpowszechnienia pierwiastków w powierzchniowej warstwie Ziemi. Zawartość wynosząca niecałe 5 tysięcznych procenta masy nie wygląda imponująco, ale ceru jest niewiele mniej niż azotu, a więcej niż licznych metali uważanych za pospolite, np. cynku, kobaltu, cyny czy też ołowiu. Nawet najrzadziej występującego tulu (pomijając promet) jest dwukrotnie więcej niż srebra, a kilka razy więcej niż złota i platynowców razem wziętych. „Rzadkość” metali grupy trzeciej spowodowana jest ich znacznym rozproszeniem w minerałach innych pierwiastków, przede wszystkim bardzo rozpowszechnionego glinu.

Minerały zawierające lantanowce dzielą się na dwie grupy: pierwsza z nich zawiera lżejsze metale, druga – itr i cięższych członków rodziny. Najważniejsze minerały z pierwszej grupy to **monacyt** (zawierający również do kilkunastu procent toru, zwykle eksploatowane

są złoża wtórne, czyli piaski monacytowe – Brazylia, Indie, USA, Cejlon) oraz znaleziony w Szwecji **bastnazyt** (występuje również w innych miejscach na świecie). Przemysłowo eksploatowane minerały z grupy drugiej to **gadolinit** (oprócz Skandynawii w większych ilościach jest spotykany w USA) i **samarskit** (największe złoża znajdują się w USA) (7). Lantanowce stanowią również domieszkę w wydobywanych na wielką skalę apatytach (surowiec do produkcji nawozów fosforowych), skąd również się je otrzymuje. Technologia przeróbki ich minerałów uzależniona jest od składu chemicznego surowca. Rozdzielanie lantanowców należy jednak do najtrudniejszych i najbardziej wymagających technologicznie procesów metalurgicznych, stąd często otrzymuje się tylko mieszaninę metali i ich związków (do wielu zastosowań nie ma potrzeby ich rozdzielania). Promieniotwórczy promet jest natomiast wyodrębniany ze zużytego paliwa do reaktorów.

Światowa produkcja metali ziem rzadkich rośnie w tempie ponad 10% rocznie i w roku 2022 osiągnęła poziom około 300 tysięcy ton (w przeliczeniu na tlenki). Monopolistą na rynku są Chiny dostarczające 210 tys. ton (w przypadku niektórych metali udział tego kraju wynosi ponad 90%). Na kolejnych miejscach znajdują się USA, Australia i Burma z udziałami rzędu 5-10%. Skupienie produkcji strategicznych metali w rękach jednego dostawcy spędza sen z powiek ekonomistów i polityków, stąd też na całym świecie trwają usilne poszukiwania nowych złóż. ■

Krzysztof Orlński



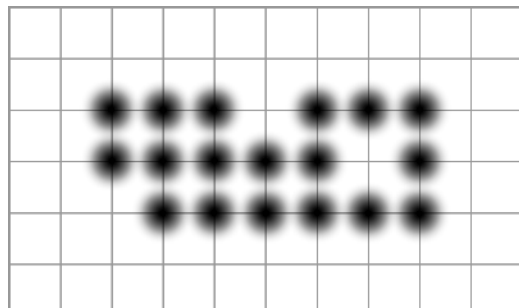
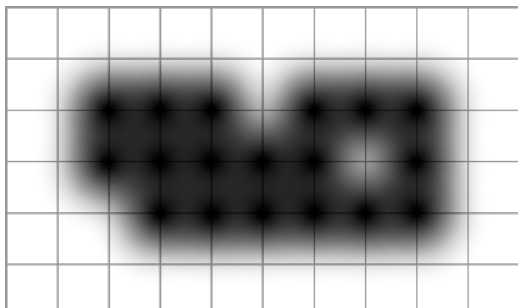
Mikroskop sił atomowych (1)

Zobaczyć atomy

Oko, zarówno w przypadku człowieka jak i zwierząt, ma bardzo podobną budowę i pozwala widzieć obrazy przedmiotów oświetlonych światłem widzialnym. Jakkolwiek nie ma górnego ograniczenia na wielkość przedmiotu jaki możemy dostrzec nieuzbrojonym okiem, problem zaczyna się pojawiać w sytuacji przedmiotów o niewielkich rozmiarach.

Ponieważ oko jest przyrządem optycznym działającym dzięki znajdującej się w nim soczewce, jak każde takie urządzenie ma ograniczoną zdolność rozdzielczą, czyli zdolność rozróżniania dwóch punktów położonych blisko siebie. Z tym mankamentem

można sobie poradzić, używając na przykład lupy lub mikroskopu optycznego, które pozwalają widzieć obserwowane obiekty w powiększeniu i z lepszą zdolnością rozdzielczą niż w przypadku oka nieuzbrojonego.



1. Obraz układu punktów na płaszczyźnie widziany przez urządzenie o niskiej zdolności rozdzielczej (lewa strona) oraz przez urządzenie o wystarczająco wysokiej zdolności rozdzielczej (prawa strona) przy tym samym powiększeniu

Sprawdź swoją wiedzę

Mamy do dyspozycji wiązkę światła widzialnego o długości fali równej 680 nm. Wybierz z poniższej listy te obiekty, które zaobserwujemy wykorzystując posiadane źródło światła:

- ☐ a) długopis o długości około 15 cm i średnicy około 1 cm;
- ☐ b) wewnętrzna struktura kryształu dla którego odległość między atomami wynosi 0,3 nm;
- ☐ c) struktura ścieżek o szerokości około 500-600 nm zapisanych na płycie CD;
- ☐ d) wirusy o rozmiarach rzędu 50 nm.

Do każdego obiektu dopasuj odpowiednią technikę obserwacji: obserwacja w świetle odbitym, obserwacja w świetle przechodzącym, rekonstrukcja budowy na podstawie obrazu interferencyjnego.

Jest jednak pewna dolna granica poniżej której urządzenia optyczne stają się zupełnie bezużyteczne. Bez problemu obserwujemy światło odbite od przedmiotów których rozmiary są większe od długości fali światła lub światło przechodzące przez materiały choćby częściowo przezroczyste (np. cienkie próbki tkanek używane jako preparaty mikroskopowe). Do tego typu przypadków stosują się prawa optyki liniowej. Można również wykorzystać światło ulegające dyfrakcji i interferencji na obiektach, których rozmiary są porównywalne z długością fali świetlnej, a następnie na podstawie uzyskanych obrazów zrekonstruować ich budowę wewnętrzną. Nie jesteśmy jednak w stanie uzyskać obrazu optycznego obiektów o rozmiarach mniejszych niż długość fali świetlnej, takich jak cząsteczki czy pojedyncze atomy. Musimy w ich przypadku zastosować inne techniki.

Odrobina historii

Pierwsze mikroskopy optyczne powstały pod koniec XVI wieku w Holandii, jednak ze względu na niewielkie powiększenie okazały się mało przydatne jako narzędzia badawcze. Dopiero w XVII wieku Antoni van Leeuwenhoek udoskonalił ich konstrukcję oraz uruchomił komercyjną produkcję tych urządzeń.

Mikroskop optyczny według jego projektu pozwalał na obserwacje szerokiej grupy tkanek roślinnych i zwierzęcych oraz organizmów jednokomórkowych.

Kolejny przełom dokonał się w roku 1931, gdy został skonstruowany mikroskop elektronowy, pozwalający zastąpić światło widzialne wiązką rozprzeczonych elektronów. Dzięki temu możliwa stała się obserwacja budowy drobnych organelli komórkowych, niedostępnych do badań wcześniejszymi metodami, jak również wewnętrznej struktury niektórych materiałów, na przykład półprzewodników.

W roku 1982 powstał pierwszy skaningowy mikroskop tunelowy, wykorzystujący zjawisko tunelowe w celu uzyskania obrazu powierzchni z rozdzielczością rzędu pojedynczego atomu. Niemniej zasada działania tego mikroskopu narzuca pewne ograniczenie: nie nadaje się on do obrazowania powierzchni izolatorów. Pomysł jako taki stał się inspiracją do skonstruowania w roku 1986 mikroskopu sił atomowych, który pozwala uzyskiwać obrazy z podobną rozdzielczością jak w przypadku skaningowego mikroskopu tunelowego. Mikroskopu sił atomowych działa jednak niezależnie od właściwości elektrycznych danego materiału. ■

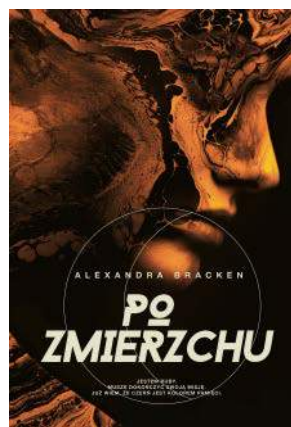
Joanna Borgensztajn

Po zmierzchu. Mroczne umysły. Tom 3

Alexandra Bracken

Wydawnictwo Jaguar, liczba stron: 592, cena: 59,90 zł

Trzeci tom zapierającego dech w piersiach cyklu Mroczne umysły. Ruby nie może oglądać się za siebie. Ona i inne dzieci, które przeżyły rządowy atak na Los Angeles, jadą na północ, by się przegrupować. Jest z nimi więzień: Clancy Gray, syn prezydenta i jedna z nielicznych osób, które mają takie zdolności jak ona. Tylko Ruby ma nad nim jakąkolwiek władzę, a jedno potknięcie może doprowadzić do katastrofy... Tymczasem tysiący takich jak oni wciąż cierpią w „obozach rehabilitacyjnych”. Ruby musi ich uwolnić...





Mysząc o informatyce możemy skupić się na klawiaturze, w którą stukając, wywołujemy pewne, oczekiwane przez nas procesy. Zachodzą one w małych urządzeniach, które dzięki konstruktorom przyjmują ergonomiczne, ekonomiczne i estetyczne kształty. By dojść do obecnego stanu rzeczy musiało upłynąć wiele lat.

Informatyka

Pierwszy komputer, który powstał w latach czterdziestych ubiegłego wieku zajmował aż 167 m². W dzisiejszych czasach również mamy odczynienia z maszynami zajmującymi znaczną przestrzeń, jak na przykład serwery GPT, ale związane jest to ze skomplikowaną architekturą i możliwościami obliczeniowymi, które przewyższają te z ubiegłego wieku. Jednak wielkość ani na początku rozwoju informatyki, ani w jej obecnym momencie nie jest istotna, a jedynie możliwości jakie dają nowe urządzenia napędzane ludzką wyobraźnią i chęcią rozwoju. Zapraszamy na najbardziej dynamicznie rozwijający się kierunek nauki. Zapraszamy na Informatykę.

Kurs vs studia

Współcześnie, praca związana z szeroko pojętą informatyką nie musi być powiązana z ukończeniem studiów, które gwarantowałyby odniesienie sukcesu w tej dziedzinie. Rozwój technologii i Internetu zwiększył dostęp do wysokiej jakości zasobów edukacyjnych online. Pojawiło się wiele platform e-learningowych, samouczków, kursów online i materiałów dostępnych dla każdego, kto chce nauczyć się programowania i innych umiejętności związanych z informatyką. Samodzielnie zdobyte doświadczenie i gotowość do realizowania różnorodnych tematycznie projektów może przyciągnąć uwagę potencjalnych pracodawców. Warto też zwrócić uwagę na fakt, że branża informatyczna rozwija się bardzo szybko, a nowe technologie i narzędzia pojawiają się na rynku regularnie. W związku z tym, wiedza uzyskana na studiach może nie być kompletna i aktualna. Chcąc jednak rozwijać swoje umiejętności, warto oprzeć się na wiedzy, którą w podstawowym zakresie mogą przekazać wszechobecne kursy programowania. Skupiają się one często na konkretnych językach programowania lub technologiach, oferując intensywne szkolenie w stosunkowo krótkim czasie. Tym samym w porównaniu do studiów, pojawia się możliwość zaoszczędzenia czasu, dzięki skupieniu się na specyficznych umiejętnościach, które są istotne dla danego stanowiska. W tym wypadku niezbędne



okazuje się posiadanie jasno określonego celu zawodowego i technologii w jakiej chciałoby się rozwijać.

Alternatywnym rozwiązaniem jest ukończenie studiów wyższych. Należy wspomnieć, że formalne wykształcenie wciąż ma swoją wartość i może zapewnić dogłębną wiedzę teoretyczną. Uzyskanie dyplomu daje szansę na pojmowanie informatyki w znacznie szerszym kontekście niż po ukończeniu kursu programowania. Otrzymuje się kompleksową wiedzę, która pomaga w zrozumieniu szerokiego spektrum zagadnień. Ponadto program studiów oferuje różnorodne specjalizacje i kursy. Studenci mają możliwość skupienia się na swoich zainteresowaniach, takich jak na przykład sztuczna inteligencja. Tym samym studia dają swego rodzaju elastyczność, umożliwiającą dopasowanie edukacji do własnych potrzeb. Warto pamiętać także o tym, że jest to przestrzeń, w której pojawia się możliwość nawiązywania kontaktów z wykładowcami i profesjonalistami z branży.

Trudne początki studiowania

Jeśli już kandydat na informatyka podejmie decyzję o studiowaniu, musi zmierzyć się z procesem rekrutacji. W przypadku studiów zaocznych wystarczy dysponować określonym budżetem na pokrycie czynszu, które będzie różnić się w zależności od wybranej uczelni. Jeśli natomiast plan studiowania obejmuje naukę w trybie stacjonarnym, należy zmierzyć się z kontrkandydatami. Już na samym początku poziom trudności drastycznie rośnie, w związku z tym, że jest to od kilku lat, najpopularniejszy kierunek studiów. Politechnika Krakowska w roku akademickim 2022/2023 na Wydziale Mechanicznym

odnotowała 14.61 kandydata na jedno miejsce, na kierunku Informatyka Stosowana, natomiast na Wydziale Informatyki i Telekomunikacji o jeden indeks na kierunku Informatyka, walczyło średnio 13.46 kandydatów. Jest to absolutny rekord, a tym samym poprzeczka postawiona przed absolwentami szkół średnich, zawiśła bardzo wysoko.

Standardy kształcenia

Informatyka to dziedzina, która odgrywa kluczową rolę we współczesnym świecie, a studiowanie tego kierunku jest niezwykle popularne i obiecujące dla przyszłych absolwentów.

Jednak studiowanie informatyki, to nie przyszłowiowa „bułka z masłem”. Wymaga ono solidnych zdolności analitycznych, logicznego myślenia i sprawnego poruszania się w obszarach związanych z matematyką. W trakcie nauki, niektóre przedmioty mogą się okazać nie lada wyzwaniem, szczególnie dla osób, które nie mają silnych podstaw w zakresie nauk ścisłych. Początkujący mogą mieć trudności z opanowaniem algorytmów, struktur danych i języków programowania. Warto także zwrócić uwagę na szybkość rozwoju technologii komputerowych, która wymaga od studentów ciągłego aktualizowania wiedzy i umiejętności. W trakcie studiów inżynierskich można spodziewać się: 45 godzin analizy matematycznej i algebry liniowej, 60 godzin metod probabilistycznych i statystyki, 60 godzin matematyki dyskretniej oraz po 45 godzin fizyki i nauk technicznych. Na studiach licencjackich rozkład godzinowy trochę się zmienia. Próżno szukać fizyki i nauk technicznych, natomiast liczba czasu, którą należy poświęcić matematyce pozostaje taka sama. Znajomość języka angielskiego jest na tym kierunku niezbędna. Dlatego trzeba sobie zarezerwować trochę czasu na dodatkowe lekcje.

Owoce pracy i możliwości zawodowe

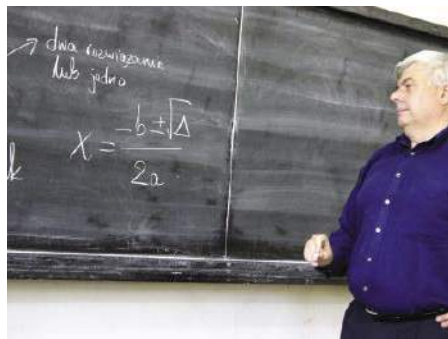
Informatyka od wielu lat jest najpopularniejszym wyborem wśród absolwentów szkół średnich. Jej renowa wpływa na wzrost konkurencji na rynku pracy. Aby wyróżnić się spośród innych kandydatów, studenci muszą wykazać się nie tylko teoretyczną wiedzą, ale również praktycznym doświadczeniem. Szybkie tempo zmian technologicznych wymaga stałego uczenia się i dostosowywania do nowych narzędzi. Mimo tych trudności, studiowanie informatyki oferuje wiele szans i perspektyw w zatrudnieniu. W dzisiejszym zglobalizowanym i z informatyzowanym świecie, popyt na specjalistów z zakresu IT jest ogromny. Praktycznie każda firma i branża korzystają z technologii informatycznych, co oznacza, że istnieje wiele

miejsz pracy dla absolwentów. Szeroki zakres możliwości zawodowych może obejmować pracę w charakterze: programisty, inżyniera programowania, analityka danych, administratora systemów, projektanta interfejsów użytkownika, specjalisty do spraw bezpieczeństwa sieciowego, inżyniera sieci. Patrząc w przyszłość, perspektywy zatrudnienia wydają się być obiecujące. Według prognoz, zapotrzebowanie na specjalistów będzie stale rosnać w najbliższych latach. Wciąż istnieje niedobór wykwalifikowanych pracowników w branży, co oznacza większe szanse na znalezienie pracy zaraz po ukończeniu studiów. Kolejnym czynnikiem wpływającym na zwiększone zapotrzebowanie na fachowców jest dynamiczny rozwój sztucznej inteligencji. Postęp w dziedzinie SI otwiera nowe możliwości i wyzwania, które wymagają specjalistycznej wiedzy i umiejętności. Pojawia się zapotrzebowanie na fachowców, którzy potrafią projektować, rozwijać i wdrażać systemy oparte na tej technologii. Firmy poszukują ekspertów do uczenia maszynowego, przetwarzania języka naturalnego, analizy danych i robotyki, aby zastosować SI w takich dziedzinach jak na przykład: medycyna, finanse, produkcja, e-commerce czy logistyka. Powstają nowe stanowiska pracy lub rozwijają się te, które do tej pory miały marginalne znaczenie. Sztuczna inteligencja jest wykorzystywana w coraz większej ilości branż, co sprawia, że specjaliści informatyki z umiejętnościami w dziedzinie SI są coraz bardziej poszukiwani. Rozwój SI prowadzi także do udoskonalania istniejących technologii informatycznych, a tym samym pojawia się zapotrzebowanie na specjalistów potrafiących stosować techniki w praktyce i optymalizować istniejące rozwiązania. Wzrost zapotrzebowania na ekspertów w zakresie SI wpływa na tworzenie nowych miejsc pracy oraz zwiększenie wynagrodzeń. Ponadto branża IT oferuje elastyczność i możliwości pracy zdalnej, co jest atrakcyjne dla wielu osób. Tym samym praca w branży informatycznej wydaje się być nie tylko pewna, ale i wyjątkowo atrakcyjna.

Wybór studiów na kierunku informatyka jest niemalże skazany na sukces, pod warunkiem, że osoba decydująca się na tę drogę posiada podstawowe zdolności analityczne, umiejętności logicznego myślenia i chęć rozwoju. Jest to jedna z najszybciej rozwijających się dziedzin nauki, a jej wpływ na nasze życie jest tak duży, że umiejętność poruszania się w jej obszarze wydaje się niezbędna. Wszystko wskazuje na to, że w kolejnych latach zainteresowanie Informatyką będzie nadal rosło, a zapotrzebowanie na fachowców nie zmaleje. Zapraszamy i polecamy. ■

Michał Pacholski

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



Felga z Bielska i powrót do przeszłości

Wrzesień to początek roku szkolnego. Patrzymy z zaciekawieniem na nowy rok, wspominamy wakacje. Chyba wszystkim kojarzą się one z wyjazdami, również samochodowymi. Posiadanie auta nie są już wielkim luksusem, jak to było za mojej młodości. A w tym roku dużo jeździłem w podgórskich okolicach.

Kontynuowałem swoje małe hobby – przypatrywałem się felgom samochodowym i szukaniu w nich interesującej geometrii. Na pierwszych trzech zdjęciach widzimy „klasykę gatunku” – proste felgi o symetrii pięciokątnej, sześciokątnej i siedmiokątnej. Te ostatnie są rzadziej spotykane, ale nie dlatego, że siedmiokąt foremny nie jest – matematycznie rzecz biorąc – konstruowalny. To może wzbudzić zdziwienie:

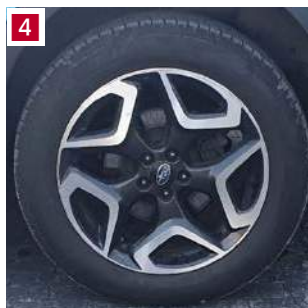
*Ruszają w dal pociągi dwa,
z miasteczka B do miasta A
i z miasta A do miasta B
i mam określić, gdzie spotkają się.
Ich prędkość znam, odległość znam
i wszystko wyznaczone mam.
Lecz to zadanie peszy mnie
I profesora oko złe.*

Z piosenki w wykonaniu Grażyny Świątły

jak to, przecież właśnie widzimy, że jest. Dla tych Czytelników, którzy byli uczniami jeszcze 30 lat temu, ta pozorna sprzeczność będzie zrozumiała. Dla pozostałych – wyjaśnię przy końcu.

Aha. Na ogół nie są to felgi, tylko plastikowe kołpaki, nałożone na nie. Ale dla prostoty będę je felgami. I oto kolega przysłał zdjęcie (4) z podpisem „Oto felga z Bielska”. Ciekawy wzór, prawda? Wszystko, co jest

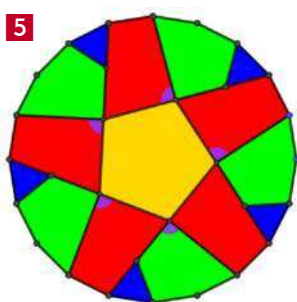




pięciokątne, jest ładne. Jaką geometrię można tu zobaczyć? Zacząłem od przerobienia tego deseni na bardziej geometryczny – w sensie szkolnej geometrii, gdzie mamy jednak proste figury. Wyszło coś takiego – jak na **rysunku 5**.

W szkole uczymy (się albo innych), jak rozwiązywać zadania. Wiele razy dawałem i uczniom, i nauczycielom, zagadnienie odwrotne: *ułożyć* zadanie na jakiś temat. Zawsze szło opornie. Nic dziwnego: O dostrzeganiu problemów, stawianiu zagadnień i inteligentnemu przypatrywaniu się światu nie mówi się w szkole prawie wcale. „Czy umiecie się dziwić?” – taki był tytuł książki wydanej dawno temu przez redakcję miesięcznika naukowego „Delta”.

Na „feldze z Bielska” zobaczyłem pięć dziobów ptasich. Przerobiłem ornament na geometrycznie prostszy. Oj, czy naprawdę prosty – to się okaże. Na **rysunku 5**



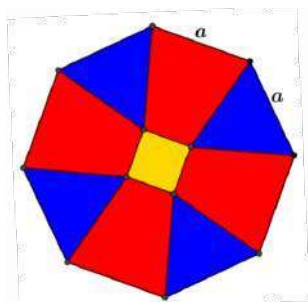
dziobom odpowiadają czerwone czworokąty z przyczepionymi niebieskimi trójkątami. Budowę figury zacząłem właśnie od narysowania pięciu trójkątów równobocznych w dwudziestokącie. Długość boku trójkątów oznaczyłem przez a . I oto główne **Zadanie 1**. Przyjmując, że bok niebieskiego trójkąta równobocznego ma długość a , obliczyć wszystkie kąty i wszystkie rozmiary figur tam widocznych.

Ponieważ zrobienie rysunku zajęło mi kilkanaście minut, to pomyślałem sobie, że obliczenia potrąją może trzy razy tyle, może cztery, może aż godzinę. Hm, nie powiem, ile siedziałem nad tym zadaniem. Zobaczymy. Już rachunek kątów zadziwia. Takich kątów nie spotyka się na szkolnych lekcjach. Idźmy systematycznie, bo się przyda. Przypomnijmy sobie lekcję geometrii.

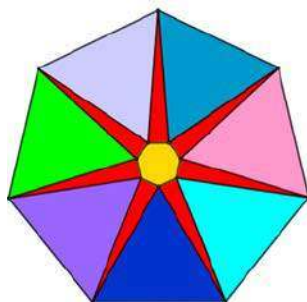
Nietrudno wyznaczyć miarę kąta wewnętrznego n -kąta foremnego. Opuszczę naprawdę proste wyprowadzenie, podam tylko końcowy wzór. Kąt w n -kącie foremnym ma miarę

$$180^\circ \left(1 - \frac{2}{n}\right)$$

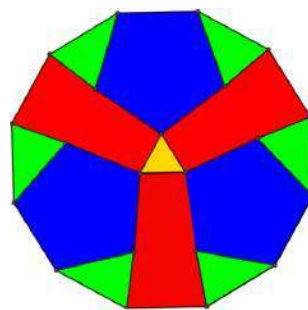
To szkolny materiał. Sprawdźmy tylko, że się zgadza. Dla kwadratu ($n=4$) istotnie dostajemy 90° , a dla trójkąta równobocznego ($n=3$) też wyjdzie, jak powinno, 60° . Pięciokąt foremny ma 108° , sześciokąt 120° ,



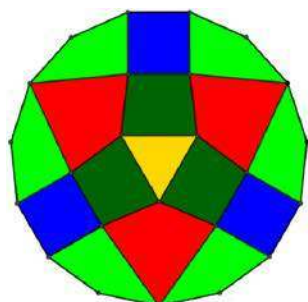
6. Felga F(8,3,4)



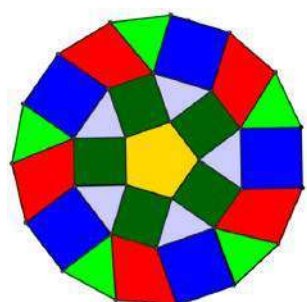
7. Felga F(7,3,7)



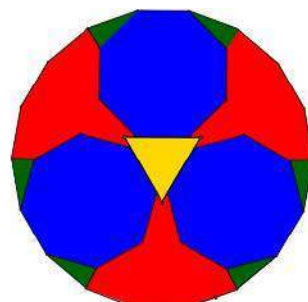
8. F(12, 5, 3)



9. F(15, 4, 3)

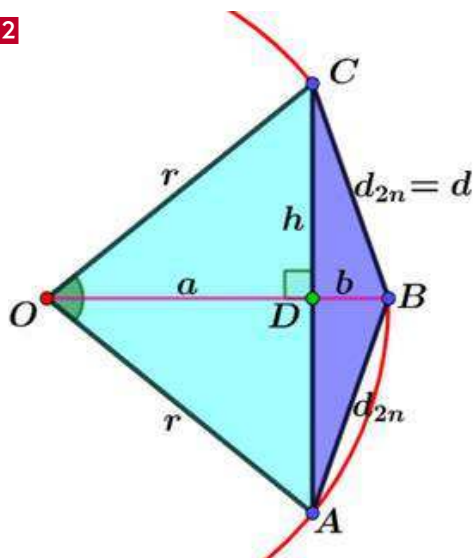


10. F(15, 4, 5)

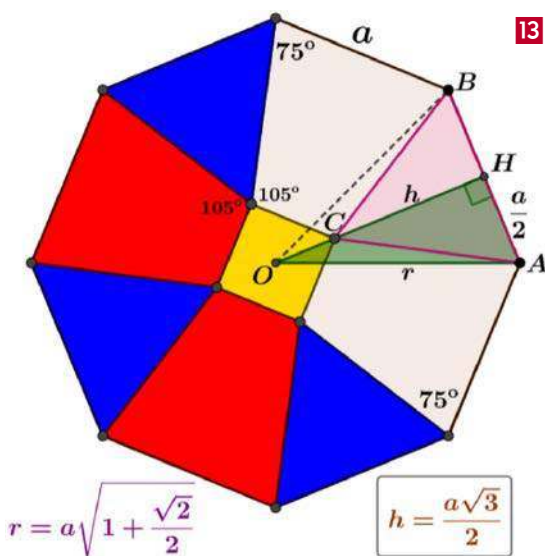


11. F(18, 6, 3)

12



13



dwunastokąt 150° , a dla dwudziestokąta otrzymujemy 162° . W mojej prywatnej, nie całkiem poważnej numeryce kojarzę to z autobusem linii 162 w Warszawie – jeździłem nim na zajęcia uniwersyteckie. Ojej, przepraszam za dygresję.

Przypominam, że jedyną daną liczbową w tym zadaniu jest długość boku trójkąta równobocznego. Można powiedzieć, że są jeszcze dwie ważne inne liczby: $n=20$ (dwudziestokąt), że zaczynamy od trójkątów (a więc $m=3$) i że jest ich pięć ($k=5$). Dlatego nazywałem tę felgę $F(20, 3, 5)$. Aby wprawić się w dobry nastrój (i dać sobie serię zadań), narysowałem kilka innych felg. Prawda, że ładne? Na **rysunku 7** mamy niezbyt ciekawą konfigurację sześciu trójkątów w sześciokącie. Na **rysunku 8** mamy ośmiokąt, trójkąty równoboczne w liczbie 4, a więc jest to $F(8, 3, 4)$. I tak dalej.

Zadanie 2. Narysuj $F(15, 3, 5)$ i $F(15, 5, 3)$.

Potrzebna nam teraz będzie ciekawa i historycznie ważna zależność między długościami boków n -kąta foremnego i $2n$ -kąta. Spójrzmy na **rysunek 12**.

Odcinek AC symbolizuje bok n -kąta, a odcinki AB i BC to boki wielokąta o dwa razy większej liczbie boków – tak, że

$$h = \frac{d_n}{2}$$

Wielokąty te są wpisane w okrąg o promieniu r . Chcemy wyznaczyć związek między h i d . Możemy założyć, że promień r jest równy 1, bo związek, którego szukamy, nie może zależeć od promienia okręgu opisanego na tych wielokątach.

Trójkąt ODC (12) jest prostokątny, a więc $OD^2 + DC^2 = OC^2$, czyli

$$a = \sqrt{1 - h^2}$$

zatem

$$b = 1 - \sqrt{1 - h^2}$$

Trójkąt BDC też jest prostokątny, a więc $BD^2 + DC^2 = BC^2$, czyli

$$d^2 = (1 - \sqrt{1 - h^2})^2 + h^2 = 1 + (1 - h^2) - 2\sqrt{1 - h^2} = 2 - 2\sqrt{1 - h^2}$$

Pamiętamy, że

$$h = \frac{d_n}{2}$$

Stąd nasz końcowy wzór

$$d_{2n} = \sqrt{2 - \sqrt{4 - d_n^2}}$$

Przyjrzyjmy mu się i wyprowadźmy kilka zależności. Bok kwadratu opisanego na okręgu o promieniu 1 ma długość $\sqrt{2}$. Dla wielokątów o 8, 16, 32 i 64, ... bokach otrzymamy estetycznie wyglądające wzory

$$d_8 = \sqrt{2 - \sqrt{2}}$$

$$d_{16} = \sqrt{2 - \sqrt{2 + \sqrt{2}}}$$

$$d_{32} = \sqrt{2 - \sqrt{2 + \sqrt{2 + \sqrt{2}}}}$$

$$d_{64} = \sqrt{2 - \sqrt{2 + \sqrt{2 + \sqrt{2 + \sqrt{2}}}}}$$

i tak dalej.

Przypomnę, że niewymierna liczba π wyraża stosunek obwodu koła do jego średnicy, a to oznaczenie pochodzi od pierwszej litery słowa

Otóż

$$\sin 27^\circ = \frac{1}{2} \sqrt{\frac{1}{2}(4 - \sqrt{2(5 - \sqrt{5})})}$$

Jak to obliczyłem? To trochę niewłaściwe pytanie – o tym przy końcu artykułu.

Mamy zatem

$$AG = \sqrt{3 + \sqrt{5} + \sqrt{\frac{1}{2}(25 + 11\sqrt{5})} \cdot \sqrt{\frac{1}{2}(4 - \sqrt{2(5 - \sqrt{5})})}}$$

co po nieco zawiłych rachunkach znacznie się upraszcza:

$$AG = a(1 + \sqrt{\frac{1}{2}(5 + \sqrt{5})})$$

Można do tego wzoru dojść prościej, wykorzystując elementy geometrii analitycznej. Ale chcę trzymać się jak najbliżej matematyki szkolnej. Trójkąt *GDA* wydaje się prostokątny, ale nie jest, choć różnica jest niewielka. Kąt *DGA* ma 48° , zatem kąt *GAD* to 42° . Zatem $AD = AG \sin 48^\circ$. Potrzebujemy wartości $\sin 48^\circ$. Oto i ona

$$\sin 48^\circ = \frac{1}{4} \sqrt{7 - \sqrt{5} + \sqrt{6(5 - \sqrt{5})}}$$

Znając *AD* i kąty trójkąta *DOA*, możemy z twierdzenia sinusów obliczyć $OD = r$. Potrzebne są jeszcze wartości sinusów kątów 21° i 36° . Otóż mamy:

$$\sin 21^\circ = \frac{1}{2\sqrt{2}} \sqrt{4 - \sqrt{7 - \sqrt{5}} + \sqrt{6(5 - \sqrt{5})}}$$

$$\sin 36^\circ = \sqrt{\frac{5}{8} - \frac{\sqrt{5}}{8}}$$

Teraz już „łatwo”:

$$\frac{AD}{\sin 36^\circ} = \frac{r}{\sin 21^\circ}$$

czyli

$$r = \frac{AD \cdot \sin 21^\circ}{\sin 36^\circ} = \dots$$

I tu „poległem”. Próbowałem napisać końcowy wzór, wyrażający *r* w zależności od *a*. To tylko kwestia cierpliwości i kaligrafii, bo przecież „wszystko wyznaczone mam” (jak w piosence, która jest mottem artykułu). Ale czy to ważne? Dla praktycznej konstrukcji nie bardzo.

Uff. Ciekawe, czy ktoś doczytał do końca. Mam jakiś dziwny sentyment do skomplikowanych, ale ładnie wyglądających wzorów. Ale w matematyce (przecież i w życiu) ważne jest też spojrzenie ogólne, z jakiejś perspektywy, dystansu na wszystko, co robimy. Spójrzmy zatem i na felgę z Bielska (mowa o Bielsku-Białej na Śląsku; przypomnę, że w czasach zaborów Bielsko należało do Prus, a Biała do Austrii, stąd dzisiejsza podwójna nazwa miasta). Zobaczmy, że miary wszystkich kątów są podzielne przez 3. To dlatego umiemy podać dokładne wartości ich funkcji trygonometrycznych. Wprawny „trygonometrysta” zrozumie, dlaczego. Na potrzeby zadania o feldzie nie obliczałem tych wartości, jak bym to czynił 60 lat temu w szkole – kliknąłem tylko raz klawiaturę. Co więcej, wszystkie takie kąty są konstruowalne za pomocą cyrkla i linijki – a inne kąty o całkowitej liczbie stopni już nie. Dlatego możemy skonstruować pięciokąt, sześciokąt, piętnastokąt, 99-kąt i 102-kąt foremny, a siedmiokąt (foremny) już nie. Oczywiście możemy go narysować dowolnie dokładnie, odmierzając starannie kąty kątomierzem, ale nie jest to konstrukcja w sensie klasycznej geometrii. Łatwo dają sobie z tym radę najprostsze programy graficzne (w tym i popularna w szkole Geogebra). Karl Friedrich Gauss (1777–1855, zwany potem „księciem matematyków”) odkrył, jak skonstruować za pomocą cyrkla i linijki 17-kąt foremny i podobno tak się zachwycił swoim odkryciem, że poświęcił się matematyce. Miał wtedy 19 lat. Konstrukcja jest zresztą pomysłowa i oparta na specjalnym „porządku Gaussa” liczb od 0 do 16.

Zadania konstrukcyjne (na budowanie figur za pomocą cyrkla i linijki) były obecne od początku istnienia gimnazjów i liceów i bardzo ważne w szkolnej matematyce. Pochodziły zresztą aż z antyku. Miały wiele zalet: uczyły myślenia i algorytmów, a do pewnego stopnia i sprawności rachunkowej, co było chyba widać. Z grubsza rzecz biorąc, konstruowalne jest to, co da się wyrazić za pomocą pierwiastków kwadratowych. Dzisiaj, jeżeli nawet uczymy myślenia w szkole (a są co do tego różne zdania), to na innych przykładach i innych zadaniach. Problemy konstrukcji za pomocą cyrkla i linijki odeszły w zapomnienie. Trochę żal, jak wszystkiego co było dobre, ale zostało przykryte czymś świeższym. ■

Michał Szurek



dr inż. Jan Sobótka
– nauczyciel akademicki,
licencjonowany instruktor
i sędzia szachowy

Benjamin Franklin (1706-1790) był nie tylko wielkim amerykańskim mężem stanu i jednym z twórców niepodległości Stanów Zjednoczonych, ale również malował obrazy, prowadził działalność publicystyczną, zajmował się filozofią i dokonał wielu odkryć naukowych. W szachy grać zaczął dopiero na początku lat 30. będąc już po dwudziestce, ale szybko stał się wielkim miłośnikiem królewskiej gry. Franklin widział w szachach zwierciadło życia i głosił ich zalety. Jest autorem świetnego eseju „Morals of chess” („Moralność gry szachowej”), który po raz pierwszy opublikował w Londynie w 1779 roku. Studium to zostało w grudniu 1786 roku reprodukowane w Filadelfijskim Columbian Magazine. Do wielu korzyści jakie niesie ze sobą gra w szachy Franklin zaliczał między innymi rozwój umiejętności obserwowania, rozważań, przezorności i zapobiegliwości. Uważał, że sposób w jaki postępujemy w czasie gry, można przełożyć na inne aspekty życia.

Szachowa dyplomacja Beniamina Franklina

Benjamin Franklin urodził się w Bostonie, w stanie Massachusetts jako jedno z siedemnastu dzieci w rodzinie biednego wytwórcy świec (1).

W wieku 12 lat zaczął pracować w drukarni starszego brata (2). Konflikty z bratem spowodowały, że w wieku siedemnastu lat Benjamin uciekł do Filadelfii. Początkowo pracował jako drukarz



1. Josiah Franklin, wytwórca świec i ojciec Beniamina Franklina, źródło: <https://tiny.pl/cc9qm>



2. Benjamin Franklin w warsztacie swojego brata w Bostonie, gdzie pracował w latach 1718-1723, źródło: <https://tiny.pl/cc9qm>



3. Portret Benjamina Franklina autorstwa Davida Martina, 1767, źródło: <https://tiny.pl/cc9qt>

a potem jako wydawca. Od 1730 zaczął wydawać „The Pennsylvania Gazette” (głos kolonialnego sprzeciwu wobec brytyjskich rządów kolonialnych), a w latach 1732–1757 poczytny zbiór maksym i porad, zatytułowany „Poor Richard’s Almanack” (Almanach

biednego Ryszarda), który oprócz korzyści materialnych, przyniósł mu także popularność.

Polityk

W 1736 został członkiem zgromadzenia Pensylwanii a w 1751 wybrano go do rady miejskiej Filadelfii. W 1753 Franklin został zastępcą generalnego naczelnika poczty, odpowiedzialnym za pocztę we wszystkich koloniach północnych. W 1757 wyjechał do Anglii by reprezentować swój stan na procesach przeciwko właścicielom kolonii (3). Z wyjątkiem dwuletniego powrotu do Filadelfii w latach 1762–64, Franklin spędził 18 lat mieszkając w Londynie. Do Ameryki wrócił w 1775 roku, tuż przed wybuchem wojny o niepodległość Stanów Zjednoczonych, został poczmistrzem generalnym Ameryki Północnej. Benjamin Franklin był współautorem i sygnatariuszem (Komitet Pięciu) uchwalonej 4 lipca 1776 Deklaracji Niepodległości Stanów Zjednoczonych (4). Jeszcze w tym samym roku ponownie wyruszył do Europy. Tym razem dotarł do Paryża, gdzie wraz z współpracownikami doprowadził do sojuszu Stanów Zjednoczonych z Francją. 6 II 1778 podpisał w imieniu Stanów Zjednoczonych traktat handlowy z Francją. W latach 1778–85 był posłem pełnomocnym Stanów Zjednoczonych we Francji. Jako członek komisji (wspólnie z J. Adamsem i J. Jayem) negocjował warunki traktatu pokojowego z Wielką Brytanią, kończącego wojnę o niepodległość Stanów Zjednoczonych, podpisanego 3 września

4. Pięcioosobowy komitet redakcyjny Deklaracji Niepodległości prezentujący swoje prace Kongresowi, źródło: <https://tiny.pl/cc9q7>





5. Benjamin Franklin z pomocą 21-letniego syna prowadzi badania nad elektrycznością,
źródło: <https://tiny.pl/cc9qr>



6. Benjamin Franklin gra w szachy z Lady Howe,
źródło: <https://tiny.pl/cc9qw>

1783 w Paryżu. W latach 1785–87 uczestniczył w pracach nad konstytucją Stanów Zjednoczonych, które ostateczny tekst został przyjęty we wrześniu 1787 r.

Naukowiec, filozof i wynalazca

Równolegle z karierą polityczną prowadził pionierskie badania naukowe nad elektrycznością. Stwierdził, że ciała naelektryzowane jednakowo odpychają się, zaś naelektryzowane różnoimiennie – przyciągają się.

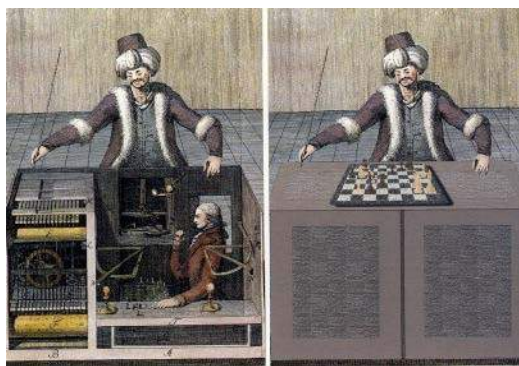
W roku 1752 złapał prąd z błyskawicy do butelki lejdeckiej po wilgotnym sznurze latawca i dowiódł, że w piorunach jest elektryczność (5). Wynalazł piorunochron, który do dziś ochrania budynki w czasie burzy. Wymyślił i wykonał tzw. dzwoneczki Franklina, które odpowiednio wcześniej sygnalizowały nadejście burzy. Wykorzystał do tego wynalazek Andrew Gordona z Niemiec z 1742 roku, tzw. lightning bells, który przekształcał energię elektryczną w mechaniczną. Uziemiając jeden z dzwonków i podłączając drugi do piorunochronu, Franklin stworzył konstrukcję ostrzegającą przed nadciągającą burzą. Jego działalność naukowa zaowocowała przyznaniem tytułów doktor honoris causa renomowanych uczelni: Uniwersytetu Yale, Uniwersytetu Harvarda (1753), College of William & Mary (1756), Uniwersytetu w St Andrews (1759) i Uniwersytetu Oksfordzkiego (1762).

Franklin zajmował się również wymyślaniem udogodnień służących w codziennym życiu, jak bezdymny piec podwyższający komfort mieszkania, fotel bujany, cewnik urologiczny, płetwy pływackie, licznik kilometrów i okulary dwuogniskowe (w roku 1785). Górna połowa szkieł służyła do widzenia dali, dolna – do czytania. Franklina męczyła konieczność noszenia kilku par okularów. Jego wszechstronność była zadziwiająca. Był organizatorem

pierwszej w Filadelfii wypożyczalni książek, zorganizował również ochotniczą straż pożarną oraz Amerykańskie Towarzystwo Filozoficzne, które istnieje do dziś. Powołał do życia i był prezesem The Academy and College of Philadelphia (później University of Pennsylvania). Oprócz prac związanych z elektrycznością napisał również traktat o pieniądzu, książkę o grze w szachy. Zajmowały go problemy związane z przyczynami trzęsień ziemi, był zwolennikiem zniesienia niewolnictwa. W roku 1775, usiłując zrozumieć, dlaczego statki z Anglii do Ameryki płyną dwa tygodnie dłużej niż odwrotnie, odkrył i opisał prąd zatokowy (Golfsztrom).

Szachista

Franklin grał w szachy od co najmniej 1732 roku. Nie zachowały się żadne zapisy partii rozgrywanych przez Franklina, ale można przypuszczać, że był ponadprzeciętnym graczem, który jednak nie osiągnął najwyższego poziomu. Wspominał o szachach w wielu swoich listach. Dostrzegał potencjał wykorzystania szachów do celów dyplomatycznych. W grudniu 1774 roku przyjechał do Londynu jako poseł negocjujący pojednanie. Tutaj był stałym bywalcem w kawiarni Old Slaughter's, gdzie grał w szachy i spotykał się m.in. z grupą artystów, którzy utworzyli Royal Society of Arts. W listopadzie 1774 roku jeden z członków stowarzyszenia przesłał wiadomość od lady Caroline Howe z zaproszeniem na partię szachów. Była to siostra lorda Richarda Howe – admirała i Sir Williama Howe – generała, który od 10 października 1775 do 20 maja 1778 był naczelnym dowódcą wojsk brytyjskich w Ameryce (Commander-in-Chief of the British Army in America). Przy okazji rozgrywania partii szachowych w rezydencji lady Caroline Howe doszło do szeregu nieoficjalnych



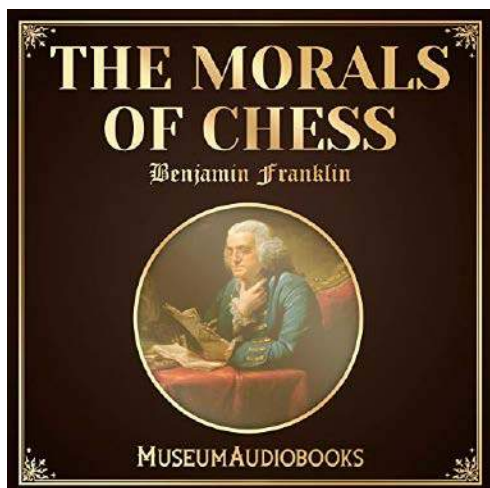
7. Rysunek mechanicznego Turka, autorstwa Karla Gottlieba von Windisch, 1783 r., źródło: Wikimedia Commons

rozmów Franklina z brytyjskim rządem, w których lord Howe pełnił rolę pośrednika. Niestety Brytyjczycy zajmowali nieprzejezdne stanowiska, doszło do wybuch wojny i nie udało się rozwiązać narastającego kryzysu przy okazji partii szachów rozgrywanych w salonie lady Caroline Howe.

W latach 70. i 80. XVIII wieku, Franklin będąc ambasadorem we Francji, doskonalił swoje umiejętności gry w szachy odwiedzając często Café de la Regence, w tamtym czasie miejsce spotkań najlepszych szachistów na świecie. Podczas swojego pobytu we Francji grał przeciwko słynnemu automatowi „Mechaniczny Turek” (7), choć wynik meczu jest nieznany.

Moralność gry szachowej

Pisanie eseju „Morals of chess” Franklin rozpoczął już około 1732 roku, ale opublikował go dopiero w 1779 roku



8. Benjamin Franklin – Moralność gry w szachy, źródło: <https://tiny.pl/cc9qc>

(8). Po przedstawieniu w prologu historii szachów porównał szachy do życia, opisał zestaw zasad moralnych, których powinien przestrzegać szachista, w tym nie oszukiwać i nie przeszkadzać przeciwnikowi. Esej jest jednym z pierwszych tekstów o szachach, jaki ukazał się w Stanach Zjednoczonych został potem przełożony na języki rosyjski, francuski, niemiecki. Franklin podkreślał w niej, że szachów nie można ograniczyć do rozrywki i zwracał uwagę na ich walory wychowawcze, kształtowanie pozytywnych cech osobowości (zdolności przewidywania i planowania, rozważ, odpowiedzialności za decyzje), a szachy porównał do życia: „Życie jest swego rodzaju grą w szachy, w której mamy często możliwość wygrania i walki z współzawodnikami i przeciwnikami, w której jest duża różnorodność dobrych i złych wydarzeń, będących w jakimś stopniu rezultatem rozsądku lub jego braku”.

Podjęcie Franklina do gry wyraźnie różniło się od większości jego poprzedników, których jedynym celem było zwycięstwo. Jednak dla Franklina gra w szachy nie była zwykłą rozrywką, ale sportem odzwierciedlającym samo życie – ponieważ życie jest rodzajem szachów, w których często zdobywamy punkty i walczymy z konkurentami lub przeciwnikami – co wymaga wykorzystania wszystkich najwspanialszych cech umysłowych i moralnych, do jakich zdolny jest człowiek. Maksymy z „Moralności gry szachowej” były szeroko cytowane w kraju i za granicą.

Franklin ułożył także własny system wartości, oparty na 13 Cnotach, które opisał w autobiografii *Żywot własny*:

1. *Umiar* – nie jedz do otyłości, nie pij do podniecenia,
2. *Milczenie* – mów tylko to, co może przynieść pożytek innym lub tobie; unikaj próżnej rozmowy,
3. *Ład* – niech wszystkie twoje rzeczy mają swoje miejsce; niech wszystkie twoje sprawy mają swój czas,
4. *Postanowienie* – postanów sobie czynić, co powinien; czyń bez zawodu to, coś postanowił sobie,
5. *Oszczędność* – nie czyń wydatków, chyba że dla dobra innych lub swego, to znaczy nic nie marnotraw,
6. *Pracowitość* – nie trać nigdy czasu, bądź zawsze zajęty czymś pożytecznym, unikaj wszelkich niepotrzebnych działań,
7. *Szczerość* – nie uciekaj się do krzywdzącego oszustwa; myśl niewinnie i sprawiedliwie, a kiedy mówisz, mów podobnie,
8. *Sprawiedliwość* – nie krzywdź nikogo, wyrządzając mu zło lub pozbawiając pożytków, jakie mu się od ciebie należą,
9. *Powściągliwość* – unikaj krańcowości, nie odczuwaj krzywd w takim stopniu, jak twoim zdaniem na to zasługują,

10. *Czystość* – nie pozwalaj na żadną nieczystość ciała, odzienia czy mieszkania,
11. *Spokój* – niech Ci nie zamącają umysłu drobności lub też zdarzenia pospolite i nieuchronne,
12. *Surowość cielesna* – rzadko używaj płci, tylko dla zdrowia albo potomstwa, nigdy do oziębnia, osłabienia lub szkody w twoim czy cudzym spokoju i opinii.
13. *Pokora* – naśladowuj Jezusa i Sokratesa

Franklin nie stworzył tej listy, by być doskonałym gentlemanem, chciał być człowiekiem ciężkiej pracy, wielkie znaczenie miały dla niego honor i zaufanie innych osób, a wysoka kultura osobista zapewniała mu prominentnych przyjaciół.

Efekt Benjamina Franklina

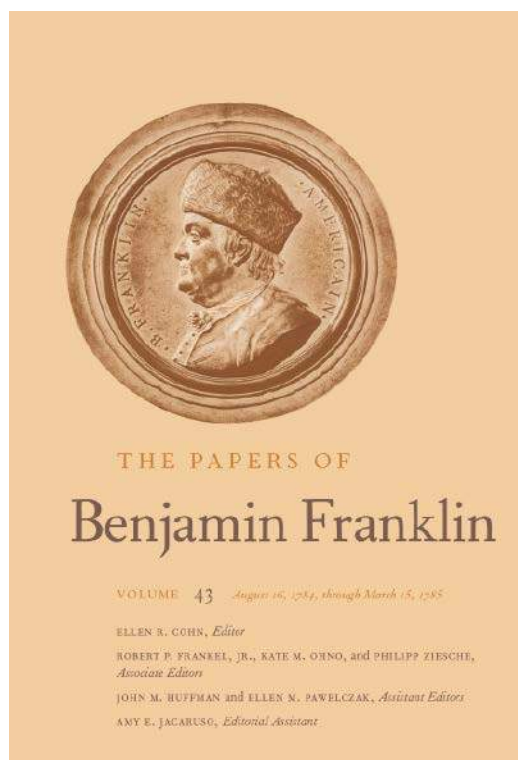
Nie tyle kochamy ludzi za dobrodziejstwa, które oni nam wyświadczyli, ile za dobrodziejstwa, któreśmy im wyświadczyli.

Efekt wyświadczonej przysługi, znany również jako efekt Benjamina Franklina to tendencja do zmiany postawy na bardziej przyjazną w stosunku do osób, dla których zrobiło się coś dobrego. Zrobienie czegoś dla drugiej osoby, szczególnie wtedy, gdy jej nie lubimy lub jest nam obojętna, może wywołać stan napięcia spowodowany dysonansem pomiędzy postawą „nie lubię tej osoby” a nieodpowiadającym jej zachowaniem „pomogłem jej”. W celu zredukowania dysonansu i obniżenia napięcia następuje zmiana postawy z „nie lubię tej osoby i jej pomogłem” na „pomogłem tej osobie, bo ją lubię”.

Benjamin Franklin opisał ten efekt w autobiografii *Żywot własny* taktykę, którą obrał w stosunku do swego przeciwnika w wyborach na pisarza Zgromadzenia Pensylwanii w 1737 roku: „Nie zamierzałem wszakże ubiegać się o jego życzliwość drogą okazywania mu służalczych względów, ale odczekawszy trochę obrałem inną metodę. Wiedząc, że ma on w swojej bibliotece pewną rzadką a interesującą książkę, wyśtoso wałem doń list, w którym wyraziłem pragnienie przejrzenia książki oraz prośbę, aby zechciał łaskawie pożyczyć mi ją na kilka dni. Przysłał mi ją natychmiast, ja zaś zwróciłem po niespełna tygodniu z nowym listem, w którym wypowiedziałem moją gorącą wdzięczność. Kiedy niebawem spotkaliśmy się w Izbie, przemówił do mnie (czego nigdy dotąd nie robił) i wdał się w nader uprzejmą rozmowę. Odtąd też okazywał zawsze gotowość dopomagania mi przy każdej sposobności i wkrótce bardzo zaprzyjaźniliśmy się, a przyjaźń ta trwała aż do jego śmierci. Wypadek ten potwierdził raz jeszcze prawdziwość starej maksymy której dawno już nauczyłem się, a która głosi,



9. Benjamin Franklin, 1778,
źródło: <https://tiny.pl/cc9qf>



10. The Papers of Benjamin Franklin,
źródło: <https://tiny.pl/cc9qf>



11. Banknot studolarowy z podobizną Benjamina Franklina

że: „Ten, kto raz ci wyrządził przysługę, będzie bardziej skory oddać ci nową niż ktoś, komuś ty usłużył”.

Przez całe swoje dorosłe życie Franklin (9) gromadził swoją korespondencję, dokumenty i inne pisma, które dziś obejmują około 30 000 pozycji (10).

Był współautorem i jednym z sygnatariuszy Deklaracji Niepodległości (1776) i Konstytucji Stanów Zjednoczonych (1787). Został pierwszym reprezentantem rządu USA we Francji (1776-1785). Podpisał też petycję o zniesieniu niewolnictwa. Chociaż nie był nigdy prezydentem USA, jego podobizna widnieje na banknocie studolarowym (11). Aż trudno uwierzyć, że zaczynał jako biedny czeladnik drukarza.

Franklin osiągnął sukces ciężką pracą, poświęceniem i szlifowaniem własnego charakteru. Pożyczał fundusze na swoje projekty i robił wszystko, by w terminie oddać je z procentem i jeszcze na nich zarobić. Starał się żyć manifestując swoją pracowitość, cnotliwość i mądrość.

Stworzył listę wartości, co wieczór robił rachunek sumienia i w specjalnym zeszycie zapisywał swoje sukcesy i porażki w opanowywaniu tych cnót. Posługiwał się również spisanim planem dnia (12). ■

RANEK
pytanie:
co dobrego
dzisiaj
uczynię?

5 } Wstanie, umycie się i odmówienie „Wszec-
moena Dobroci”. Rozważenie spraw dnia i po-
6 } wzięcie odpowiednich postanowień na cały
7 } dzień. Przestudiowanie tabelki. Śniadanie.

8 }
9 } Praca.
10 }
11 }

PÓŁDNIEM

12 } Lekcja lub przejrzenie rachunków i
1 } obiad.

2 }
3 } Praca.
4 }
5 }

WIECZÓR
pytanie:
co dobrego
dzisiaj
spełniłem?

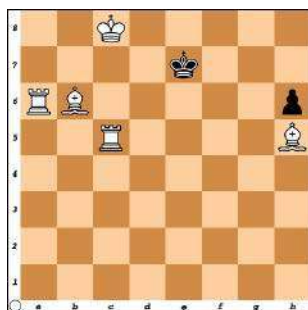
6 } Ułożenie rzeczy na właściwych miejscach. Ko-
7 } lacja. Muzyka lub rozrywka, albo rozmowa.
8 } Przebicie dnia.
9 }

10 }
11 }
12 } Sen.
1 }
2 }
3 }
4 }

NOC

12. Plan dnia Benjamina Franklina, źródło: <https://tiny.pl/cc9q5>

Zadania do samodzielnego rozwiązania



Zadanie 1. 13. Víctor Baja
Mat w 2 posunięciach



Zadanie 2. 14. G. Heathcote,
„American Chess Bulletin”,
1911. Mat w 2 posunięciach

Rozwiązanie zadań z MT 8/2023

Zadanie 1
Wilhelm Steinitz – Herbert Trenchard,
Nowy Jork, 1887
Mat w 2 posunięciach
Rozwiązanie: 1.Se7

Zadanie 2
Niels Høeg, 1905
Mat w 3 posunięciach
Rozwiązanie: 1.f7
1...e4 2.f8=H lub 2.f8=W, 1...e:d4 2.f8=G,
1...e:f4 2.f8=W, 1...Kd6 2.f8=H+, 1...Kf6
2.f8=S



1. Piloty zdalnego sterowania

Piloty do TV i innych urządzeń Czy zmienimy władzę?

Czy pilot do telewizora (1), już wkrótce zniknie z naszego domu i naszego życia? To prawdopodobne. Należy przypuszczać, że w niebyt pójdą za nim inne piloty, do wież stereo, do klimatyzacji i inne. Już dawno wiadomo, że dobie smartfonów nie potrzeba dodatkowych urządzeń do sterowania. Potem okazało się, że nie potrzeba do tego nawet smartfona – to kolejny etap.

W współpracy z aplikacjami, smartfony ale również tablety, stopniowo przejmują zadania pilotów sterujących domową rozrywką, domowym oświetleniem, klimatyzacją i systemami bezpieczeństwa. Może to być jednak jedynie faza przejściowa, gdyż, jak się okazuje, urządzenia w coraz większym stopniu potrafią w słuchać i reagować bezpośrednio na ludzki głos a nawet na ruchy ciała, gesty, mimikę.

Telefony komórkowe w roli pilotów (2) mają pewną poważną wadę – są urządzeniami osobistymi, a telewizja jest nadal przeważnie aktywnością grupową. Dlatego uważa się, że ostatecznie rolę przestarzałych pilotów przejmą urządzenia głosowe, takiego rodzaju jak Alexa. Inteligentne telewizory od dawna obsługują sterowanie głosem, choć wydaje się, że ten rodzaj interakcji wciąż jeszcze nie jest bardzo popularny.



2. Smartfon jako pilot

Kwestia przyzwyczajenia

Poszukuje się rozwiązań, które pozwoliłyby posługiwać się sterowaniem głosowym w sposób jak najbardziej naturalny. Przykładem takich wysiłków jest brytyjska platforma smart TV YouView, która pracuje nad integracją aktywowanego głosem wirtualnego asystenta Amazon Alexa ze swoimi systemami. W trakcie różnego rodzaju eksperymentów YouView odkryło m.in., że proces wyszukiwania w interfejsie telewizyjnym zmienia się diametralnie, gdy ludziom zabierze się pilot. Gdy mają do dyspozycji obie opcje, to głos jest używany tylko eksplorowania nieznanych zasobów, np. bibliotek filmowych, zaś pilot pozostaje narzędziem do zarządzania rutynową aktywnością telewizyjną.

Specjaliści YouView zauważyli także, że przyjęcie wyszukiwania głosowego w telewizji nie jest u widzów natychmiastowe, wymaga przyzwyczajenia. Kiedy jednak widzowie przechodzą podejście głosowe, mają tendencję do przełączania się na taki tryb na stałe, tylko sporadycznie używając pilota.

Nie wszyscy zgadzają się, że piloty znikną. Peter Docherty z firmy ThinkAnalytics, zajmującej się dostarczaniem spersonalizowanych rekomendacji treści, twierdził w rozmowie z „Forbesem” w 2018 r., że pilot nie stanie się przestarzały, ponieważ różni ludzie lubią korzystać z technologii na różne sposoby i zawsze tak będzie. „Nie ma jednego najlepszego sposobu na zrobienie wszystkiego”, mówi.

„Pilot pozwala na skakanie po kanałach i sterowanie funkcjami za pomocą przycisków skrótów. W końcu do nowoczesnych pilotów można nawet mówić. Głos, tablety i telefony komórkowe oferują różne korzyści. Ale nie chciałbyś robić wszystkiego za pomocą głosu, nawet gdybyś mógł. A konieczność wzięcia telefonu

komórkowego, odblokowania go i uruchomienia aplikacji telewizyjnej tylko po to, aby zmienić kanał lub dostosować głośność, to wcale nie jest najlepsze doświadczenie użytkownika”.

Odchodzą, ale może nie całkiem

Niemniej jest obecnie sporo przesłanek wskazujących na, jeśli nie całkowite pożegnanie z pilotem, to spadek znaczenia tego rodzaju urządzeń. Zasilani coraz bardziej wyrafinowanymi modelami AI asystenci głosowi, tacy jak Amazon Alexa, Google Assistant i Apple Siri robią coraz większe postępy i mogą sterować telewizorami oraz innymi urządzeniami za pomocą poleceń głosowych. Użytkownicy mogą wydawać polecenia w rodzaju „włącz telewizor”, „przełącz kanał” lub „podgłośnij” zamiast używać pilota. Większość nowoczesnych telewizorów i dekodów oferuje aplikacje na smartfony i tablety do zdalnego sterowania tymi urządzeniami. Użytkownicy mogą włączyć/wyłączyć telewizor, zmieniać kanały i regulować głośność bezpośrednio ze swoich telefonów. Niektóre nowoczesne telewizory obsługują sterowanie poprzez gesty dłoni. Użytkownicy mogą zmieniać ustawienia lub wykonywać podstawowe funkcje telewizora wykonując gesty takie jak przesunięcie w górę/w dół dłoni.

Coraz więcej telewizorów oferuje funkcje Smart TV z dostępem do internetu. Użytkownicy mogą przeglądać treści wideo, obrazki i informacje bezpośrednio na ekranie telewizora za pomocą jego menu i przycisków sterowania. Nie potrzebują pilota do poruszania się po opcjach Smart TV. Nowoczesne telewizory stają się częścią szerszych ekosystemów automatyki domowej, które można kontrolować z jednej aplikacji. Na przykład właściciele produktów SmartThings lub Nest mogą sterować telewizorem za pomocą tych

samych aplikacji, których używają do kontroli ogrzewania, oświetlenia, zabezpieczeń itp. W tym modelu pilot schodzi na dalszy plan.

Niektórzy dostawcy usług wideo streamingowych oferują własne moduły i przystawki, które podłącza się do telewizora. Służą one do przeglądania biblioteki treści dostawcy, na przykład Roku lub Fire TV firmy Amazon. Takie moduły zazwyczaj mają własne przyciski do sterowania i przeglądania, więc pilot do telewizora nie jest potrzebny.

Widać wyraźnie, że tradycyjne piloty telewizyjne powoli odchodzą w przeszłość wraz z pojawianiem się nowszych i wygodniejszych technologii, które je zastępują. Jednak ciągle jeszcze pozostaną używane przez starsze i bardziej konwencjonalne telewizory. Ich dominacja dobiega jednak końca.

Ekspertki kreslą kilka różnych scenariuszy rozwoju wydarzeń w tej dziedzinie. W jednym z nich telewizory mogą być wyposażone jedynie w minimalną liczbę przycisków sterowania na obudowie, na przykład przyciski włączania/wyłączania, regulacji głośności i może kilka innych. Cała reszta kontroli odbywałaby się za pośrednictwem głosu, aplikacji lub technologii gestów. Taki model mógłby całkowicie wyeliminować potrzebę pilota.

W innym możliwym wariantcie piloty mogłyby pozostać, ale stać się bardziej uniwersalne i współpracować z różnymi urządzeniami smart home, nie tylko z telewizorami. Na przykład jeden pilot Logitech Harmony może służyć do sterowania telewizorem, oświetleniem, klimatyzacją i zabezpieczeniami w domu. Takie zaawansowane uniwersalne piloty mogą utrzymać się na rynku dłużej. Piloty mogą ewoluować i oferować dodatkowe funkcje, takie jak wbudowane mikrofony do sterowania głosowego, ekrany dotykowe, czytelniki linii papilarnych do uwierzytelniania użytkownika itp. Im bardziej piloty staną się „inteligentne” i zintegrowane z systemami Smart Home, tym większe mają szanse na przetrwanie.

Plusy i minusy

Wśród możliwych pozytywnych konsekwencji zaniku pilotów, wymienia się wzrost poziomu wygody, gdyż sterowanie telewizorem i systemem smart home za pomocą głosu, aplikacji i gestów uchodzi za wygodniejsze dla wielu osób niż używanie tradycyjnego pilota. Niewątpliwą zaletą jest też brak problemów związanych z gubieniem pilotów i zmniejszenie liczby elektronicznych odpadów.

Przewiduje się też możliwe negatywne skutki, np. wyższy koszt, gdyż rozwiązania takie jak rozpoznawanie mowy, gestów czy aplikacje mobilne mogą podnieść cenę telewizorów i urządzeń smart home.



3. Jedno z rozwiązań panelu kontrolnego smart home

Koszty te z czasem prawdopodobnie spadną, ale na początku mogą stanowić barierę dla niektórych klientów. Nowe systemy sterowania mogą być też podatne na błędy i awarie, które uniemożliwią kontrolę nad telewizorem. W przypadku braku pilota, może to powodować frustrację.

Inna kwestią są bariery dla niektórych grup na przykład osoby starsze, z ograniczoną mobilnością lub percepcją mogą mieć problem ze sterowaniem telewizorem za pomocą głosu lub gestów. Dla nich prosty pilot może być wciąż najwygodniejszym rozwiązaniem.

I w końcu trzeba zwrócić uwagę na ograniczoną kompatybilność. Różni producenci telewizorów i urządzeń smart home mogą oferować niezgodne ze sobą systemy kontroli, co utrudni zarządzanie całym domem. Nie jest to więc szczególnie postęp w porównaniu z pilotami. Jeśli przypomnimy sobie o tzw. „uniwersalnych pilotach”, to przewaga znika całkiem a nawet pilot może wyjść na prowadzenie.

Dyrektor wykonawczy Microsoftu James A. Baldwin z Microsoftu, już w 2010 r., podczas konferencji technicznej i wystawy Society of Motion Picture and Television Engineers w Hollywood & Highland przewidywał koniec ery pilota telewizyjnego. Minęło 13 lat i pojawiły nowe rozwiązania, których wtedy jeszcze nie znano, jednak pilot wciąż odgrywa swoją ważną rolę. I kto wie, czy pomimo wielu znaków jego zmierzchu, w kolejnej dekadzie ciągle będzie z nami. W końcu jest symbolem władzy. ■

Mirosław Usidus



Szkoła Wynalazców

dozwolone do lat 15

Mieścieście zadanie: *jak zmusić nieprzyjacielskie wojsko do ujawnienia swojego dowódcy, nieczym nie wyróżniającego się od reszty żołnierzy?*

Dowódca prawdopodobnie nie walczył w pierwszym szeregu, gdyż narażałby nie tylko siebie, ale i armię, która mogłaby stracić dowodzenie. Jednym z ostatnich takich wodzów był Ulryk von Jungingen – wielki mistrz krzyżacki, który pod Grunwaldem kilkakrotnie prowadził, jadąc na czele, chorągwie krzyżackie do ataku. W ostatniej takiej szarży zginął. Znamienne było zachowanie Władysława Jagiełły, który zajął stanowisko dowodzenia na pagórku, z dala od centrum bitwy, skąd mógł mieć wgląd na jej przebieg. Oczywiście jest, że łatwo było zauważyć gońców biegnących do i od króla z rozkazami. Jeżeli dowódca w omawianym przypadku maskował się, unikając rozpoznania, to jednak czynności dowódcze musiał wykonywać. Żeby ujawnić centrum dowodzenia i m.in. głównodowodzącego, można było wystrzelić z łuku strzałę z dowiązaną ważną informacją, która powinna trafić natychmiast do dowódcy. Śledząc kto i gdzie pobiegł z wiadomością, można było zorientować się kto dowodzi. Oczywiście zadanie to z dzisiejszego punktu widzenia ma charakter nieco bajkowy, ale obyczaje rycerskie z latami zmieniały się i taka sytuacja była niekiedy możliwa. Istotne było tu znalezienie logicznego rozwiązania.

Wacław Krzanowski można było wysłać wiadomość, wymagającą podjęcia decyzji dowódcy. Mogła to być np. propozycja poddania się i zaproszenie do uzgodnienia warunków. Mogło też być odwrotnie: wysyłamy wiadomość, że to my chcemy się poddać i prosimy o przysłanie parlamentariuszy. Zarówno w pierwszym jak i drugim przypadku decyzję musiał podjąć dowódca. Śledząc, gdzie pobiegł żołnierz ze strzałą z przyczepioną do niej wiadomością, łatwo było już ustalić kto dowodzi.

Istotnie, najprostszy i zarazem najlepszy sposób. Co prawda nie wiemy w jakiej odległości stały naprzeciw siebie obie armie, ale jeśli założyć, że w zasięgu strzału z łuku, to mogło to być 200 – 300 m. W tych dawnych czasach nie było smartfonów, telewizji i ekranów komputerów: ludzie mieli zdrowe oczy i odległość nawet 400 m nie byłby przeszkodą w dostrzeżeniu komu żołnierz podał strzałę z „depesz”.

Grzegorz Wiejas w dawnych czasach bywało, że gdy dwie armie stały naprzeciw siebie, można było rozstrzygnąć o zwycięstwie metodą pojedynku wodzów. Ostatni taki znany przykład dotyczy... bitwy pod Grunwaldem (!). Wielki Mistrz – Ulryk von Jungingen rzucił wyzwanie królowi Władysławowi Jagiellie i księciu Witoldowi. Jagiełło odrzucił wyzwanie, motywując to względami religijnymi, a Witold przyjął, ale postawił warunek, że potykać się będą bez zbroi i żadnego uzbrojenia, na co Mistrz Ulryk von Jungingen się nie zgodził i w rezultacie do pojedynku wodzów nie doszło. Gdyby moment wyzwania miał miejsce tuż przed bitwą, gdy obie armie już stały naprzeciw, wymagałoby to wysłania heroldów, którzy oczywiście dokładnie obejrzeliby sobie dowódcę, a poza tym jednym z heroldów mógł być łucznik, który miał zadanie zabić dowódcę.

Bardzo dobry sposób. Łucznik w bezpośrednim kontakcie z dowódcą dobrze go sobie zapamięta i później wie do kogo ma strzelać! Sposób ten ma jednak jedną istotną wadę. W dawnych czasach niekiedy (a nawet dość często) obu heroldów ścinano, a głowy odsyłano do ich mocodawców, co oznaczało odrzucenie wyzwania.

Obu kolegom gratuluję i zapraszam do nowych zadań.

Nowe zadanie:

Wszyscy mamy świadomość, że życie na statku odbiega mocno od wzorów lądowych. Nawet najprostsze czynności wymagają odpowiedniego podejścia. Najprostszą sprawą jest jedzenie przy stole. Statek to nie wagon kolejowy i oczywiście kołysze się na fali, Przy czym im mniejszy, tym mocniej. Problemem staje się jedzenie zupy, Przy kołysaniu statku talerz może zjechać: to w jedną to w drugą stronę, a nawet wylać zawartość na spodnie głodnego matrosa. Oczywiście są na to sposoby, od wgłębień w stole zaczynając, po zawieszenie Cardana. Nie są to jednak sposoby w pełni załatwiające problem. Jednocześnie interesują na jednostki raczej niewielkie, które nie mają kompensatorów przechyłu. Trzeba coś wymyślić, co ułatwiłoby życie marynarzom, ale nie będzie wymagało jakichś dużych nakładów. Jak zwykle preferujemy rozwiązania najprostsze. Waszym zadaniem będzie:

Zaproponować jakiś sposób na to, żeby na kołyszącym się statku dało się zjeść zupełnie, bez większych problemów.

Wspomniane i znane rozwiązania: wgłębenia w blacie stołu lub zawieszenia Cardano, to w sumie raczej teoretyczne propozycje niż praktyczne sposoby. Nam chodzi tym razem o coś, co być może nie

będzie w 100% idealne, ale pozwoli zjeść zupełnie bez żadnej „awarii”.

Wszystkim życzymy dobrych pomysłów i oczekujemy na propozycje rozwiązania problemu do końca października br.

Klub Wynalazców

bez ograniczeń wieku

Zadaniem waszym było: do niedawna jeszcze hinduskie wdowy praktykowały samospalenie w przypadku śmierci męża. Ta tradycja była dla wielu odrażająca. W plemieniu Newarów, mieszkającym w górach Nepalu, a także praktykującym hinduizm, znaleźli sposób na uratowanie swoich kobiet, które straciły mężów, przed koniecznością popełnienia samobójstwa, bez naruszania przykazań religijnych.

Mateusz Frankowski przeprowadził obszerne studium tematu, ale okazał się on zbyt peryferyjny. Natomiast skorzystał ze „sztucznej inteligencji” Bing i ustalił, że Newarowie obchodzą raz do roku święto Gai Jatra – „festiwal krowy” obchodzone w sierpniu, podczas którego rodziny, które straciły kogoś w ciągu roku, prowadzą procesję z krowami lub dziećmi przebranymi za krowy, aby ułatwić drogę duszy zmarłego do nieba. Mateusz podejrzewa, że to święto, ułatwiające duszy zmarłego męża przejście do nieba,

może zwalniać wdowę od popełnienia sati. Poza tym Newarowie zniesli sati pod wpływem buddyzmu, który głosił szacunek dla życia i sprzeciwiał się przemocy.

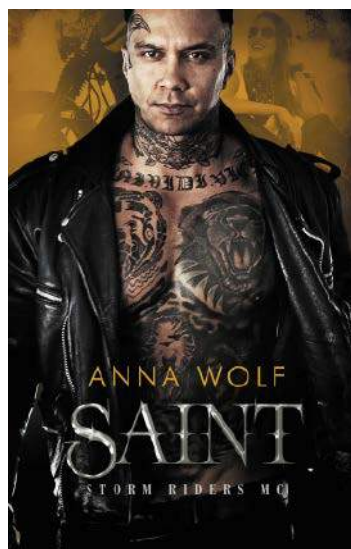
Usługa Bing obsługiwana przez sztuczną inteligencję rozwija się i kto wie jakie będą dalsze losy TRIZ? Na razie przetestowałem kilka typowych zadań trizowskich i te prostsze Bing rozwiązywał, ale takie bardziej „rasowe” to już nie. Religie hinduskie wraz z odmianami są bardzo skomplikowane i niekiedy są zupełnie niezrozumiałe dla Europejczyków. Informacja o święcie

Saint

Anna Wolf

Wydawnictwo Akurat, cykl: Storm Riders MC (tom 5), liczba stron: 320, cena: 42,90 zł

Klub i bracia to wszystko, czego chce Saint, tak mu się wydaje, dopóki nie poznał jej... Saint dawno temu podjął decyzję, że już nigdy nie zwiąże się z nikim na stałe. Przygodny seks to wszystko, czego mu potrzeba. Tak właśnie poznaje Hazel. Jednak przybierająca na sile wojna klubów sprawia, że pewnego dnia w tragicznych okolicznościach kobieta trafia do siedziby Storm Riders MC. W obawie o życie dziewczyny, Saint postanawia objąć ją opieką. Wbrew rozsądkowi oraz złożonym obietnicom, Saint robi wszystko, żeby Hazel z nim została. Tylko, że ona ma zupełnie inny plan na życie. Sytuacja znacznie się komplikuje, kiedy wychodzą na jaw jej powiązania z wrogim gangiem motocyklowym, który jest w stanie posunąć się do wszystkiego, żeby osiągnąć swój cel. Jak potoczy się historia Hazel i Saint? Czy tych dwoje dojdzie do porozumienia i da sobie drugą szansę?





Gai Jatra, które mogło zwalniać wdowy od obowiązku sati, wygląda prawdopodobnie. Zwyczaj sati był wielokrotnie zwalczany zarówno przez Anglików, jak i przez główne odłamy religii hinduskich. I mimo to jeszcze w roku 1987 w stanie Radżastan sati popełniła 18 – letnia Roop Kanwar.

Zachęcam do korzystania z usługi Bing. Daje ona jakies wstępne wejście w temat i może pomóc w korzystaniu z TRIZ.

Zbigniew Przygodzki słyszał gdzieś, kiedyś, że w podobnej sytuacji wystarczy spalić wizerunek wdowy. W dzisiejszych czasach można oczywiście posłużyć się zdjęciem. Zbigniew nie zna oczywiście teologii i obrzędów hinduskich, ale uważa, że prawdopodobnie da się temu zdjęciu nadać charakter osobowy, tzn. jakiś ceremoniał nadawałby zdjęciu „duszę”, co oznaczałoby, że spalenie zdjęcia jest równoważne spaleniu wdowy. Być może dałoby się posłużyć rytuałem voodoo i spalić laleczkę, oczywiście wykonaną wg rytuałów voodoo.

Ciekawa i realna koncepcja. Jednakże nie wiadomo, czy teologia hinduskich wietrzeń dopuściłaby takie „zastępcze sati”. Pomysł kolegi jednakże realistyczny i chyba dobry.

Obu kolegom gratuluję i zachęcam – wzorem kolegi Mateusza do korzystania z usługi Bing, która z czasem będzie coraz bardziej skuteczna.

Nowe zadanie:

Nowoczesność niekiedy spotyka się z problemami „na oko” prymitywnymi, ale których załatwienie nie jest proste. Jednym, z takich problemów jest sprawa oczyszczenia wnętrza rury, w której ściankach są wycinane różne otwory kształtowe. Samo wycinanie

laserem lub aparatem tlenowym to dzisiaj nie jest problemem. Odpowiednie oprogramowanie pozwala na wycinanie „najdzijszych” kształtów. Niestety podczas wycinania część żużla i nadtopionego materiału dostaje się do wnętrza rury, a także na krawędziach wyciętego otworu powstają nadtopione gruzelki spalonego w zasadzie metalu. Oczywiście można to usunąć metodą wycioru armatniego, ale to się może udać w przypadku rury prostoliniowej. A co zrobić gdy mamy do czynienia z rurą powyginaną? I to jest właśnie zadanie dla was:

Zaproponować metodę oczyszczania wnętrza rury, w której pozostają ponaklejane cząstki żużla powstałe podczas cięcia laserem lub tlenem. Rozważyć problem rury powyginanej.

Kiedyś podczas egzaminu ustnego z inżynierii sanitarnej, student na pytanie: jakie zna metody oczyszczania wnętrza kanałów ściekowych, usłyszał „kątem ucha” żartobliwą odpowiedź kolegi: metodą pudła! W stresie egzaminacyjnym student nie wyczuł żartu i głośno oświadczył: „metodą pudła”. To oczywiście stary dowcip z długą brodą, ale jest w nim idea: do kanałów o skomplikowanym przebiegu należy użyć metody sprzętu autonomicznego lub niezależnego od trasy kanału, czyli, albo samobieżny sprzęt, samodzielnie wędrujący we wnętrzu krzywoliniowej rury, albo coś w rodzaju węża. To ogólne idee, ale trzeba im nadać inżynierską postać. Nie oczekujemy oczywiście dokumentacji, ale dość dokładnie sprecyzowanej zasady działania.

Wszystkim życzymy fantazji, wyobraźni i inżynierskiej kreatywności. Przypominam o terminie nadsyłania propozycji rozwiązań: do końca października br.

Vademecum Młodego Wynalazcy

Wracamy do systemu standardów, tym razem klasa 3, chyba najciekawsza i najprzyjemniejsza.

Klasa 3. Przejście do nadsystemu i na mikropoziom

3.1. Przejście do bisystemów i polisystemów

Równolegle z „wewnątrz-systemowym” ulepszaniem (linia standardów klasy 2) istnieje linia „zewnątrzsystemowego” rozwoju: na dowolnym etapie wewnętrznego rozwoju system może być połączony z innym systemem w nadsystem o nowych właściwościach.

3.1.1. Przejście do bisystemów i polisystemów

Efektywność systemu na dowolnym etapie jego rozwoju może być podniesiona systemowym przejściem czyli połączeniem systemu z innym systemem (lub systemami) w bardziej złożony bisystem lub polisystem.

Przykład 1.

Sposób transportu gorących „slabów” (produkty walcarki wstępnej tzw. „slabingu”) ze slabingu do odbiorczego transportera rolkowego walcarki, przy czym słaby są cięte i następnie przemieszczane transporterem rolkowym. Znamienny tym, że w celu zmniejszenia strat ciepła slabów przez zmniejszenie powierzchni schładzania każdego slabu, przemieszczanie ich odbywa się pakietami,

złożonymi z co najmniej dwóch słabów i cięciem bezpośrednim przed podaniem ich do klatki walcarki.

Wyjaśnienia

1. Do stworzenia bisystemu i polisystemu w najprostszym przypadku łączymy dwie substancje S1 i S2 (bisubstancjalne i polisubstancjalne wepola).

2. Przytoczony wyżej standard 3.1.1. także można rozpatrywać jako przejście do polisystemu (choć dokładniej należy go uważać za powiększenie stopnia polisystemowości) Jedność przeciwstawień: rozdzielenie i połączenie prowadzi do jednego i tego samego: tworzą się bisystemy i polisystemy.

Przykład 2.

Dla otrzymania detali z cieniutkich szklanych płytek, półfabrykaty są sklejane w blok, obrabiane maszynowo, a później rozklejane.

Przykład dobrze pokazuje jedną z głównych cech polisystemu: przy jego tworzeniu po-wstaje wewnętrzne środowisko (lub tworzą się warunki dla jego powstania) o szczególnych właściwościach. W danym przypadku pojawiła się możliwość wprowadzenia do wewnętrznego środowiska kleju i otrzymania nie tylko zwykłej „sumy szkielek”, a jednego bloku, o nowych właściwościach. Posmarowanie klejem jednej płytki nic by nie dało. Wytrzymałość pojedynczej płytki można by podnieść, oblewając płytkę dużą ilością kleju (standard 1.1.3.) i tworząc bryłę. Wtedy jednak koszt obróbki byłby znacznie większy.

Druga charakterystyczna cecha bisystemu i polisystemu to „efekt wielostopniowości”

Przykład 3.

Sposób podnoszenie prędkości obrotowej turbowiertła, znamienny tym, że w celu powiększenia liczby obrotów wirnika turbiny, przy zachowaniu dopuszczanej

prędkości roboczego płynu, turbowiertło składa się z kilku sekcji tak, że wał wirnika turbiny pierwszej sekcji jest połączony z korpusem turbiny drugiej sekcji itd, dzięki czemu prędkość obrotów wirnika wzrasta stopniowo od pierwszego do następnych. Względne prędkości kolejnych wirników pozostają w granicach prędkości dopuszczalnych, prędkość ostatniego wirnika jest sumą prędkości wszystkich stopni.

Wyjaśnienia

3. Możliwe jest tworzenie bipolowych i polipolowych, a także wepolowych systemów. W których jednocześnie zachodzi multiplikacja substancji i pól. Niekiedy multiplikuje się parę (P – S) lub wepole w całości.

Przykład 4.

Zadanie wykonania elektrochemiczną metodą otworu, który posiada rozszerzenie w środku swej długości. Elektrode podzielono (w długości) na trzy części i na każdą z nich podano prąd pod innym napięciem.

Wyjaśnienia

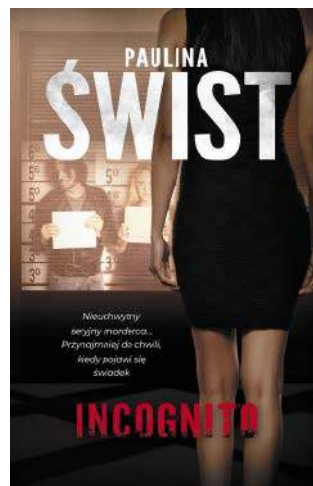
W omawianych opracowaniach, opartych na systemie standardów, przejście do nadsystemu rozpatrywano jako ostateczny etap rozwoju systemu. Zakładano, że system najpierw powinien wyczerpać rezerwy rozwoju „na swoim poziomie”, a potem dopiero przejść do nadsystemu. Jednakowoż zgromadzono dużą ilość informacji, świadczących o tym, że takie przejście może być realizowane na dowolnym etapie rozwoju systemu. Przy tym dalszy rozwój idzie dwoma torami: ulepsza się tworzący się nadsystem i jednocześnie trwa kontynuacja systemu początkowego (wyjściowego). Nieco podobna sytuacja zachodzi w chemii: bardziej złożone chemiczne substancje tworzą się nie tylko metodą dobudowywania nowych

Incognito

Paulina Świst

Wydawnictwo Akurat, liczba stron: 320, cena: 42,90 zł

Nieuchwytny seryjny morderca... Przynajmniej do chwili, kiedy pojawi się świadek incognito. Artur Cienowski i Anka Sawicka znają się od lat. Choć prokurator rejonowy i dziennikarka śledcza nieczęsto stoją po tej samej stronie barykady, to akurat oni zawsze mogą na siebie liczyć. Już niedługo przyjdzie im się o tym dobitnie przekonać. Jak schwycić mordercę, który doskonale zna wszystkie twoje tajemnice? Który podąża za tobą jak cień, a chwilami wydaje się nawet być krok przed tobą? A co, jeśli na dokładkę z niektórymi osobami prowadzącymi śledztwo wiąże cię coś więcej niż tylko profesjonalna solidarność, sympatia albo przyjaźń? Znaków zapytania jest wiele, odpowiedź może być tylko jedna. Całkowicie oczywista i właśnie dlatego kompletnie nieoczekiwana.





orbit elektronów ale i kosztem nie zakończonych orbit wewnętrznych.

3.1.2. Rozwój związków w bisystemach i polisystemach

Podniesienie efektywności bisystemów i polisystemów osiąga się przede wszystkim metodą rozwoju powiązań elementów w tych systemach.

Nowoutworzone bisystemy i polisystemy często mają „zerowe powiązania” (termin zaproponowany przez A. Timoszczuka) t.j. przedstawiają sobą po prostu „zbiór” elementów. Rozwój idzie więc w kierunku wzmacniania międzyelementowych powiązań.

Z drugiej strony, elementy nowoutworzonych systemów niekiedy bywają połączone „sztywnymi” wiązaniami. W takich przypadkach rozwój idzie w kierunku wzrostu stopnia dynamizacji powiązań.

Przykład 4.

„Usztywnienia powiązań”. Podczas grupowego wykorzystania dźwigów (trzy dźwigi o udźwigu po 60 T, podnoszą ciężar 150 T) trudno uzyskać synchronizację ich pracy. W patencie nr. 742372 zaproponowano urządzenie (sztywny wielobok) wiążące ramiona dźwigów.

Przykład dynamizacji powiązań. Katamarany początkowo miały korpusy sztywno połączone ze sobą. Później wprowadzono ruchome łączniki, pozwalające na zmianę odległości pomiędzy kadłubami.

3.1.3. Wzmocnienie różnicy między elementami bisystemu i polisystemu

Efektywność bisystemów i polisystemów zwiększa się w wyniku powiększenia zróżnicowania elementów systemów (przejście systemowe 1-b) np.

- od jednakowych elementów (komplet ołówków) do:
- elementów ze zmienionymi charakterystykami (komplet kolorowych kredek) i wreszcie do:
- całkiem różnych elementów (przybornik kreślarski),
- a także w wyniku inwersji połączenia typu: „element i antyelement” (ołówek i gumka)

Przykład 5.

„Dwupiętrowa piła”, w której dolne zęby są rozwiedzione bardziej niż górne – czysto tnie włókniste materiały.

Przykład 6.

Podczas spawania grubych stalowych blach, elektrody układane są jedna za drugą, przy czym prąd spawania dla każdej następnej elektrody i głębokość jej wsunięcia w szczelinę pomiędzy blachami jest większa niż u poprzedniej. (Typowy polisystem ze zwiniętymi charakterystykami. Efekt osiągnięto, w gruncie rzeczy kosztem przejścia od zwykłego polisystemu, do polisystemu ze zwiniętymi charakterystykami.)

Przykład 7.

Urządzenie do mocowania detali na wewnętrznej powierzchni, zawierające rozcięty sprężysty element, znamienne tym, że w celu podniesienia dokładności zacisku i rozszerzenia możliwości technologicznych urządzenia sprężysty element wykonano w formie dwóch połączonych ze sobą pierścieni z materiałów o różnych współczynnikach rozszerzalności liniowej.

Przykład 8.

Spawalniczy generator mechanicznych drgań, zawierający wykonany w postaci rolki cierny element roboczy, mający możliwość cierno – poślizgowej współpracy z obrabianym obiektem i wykonujący ruch obrotowy. Znamienne tym, że w celu poprawienia jakości spawania przez powiększenie amplitudy drgań, rolkę wykonano jako wielosekcyjną z materiałów o różnych współczynnikach tarcia.

Przykład 9.

Metoda otrzymywania zawieszin sposobem oddziaływania wibracjami na płyn w warunkach wibroturbulencji, przez wprowadzenie do naczynia z płynem sprężystego rezonatora i jednoczesnego działania na naczynie drganiami o częstotliwości rezonansowej. Znamienne tym, że w celu podniesienia ekonomii i wydajności procesu do naczynia z płynem wprowadzono kilka sprężystych rezonatorów o różnych częstotliwościach drgań własnych.

Zadanie 10.

Obiekt – Ciepłarnia. Przedstawić prognozę dalszego rozwoju.

Rozwiązanie zadania 10 wg. standardów 3.1.1 i 3.1.3:

Wg standardu 3.1.1. należy przejść do bisystemu (np. podwójna ciepłarnia). Żeby uzyskać przy tym jakąś nową jakość, trzeba zapewnić wzajemne oddziaływanie między częściami „bicieplarni” lub pomiędzy znajdującymi się w „bicieplarni” roślinami. Maksimum wzajemnego oddziaływania może nastąpić wtedy, gdy rośliny są pod jakimś względem przeciwstawne. Odpowiedź: wg patentu nr 950241 w jednym przedziale rośliny zużywające dwutlenek węgla, a wydzielające tlen, a w drugim rośliny zużywające tlen i wydzielające dwutlenek węgla.

Symbolem klasy 3 standardów, jest rozwój broni strzeleckiej: od jednorurki, poprzez dubeltówkę, następnie dryling, później – już w innym zastosowaniu: poczwórnie sprzężone działka przeciwlotnicze. Ten system został następnie zwinięty do pojedynczego systemu: wielolufowy karabin maszynowy dla samolotów myśliwskich. ■

Prezes Klubu Wynalazców
Champion TRIZ
Jan Boratyński

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

Pomysł miesiąca 9/2023

Interesującym tokiem myślenia wydaje się pomysł ulepszenia mocowania zewnętrznych lusterek samochodowych. Kto wie czy nie doprowadzi to do przeobrażenia tego elementu w część miękką i podatną, o którą nie trzeba się martwić w kontekście kolizji.

Autorem pomysłu jest Ryszard Mijas

1 Jerzy Durlej uważa, że ciągly kłopot z podładowaniem komórki powinien być rozwiązany z wykorzystaniem okoliczności i możliwości miejsca, w którym przebywa turysta. Podczas wędrówek górskich biwakujemy najczęściej w pobliżu strumienia. Należałoby opracować i wyprodukować małą turbinkę – hydrogenerator, która wrzucona do strumienia i zakotwiczona, dawałaby prąd dla komórki.

Wobec ogromnego postępu w dziedzinie budowy małych maszyn elektrycznych (dzięki magnesom neodymowym) pomysł ma realne podstawy. Nowoczesny hydrogenerator może być już na tyle lekki, że jego noszenie nie powinno być problemem.

2 Marek Oskarbski nadesłał pomysł niejako konkurencyjny do pomysłu kolegi Jerzego. Proponuje on pomalowanie płócien namiotowych słynnym perowskitem, co powinno dać prąd wystarczający dla komórki i oświetlenia wnętrza namiotu z pomocą wysokosprawnych diod LED. Jeżeli istnieją działające powłoki z perowskitu, nanoszone na pokrywę smartfona, to o wiele większa powierzchnia namiotu powinna dać zadowalające rezultaty.

Nie wiemy dokładnie jaka jest wydajność powłoki pomalowanej perowskitem, nie wiemy też, co się dzieje gdy poranny deszcz zmoczy tkaninę: czy perowskit będzie działał? Jednakże pomysł jest dobry i firma pani Olgi Malinkiewicz na pewno już o czymś takim myśli.

3 Ryszard Mijas proponuje ulepszyć lusterka nie całkiem współczesnych samochodów. Wiadomo, że na parkingach jest tłok, a nawet na wąskich uliczkach zastawionych obustronnie zaparkowanymi samochodami, trzeba bardzo uważać i mijać się z pojazdem jadącym z przeciwną dosłownie na centymetry. Najnowsze samochody mają lusterka elektrycznie składane, ale nie są one przydatne w ruchu, ułatwiają parkowanie, ale mijania – nie. Ryszard proponuje wyjście najprostsze: zaopatrzenie lusterek w poduszki z gąbczastego materiału i skorygowanie mechanizmu składania tak, żeby się łatwiej zamykał.

Tzw. „święta prawda”! Lusterka obecnie produkowane powiększają szerokość samochodu o ok. 42 cm. Dzisiaj to bardzo dużo. Lusterka składane elektrycznie składają się zbyt powoli, żeby zdążyć je zamknąć przed nadjeżdżającym z przeciwną samochodem. Pomysł kolegi wydaje się dobry, oczywiście trzeba go dopracować, ale mógłby stanowić istotne ułatwienie w ruchu, zwłaszcza na wąskich uliczkach starych centrów miast i osiedli.

4 Karol Zdyb proponuje opracowanie i wdrożenie do komputerów i smartfonów aplikacji, która selekcjonowałaby całe morze informacji potrzebnych i niepotrzebnych i wycinałaby wszelkie teksty, noszące znamiona hejtu. Aplikacja powinna działać z wykorzystaniem sztucznej inteligencji i stanowiłaby ogromne ułatwienie życia wszelkich „komputerowców”.

Faktem jest, że obserwuje się kryzys nadmiaru informacji. Łatwość dostępu do niemal wszelkich informacji, powoduje swoiste rozleniwienie i znudzenie. Przypomina to trochę nowelkę Stanisława Lema o dyplomowanym kosmicznym zbójcu Gębonie. (Stanisław Lem – „Cyberiada”, Wydawnictwo Literackie, Kraków, 1978). Gdy dwaj znamienici konstruktorzy – Trurl i Kłapaucjusz – wpadli w jego łapy i chcieli się wykupić, ofiarowali Gębonowi „demonów II rodzaju”, który z ruchów Browna częsteczek gazu, wydobywającego się ze starej beczki, odfiltrowywał sensowne treści, a więc Gębon dowiedział się: jak się wije arlebardzkie wiją i że córka króla Petrycego z Labaudii zwała się Garbunda, i co jadł na drugie śniadanie Fryderyk II, król bladawców, nim wojnę wypowiedział Gwendolinom, o czym myślałby stół gdyby myślał, itp. Gębon szybko zorientował się, że wpadł w pułapkę nadmiaru informacji niepotrzebnych i bzdurnych. Niestety było już za późno i padł przygnieciony górą zwojów papierowej wstęgi, na której rysik zapisywał wciąż nowe informacje.

5 Krzysztof Dudzik zatrudniany często przez „władze domowe” do krojenia warzyw na sałatkę, przy czym wymagane są równe sześcianiki marchewki, pietruszki i rzepy, zauważył wyzwolenie!! Na wyposażeniu kuchni, mama Krzysztofa ma maszynkę do krojenia ziemniaków na frytki. Z maszynki tej jednym ruchem otrzymuje się „ślupki” ziemniaczane i wystarczyłoby pokroić je w poprzek i gotowe, Krzysztof proponuje więc uzupełnienie maszynki do frytek w poprzeczny nóż, sprzężony z ruchem kratki wieloostrowej. W ten sposób jednym ruchem uzyskiwałoby się garść sześcianików na sałatkę.

Faktem jest, że kandydaci na wynalazców nie znoszą tych monottonnych i prymitywnych czynności kuchennych, jednakże pamiętamy stare powiedzenie: „lenistwo jest rękojmą postępu”, które tak pięknie się zrealizowało w przypadku Krzysztofa. Fakt, faktem: nikt nie lubi monottonnych i prostackich czynności i rzeczywiście stanowi to motywację dla wielu wynalazków. ■



W naszej rubryce „Elektronika dla Ciebie” zachęcamy Cię, drogi Czytelniku, do wykonywania prostych projektów – zabawek, gadżetów itp. Każdy to potrafi. Opis jest zawsze zrozumiały dla nieelektroników, a montaż niemal intuicyjny. A jeśli złapiesz bakcyla pasji elektronicznej, na co liczymy, to podstawy elektroniki przyswoisz sobie z łatwością za pomocą naszego „Praktycznego Kursu Elektroniki” (dostępnego pod adresem: <http://bit.ly/2ThcNxU>).

Płynny regulator prędkości i kierunku obrotów silnika 12 V

Regulatory PWM do sterowania prędkością obrotową silników prądu stałego są dostępne w wielu wykonaniach. Część z nich ma wbudowaną możliwość zmiany kierunku obrotów wału silnika, co odbywa się poprzez przełączenie oddzielnego przycisku lub przestawienie dźwigni. Zaprezentowany układ pozwala zmieniać zarówno szybkość obrotową, jak i kierunek obrotów przy użyciu tylko jednego pokręćła.

Opis układu

Obsługa suwnicy, wciągarki, ramienia czy innej maszyny tego typu wymaga od jej operatora precyzyjnego sterowania elementem napędowym w obie strony. Musi mieć możliwość jego przyspieszenia, zwolnienia oraz płynnej zmiany kierunku obrotów silnika, który jest elementem wykonawczym. Dlatego regulator PWM do silników prądu stałego wydaje się być idealny.

Ale obsługa maszyny, w której jedno pokręćło służy jako regulator prędkości obrotowej, a drugi podzespół (np. przełącznik) odpowiada za zmianę kierunku obrotów, jest niewygodne i nieintuicyjne. Lepiej byłoby mieć tylko jedno pokręćło, które przekręcamy w lewo lub w prawo – im większy kąt obrotu, tym szybciej porusza się nasz mechanizm. Powrót pokręćła do położenia neutralnego powoduje zatrzymanie silnika.

Właśnie do realizacji tego typu sterowania służy opisany układ. Operator trzyma w dłoni tylko jeden element, jakim jest pokręćło potencjometru, a elektronika odpowiada za zmianę kierunku obrotów silnika oraz jego przyspieszanie i zwalnianie. Zasadę



działania tego układu, a dokładniej sposób reagowania na kąt obrotu osi potencjometru, został pokazany na **rysunku 1**. Schemat ideowy regulatora został pokazany na **rysunku 2**. Główną rolę odgrywa w nim mikrokontroler ATtiny13.

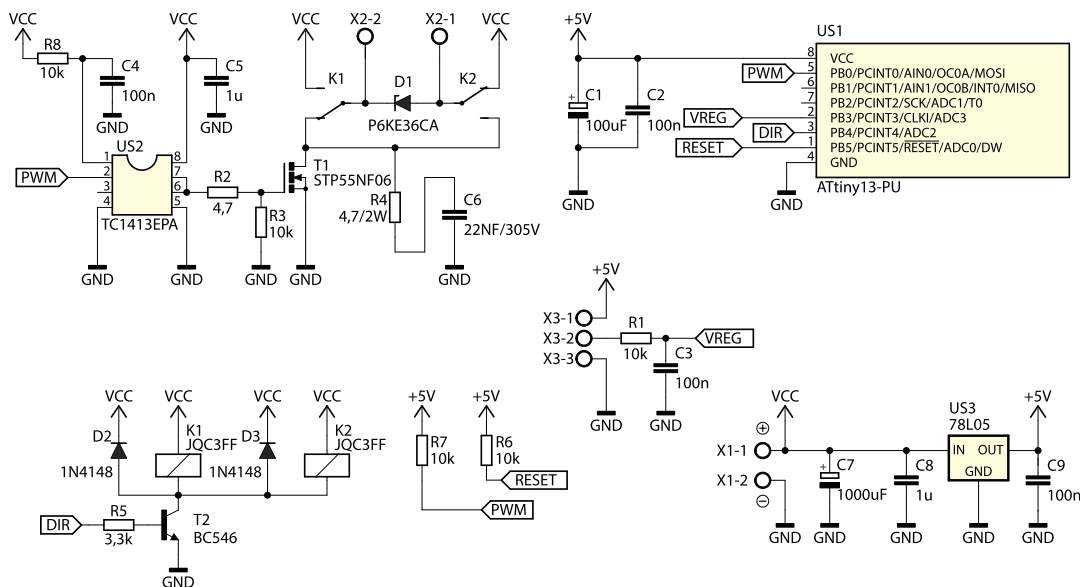
Do odczytu kąta obrotu potencjometru służy przetwornik analogowo-cyfrowy mikrokontrolera, który mierzy pojawiające się na ślizgaczu napięcie. Sam potencjometr został włączony do układu jako dzielnik

Właściwości

- regulacja PWM od 0 do 100%
- automatyczna zmiana polaryzacji napięcia wyjściowego
- obsługa jednym pokręćłem: położenie środkowe wyłącza silnik, przekręcenie powoduje ruch w zadanym kierunku
- częstotliwość sygnału PWM około 600 Hz
- maksymalny prąd wyjściowy 6 A
- zasilanie napięciem stałym 9...18 V, typowo 12 V



1. Reakcja układu na kąt obrotu osi potencjometru



2. Schemat ideowy płynnego regulatora obrotów

napięcia zasilającego mikrokontroler, które wynosi około 5 V. Jednak potencjał ślizgacza może ulegać pewnym fluktuacjom względem masy układu, do czego mogą przyczyniać się zarówno zakłócenia elektromagnetyczne, jak i niedoskonałości samego elementu mechanicznego. Z tego powodu między potencjometrem, a wejściem przetwornika został włączony bardzo prosty filtr RC, złożony z rezystora R1 i kondensatora C3, który tę niepożądaną składową tłumi.

Regulacja PWM sygnałem wygenerowanym przez mikrokontroler odbywa się poprzez kluczkowanie tranzystora MOSFET z kanałem typu N. Do sterowania jego bramką został użyty specjalizowany driver. Rezystor R2 ogranicza prąd bramki w momencie przeładowywania oraz nieco wydłuża ten proces, aby ograniczyć natężenie zakłóceń elektromagnetycznych wywoływanych przez przełączaną indukcyjność uzwojenia silnika. Temu samemu służy gasik, w skład którego wchodzi elementy R4 i C6.

Zmiana kierunku obrotów silnika odbywa się poprzez przełączenie styków przełączników K1 i K2. W położeniu spoczynkowym prawy zacisk złącza X3 ma potencjał zerowy, a lewy dodatni. Przetawienie

Wykaz elementów

Rezystory:

R1, R3, R6...R8: 10 kΩ
R2: 4,7 Ω
R4: 4,7 Ω/2 W
R5: 3,3 kΩ
P1: potencjometr obrotowy 10 kΩ

Kondensatory:

C1: 100 μF/25 V
C2-C4, C9: 100 nF
C5, C8: 1 μF
C6: MKP X2 22NF/305V AC RM10
C7: 1000 μF/25 V

Półprzewodniki:

D1: P6KE36CA
D2, D3: 1N4148
T1: STP55NF06 (lub podobny)
T2: BC546 lub podobny
US1: ATtiny13-PU DIP8
US2: TC1413EPA DIP8
US3: 78L05 TO92

Inne:

X1, X2: ARK2/500
X3: ARK3/500
K1, K2: JQC3FF 12 V SPDT
Radiator

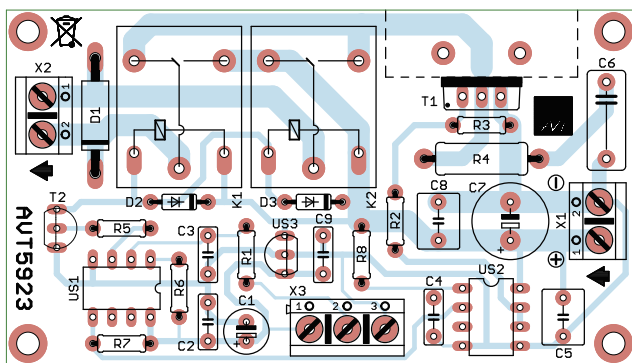


obu zestyków jednocześnie zmienia biegunowość. To bardzo proste rozwiązanie niweluje konieczność stosowania w układzie pełnego mostka H, który wprowadzałby więcej strat i generował problemy ze sterowaniem wszystkich czterech tranzystorów chodzących w jego skład. Tutaj pojedynczy tranzystor steruje mocą silnika (od 0 do 100%), a przełączniki zmieniają kierunek obrotów silnika. Dioda D1 ogranicza amplitudę powstających przebiegów, aby nie doszło do uszkodzenia tranzystora T1.

Cewki przełączników K1 i K2 nie wymagają szczególnego sterowania. Wystarczy zwykły klucz na niemal dowolnym tranzystorze bipolarnym, jak BC546, aby mikrokontroler mógł je załączać. Diody zabezpieczające, D2 i D3, zostały umieszczone przy każdej cewce przełącznika, aby długość drogi prądu wynikającego z samoindukcji była możliwie krótka. Silnik, driver tranzystora MOSFET i cewki przełączników są zasilane wprost z napięcia zasilającego cały układ.

Montaż i uruchomienie

Układ został zmontowany na jednostronnej płytce drukowanej o wymiarach 83×47 mm. Jej schemat został pokazany na **rysunku 3**. Montaż należy rozpocząć od elementów o najmniejszej wysokości obudowy. Zmontowana płytka jest widoczna na fotografii tytułowej. Maksymalny prąd pobierany przez silnik może wynosić około 6 A i to ograniczenie wynika z szerokości ścieżek znajdujących się na powierzchni obwodu drukowanego. Zasilanie układu wymaga podania napięcia stałego o wartości około 12 V (lub z przedziału 9...18 V) do złącza X1. Dolna granica wynika z konieczności prawidłowego sterowania tranzystorem MOSFET. Górny limit napięcia zasilającego układ to efekt troski o izolację podbramkową tego elementu oraz wytrzymałość cewek przełączników. Tym samym napięciem jest też zasilany sterowany silnik. Pobór prądu przez układ nie przekracza 60 mA przy zasilaniu napięciem 12 V.



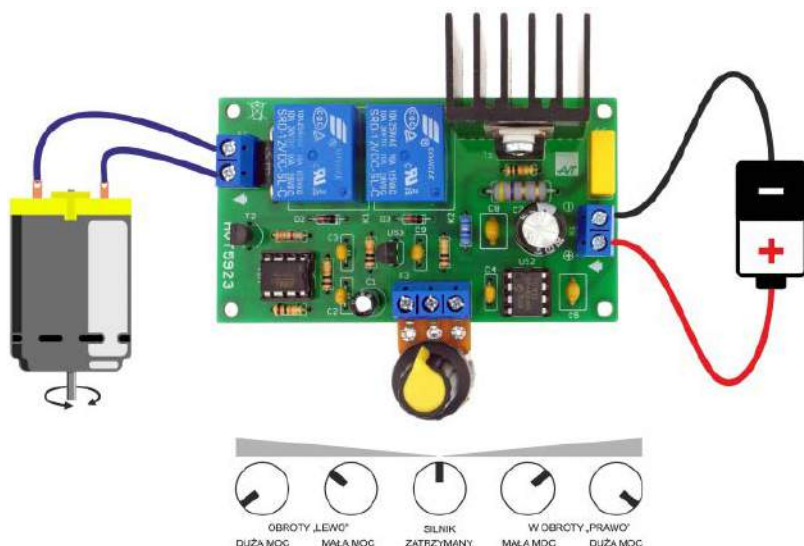
3. Rozmieszczenie elementów na płytce drukowanej



Wszystkie niezbędne części do tego projektu zawiera kit AVT5923, w cenie 60 zł, dostępny pod adresem: <https://sklep.avt.pl/avt5923.html>



Potencjometr nie musi być obrotowy, można z równie dobrym skutkiem użyć suwakowego. Częstotliwość sygnału PWM, który steruje mocą dostarczaną do silnika, wynosi około 600 Hz. Z tego powodu może być słyszalne „piszczenie” uzwojeń silnika w trakcie jego pracy, zwłaszcza z niskim wypełnieniem (przy małej mocy). ■



4. Widok zmontowanego układu



Konkurs „AS Majsterkowania”

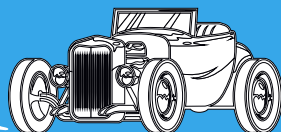
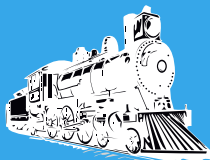
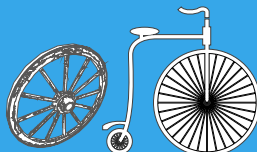
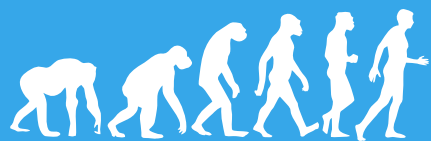
W grudniu tego roku będziemy obchodzić stulecie urodzin Adama Słodowego. By uświetnić tę wyjątkową rocznicę, popularyzując tak bliską Mistrzowi i najslynniejszemu polskiemu popularyzatorowi idei „Zrób to sam”, a przy tym umiejętności konstruktywnego radzenia nie tylko z technicznymi wyzwaniami, zapraszamy Czytelników „Młodego Technika” (i nie tylko) – tych młodych wiekiem, i tych młodych ciekawością świata – do jubileuszowego konkursu pt. AS Majsterkowania!

Zadaniem konkursowym dla uczestników (niepełnoletnich lub dwuosobowych zespołów, w których co najmniej jedna osoba będzie niepełnoletnią) jest samodzielne wykonanie modelu, zabawki albo usprawnienia technicznego inspirowanego dowolną publikacją Mistrza Adama i przesłanie fotorelacji z tej pracy (zawierającej również ilustrację materiału będącego bezpośrednią inspiracją projektu konkursowego) na adres redakcji do końca września tego roku.

Dla młodych laureatów konkursu przewidujemy atrakcyjne nagrody rzeczowe, dyplom uczestnictwa i upominek od Pana Dobromira dla pierwszych stu zgłaszających (zawodników lub drużyn).

Prawdopodobnie dorosłych członków zwycięskich ekip zapewne najbardziej ucieszą kopie mistrzowskiego wskaźnika z wygrawerowanym podpisem samego Mistrza. ■





Śruby i gwinty

około 400 p.n.e.

Grecki filozof Archytas z Tarentu wspomina w swoich pismach o spiralnej konstrukcji umieszczonej w drewnianym cylindrze, co uznaje się za pierwszy historycznie opis śruby, która jak się podejrzewa była używana już wcześniej przez starożytnych Egipcjan.

III w. p.n.e.

Pierwowzorem mechanizmu śrubowego była tzw. śruba Archimedesza (1), nie będąca elementem złącznym, a urządzeniem do podnoszenia poziomu wody. Choć wynalazek ten jest powszechnie przypisywany Archimedesowi (ok. 287–212 r. p.n.e.), to najprawdopodobniej nie wymyślił go sam, lecz opisał znane wcześniej Egipcjanom i Babilończykom rozwiązanie. Śruba Archimedesza używana jest od czasów starożytnych do nawadniania kanałów irygacyjnych. W Holandii służyła do osuszania terenów położonych poniżej poziomu morza. Obecnie ze względu na takie zalety jak nieczułość na zanieczyszczenia, odporność na niskie temperatury oraz niezawodność działania, coraz częściej wykorzystuje się ten mechanizm do transportu ścieków.

starożytność
i średniowiecze

Śruby są używane w Europie powszechnie w prasach do wina i oliwy. Ich działanie polegało na obracaniu śruby, która dociskała owoc do perforowanego kosza, wyciskając sok lub olej.

XV w.

Najwcześniejsze udokumentowane wkręta (śrubokręty) były używane w późnym średniowieczu. Zostały one prawdopodobnie wynalezione pod koniec XV wieku w Niemczech lub we Francji. Oryginalne nazwy narzędzia w języku niemieckim i francuskim brzmiały odpowiednio Schraubenzieher i tournevis. Pierwsza dokumentacja tego narzędzia znajduje się w średniowiecznej księdze domowej zamku Wolfegg, manuskrypcie napisanej między 1475 a 1490 r. Te najwcześniejsze miały uchwyty w kształcie gruszki i były przeznaczone do śrub z rowkiem.

1564

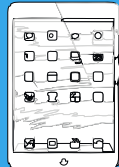
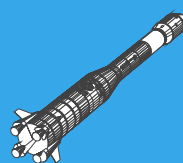
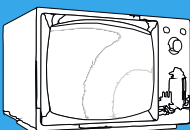
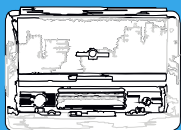
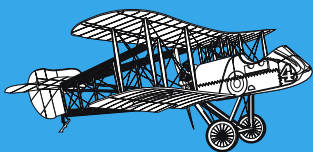
Pierwsze wzmianki o wkrętach, wykonanych z drewna, jako elementach łączących. Opisał je francuski architekt Philibert Delorme. Miały gwintowaną część na końcu i służyły do łączenia belek.

1568

Prawdopodobnie pierwsza maszyna używana do produkcji śrub i wkrętów została wyprodukowana przez Jacquesa Bessona we Francji, który opisał ją podobnie jak wiele innych skonstruowanych przez siebie urządzeń w dziele pt. „Theatrum Instrumentorum et Machinarum” (2). Stworzył on również płytę do wycinania gwintów współpracującą z tokarkami, która została później udoskonalona i wprowadzona do szerszego obiegu przez angielską firmę Hindley of York.

XVII w.

Szereg nowych zastosowań konstrukcji śrubowych. W 1629 roku śruby zostały po raz pierwszy zastosowane w prasie drukarskiej Holendra Willema Blaeu. W latach 1664–65 Christiaan Huygens skonstruował pierwsze zegarki z regulowaną sprężyną balansową wykorzystującą mechanizm śrubowy do naciągania.



XVIII w.

1710

1750

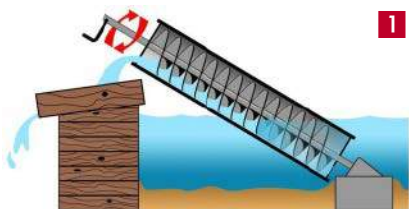
Pomysł na wykorzystanie śruby do łączenia poszczególnych elementów na skalę przemysłową datuje się dopiero na osiemnasty wiek. Wcześniej do spajania elementów używano drewnianych kołków lub metalowych gwoździ, prostych lub zakrzywionych. I to te ostatnie wskazały ścieżkę dla zastosowania śruby. Nie oznacza to, że już wcześniej, w czasach nowożytnych, nie próbowano wykorzystać koncepcji łączenia elementów za pomocą śrub, ale były to przypadki odosobnione, które się nie upowszechniały. Ze śrubą eksperymentował Brunelleschi czy Leonardo da Vinci. Opracowywane przez nich konstrukcje maszyn budowlanych lub wojennych, obłąnne były zapowiedzią przyszłości.

Christopher Polhem ze Szwecji opracowuje śrubę pociągową, która zamienia ruch obrotowy na ruch postępowy i zapewnia napęd dla przesuwania części lub suportu obrabiarki.

Antoine Thiout, francuski zegarmistrz, wprowadza innowację polegającą na wyposażeniu tokarki w napęd śrubowy umożliwiający półautomatyczne przesuwanie wzdłużne wózka narzędziowego (3). Śruby o drobnym skoku są niezbędne w wielu różnych instrumentach, takich jak mikrometry. Do wykonania takiego gwintu niezbędna była tokarka.

1760

Anglicy, bracia Hiob i William Wyatt, opatentowali maszynę do pierwszej przemysłowej produkcji śrub (4). Urządzenie pozwalało wykonać dziesięć śrub na minutę. Ich maszyna wykorzystywała ona śrubę pociągową do prowadzenia frezu w celu uzyskania pożądanego skoku, a szczelina była wycinana za pomocą pilnika obrotowego, podczas gdy główne wrzeciono pozostawało w bezruchu. Dopiero w 1776 roku bracia Wyatt uruchomili fabrykę wkrętów do drewna. Ich przedsięwzięcie nie powiodło się, ale nowi właściciele wyprowadzili je na prostą, a w latach 80. XVIII wieku produkowali 16 tysięcy wkrętów dziennie, zatrudniając zaledwie trzydziestu pracowników.

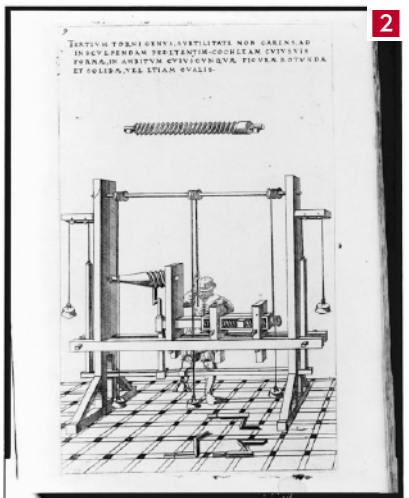


1

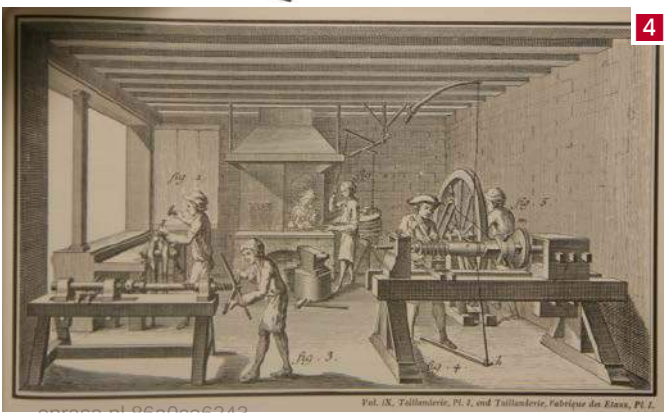


3

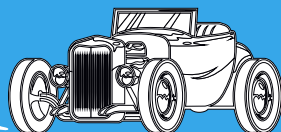
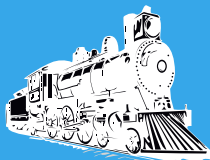
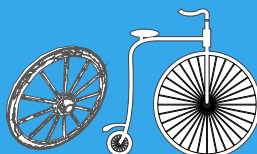
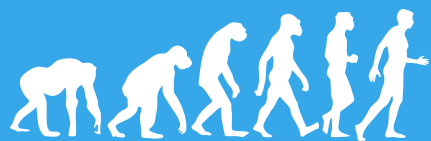
1. Śruba Archimedesa, 2. Ilustracja pochodząca z dzieła „Theatrum Instrumentorum et Machinarum”, 3. Urządzenie zaprojektowane przez Antoine Thiout, 4. Osiemnastowieczna manufaktura produkująca śruby



2



4



1777

Angielski fabrykant Jesse Ramsden (1735–1800) zajmował się produkcją narzędzi i instrumentów, a w 1777 roku wynalazł pierwszą zadowalającą tokarkę do cięcia śrub.

ok. 1795

Pierwszą tokarkę do wytwarzania śrub skonstruował Anglik Henry Maudslay (5), zaś rok później podobne urządzenie zaprezentował amerykański wynalazca David Wilkinson. Wynalazek Brytyjczyka to pierwsza tokarka do metali z suportem krzyżowym napędzanym śrubą pociągową, co pozwalało precyzyjne nacinanie gwintów. Urządzenie przyczyniło się do standaryzacji produkcji śrub i zapoczątkowało wielkoprzemysłową ich produkcję, a konstruktor założył własną fabrykę tokarek, napędzanych silnikami parowymi.

XIX w.

Pojawienie się i rozwój tokarki rewolwerowej i wywodzących się z niej automatów śrubowych (lata siedemdziesiąte XIX wieku) skokowo obniżył koszt jednostkowy gwintowanych elementów złącznych poprzez coraz większą automatyzację sterowania narzędziem maszynowym. Redukcja kosztów spowodowała coraz większe wykorzystanie śrub.

1836

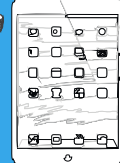
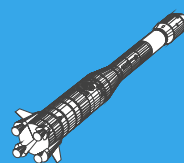
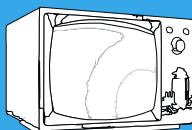
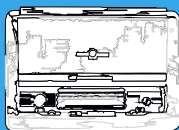
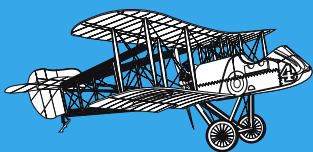
Joseph Whitworth patentuje narzynkę do wykonywania gwintów zewnętrznych. W połączeniu z jego gwintownikiem umożliwiła ona produkcję precyzyjnych gwintów na skalę masową. Whitworth jako pierwszy zastosował narzynki 3-częściowe.

1841–64

Brak standaryzacji gwintów sprawił, że wymiennosc elementów złącznych stała się problematyczna. Aby przezwyciężyć te problemy, Joseph Whitworth zebrał próbki śrub z wielu brytyjskich warsztatów i przedstawił dwie propozycje: kąt nachylenia boków gwintu powinien być znormalizowany i wynosić 55 stopni a liczba nacięć gwintu na cal powinna być znormalizowana dla różnych średnic. Jego propozycje stały się standardową praktyką w Wielkiej Brytanii w latach 60. XIX wieku. W 1864 roku w Ameryce William Sellers niezależnie zaproponował inny standard oparty na nachyleniu gwintu 60 stopni i różnych skokach gwintu dla różnych średnic. Został on przyjęty jako standard amerykański, a następnie przekształcony w amerykańską standardową serię grubozwojną (NC) i serię drobnozwojną (NF). Ten gwint miał płaskie dna bruzdy i szczyty, co sprawiało, że śruba była łatwiejsza do wykonania niż standard Whitwortha, który miał zaokrąglone dna i szczyty. Mniej więcej w tym samym czasie w Europie kontynentalnej zaczęto przyjmować metryczne standardy gwintów z wieloma różnymi kątami zarysu gwintu. Na przykład niemiecki gwint Loewenherz miał kąt boku gwintu 53 stopnie i 8 minut, a szwajcarski gwint Thury kąt 47,5 stopnia. Standardowy międzynarodowy gwint metryczny ostatecznie wyewoluował z niemieckich i francuskich standardów metrycznych opartych na kącie boku 60 stopni. Potem Wielka Brytania, Kanada i Stany Zjednoczone wprowadziły Zunifikowany Gwint Śrubowy (Unified Screw Thread) używany do dzisiaj. Ta koncepcja z pewnymi modyfikacjami została przyjęta przez Międzynarodową Organizację Normalizacyjną (International Standards Organization) przy standaryzacji gwintów metrycznych o różnych zarysach (6).

1870

Christopher M. Spencer patentuje automat tokarski, który pozwalał gwintować na potrzeby produkcji masowej, napędzany za pomocą wału. Stosowany był m.in. w fabryce maszyn do szycia Singera.



1911

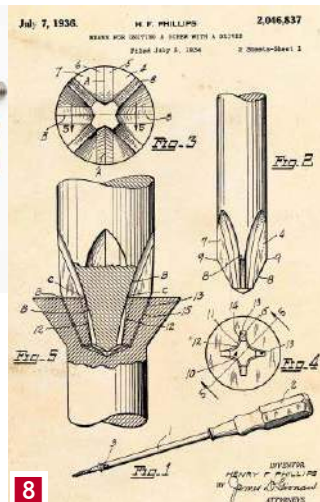
Przez cały XIX wiek w łbach śrub stosowano najczęściej proste rowki do płaskiego wkrętaka oraz kwadraty do pracy z kluczem zewnętrznym. Były one łatwe w obróbce i dobrze sprawdzały się w większości zastosowań. Gdy do ręcznych śrubokrętów dołączyły wkrętarki ułatwiające szybkie wkręcanie śrub, proste nacięcia w łbkach stały się problemem. Aby wkręcić taką śrubę najpierw trzeba było odpowiednio ustawić wkrętarke, aby w ogóle była w stanie wkręcić śrubę oraz, aby nie zniszczyć urządzenia i samej śruby, co przy masowej produkcji było na swój sposób czasochłonne. Wiele konstruktorów próbowało rozwiązać ten problem, a przełom zapoczątkował Peter L. Robertson, który zaprojektował śrubę z kwadratowym wycięciem na głowce (7). Oczywiście wymagała ona zastosowania odpowiedniego śrubokrętu, innego niż dotychczasowe, ale znacząco ułatwiała pracę. Wynalazek Robertsona został opatentowany i stopniowo zaczął zyskiwać popularność.

1933–36

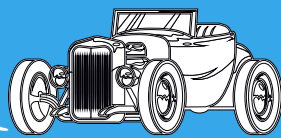
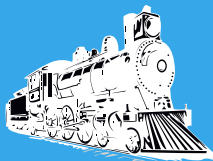
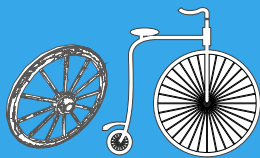
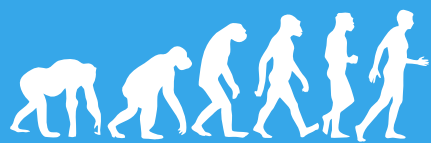
Jedną z wad rozwiązania Robertsona było dość częste zjawisko wypadania wkrętarke, a dodatkowo brak możliwości stosowania standardowych, płaskich śrubokrętów, które były już bardzo popularne. Już po wygaśnięciu patentu Robertsona (w 1928 roku), inny konstruktor, John P. Thompson zaprojektował śrubę, w której wycięcie na łbku miało kształt krzyża i odpowiedni, pasujący do tego wkrętak. Początkowo pomysł Thompsona nie zyskał popularności, ponieważ wielu producentów uważało, że technicznie nie było możliwości produkowania na masową skalę takich śrub. Dopiero Henry F. Phillips dostrzegł potencjał w projekcie swojego przyjaciela i postanowił go spopularyzować. Po wprowadzeniu drobnych poprawek, w 1934 roku założył firmę Phillips Screw Company. Po odkupieniu w 1935 roku patentu od Thompsona, rozpoczął produkcję nowych śrub w 1936 roku na podstawie nowego patentu (8). Krzyżowe nacięcie ułatwiała wkręcanie i wykręcanie. Ciekawostką jest to, że już około sześćdziesiąt lat wcześniej angielski wynalazca John Frearson opatentował śrubę z „otworem w kształcie krzyża”, jednak wtedy ta konstrukcja się nie przyjęła.

1936

W Niemczech powstają śruby z nacięciem w formie sześciokąta. Zaprojektowano je w firmie Bauer & Schaurte Karche z Neuss. Po niemiecku nazywano je Innensechskantschraube Bauer und Schaurte – po polsku imbus. Dalszym rozwinięciem imbusów było powstanie systemu Torx, który dzięki zastosowaniu mocniejszych materiałów i nacięć w formie sześcioramiennej gwiazdy pozwala na przenoszenie większej siły niż imbus. W późniejszych latach i dekadach eksperymentowano i wprowadzano kolejne warianty nacięć (9).



5. Henry Maudslay,
6. Normy ISO dotyczące śrub i gwintów,
7. Wkręty z rowkiem według patentu Robertsona, 8. Ilustracja z opisu patentu Phillipsa,
9. Nacięcia w łbkach współczesnych śrub



Rodzaje śrub i gwintów

Śruby to maszyny proste służące do łączenia ze sobą różnych elementów. Gwinty to definicyjnie część konstrukcji śruby.

Jeśli chodzi o śruby to mamy do czynienia z wielką różnorodnością konstrukcji, które trudno ująć w pełnej klasyfikacji. Na rynku występuje wiele rodzajów śrub, różniących się pomiędzy sobą rozmiarem (średnicą oraz długością trzpienia), kształtem łba, zakończeniem, rodzajem gwintu oraz długością gwintu w stosunku do długości trzpienia.

Ze względu na kształt łbów możemy wyróżnić kilkanaście rodzajów, wśród których najbardziej rozpoznawalnym jest śruba z łbem sześciokątnym. Jest bardzo wytrzymała i stosuje się ją przy łączeniach, narażonych na duże obciążenia. Śruby skrzydełkowe (motylkowe) montowane są w elementach, wymagających regulacji, np. teleskopowych. Śruby z oczkiem lub uchem służą z kolei do mocowania lin, holowania i przenoszenia elementów o niewielkiej wadze. Śruby z łbem stożkowym natomiast najczęściej wykorzystywane są w branży meblarskiej – przy skręcaniu desek czy mocowaniu zawiasów. Dzięki specyficznej budowie, śruba ta nie wystaje ponad płaszczyznę materiału. Poza wymienionymi, występują również śruby z łbem czworokątnym i trójkątnym, a także grzybkowym, kulistym, walcowym, młoteczkowym, a nawet śruby bez łba – z gniazdem.

Śruba może występować z gwintem na całej długości lub tylko na części. Może również posiadać dodatkowe elementy takie, jak nosek czy wieniec. Typy śrub można ponad to klasyfikować w oparciu o kryterium zakończenia lub gniazda. Najczęściej stosowanym gniazdem jest gniazdo sześciokątne, ale można również spotkać w jego miejscu wgłębienie krzyżowe lub rowek. Z kolei najczęstsze zakończenia śrub to: płaskie, soczewkowe, ścięte, stożkowe i wgłębione.

Gwinty

Gwinty przede wszystkim klasyfikuje się ze względu na system miar – na metryczne, w przypadku, gdy kąt zarysu gwintów wynosi 60° , a wymiar skoku jest podany w milimetrach i calowe, jeśli kąt zarysu gwintów wynosi 60° albo 55° , a wymiar skoku jest podawany w ilości zwoi gwintu na jeden cal.

Gwinty dzielone są także ze względu na ich kształt. Ty wyróżnia się m.in. gwint do drewna, gwint okrągły, gwint prostokątny, gwint rowerowy, gwint stożkowy, gwint toczny, gwint trapezowy niesymetryczny, gwint trapezowy symetryczny, gwint trójkątny, gwint



walcowy. Nazwy – gwintów trójkątnych, stożkowych, trapezowych i innych wywodzą się z kształtu zarysu.

Innym podziałem jest ten, w którym wyróżniamy gwint wewnętrzny i zewnętrzny. Gwinty dzielimy też ze względu na kierunek obrotu: prawy (gwint śruby należy wkręcać w prawo, zgodnie z ruchem wskazówek zegara) i lewy (gwint śruby trzeba wkręcać w lewo, przeciwnie do ruchu wskazówek zegara).

Jest też klasyfikacja uwzględniająca stosunek podziałki do średnicy. W niej wyróżnia się takie typy gwintów:

- gwint zwykły, o wartości skoku dla danej średnicy zgodnej z Polską Normą,
- gwint drobny, o wartości skoku dla danej średnicy mniejszej niż dla gwintu zwykłego,
- gwint gruby, o wartości skoku dla danej średnicy większej niż dla gwintu zwykłego. ■

M.U.



Gorące czary klasy A

Musical Fidelity A1

Do sprzedaży wrócił słynny wzmacniacz A1 firmy Musical Fidelity. Swoją karierę rozpoczął 40 lat temu, szybko stał się bardzo popularny, bowiem nie był specjalnie drogi – kosztował ok. 300 funtów, ale sławę zawdzięczał niezwyklej konstrukcji, parametrom i brzmieniu. Był kilkakrotnie modyfikowany, jednak na początku XXI wieku zniknął, więc jego powrót jest w audiofilskim świecie sensacją. Dokładny test najnowszej wersji ukazał się w numerze 9/2023 AUDIO, którego redakcja przygotowała dla MT specjalne, esencjonalne opracowanie tego tematu.

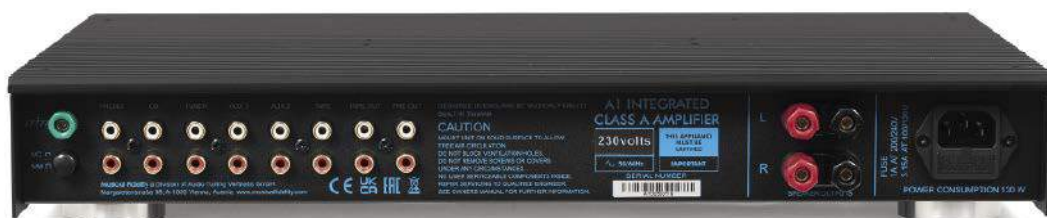
Wzmacniacze w klasie A są zwykle duże, nawet bardzo duże, chociaż wcale nie wykazują się wysoką mocą. Dlaczego – wiedzą wszyscy, którzy chociaż trochę liźnęli techniki. Wysoki prąd spoczynkowy, polaryzujący tranzystory, oznacza duży pobór energii również wtedy, gdy wzmacniacz gra bardzo cicho... a nawet nie gra wcale. Oczywiście pobierana z sieci zasilającej energia elektryczna, jeżeli nie płynie do głośników, musi zamienić się w inną formę energii, i jest to zwykłe ciepło. Sprawność energetyczna – stosunek mocy elektrycznej oddawanej (użytecznej) do pobieranej – układu w klasie A jest więc bardzo niska. Oznacza to nie tylko wysoki koszt energii, ale sam fakt, że w większości zamienia się w ciepło (w epoce ocieplenia klimatu!) wywołuje przykre skutki – wymaga zbudowania urządzenia, które będzie to ciepło zdolne efektywnie oddawać (a to kolejne koszty), jak też powoduje szybsze zużycie się elementów wrażliwych na wysokie temperatury (głównie kondensatorów). Mimo to są chętni – najbardziej wymagający i dostrzegający brzmieniowe zalety klasy A (lub choćby wierzący, że klasa A z założenia brzmi lepiej).

Wzmacniacz A1 nie był i nie jest jedynym wzmacniaczem w klasie A, ale jego wyjątkowość polega na umiarkowanej wielkości. Mogłoby to rodzić podejrzenia, czy rzeczywiście wzmacniacz pracuje w klasie A, bowiem zdarzały się konstrukcje, które tylko jako A-klasowe

przedstawiano, a w rzeczywistości pracowały w klasie AB, jednak w tym przypadku co do tego nie ma wątpliwości – gruntownie przebadany został zarówno oryginalny A1, jak też jego najnowsza wersja. Laboratorium AUDIO porównało obydwa egzemplarze.

Jak udało się zaprojektować „kompaktowy” wzmacniacz w klasie A? Po pierwsze moc wyjściowa jest umiarkowana, deklarowana przez producenta to 2×25 W, podczas gdy tej wielkości wzmacniacz w klasie AB miałby kilka razy większą, a nowoczesny wzmacniacz w klasie D – setki watów. Mimo to potrzebne było specjalne rozwiązanie. Otóż radiatorem odprowadzającym ciepło jest cała górna ścianka, która już niedługo po podłączeniu do sieci nagrzewa się do temperatury 70°C ; ponieważ jednak wzmacniacz stale pobiera z sieci ok. 100 W, więc jeżeli część tej mocy zamieniamy na elektryczną moc wyjściową („gramy”), to nieco mniej niż początkowo zamienia się w ciepło... i radiator jest odrobinę chłodniejszy!

Mimo to więcej niż 27 W (z każdego kanału) nie wy ciśniemy, a i to pod warunkiem obciążenia impedancją $8\ \Omega$. Przy impedancji $4\ \Omega$ moc spada do 17 W, co jest nietypowe dla „porządnych” wzmacniaczy tranzystorowych, które przy niższych impedancjach moc zwiększają (znacznie zwiększając prąd nawet przy pewnym spadku napięcia), ale tutaj większy prąd oznaczałby jeszcze większą temperaturę, czemu trzeba było zapobiec.



A1 trzyma się funkcjonalności typowej dla sprzętu sprzed kilkadziesiąt lat

Charakterystyki częstotliwościowe swobodnie pokrywają pasmo akustyczne, ze spadkiem -3 dB daleko poza nim, przy 64 kHz.

Natomiast spektrum zniekształceń harmonicznym, które w dużym stopniu określa charakter brzmienia, jest dość emocjonujące... Pojawia się wysoka (-57 dB) druga harmoniczna, czym jeszcze się nie martwimy, wiedząc że parzyste harmoniczne mniej szkodzą niż nieparzyste, ale trzecia też jest wysoka (-62 dB), a powyżej -90 dB wychodzi jeszcze piąta. To sytuacja często spotykana we wzmacniaczach... lampowych, które wcale nie zawsze, wbrew teorii, promują tylko parzyste.

Charakterystyka THD+N w funkcji zniekształceń wygląda nietypowo (dla ogółu wzmacniaczy tranzystorowych, a znowu podobnie do wzmacniaczy lampowych), bowiem minimalna wartość nie pojawia się przed przesterowaniem, ale znacznie niżej, a powyżej powoli rośnie. To jednak jest zjawiskiem raczej korzystnym.

Przy tak niskiej mocy na bardzo wysoką dynamikę nie ma co liczyć, ale może chociaż odstęp od szumu jest znośny? Jeżeli również pod tym względem A1 upodobniłby się do wzmacniaczy lampowych, raczej by to nas nie ucieszyło. Na szczęście jest inaczej, chociaż znowu ciekawie. Wzmacniacz ma dwa tryby czułości – w trybie Normal wynosi ona 250 mV, jest

więc bliska dawnemu standardowi, a w trybie Direct (wówczas omijamy jeden ze stopni wzmacnienia) spada do 900 mV – to jednak wartość bliska nowoczesnemu trendowi w projektowaniu wzmacniaczy, który wynika z wysokiego (2 V) poziomu napięcia wyjściowego współczesnych urządzeń źródłowych. Ustalenie niskiej czułości niemal zawsze obniża poziom szumów wzmacniacza, dzięki temu w trybie Direct wynosi doskonale 90 dB, ale i w trybie Normal całkiem nieźle 84 dB; dlatego dynamika wyniosła odpowiednio 104 dB i 98 dB.

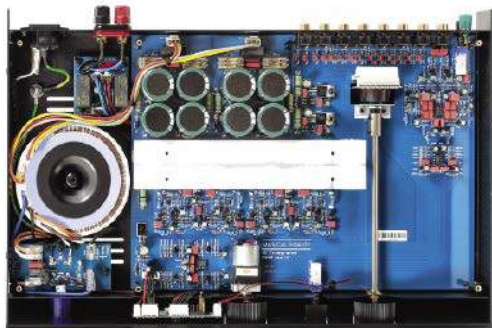
Współczynnik tłumienia znowu przynosi reminiscencje wzmacniacza lampowego... wynosi 32 w odniesieniu do 8 omów, a więc tylko 16 do 4 omów.

Oryginalny, ale lekko wyremontowany A1 (wymienione kondensatory i potencjometr głośności) miał bardzo podobną moc, szerszą charakterystykę przenoszenia (-3 dB przy 92 kHz), wyższe harmoniczne... ale przede wszystkim parzyste, i podobnie niski współczynnik tłumienia. Jego słabością był niski odstęp od szumu (70 dB) powodowany zarówno wysoką czułością (0,14 V) jak też dłuższą ścieżką sygnału (na skutek niezbyt szczęśliwie zaprojektowanego przełącznika tzw. pętli magnetofonowej „tape monitor”, który w nowej wersji zamienił się w przełącznik Normal/Direct).

Mimo tych różnic można stwierdzić, że udało się „odtworzyć” praktycznie wszystkie ważne cechy oryginalnego A1, niektóre parametry poprawić, to jednak sukces wciąż względny i kontrowersyjny, określający wzmacniacz mało uniwersalny (niska moc, problemy z obciążeniem 4-omowym, niski współczynnik tłumienia), do którego trzeba starannie dobrać kolumny. Pod względem otwartości na źródła sygnału, A1 pozostaje wzmacniaczem czysto analogowym (żadnego DAC-a, a tym bardziej odtwarzacza strumieniowego), ale oprócz wejść liniowych ma też cenione dzisiaj wejście phono (dla gramofonu, i to z korekcją zarówno MM, jak i MC). Ponadto został doposażony w zdalne sterowanie. ■

Andrzej Kisiel

Zachowano najważniejsze cechy oryginalnego układu, ale niektóre elementy trzeba było wymienić, a niektóre rozwiązania poprawić





Nadmorska impresja

To zdjęcie zostało wykonane w dwóch etapach. Podczas naświetlania przez 10 s pierwsza część ekspozycji została wykonana z aparatem umieszczonym na statywie, aby uchwycić ostrą scenę; następnie podczas drugiej połowy ekspozycji statyw został przesunięty, aby uzyskać impresjonistyczne rozmycie.



Wyjdź poza sztywne klasyczne ramy

Celowe poruszanie aparatem to technika stosowana podczas długiego czasu naświetlania w celu uzyskania charakterystycznych efektów rozmycia, które dają bardziej abstrakcyjny obraz wybrzeża. Najpopularniejszy sposób polega na poziomym przesuwaniu horyzontu, ale technika, którą tu omówimy, tworzy bardziej artystyczne obrazy.

Jak dobrze poruszyć?

1. Połącz stabilność statywu z celowym ruchem aparatu.
2. Ustawienia aparatu. Fotografuj z priorytetem przysłony, tak by uzyskać czas od 2 do 10 s przy ISO 100 i $f/11$ –16. Zwolnij migawkę i w połowie ekspozycji podnieś aparat ze statywem.
3. Skomponuj i działaj. Skomponuj zdjęcie tak, jakby było to „standardowe” ujęcie, upewniając się, że horyzont jest równy. Załóż wszelkie filtry niezbędne do kontroli ekspozycji.

KREATYWNE fotografowanie



Trzymaj się zasad

Stosuj reguły kompozycyjne, aby dodać wizualną równowagę i poczucie dynamiki do swoich zdjęć morskich krajobrazów

Sposób, w jaki rozmieszczasz elementy sceny, jest niezwykle ważny, bo to wpływa na równowagę wizualną i wciąga widza w obraz. Podniesienie aparatu do oka i robienie zdjęć bez przemyślenia kompozycji zdecydowanie mija się z celem. Przed zamocowaniem aparatu na statywie spróbuj wykonać zdjęcia z ręki, z różnych wysokości i staraj się zmieniać miejsce, z którego fotografujesz, chodząc wokół tematu, aby znaleźć najlepszy kąt. Zwracaj przy tym uwagę na to, jak poszczególne elementy sceny równoważą się wzajemnie w kadrze.



1. Centralna kompozycja

Często mówi się, że nigdy nie należy komponować krajobrazu z horyzontem lub obiektem w centrum kadru, jednak w przypadku niektórych pejzaży morskich zapewnia to większą równowagę niż zasada trójpodziału.



2. Przestrzeń negatywna

Duża ilość pustej przestrzeni z wyróżniającym się punktem głównym może stworzyć wyjątkowo mocny kadr. W tym przypadku dwie łodzie są kluczowymi punktami kompozycji, a niebo i woda współgrają, tworząc dużo przestrzeni negatywnej.



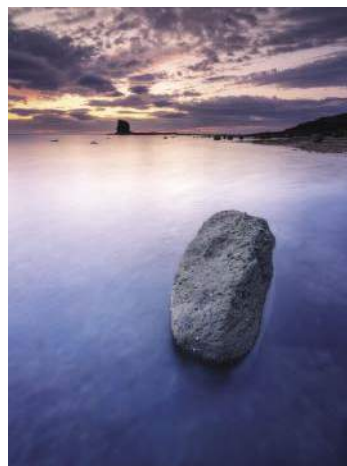
3. Pierwszy plan

Obiekt pierwszego planu to po prostu interesujący temat umieszczony w dolnej trzeciej części obrazu, który ma związek ze sceną i stanowi wizualną bramę, przez którą wchodzimy w głąb kadru. W połączeniu z zasadą trójpodziału jest to prosty sposób tworzenia mocnych kadrów.



4. Linie wiodące

To elementy liniowe w obrębie sceny, takie jak falochrony, linie w piasku lub nawet rząd skał, które tworzą linię prowadzącą od pierwszego planu do środka lub tła. Skuteczne linie prowadzące nigdy nie są poziome. Jest to świetna technika, która pozwala przyciągnąć widza do głównego tematu.



5. Trójpodział

Wyobraźmy sobie, że kadr jest podzielony równo przez dwie linie poziome i dwie pionowe. Aby uzyskać równowagę wizualną, obiekt główny należy umieścić w jednym z czterech punktów przecięcia tych linii, a horyzont – wzdłuż linii podziału. Tę zasadę warto stosować nie tylko w pejzażach.

*** Pisownia oryginalna ***

PRZEGLĄD TECHNICZNY Samoczynne smarowanie obrzeży kół taboru kolejowego

Na kolei Montreux-Berne zastosowano smarowanie samoczynne obrzeży kół wagonów i łbów szyn, przyczem osiągnięto wyniki pomyślne. Badania, przeprowadzone w Ameryce, wykazały również, że prawidłowe i stałe smarowanie tych powierzchni zmniejsza 4–5-krotnie ich zużycie. Przyrząd do smarowania, zastosowany na kolei Montreux-Berne, działa tylko podczas ruchu pociągu; na postoju natomiast wyłącza się samoczynnie. Na licznych zaokrągleniach o promieniu 40 do 80 m ścieranie łba szyny zostało prawie całkowicie usunięte. Do smarowania używano stary smar, już zużyty poprzednio do innych celów.

4 września 1923

Aluminyjne lane drzwi do wagonów żelaznych

Od pewnego czasu zaczęto wprowadzać na kolejach amerykańskich St. Zj. w żelaznych wagonach osobowych drzwi metalowe, odlewane z gliny, zamiast dotychczasowych, wykonywanych z blachy żelaznej. Te ostatnie miały bowiem tę wadę, że woda, skraplająca się pomiędzy podwójnymi ich ściankami blaszanymi, wywoływała rdzewienie blachy. Natomiast drzwi aluminiowe, prócz odporności na rdzę, mają też jeszcze zaletę, że są lżejsze i w razie uszkodzenia mogą być przetopione. Ponieważ metal ten znacznie się kurczy przy odlewie, powstają trudności z wykonaniem dużych i cienkich płyt. Formowanie odbywa się za pomocą modeli aluminiowych, mających wymiary zwiększone o długość 35 mm, licząc na skurcz odlewu. Obie skrzynie formierskie (górna i dolna) muszą być mocno do siebie przyciśnięte ciężkimi szynami, dla otrzymania ścianki odlewu o grubości jednostajnej i zgodnej z przepisami.

4 września 1923

Zastosowanie masek gazowych w kolejnictwie

Dla wytworzenia dogodniejszych warunków pracy dla personelu kolejowego, szczególnie dla obsługujących parowóz, wprowadzono w St. Zjedn. maski gazowe do użytku podczas przejazdu przez tunele. Chodzi tu o pochłanianie kwasu węglowego i tlenku węgla oraz kwasu siarkowego. Badania zapomocą masek z czasów wojny oraz nowych modeli wykazały, że najpraktyczniejszą jest maska nowego modelu formatu kieszonkowego. Maska

posiada ustnik kauczukowy, przez który należy wciągać powietrze przy oddychaniu, nos zaś zaciska się odpow. sprężynkami, jak przy binoklach. Środkiem pochłaniającym jest proszek z tlenków metali (miedzi

i manganu), który przetwarza CO na CO₂, oraz mieszanina z proszku węgla drzewnego i wapna, pochłaniająca CO₂ i SO₂. Wyniki stosowania tych masek okazały się b. zadowalające, zarówno w licznych tunelach na kol. Baltimore and Ohio Railroad, jednotorowych, gdzie chodzą pociągi o 2-ch parowozach, jak też w tunelu kol. Pennsylvania Railroad, gdzie używa się po 3 parowozy do przewożenia pociągów.

4 września 1923

Komunikacja radiotelefoniczna w tramwajach

Third Avenue Railway Co wykonała próby wprowadzenia komunikacji telefonicznej na swych liniach kolei elektrycznych. Jak podaje Electric Railway Journal (...), urządzenie polegało na nათოზიუნი na przed trakcyjny, płynący przez przewodnik linii, prądu zmiennego wielkiej częstotliwości, nie oddziaływującego na prąd pierwszy. Ten prąd szybkozmienny użyto do komunikacji telefonicznej obustronnej – pomiędzy stacją centralną a obsługą wagonu. Sieć tramwajowa służyła tu jako antena. Zauważono, że pomiędzy podstacją, zaopatrzoną

w aparat nadawczo-odbiorczy, a wagonem, posiadającym także urządzenie, można było podtrzymać stałą rozmowę wzajemną, zarówno podczas jazdy, jak na postoju. Przyrząd nadawczy na podstacji składał się z 3-ch rurek (lamp) trójelektrodowych po 50 wolt, służących: pierwsza, jako oscylator, druga – amplifikator i trzecia – modulator. Jako przyrząd odbiorczy, dodano do tego jedną rurkę, działającą jednocześnie jako detektor i amplifikator. W wagonie układ aparatów był mniej więcej ten sam; tylko dla ułatwienia słuchania podczas huku w czasie jazdy, dodano jeszcze jedną lampę do aparatu odbiorczego. Energia, niezbędna do wysłania (emisji), była doprowadzana do stacji nadawczej (czyli do powyższych 3 lamp kat.) zapomocą grupy generatorów 1000-woltowych oraz baterii akumulatorów z 12 ogniw, czyli 24-woltowej. Dla obustronnej rozmowy, czyli podwójnego działania radiotelefonu, naogół posługują się obecnie przekładnikami (relais), pozwalającymi ustawić każdą stację

na odbiór, albo na nadawanie. Przy próbach w Nowym Yorku utrudenienie to usunięto przez zastosowanie różnych częstotliwości dla prądów do nadawania i do odbioru, mianowicie, 73 000 okresów na sek. w pierwszym i 49 000 – w drugim wypadku.

18 września 1923

Rozwój samochodów elektrycznych

Od kilku lat w technice samochodowej zauważa się ponowny zwrot ku samochodom akumulatorowym, które szczególnie licznie mają zwolenników w Stanach Zjednoczonych Am. Półn. Bulletin of the NELA (National Electric Light Association) oblicza ilość takich samochodów w samym New-Yorku na 5000, zaznaczając, że liczba ta wzrasta o 500 sztuk rocznie. Również w Europie coraz więcej się poświęca uwagi temu typowi samochodów, czego dowodem służy organizowany w Paryżu w mies. bież. konkurs samochodów elektrycznych, o którym donosi „Le Génie Civil” N° 9 r. b., przytaczając jednocześnie dane o samochodach akumulatorowych angielskich (w szczególności z praktyki komunikacji miejskiej w Glasgow, podł. Electrical Review). Władze miejskie tego miasta przeszły w r. 1913 do obsługi komunikacji na miejskiej sieci elektrycznej zapomocą własnych samochodów elektrycznych, zastępując dawnych przedsiębiorców prywatnych, w których ręku była ta obsługa. Obecnie miasto posiada kilkadziesiąt samochodów (1, 2, 3½ i 5-tonnowych) tego typu, a z nich 21 – tylko do służby na sieci elektrycznej (ciężarowe i osobowe), reszta do użytku publicznego. Jak się okazało, amerykańskie samochody z ogniwami Edisona (niklowymi), są lepsze od angielskich, wyposażonych w akumulatory ołowiane. Wyjaśniło się także, że konstruktorzy angielscy mają mniejsze doświadczenie i niektóre części mechanizmu (hamulce, łożyska kulkowe, łańcuchy napędne) niezbyt sprawnie działały. Obecnie braki te usunięto. Samochody posiadają 1, 2 lub 4 motory z przekładnią łańcuchową albo zębatą. Zasobniki są ładowane na 3-ch stacjach ładunkowych (2 po 75 kW i jedna – 150 kW), które pracują bez przerwy, ładując w nocy zasobniki do jazdy dziennej, a w dzień do samochodów, wywożących odpadki i śmiecie w nocy. Obok stacji mieszczą się remizy na 48 samochodów. Wypadków zepsucia samochodów nie zdarza się; tylko

co 2 lata wymagają one rewizji. Co zaś do zasobników, to jakkolwiek Edisonowskie są uważane za lepsze niż ołowiane, jednak, zadowolając ostatnim ulepszeniem, te ostatnie zaczynają skutecznie współzawodniczyć z pierwszymi. Zaletą ogniw Edisona jest ich służba 4–5-letnia, a nawet czasem aż 8-letnia, wówczas gdy zasobniki ołowiane mogą pracować 20–24 miesiące, względnie jeszcze o 20 mies. więcej, w razie zamiany płyty dodatkowej. Na próbnej jeździe na szlaku 32 km, samochód 1½ t-owy zużył 143 ampero-godzin przy napięciu 80 woltów. (4½ amp.-godz. na 1 km, wzgl. ok. 352 watt-godz. na 1 km). Średnia prędkość wyniosła 15 km/godz.: kolektor na grzewał się nie wyżej jak do 20°, zasobnik zaś – do 29° C. (...)

25 września 1923

PRZEGLĄD ELEKTROTECHNICZNY Nowy przekładnik elektryczny

Nowy typ przekładnika zapobiega odkształcaniu dźwięków. To częścią ruchomą jest stożek, owinięty cienkim jedwabiem, dokoła którego znajduje się jedno- lub parowarstwowa spirala z cienkiego drutu glinowego. Stożek ten jest umieszczony pomiędzy biegunem elektromagnesu. Przy przepływie prądu (telefonicznego) przez drut glinowy wzbudza się w nim siła elektromotoryczna na skutek oddziaływania pola magnetycznego elektromagnesu. Ponieważ spirala nie ma swej własnej częstotliwości, przeto, odtwarzając dźwięki, nie odkształca ich. Osiąga się przeto wyraźne i silne dźwięki.

1 września 1923

Telefon głośnomówiący w szpitalu

W jednej z francuskich klinik zastosowano telefon głośnomówiący w charakterze pomocy naukowej do nauczania studentów-medyków podczas operacji chirurgicznych. Urządzenie polega na tem, że nad salą operacyjną znajduje się szklany sufit, przez który zebrani studenci obserwują każdy ruch chirurga przy operacji; profesor natomiast, dokonując swych czynności, mówi przez mikrofonem. Wzbudzony w mikrofonie prąd amplifikuje się za pomocą baterji z 10 lamp katodowych i zasila kilka telefonów głośnomówiących na sali górnej. W ten sposób studenci otrzymują wszystkie objaśnienia podczas wykładu pokazowego.

1 września 1923

Poznaj całą serię

Zupełnie nowa edukacyjna seria kitów AVTEDU. Wypróbuj je wszystkie i zostań mistrzem lutownicy, poznaj świat elektroniki i zgłębiaj go razem z nami

#AVTEDU #NaukaLutowania #KityAVT

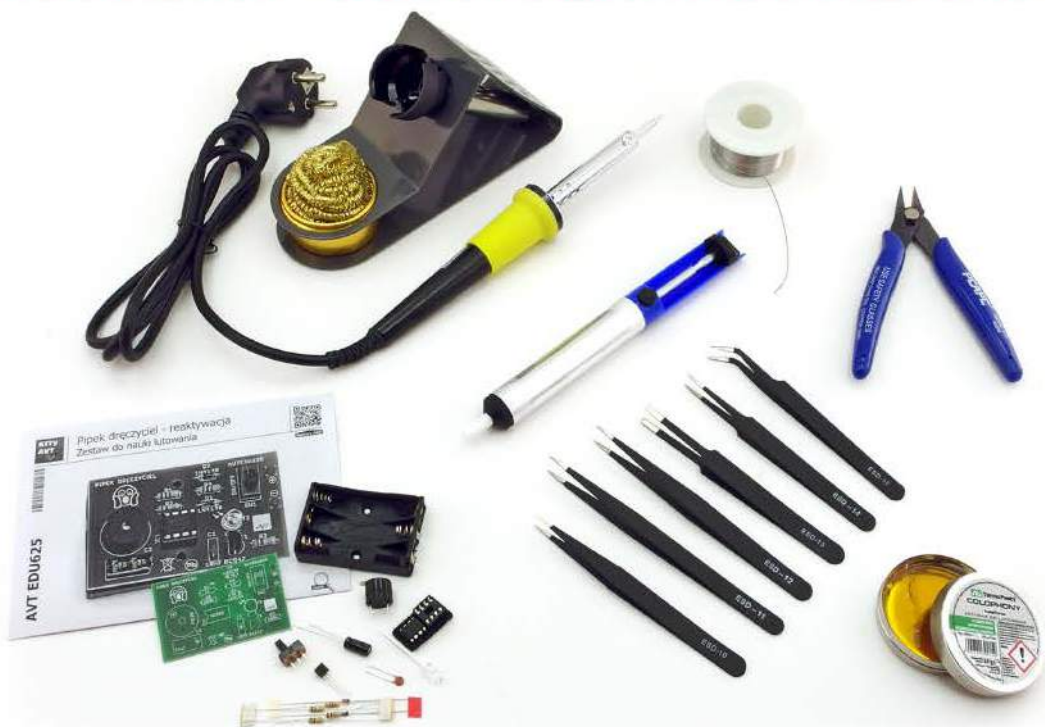
Zestaw umożliwiający rozpoczęcie nauki techniki lutowania elementów elektronicznych. Wraz z serią kitów AVTEDU tworzy idealne uzupełnienie zagadnienia montażu prostych urządzeń elektronicznych.

Zestaw zawiera **lutownicę**, wysokiej jakości **podstawkę** z czyścikiem, **cynę** z topnikiem, **kalafonię**, **pęsety**, **odsysacz** do cyny oraz **szczypce** tnące boczne.

W komplecie na dobry początek znajduje się również **zestaw AVTEDU do zlutowania**.



AVTEDUSTART - zestaw narzędzi do nauki lutowania



sklep.avt.pl

AVT SPV Sp. z o.o., 03-197 Warszawa, ul. Leszczynowa 11
tel.: (22) 257 84 51 e-mail: handlowy@avt.pl

tel.: (22) 86406243