



nr 8. sierpień 2024

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

TAJEMNICE I ZAGADKI

Na co nauka nie ma odpowiedzi



ISSN 0462-9760 Indeks 365408

9 47704621976243

cena: **14,90 zł** (w tym 8% VAT)

Górnictwo pozaziemskie

Kopiąc wśród gwiazd



Prenumerata

oszczędzasz 20% • cieszysz się darmową dostawą

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

Zamów prenumeratę na www.UlubionyKiosk.pl



Temat okładkowy

Historia matematyki, fizyki innych nauk ścisłych to historia ciężkich zmagania z problemami, których od wieków nie da się rozgryźć. Mimo wielkich postępów, jakie poczyniono, nadal sobie z tymi największymi zagadkami nie radzimy.

Czy liczby pierwsze będą ostatnimi?

Zaszczyśliśmy z naszą nauką i techniką daleko. Niedługo zamykamy pierwsze ćwierćwiecze XXI wieku. Jednak wciąż nie umiemy policzyć tak wielu liczb, rozwiązać tylu fundamentalnych zagadek fizyki czy biologii, odpowiedzieć na te wszystkie pytania bez odpowiedzi, niekiedy od wieków. Ogromny obszar rzeczywistości jest dla nas wciąż tajemnicą. Jak jest ogromny, też zresztą nie wiemy.

Czy, na przykład, liczby pierwsze będą ostatnimi, których tajemnicę rozwikłamy? Nie rozgryźliśmy jeszcze wielu uchodzących za łatwiejsze problemów matematycznych, z którymi warto byłoby sobie poradzić, zanim zmierzmy się z zagadką liczb pierwszych. Jednak człowiek ma taką naturę, że zabiera się do wszystkich rzeczy na raz, od poszukiwania wszystkich miejsc po przecinku w π po enigmatyczną hipotezę Riemanna.

Niektóre zagadki są z nami od wieków i wciąż nie znamy odpowiedzi

Fizyka może się „pochwalić” (w cudzysłowie, bo to ironiczne w tym kontekście słowo) bardzo obszernym katalogiem nierozwiązanych problemów. Modny dzięki książce i serialowej ekranizacji „pro-

blem trzech ciał” kogoś mniej wtajemniczonego może zaskakiwać, bo w potocznym rozumieniu jest to fizyka newtonowska, stara i, wydawałoby się, opanowana. Przykład ten wskazuje, że nie trzeba sięgać do problematyki kosmologicznej, ciemnych energii i materii czy unifikacji fundamentalnych oddziaływań, by znaleźć nierozwiązane fizyczne zagadki.

Piszemy w tym numerze „Młodego Technika” o sztucznej inteligencji, która od kilku lat zмага się z wyzwaniem naukowymi. Ma pewne sukcesy. Program AlphaFold firmy DeepMind np. świetnie radzi sobie z trudnymi problemami modelowania struktur protein. Od czasu do czasu pojawiają się doniesienia o rozwiązywaniu niektórych problemów matematycznych czy o odkryciach w mechanice kwantowej z udziałem AI. Nic jednak nie słyhać o jakichkolwiek sukcesach algorytmów w zmaganiach z fundamentalnymi problemami klasy liczb pierwszych itp.

Może to kwestia rozwoju techniki AI, a może powinniśmy mimo wszystko wciąż liczyć na geniuszy w ludzkiej skórze, a nie na maszyny. Tak czy inaczej, musimy czekać, jak czekamy od pokoleń.

Mirosław Usidus

Spis treści

Temat numeru: Tajemnice i zagadki. Na co nauka nie ma odpowiedzi

- 24 • Czy AI rozwiąże nierozwiązywalne?
I czy ta odpowiedź nam cokolwiek da?
Algorytmy coraz bardziej pomocne,
ale z Noblami jeszcze się wstrzymajmy
- 30 • Większość liczb rzeczywistych jest nieznana
– nawet matematykom.
Czy coś, czego nie da się policzyć, istnieje?
- 34 • Lista nieznaną od dawna rozwiązania
zagadek fizyki jest bardzo długa.
Nie tylko problem trzech ciał
- 39 • Wytoniło się i wyłania nadal, a my nie mamy
pojęcia – jak i dlaczego? Tajemnica życia

Technika

- 8 Info Zoom
- 16 Dodaj do obserwowanych
Horyzonty mgłą spowite
- 17 • Nanotechnologie w fotowoltaice. Jak pomóc
elektronom
- 20 • Świat kosmicznych dziwów nie przestaje zadziwiać.
Matuzalem, wielkie bum i zagadkowa pustka
- 22 • Hakowanie nowoczesnych samochodów dla funkcji
premium i nie tylko. Włamać się do auta, by
posiedzieć na podgrzewanym fotelu
- 44 Czekając na aerostrady przyszłości. Rozmowa
z Tomaszem Patanem, założycielem, szefem
ds. technologii, i wynalazcą Jetson ONE
- 50 Górnictwo pozaziemskie. Kopać wśród gwiazd
- 54 Nasi idole – liderzy innowacji: Z Tajwanu na podbój
świata – Jen-Hsun „Jensen” Huang

Powrót do przyszłości

- 57 Fantastyka naukowa znów w „Młodym Techniku”
- 58 Nowy dom

m.technik

- 60 e-Technologie: Ataki bez plików i inne nowe
cyberzagrożenia. Wiedza chroni najlepiej

Szkoła

- 63 Edukacja przez szachy: Dawid Przepiórka
- 68 Fizyka bez granic: Światło elektryczne.
Dlaczego żarówka świeci?
- 70 MT studiuje: Cyberbezpieczeństwo
- 72 Chemia inna niż w szkole: Błędne ścieżki prowadzą
do celu (2)
- 76 Matematyka z ludzką twarzą: Twierdzenie Halla
o małżeństwach
- 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: Budujemy wyścigowe modele
„dragsterów” z napędem sprężynowym
- Odkryj historię wynalazków
- 92 • Lody
- 96 • Typy łodów i desrów lodowych

Hobby

- 97 Akademia audio: Klipsch JUBILEE (2)

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



44

Rozmowa z Tomaszem Patanem, założycielem, szefem ds. technologii, i wynalazcą Jetson ONE

W tym wydaniu MT m.in.:

- **Horyzonty mgłą spowite: Nanotechnologie w fotowoltaice**
Światło słoneczne dociera na Ziemię silnie „rozcięćzone”, więc energia padająca na jednostkę powierzchni ogniwa jest stosunkowo niska. Nanotechnologie zwiększają ich sprawność
- **Nasi idole: Jen-Hsun „Jensen” Huang**
Jako współzałożyciel i wieloletni szef firmy NVIDIA, jest jednym z najważniejszych „sprawców” rewolucji w przetwarzaniu grafiki komputerowej, a od niedawna także w dziedzinie AI.
- **eTechnologie: Ataki bez plików i inne nowe cyberzagrożenia**

• 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 Reklamy:
e-mail: reklama@mt.com.pl

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



24

Tajemnice i zagadki. Na co nauka nie ma odpowiedzi

W nauce, matematyce i fizyce, mamy wiele problemów, które pozostają nierozwiązane od dziesięcioleci, a czasem nawet stuleci. Problemy te są często tak złożone, że nawet najbystrzejsze umysły mają trudności ze znalezieniem rozwiązań. Obecnie próbujemy zaprząć do ich rozwiązania sztuczną inteligencję. Czy AI da radę? To się okaże.

List miesiąca

Kontrowersyjne metody CRISPR

W odniesieniu do artykułu, który ukazał się w czerwcowym wydaniu „Młodego Technika”, pozwalam sobie przedstawić Państwu szersze spojrzenie na technologię edycji genów CRISPR (Clustered Regularly Interspaced Short Palindromic Repeats), która w ostatnich latach zrewolucjonizowała badania genetyczne i biotechnologię, jednocześnie budząc wiele kontrowersji natury etycznej i społecznej. Jest to niezwykle istotny temat naukowy i społeczny, który zasługuje na pogłębioną dyskusję.

CRISPR to rewolucyjna technologia, która umożliwia precyzyjną modyfikację sekwencji DNA. Od momentu jej odkrycia w 2013 roku zyskała ona ogromne zainteresowanie w środowisku naukowym, otwierając nowe horyzonty w takich obszarach jak terapie genowe, inżynieria tkankowa, hodowla roślin i zwierząt, a także badania podstawowe.

CRISPR umożliwia naukowcom precyzyjne celowane wprowadzanie zmian w genach organizmów. Technika ta opiera się na wykorzystaniu białka Cas9 (lub jego wariantów), które działa jak „nożyczki molekularne”, oraz krótkiego fragmentu RNA, który naprowadza białko na określoną sekwencję DNA. CRISPR pozwala na usuwanie, modyfikowanie lub wstawianie fragmentów DNA z niespotykaną dotąd precyzją i efektywnością.

Kluczową zaletą CRISPR jest jej wysoka skuteczność i stosunkowo niska złożoność w porównaniu do wcześniejszych metod inżynierii genetycznej. Umożliwia ona dokonywanie precyzyjnych zmian w obrębie genomu, co stwarza ogromne możliwości terapeutyczne. Już teraz CRISPR jest z powodzeniem wykorzystywana w leczeniu takich schorzeń jak mukowiscydoza, hemofilia czy dystrofia mięśniowa Duchenne'a. W rolnictwie CRISPR może pomóc w tworzeniu roślin odpornych na szkodniki czy zmiany klimatu. W badaniach podstawowych umożliwia głębsze zrozumienie funkcji genów i mechanizmów biologicznych.

Jednocześnie technologia ta budzi ogromne kontrowersje natury etycznej, prawnej i społecznej. Wśród nich znajdują się obawy dotyczące potencjalnych skutków modyfikacji linii zarodkowej człowieka, możliwości nadużyć i nierównego dostępu do terapii genowych. Pojawiają się także pytania o granice ingerencji w ludzki genom i konsekwencje takich działań dla przyszłych pokoleń.

Oto lista najczęściej podnoszonych zastrzeżeń i wątpliwości wobec tej metody:

- **Kwestia bezpieczeństwa i niezamierzonych skutków.** Mimo wysokiej precyzji CRISPR, istnieje ryzyko wystąpienia nieplanowanych mutacji poza docelowym miejscem w genomie. Długoterminowe konsekwencje takich zmian są trudne do przewidzenia.
- **Modyfikacje linii zarodkowej.** Edycja genów w komórkach rozrodczych lub wczesnych zarodkach budzi szczególne kontrowersje, gdyż zmiany te byłyby dziedziczone przez przyszłe pokolenia.
- **Ulepszanie człowieka.** Możliwość modyfikacji cech ludzkich, takich jak inteligencja czy wygląd, rodzi obawy o tworzenie „zaprojektowanych dzieci” i pogłębianie nierówności społecznych.
- **Kwestie bioetyczne.** Manipulacja materiałem genetycznym stawia fundamentalne pytania o granice ingerencji człowieka w naturę i definicję życia.
- **Regulacje prawne.** Brak jednolitych, międzynarodowych regulacji dotyczących stosowania CRISPR stwarza ryzyko niekontrolowanego rozwoju tej technologii.
- **Dostępność i patenty.** Istnieją obawy, że patenty na technologię CRISPR mogą ograniczyć jej dostępność i zastosowanie w badaniach naukowych.
- **Biobezpieczeństwo.** Technika CRISPR może potencjalnie zostać wykorzystana do tworzenia broni biologicznej lub modyfikowania organizmów w sposób zagrażający ekosystemom.



W świetle tych kontrowersji kluczowe jest prowadzenie szerokiej debaty społecznej na temat etycznych aspektów edycji genów. Niezbędne jest wypracowanie międzynarodowych standardów i regulacji prawnych, które zapewnią odpowiedzialne wykorzystanie tej technologii. Jednocześnie ważne jest, aby obawy i kontrowersje nie hamowały rozwoju badań nad CRISPR. Technika ta ma ogromny potencjał w zakresie poprawy ludzkiego zdrowia, rozwiązywania problemów środowiskowych czy zwiększania bezpieczeństwa żywnościowego.

Jerzy Zwierski, Siemianowice Śląskie

Kosmologia

Dziękuję za poruszenie na łamach waszego pisma kwestii problemów, z jakimi boryka się współczesna kosmologia. Ostatnie odkrycia dokonane dzięki Kosmicznemu Teleskopowi Jamesa Webba (JWST) rzucają nowe światło na nasze rozumienie wczesnego Wszechświata, stawiając pod znakiem zapytania niektóre z fundamentalnych założeń dotychczas akceptowanego modelu kosmologicznego.

Standardowy model kosmologiczny, znany jako model Lambda-CDM (Lambda Cold Dark Matter), przez dekady z powodzeniem opisywał obserwowane właściwości Wszechświata. Jednakże dane z teleskopu Webba sugerują, że może on wymagać istotnych modyfikacji lub nawet całkowitego przemyślenia. Jednym z kluczowych problemów jest kwestia wczesnego formowania się galaktyk. Zgodnie z dotychczasowym modelem, pierwsze galaktyki powinny być stosunkowo małe i młode w odległych, wczesnych epokach Wszechświata. Tymczasem JWST odkrył zaskakująco masywne i dojrzałe galaktyki w okresie zaledwie kilkuset milionów lat po Wielkim Wybuchu. Te obserwacje sugerują, że proces formowania się galaktyk mógł przebiegać znacznie szybciej i efektywniej, niż przewidywały to nasze modele. Odkrycie to sugeruje też, że albo początkowe fluktuacje gęstości we wczesnym Wszechświecie były znacznie większe, niż dotychczas sądziliśmy, albo procesy formowania się gwiazd i galaktyk przebiegały o wiele efektywniej. Obie te możliwości wymagają istotnych modyfikacji naszych modeli ewolucji kosmicznej.

Wielkim wyzwaniem jest również problem ciemnej materii. Standardowy model zakłada istnienie znacznych ilości zimnej ciemnej materii, która miała odgrywać kluczową rolę w formowaniu się struktur kosmicznych. Jednakże obserwacje JWST wskazują na możliwość, że rola ciemnej materii mogła być przeceniana. Niektóre z obserwowanych galaktyk wydają się być bardziej „czyste”, bardziej pozbawione ciemnej materii, niż przewidywano. Obserwacje sugerują również, że albo ciemna materia zachowuje się inaczej, niż przewidywaliśmy (np. może być „cieplejsza” lub bardziej interaktywna), albo jej rola w formowaniu się struktur kosmicznych jest mniejsza,

niż sądziliśmy. To z kolei otwiera drogę dla alternatywnych teorii, takich jak zmodyfikowana grawitacja (MOND) czy inne formy egzotycznej materii.

Ponadto teleskop Webba dostarczył danych sugerujących możliwość istnienia alternatywnych scenariuszy dla wczesnej inflacji kosmicznej. Obserwacje struktury i rozkładu materii w najwcześniejszych dostępnych nam epokach kosmicznych nie w pełni odpowiadają przewidywaniom teorii inflacji w jej obecnej formie.

Warto dodać, że obserwacje JWST mogą mieć również wpływ na tzw. problem stałej Hubble’a. Od lat obserwujemy rozbieżność między wartościami stałej Hubble’a (opisującej tempo ekspansji Wszechświata) mierzonymi różnymi metodami. Dane z JWST mogą pomóc w rozwiązaniu tego problemu, dostarczając precyzyjnych pomiarów dla bardzo odległych obiektów. Jednakże, jeśli rozbieżność utrzyma się lub nawet pogłębi, może to wskazywać na fundamentalne problemy w naszym rozumieniu ekspansji Wszechświata.

Te odkrycia stawiają przed nami fascynujące, ale i trudne pytania. Czy nasze fundamentalne założenia dotyczące natury ciemnej materii i ciemnej energii są poprawne? Czy potrzebujemy nowej fizyki, aby wyjaśnić obserwowane zjawiska? Jak moglibyśmy zmodyfikować nasz model kosmologiczny, aby uwzględnić te nowe dane?

Warto podkreślić, że sytuacja ta nie jest kryzysem, lecz ekscytującą możliwością rozwoju naszego rozumienia Wszechświata. Historia nauki pokazuje, że takie momenty często prowadzą do przełomowych odkryć i nowych teorii.

Potrzebujemy nowych modeli teoretycznych, które będą w stanie wyjaśnić zarówno dotychczasowe obserwacje, jak i nowe dane z JWST. To może wymagać fundamentalnego przemyślenia niektórych aspektów fizyki, od skali kwantowej po kosmologiczną. Jednocześnie musimy pamiętać o zachowaniu ostrożności i rygorystycznego podejścia naukowego. Pojedyncze obserwacje, nawet tak przełomowe jak te z teleskopu Webba, muszą być potwierdzone i zweryfikowane przez niezależne badania i instrumenty.

Tadeusz Huk z Bydgoszczy



AKWANAUTYKA

Morski dron DARPA jak płaszczka

Pierwsze testy w morzu ma za sobą prototyp zrobotyzowanej łodzi podwodnej, a mówiąc bardziej technicznie, bezzałogowego pojazdu podwodnego (UUV), opracowany przez amerykańską agencję DARPA. Pojazd nazwany Manta Ray przeszedł pierwsze podwodne próby w pobliżu pacyficznych wybrzeży południowej Kalifornii.

Projekt jest dość tajemniczy. Niewiele wiadomo na temat szczegółów konstrukcyjnych i specyfikacji Manta Ray. Ze zdjęć udostępnionych z miejsca, w którym odbywały się testy, wynika, że prototyp ma kilka metrów długości i szerokości. Ma płaski kształt nawiązujący do płaszczkowatego stworzenia morskiego, do którego nawiązuje jego nazwa. Systemy napędowe są, jak wynika z relacji, schowane w korpusie konstrukcji.

Według głównego wykonawcy prototypu, firmy Northrop Grumman, Manta Ray ma modułową konstrukcję i energooszczędne systemy, w tym zdolność do zakotwiczenia się na dnie morskim i przejścia w tryb hibernacji. Na pokładzie znajdują się podwodne systemy napędowe o niskim poborze mocy, nowe czujniki do podwodnego wykrywania i klasyfikowania zagrożeń i niebezpieczeństw, wysokowydajne systemy nawigacji i dowodzenia oraz kontroli, a także rozwiązania zapobiegające porastaniu i degradacji materiałów. ■



Czwarty testowy lot wielkiej rakiety Starship firmy SpaceX miał po raz pierwszy jednoznacznie udany przebieg od początku, czyli od startu, przez lot na orbitę, aż po powrót z kontrolowanym lądowaniem głównych członów w oceanicznych wodach. Firma osiągnęła kamień milowy w postaci powrotu na Ziemię silnika i górnego stopnia. Elon Musk, szef SpaceX, powiedział: „Pomimo utraty wielu płytek [osłonowych – red.] i uszkodzonej klapy, Starship zdołał miękko wylądować”. Niestety nie można tego samego powiedzieć o misji kapsuły Starliner firmy Boeing, która, choć w końcu wystartowała i dowiozła astronautów do Międzynarodowej Stacji Kosmicznej, to ze względu na problemy techniczne nie mogła wrócić w zaplanowanym czasie.

Po wielu problemach i wielokrotnym przekładaniu lotu statek kosmiczny Starliner z dwójką astronautów, Butchem Wilmore i Suni Williams, na pokładzie wystartował. Próba od momentu startu przebiegała pomyślnie i kapsuła znalazła się na orbicie okołozemskiej, na trajektorii w kierunku Międzynarodowej Stacji Kosmicznej, gdzie po zadokowaniu załoga



EKSPLORACJA KOSMOSU

Udany test Starshipa i problemy kapsuły Boeinga

Starlinera miała zaplanowany tygodniowy pobyt. Astronauci informowali podczas lotu o dwóch nowych wyciekach helu z układów paliwowych statku. Jednak uznano, że w danym momencie nie zagrażają one misji. Potem jednak pojawiły się kolejne doniesienia nie tylko o wyciekach, ale również o kłopotach z pędnikiem. Ostatecznie NASA przełożyła powrót astronautów na czas nieokreślony, a media pisały, że utknęli oni na ISS. Do momentu zamknięcia niniejszego wydania MT sytuacja nie została wyjaśniona.

Nieco wcześniej chińska bezzałogowa sonda Chang'e-6 wylądowała po niewidocznej z Ziemi stronie Księżyca. Jej misja miała również charakter przełomowy. Pobrała pierwsze na świecie próbki

skał i gleby z nigdy nieoglądanej z Ziemi strony naszego satelity. Statek Chang'e-6, wyposażony w szereg instrumentów i własny system startowy, wylądował w gigantycznym kraterze uderzeniowym Aitkena. Używając czerpaka i wiertła, Chang'e-6 miał za zadanie zebrać 2 kg materiału księżycowego w ciągu dwóch dni i przywieźć go z powrotem na Ziemię. Misja się powiodła i statek powrócił na Ziemię z pobranymi próbkami. Chińscy naukowcy przewidują, że przywieziony materiał będzie zawierać skały wulkaniczne sprzed 2,5 miliona lat, które pozwolą wyjaśnić różnice geologiczne pomiędzy widoczną i niewidoczną stroną Księżyca. ■



TELEKOMUNIKACJA

Łączność satelitarna przez Bluetooth

Hubble Network, start-up z siedzibą w Seattle, nawiązał pierwsze w historii połączenie Bluetooth bezpośrednio z satelitą. Potwierdzono, że pokładowe chipy Bluetooth 3,5 mm, umieszczone w dwóch satelitach firmy wystrzelonych na rakiecie SpaceX, odbierają sygnał z odległości ponad sześciuset kilometrów.

Firma stosuje na pokładzie satelitów opatentowaną antenę typu phased array, która może zostać umieszczona na małym satelicie. Anteny te umożliwiają komunikację zwykłego chipu Bluetooth z satelitą Hubble. Trzeba było również pokonać problemy związane z efektami dopplerowskimi, czyli niedopasowaniem częstotliwości występującym między szybko poruszającymi się obiektami wymieniającymi dane za pośrednictwem fal radiowych.

Według firmy, podłączenie dowolnego gotowego urządzenia Bluetooth do sieci satelitarnej to potencjalnie sieć o globalnym zasięgu przy dwudziestokrotnie mniejszym zużyciu baterii i pięćdziesięciokrotnie niższych kosztach operacyjnych niż inne sieci satelitarne w tym Starlink. Hubble Network informuje, że współpracuje już we wdrażaniu pilotażowych usług z klientami z różnych sektorów, w tym urządzeń konsumenckich, budownictwa, infrastruktury, łańcucha dostaw, logistyki, ropy i gazu oraz obronności. Start-up planuje wystrzelić jeszcze w tym roku kolejne dwa, a w czwartym kwartale 2025 r. lub na początku 2026 r., jeszcze trzydzieści dwa, co ma utworzyć konstelację nowej sieci. ■



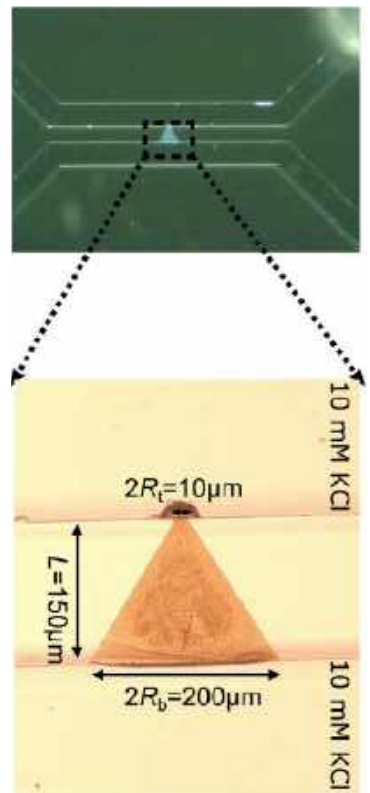
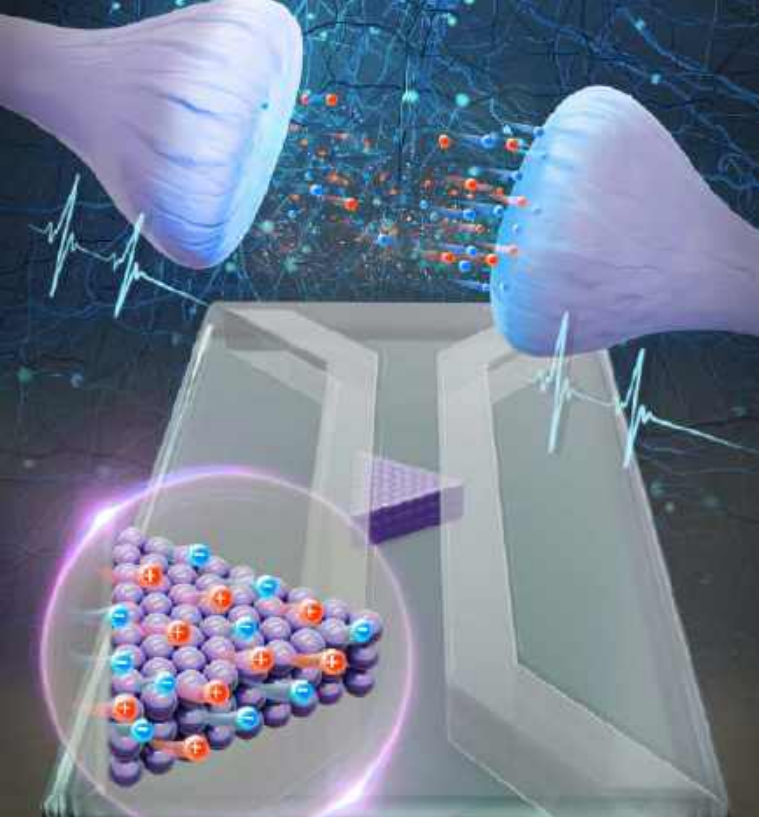
TECHNIKA MATERIAŁOWA

Sztuczne diamenty produkowane bez wielkich ciśnień i temperatur

Naukowcy z Instytutu Nauk Podstawowych w Korei Południowej opracowali nową technikę syntezy sztucznych diamentów, która nie wymaga ekstremalnych ciśnień i wysokich temperatur, znanych z procesów produkcji tych szlachetnych kamieni obecnie. Według publikacji w „Nature”, wytwarzanie drogocennych kamieni odbywa się przy ciśnieniu atmosferycznym i bez konieczności użycia diamentu załączkowego.

Proces zastosowany przez badaczy z Korei zaczyna się od elektrycznego podgrzania galu z dodatkiem krzemu w grafitowym tyglu. Gal, jak wynikało z poprzednich badań, może katalizować tworzenie grafenu z metanu. Grafen, podobnie jak diament, jest czystym węglem. Naukowcy umieścili tygiel w komorze, przez którą przepływa metan, gaz będący związkiem węgla z wodorem. Doszli do wniosku, że mieszanina galu, niklu i żelaza, w połączeniu ze szczyptą krzemu bardzo dobrze katalizuje wzrost diamentów. W praktyce założenia się potwierdziły i zespół uzyskał diamenty w tyglu już po 15 minutach.

Mechanizm, który powoduje formowanie diamentów, nie jest jeszcze dokładnie wyjaśniony. Badacze uważają, że przyczyną przechodzenia węgla z metanu w postać kryształów diamentowych jest spadek temperatury. Diamenty nie powstają bez krzemu, więc naukowcy sądzą, że może on działać jako załączek, wokół którego krystalizuje węgiel tworzący kryształy diamentu. Nie trzeba stosować wysokich ciśnień ani temperatur, jednak w tej chwili powstające tą drogą kamienie szlachetne są bardzo drobne i zawierają krzemowe zanieczyszczenia. Jednak gdy mówimy o diamencie w technicznych zastosowaniach, niekoniecznie są to duże problemy. ■



MÓZG

Symulowane synapsy mózgowe na bazie roztworu wody i soli

Naukownicy z uniwersytetów w Utrechcie w Holandii i Sogang w Korei Południowej stworzyli funkcjonującą komórkę mózgową na bazie mieszanki tych samych składników, substancji opartych na wodzie i soli, których używa mózg. W ten sposób symulowane są synapsy mózgowe. System nazwany „jonotronicznym” memrystorem „zapamiętuje”, ile ładunku elektrycznego wcześniej przez niego przepłynęło, co zdaniem uczonych jest bliskie sposobowi, w jaki funkcjonuje nasz mózg.

Jak wynika z opisu, który ukazał się w czasopiśmie „PNAS”, mający kształt stożka z roztworem wody i soli wewnątrz, jonotroniczny memrystor ma wymiary 150 na 200 mikrometrów. Impulsy elektryczne powodują ruch jonów przez stożkowy kanał, przy czym zmiany ładunku elektrycznego prowadzą do zmian w ruchu jonów. Zmiany w sposobie przewodzenia

przez synapsę energii elektrycznej mogą być mierzone i dekodowane w celu odtworzenia sygnału wejściowego, co stanowi rodzaj pamięci.

Badacze podkreślają, że wciąż są na bardzo wczesnym etapie, jeśli chodzi o konstrukcję i działanie tego urządzenia. Także ogólnie cała dziedzina jonotroniki stawia dopiero pierwsze kroki. Jednak charakterystyka wpływu długości kanału na czasową trwałość pamięci w memrystorze sugeruje, że kanały mogą być dostosowane do określonych zadań, czyli mogą być elastyczne, podobnie jak synapsy i neurony w mózgu. Fizyków czeka również praca nad tym, w jaki sposób te syntetyczne synapsy można by łączyć. System jest stosunkowo szybki i tani w produkcji i teoretycznie można go skalować w produkcji, pod warunkiem że jego przydatność zostanie udowodniona w dalszych testach i badaniach naukowych. ■



AERONAUTYKA

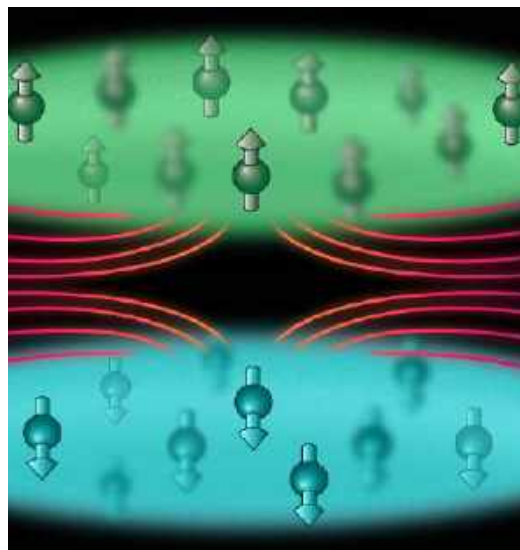
Samoloty pokryte skórą rekina oszczędzają paliwo

Linie lotnicze SWISS poinformowały, że ukończyły projekt wyposażenia zewnętrznych części kadłubów swojej floty samolotów Boeing 777-300ER w oszczędzającą paliwo folię imitującą skórę rekina. Ta, nazwana AeroSHARK, technika została zaprojektowana w celu odtworzenia hydrodynamicznych właściwości skóry tej ryby.

Umieszczana na poszyciu przezroczysta folia ma „żeberka” o mikrometrowych rozmiarach, ułożone zgodnie z kierunkiem przepływu powietrza. Jak twierdzą przedstawiciele firmy SWISS, zastosowanie tej warstwy nakładanej na poszycie samolotów obniżyło zużycie paliwa o ponad dwa tysiące ton w ubiegłym roku, co miało zmniejszyć emisję CO₂ o 7100 ton.

Nałożenie warstw folii AeroSHARK w określonych, precyzyjnie wybranych miejscach na powierzchni samolotu, zajmuje specjalnie przygotowanemu personelowi ok. tygodnia. Folia została opracowana wspólnie przez Lufthansa Technik i koncern BASF. SWISS rozważa rozszerzenie jej zastosowania na kolejne samoloty długodystansowe. ■

308 obrotów na minutę (ponad 5 na sekundę) – z taką prędkością wirowała rekordzistka tyżwiarzkiego piruetu, Natalia Kanounnikowa.



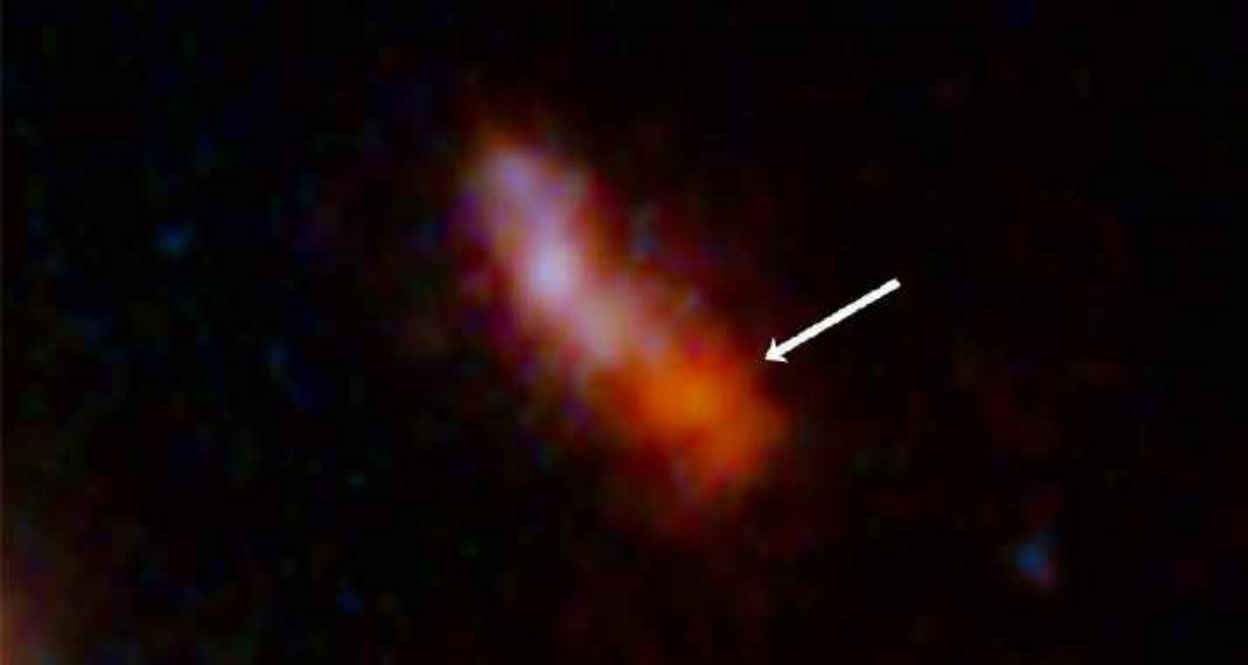
FIZYKA

Atomy w nieznanej dotąd bliskości

Naukowcy wymusili zbliżenie dwu warstw ultrazimnych atomów magnetycznych na odległość pięćdziesięciu nanometrów – to dziesięć razy bliżej, niż kiedykolwiek udało się poprzednio. Według publikacji na temat tego eksperymentu, która ukazała się w „Science”, pojawiły się niezwykle efekty kwantowe.

Badacze byli zainteresowani zachowaniem atomów dysprozu, które są wyjątkowe na tle innych metali w tym sensie, że mogą oddziaływać ze sobą na duże odległości jako dipol-dipol (słabe siły przyciągające między cząstkowymi ładunkami na sąsiednich atomach). Po utworzeniu dwuwarstwowego układu atomów tego pierwiastka, zespół badawczy za pomocą wiązek laserowych podgrzał jedną z warstw, oddzieloną od drugiej szczeliną próżniową. Zaobserwowano przenoszenie ciepła do drugiej warstwy przez pustą przestrzeń. Podejrzewa się, że przenoszenie ciepła ma związek z kwantową interakcją dipoli.

Transfer ciepła to tylko jeden z interesujących efektów. Uczeń zainteresowali się też innym obserwowanym efektem kwantowym, zwanym parowaniem Bardeena-Coopera-Schrieffera (BCS), którego doświadczają cząstki subatomowe zwane fermionami w niskich temperaturach. Efekty też są związane ze zjawiskami nadprzewodnictwa elektrycznego materiałów, zatem ich dokładniejsze poznanie może dać wiele w tej dziedzinie. ■



ASTRONOMIA

Najodleglejsza i najstarsza odkryta galaktyka burzy ustalone poglądy

Jak podała agencja NASA, Kosmiczny Teleskop Jamesa Webba odkrył najbardziej odległą z dotąd znanych galaktykę we Wszechświecie. Obiekt nazwany JADES-GS-z14-0 uformował się najpóźniej 290 milionów lat po Wielkim Wybuchu. Jak na tak wczesną datę powstania, galaktyka jest stosunkowo duża – ma 1600 lat świetlnych średnicy. Jest też niezwykle jasna, co stanowi kolejną anomalię z punktu widzenia znanych teorii naukowych.

Długości fal światła emitowanego przez JADES-GS-z14-0, zarejestrowane przez instrument MIRI (Mid-Infrared Instrument) na pokładzie JWST, wskazują na obecność silnej emisji zjonizowanego gazu, co wynika prawdopodobnie z obfitości wodoru i tlenu.

Zdaniem naukowców to też dziwne, ponieważ tlen nie występuje w takich ilościach na wczesnym etapie życia galaktyki. Sugeruje to, istnienie pokoleń bardzo masywnych gwiazd. To nie jest zgodne z modelami i teoriami, które zostały wcześniej opracowane.

„Wszystkie te obserwacje zebrane razem mówią nam, że JADES-GS-z14-0 nie przypomina galaktyk, których istnienie w bardzo wczesnym Wszechświecie przewidywały modele teoretyczne i symulacje komputerowe”, piszą naukowcy w komunikacie NASA. „Jest prawdopodobne, że w ciągu następnej dekady astronomowie za pomocą Webba znajdą wiele takich silnie świecących galaktyk, być może nawet z jeszcze wcześniejszych okresów”. ■

30000 km/s czyli ok. 1/10 prędkości światła
wynosi prędkość orbitalna gwiazdy S62 w gromadzie
okrążającej Sagittarius A*, supermasywną czarną
dziurę w centrum Drogi Mlecznej.



NOWE MATERIAŁY

Beton wzmocniany tym, co pozostaje po kawie

Fusy po kawie nie tylko z powodzeniem można dodawać do cementu, ale poprawiają one właściwości wyjściowego materiału – ogłosili po serii eksperymentów badacze z australijskiego Uniwersytetu RMIT. Rajeev Roychand, główny autor badań podaje, że to wzmocnia beton o 30 proc. i to przy wykorzystaniu niskoenergetycznego beztlenowego procesu w temperaturze 350 stopni Celsjusza.

Według pracy na ten temat, opublikowanej w „Journal of Cleaner Production”, fusy nie mogą być dodawane do betonu bez przetworzenia, gdyż wydzielają substancje szkodliwe dla jakości materiału budowlanego. Dlatego potrzebny jest beztlenowy proces pirolizy, który rozkłada cząsteczki organiczne, w wyniku czego powstaje porowaty materiał podobny do węgla drzewnego, który może tworzyć wiązania, a tym samym wbudowywać się w strukturę cementu.

Jak się szacuje, na świecie powstaje w wyniku konsumpcji kawy ok. dziesięciu milionów ton fusów rocznie. Fusy oczywiście nie są w stanie zastąpić piasku w produkcji betonu, którego zużywa się na świecie w celach budowlanych pięćdziesiąt miliardów ton rocznie, jednak jeśli wzmocniające materiał ich właściwości oraz trwałość takiego produktu zostanie potwierdzona w praktycznych zastosowaniach, to może być cenny dodatek. ■

REKREACJA

Rower dla pływaków i nurków

Francuska firma Seabike opracowała urządzenie wspomagające pływaków, które wykorzystuje siłę ich nóg do przyspieszania w wodzie ze znacznie większą prędkością, niż jest w stanie w wodzie osiągnąć pływający człowiek.

Napęd w tym wynalazku pochodzi z układu, w którym człowiek pedałuje w sposób zbliżony do ruchów rowerzysty.

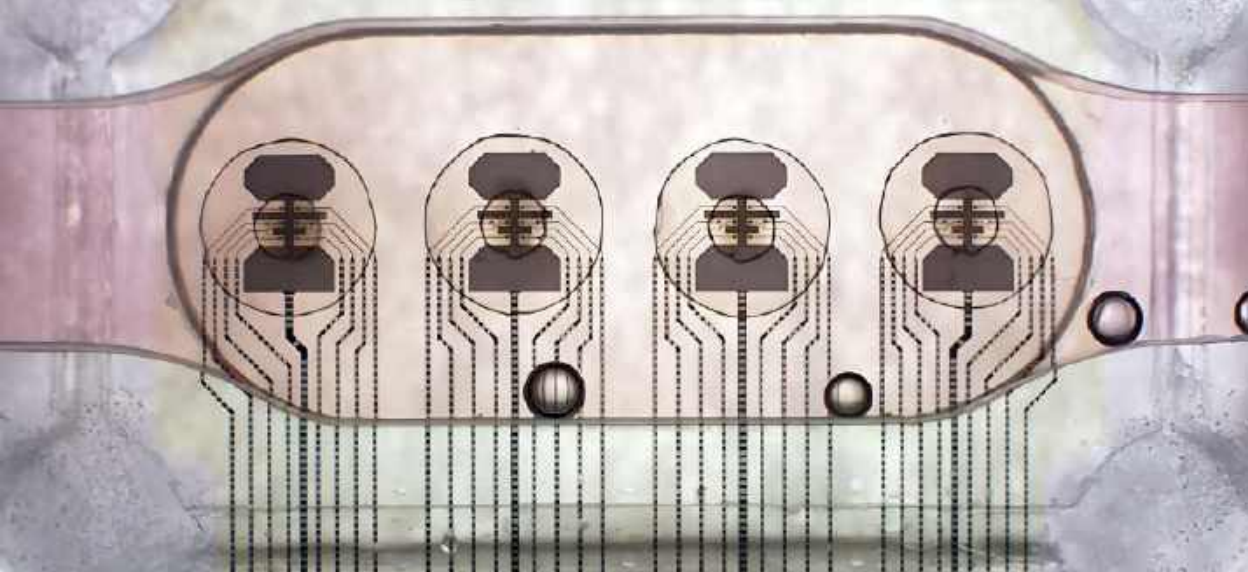
Urządzenie składa się z czegoś, co można by nazwać „dyszlem”, który do pływającego mocowany jest za pomocą pasa powyżej bioder. Na drugim końcu dyszla, za pływakiem, znajduje się śruba napędowa o średnicy 38 centymetrów. Ta jest napędzana przez korbę obracaną za pomocą dwóch pedałów przez osobę zanurzoną w wodzie. Śruba działa w obie strony – można więc odwrócić urządzenie i trzymać śmigło przed sobą, zamiast pedałów przymocować uchwyty na dłonie i napędzać urządzenie ręcznie, a nie nożnie.

Choć informacje o tym sprzęcie pojawiły się w mediach dopiero teraz, Seabike sprzedaje go już od około roku. Używa też tego rodzaju wspomagania dla uczestników imprez organizowanych dla nurków na południowym wybrzeżu Francji. W sieciach społecznościowych, np. na Instagramie, można obejrzeć już sporo filmów demonstrujących działanie urządzenia pod wodą. ■



Prezentacja urządzenia
wspomagającego pływanie:
<https://youtu.be/6btHaTwFKw>





BIOELEKTRONIKA

Oparty na komórkach mózgowych bioprocessor podłączony do internetu

FinalSpark, szwajcarski start-up z branży biokomputerów, uruchomił podłączoną do internetu platformę, która zapewnia zdalny dostęp do szesnastu organoidów ludzkiego mózgu. To pierwsze na świecie rozwiązanie zapewniające dostęp do biologicznych neuronów powstałych techniką in vitro. Według komunikatu, bioprocessory te „zużywają milion razy mniej energii niż tradycyjne jednostki cyfrowe”.

Główną innowacją jest w tym rozwiązaniu wykorzystanie czterech macierzy wieloelektrodowych (MEA), w których znajduje się żywa tkanka, organoidy będące trójwymiarowymi zgrupowaniami komórek tkanki mózgowej. Każda MEA zawiera cztery takie organoidy, połączone ośmioma elektrodami używanymi zarówno do stymulacji, jak i nagrywania. Dane są przesyłane tam i z powrotem za pośrednictwem przetworników cyfrowo-analogowych (kontroler Intan RHS 32)

o częstotliwości próbkowania 30 kHz i rozdzielczości 16 bitów. System wspierany jest przez mikroprzepływowy system podtrzymywania życia komórek i kamery monitorujące. Jest jeszcze specjalne oprogramowanie umożliwiające badaczom wprowadzanie zmiennych danych, a następnie odczytywanie i interpretowanie danych wyjściowych procesora.

Jak opisuje FinalSpark, jednostka nazywana Neuroplatform opiera się na architekturze, którą można sklasyfikować jako „wetware”, będący połączeniem sprzętu, oprogramowania i biologii. FinalSpark udostępnił swoją zdalną platformę obliczeniową dziewięciu instytucjom, a celem jest wsparcie dla badań i rozwoju bioprocessorów. Dzięki współpracy tych instytucji, FinalSpark ma nadzieję stworzyć pierwszy na świecie żywy procesor w pełnym tego słowa znaczeniu. ■

1 888 529 cyfr ma największa znana liczba pierwsza palindromiczna, czyli taka, która nie zmienia się, gdy jej cyfry dziesiętne zapiszemy w odwrotnej kolejności.



ENERGIA JĄDROWA

♦ Władze kanadyjskiej prowincji Saskatchewan zapowiadają, że do 2029 r. na ich terenie ma zostać uruchomiony kompaktowy reaktor jądrowy eVinci firmy Westinghouse o mocy 13 megawatów, zdolny do pracy przez osiem lat bez korzystania z wody do chłodzenia, gdyż stosowana jest w nim innowacyjna technologia „rurek cieplnych”, która, według konstruktorów, eliminuje potrzebę stosowania wody do chłodzenia systemu. ♦ Szwajcarska firma Transmutex opracowała technikę redukcji odpadów z elektrowni jądrowych nawet o 80%, która przez wykorzystanie specjalnie skonstruowanego akceleratora połączonego z reaktorem wytwarza uran z toru, przy czym to paliwo jądrowe wykorzystane następnie w reaktorze nie pozostawia za sobą wielu kłopotliwych odpadów, np. plutonu. ♦

SAMOCZODY ELEKTRYCZNE

♦ Chiński producent aut elektrycznych, BYD, zaprezentował hybrydową konstrukcję pojazdu z nowym typem układu napędowego, która, jak zapewnia firma, ma pozwolić na przejechanie prawie dwóch tysięcy kilometrów (z Warszawy np. na południe Włoch) bez ładowania lub tankowania. ♦ Naukowcy z Uniwersytetu Stanforda zaprezentowali na łamach „Nature” nową technikę eksploatacji akumulatorów litowo-metalowych polegającą na stosunkowo prostych zabiegach, takich jak rozładowanie baterii i pozostawienie jej na kilka godzin w stanie spoczynku, która może zwiększyć zasięgi aut elektrycznych do 800...1100 kilometrów na jednym ładowaniu. ♦

ELEKTRONIKA

♦ Znana z konstrukcji superpotężnych procesorów firma Ampere zapowiedziała kolejną generację CPU o nazwie AmpereOne, która ma pojawić się na rynku w 2025 r. i będzie wytwarzana w procesie 3 nm (poprzednik tego układu – w 5 nm) i wyposażona w maksymalnie 256 rdzeni oraz dwunastokanałową pamięć DDR5. ♦ Zespół badaczy z Uniwersytetu Kalifornijskiego w Riverside przeprowadził

testy nowej techniki obliczeniowej określanej jako przetwarzanie wielowątkowe jednocześnie i heterogeniczne (SHMT), używając układu składającego się z procesorów ARM Cortex-A57, procesora graficznego Nvidia i Google Edge TPU, i dowodząc, że wykonanie przykładowego kodu było 1,95 razy szybsze, a zużycie energii zostało zmniejszone o 51% w porównaniu z tradycyjnymi układami przetwarzającymi, co zostało zaprezentowane na dorocznym międzynarodowym sympozjum IEEE/ACM na temat mikroarchitektury w Toronto, Kanada. ♦



FIZYKA

♦ Według publikacji w „Physical Review Letters”, naukowcy z eksperymentu BESIII w Wielkim Zderzaczu Elektronów i Pozytonów 2 w Pekinie w Chinach dokonali odkrycia nowej cząstki znanej jako $X(2370)$, która, ich zdaniem, może być długo poszukiwaną „kulką gluonową”, elementarną cząstką w odróżnieniu od hadronów składających się z kwarków i gluonów pozbawioną kwarków i składającą się wyłącznie z gluonów, będącą przedstawicielem zupełnie nowej klasy cząstek elementarnych. ♦ Fizycy, Jim Al-Khalili z uniwersytetu w Surrey i Eddy Keming Chen z Uniwersytetu Kalifornijskiego w San Diego opublikowali studium, z którego wynika, że na wczesnym etapie istnienia Wszechświata nie istniało splątanie kwantowe, zjawisko kluczowe dla współczesnej fizyki – pojawiło się i narastało później w miarę rozwoju Uniwersum. ♦ Jak podało czasopismo „Science”, naukowcom z uniwersytetu technicznego ETH w Zurychu, pod kierownictwem Christiana Degena, udało się za pomocą czujnika pola magnetycznego o wysokiej rozdzielczości po raz pierwszy bezpośrednio wykryć wiry elektronowe w grafenie. ■ M.U.



Nanotechnologie w fotowoltaice

Jak pomóc elektronom

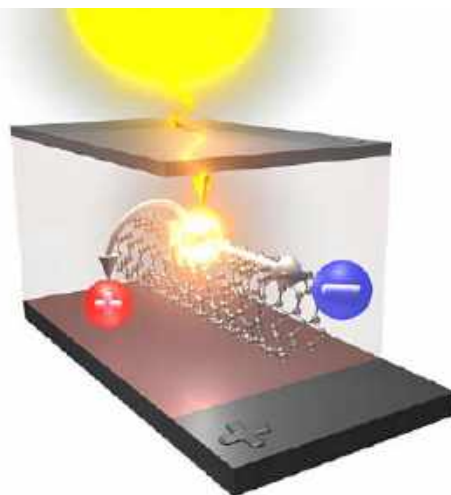
Światło słoneczne dociera na Ziemię silnie „rozcieńczone”, więc energia padająca na jednostkę powierzchni jest stosunkowo niska. Farmy słoneczne pokryte panelami z ogniwami fotowoltaicznymi kompensują to, pokrywając relatywnie duże obszary. Jest to rozwiązanie dalekie od efektywności.

Od dekad trwają prace nad zwiększeniem sprawności techniki fotowoltaicznej, co w praktyce sprowadza się do dążenia do zwiększenia produkcji użytkowej energii na jednostkę powierzchni ogniwa. Nanotechnologia nie jest jedynym polem poszukiwań, ale w ostatnich latach sprawia wrażenie techniki rozwijającej się szczególnie dynamicznie.

Obecne ogniwa słoneczne nie są w stanie przekształcić całego wpadającego światła w użyteczną energię, ponieważ część światła może uciec z powrotem z ogniwa do powietrza. Ponadto widmo światła słonecznego dzieli się na różne kolory odpowiadające zakresom fal. Ogniwo może być bardziej wydajne w przekształcaniu na energię użytkową światła w zakresach bliskich niebieskiej części spektrum, a mniej wydajne w przekształcaniu światła czerwonego. Promieniowanie świetlne o niższej energii przechodzi przez ogniwo niewykorzystane. Światło o wyższej energii wzbudza elektrony do pasma przewodnictwa, ale energia wykraczająca wartością poza energię przerwy pasmowej jest tracona jako ciepło. Jeśli te wzbudzone elektrony nie zostaną przechwycone i przekierowane, będą spontanicznie rekombinować z powstałymi dziurami, a energia zostanie utracona w postaci ciepła lub światła.

Nanorurki i inne przyspieszacze

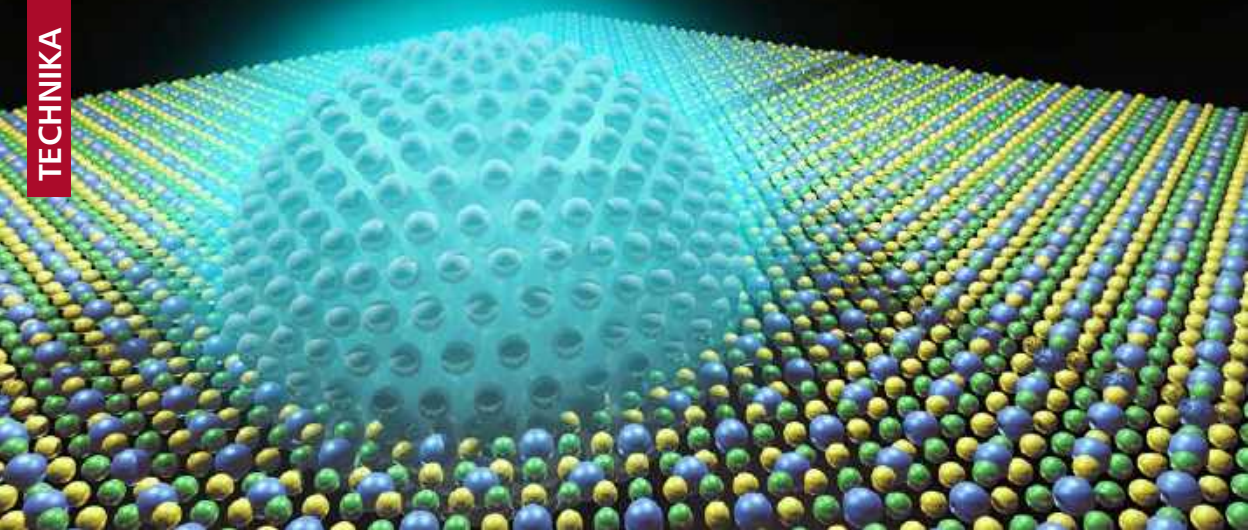
Nanomateriały to materiały, których rozmiar fizyczny mieści się w przedziale 1 nm...100 nm. Materiały te mogą występować w postaci nanoprętów, nanorurek, nanocząstek, nanosfer lub nanowłókien. Wchodzą ze sobą we wzajemne wiązania, co prowadzi do niskiej temperatury topnienia, wysokiej rozpuszczalności i wysokiej stabilności chemicznej. Efekt zwany tunelowaniem kwantowym prowadzi w nich w konsekwencji do możliwości optymalizacji lub zmiany właściwości optycznych i wewnętrznych właściwości materiału, takich jak promień Bohra i zmniejszona odległość migracji. Nanomateriały mogą obniżyć



1. Wizualizacja nanorurki węglowej w ogniwie słonecznym

koszty produkcji materiału ze względu na możliwość modyfikacji tańszych materiałów w celu dostosowania ich do celów danego zastosowania.

Nanotechnologia pokazała już ogromne możliwości w dziedzinie energii słonecznej; np. integracja wysokiej jakości warstwy nanocząstek krzemu o rozmiarze jednego nanometra bezpośrednio z krzemowymi ogniwami słonecznymi poprawia, według zwolenników tej techniki, wydajność energetyczną nawet o 60 proc. w zakresie widma ultrafioletowego. W innych zakresach mniej, ale też jest wzrost. Można zastosować też nanorurki węglowe (1) na bazie folii z nanocząstek dwutlenku tytanu, podwajając wydajność przekształcania światła ultrafioletowego w elektrony w porównaniu z wydajnością samych nanocząstek. Bez nanorurek węglowych elektrony generowane podczas pochłaniania światła przez cząsteczki tlenku tytanu muszą przeskakiwać z cząsteczki na cząsteczkę, aby dotrzeć do elektrody. Nanorurki węglowe „zbierają”



2. Wizualizacja kropki kwantowej na powierzchni atomowej sieci krystalicznej

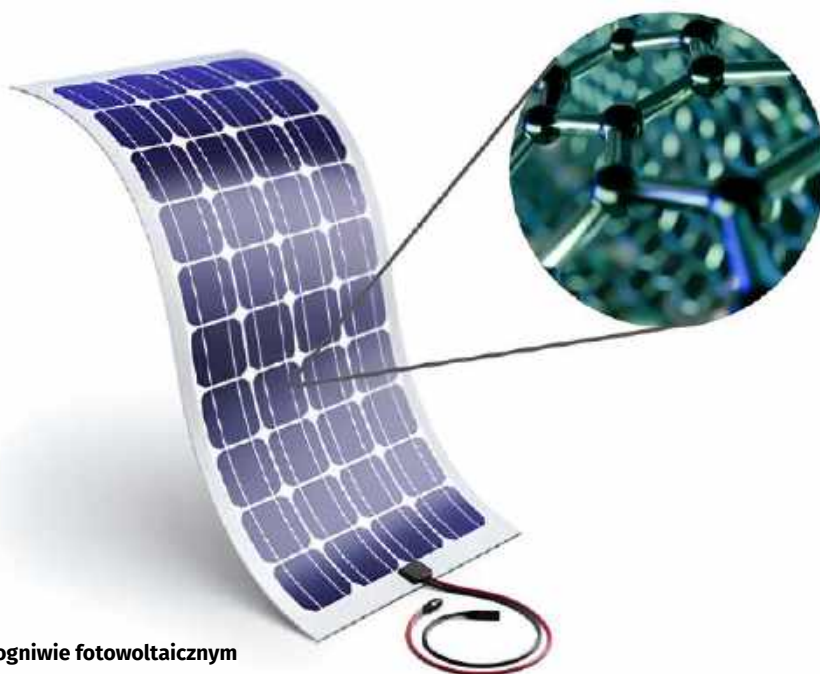
elektrony i zapewniają bezpośredni kanał dotarcia do elektrody i wysoką gęstość prądu na powierzchni ogniwa słonecznego bez większych strat. Innym pomysłem na zwiększenie wydajności konwersji ogniw słonecznych jest zastosowanie półprzewodnikowych kropek kwantowych (2), chociaż ogniwa słoneczne z kropkami kwantowymi są nadal przedmiotem badań. Rozważane są półprzewodniki III/V i inne kombinacje materiałów, takie jak Si/Ge lub Si/Be Te/Se. Potencjalne zalety tych ogniw to wyższa absorpcja światła w szczególności w podczerwonym obszarze widmowym, kompatybilność ze standardowymi procesami produkcji krzemowych ogniw słonecznych, wyższe parametry prądu w wyższych temperaturach, lepsza odporność na promieniowanie w porównaniu z konwencjonalnymi ogniwami słonecznymi. W przeciwieństwie do konwencjonalnych materiałów, w których jeden foton generuje tylko jeden elektron, kropki kwantowe mogą przekształcać wysokoenergetyczne fotony w wiele elektronów. Elektrony przemieszczają się z pasma walencyjnego do pasma przewodzenia. Kropki wychwytyują również więcej zakresów fal światła słonecznego, zwiększając w ten sposób wydajność konwersji nawet do 65 proc.

Obecnie dostępne nanotechnologiczne ogniwa słoneczne na ogół wciąż nie są tak wydajne jak tradycyjne, jednak rekompensuje to ich niższy koszt. Maksymalna sprawność konwencjonalnego krzemowego ogniwa słonecznego wynosi 33,7 proc. W dłuższej perspektywie wersje nanotechnologiczne powinny być zarówno tańsze, jak i, przy użyciu kropek kwantowych, powinny być w stanie osiągnąć wyższy poziom wydajności niż konwencjonalne.

Plazmonika i supermateriały grafenowe

Niedawno w „Nature Catalysis” zespół naukowców hiszpańskich i niemieckich przedstawił koncepcję nanoskalowego dwuwymiarowego superkryształu wykorzystującego dwa różne metale. „Najpierw tworzymy cząsteczki w zakresie 10...200 nanometrów z plazmonicznego (cechującego się występowaniem oscylującego gazu elektronowego – przyp. red.) metalu – w naszym przypadku jest to złoto”, piszą uczeni. „W tej skali światło widzialne bardzo silnie oddziałuje z elektronami metalu, powodując ich rezonansowe oscylacje”. Oznacza to, że elektrony poruszają się bardzo szybko z jednej strony nanocząstki na drugą, tworząc rodzaj minimagnesu. Uczeni nazywają to momentem dipolowym. „Można myśleć o tym procesie jako o supersoczewce koncentrującej energię. Nasze nanomateriały robią to samo co koncentratory, ale w skali molekularnej”, wyjaśniają badacze. Pozwala to nanocząsteczkom na wychwytywanie większej ilości światła słonecznego i przekształcanie go w elektrony o bardzo wysokiej energii. Te z kolei pomagają napędzać reakcje chemiczne, np. produkcję wodoru przy użyciu światła słonecznego. „Łącząc metale plazmoniczne i katalityczne, przyspieszamy rozwój silnych fotokatalizatorów do zastosowań przemysłowych. Jest to nowy sposób wykorzystania światła słonecznego, który oferuje potencjał dla innych reakcji, takich jak konwersja CO₂ w substancję użytkową”, wyjaśniają naukowcy.

W 2023 roku zademonstrowano z kolei innowacyjny nanokoncentrator promieniowania słonecznego, który wykorzystuje metamateriał grafenowy (3). Składa się on z układu o rozmiarach $\sim 1,5 \times 1,5 \mu\text{m}$



3. Grafen w ogniwie fotowoltaicznym

złotych komórek koncentratora, które mogą ogniskować promieniowanie słoneczne do plamki o rozmiarze 60 nm na aktywnym obszarze ogniwa słonecznego. Może działać w całym paśmie energii słonecznej do długości fali 2400 nm, co stanowi rozszerzenie zakresu długości fali konwencjonalnych ogniw słonecznych. Nanokoncentrator może zbierać i wykorzystywać również słabe promieniowanie słoneczne, które występuje między 1800 nm a 2400 nm. Grafen cechuje przezroczystość optyczna, wysoka ruchliwość elektronów, wysoka przewodność elektryczna, szeroko-pasmowe widmo optyczne, zdolność do domieszkowania elektrostatycznego i wysoka przewodność cieplna. Może być stosowany w ogniwach słonecznych jako powłoka antyrefleksyjna w celu maksymalizacji absorpcji energii słonecznej. Warto jednak wspomnieć, że grafen nadal ma dobrze znaną wadę – niską absorpcję optyczną na poziomie ~2,3 proc.

Warto tu wyjaśnić, że plazmonika to dziedzina związana ze zbiorowymi oscylacjami wolnych elektronów, które występują na powierzchni metali szlachetnych, takich jak złoto czy platyna. Zwykle towarzyszą im pola elektromagnetyczne o długości poniżej fali i koncentrują się z dużą intensywnością na krawędziach metalu. Dlatego połączenie metalu szlachetnego z grafenem w strukturze nanomateriału może przewyższyć niską absorpcję optyczną grafenu i zwiększyć absorpcję energii.

Metamateriał grafenowy składa się z dziesięciu koncentrycznych warstw powłoki grafenowej, które są połączone krzyżowo przez pięć poziomych pierścieni grafenowych (struktury w kształcie pierścienia). Konstrukcja koncentratora wraz

z konstrukcją metamateriału maksymalizuje wydajność absorpcji optycznej światła słonecznego w metamateriale grafenowym. Zmierzona absorpcja optyczna grafenu została zwiększona do 27 razy w stosunku do pierwotnej wartości, z maksymalną wartością 62 proc. i obejmowała bardzo szerokie pasmo słoneczne. Zwiększona absorpcja optyczna słabego promieniowania słonecznego odbieranego w rozszerzonym paśmie (1800...2400 nm) może sięgać nawet 50 proc. Testowane urządzenie wykazuje działanie niewrażliwe na polaryzację, co pozwala na zbieranie całego spolaryzowanego i niespolaryzowanego promieniowania słonecznego. Ma również dobre pole widzenia wynoszące 120°, co pozwala na zbieranie promieniowania słonecznego bez potrzeby stosowania dwuosiowych mechanicznych trackerów.

Nanotechnologia w ogniwach słonecznych zainteresowała m.in. wojsko. Armia Stanów Zjednoczonych zatrudniła już firmę Konarka Technologies, by pomogła zaprojektować lepszy sposób zasilania urządzeń elektrycznych żołnierzy. Według Daniela McGahna, wiceprezesa wykonawczego firmy Konarka, „zwykły żołnierz połowy nosi obecnie 0,7 kg baterii. Żołnierz biorący udział w operacjach specjalnych musi nosić ze sobą 70 kilogramów balastu, z czego połowa to baterie”. Nanotechnologia wykorzystana w ogniwach mogłaby znacznie poprawić mobilność żołnierzy, dając im wydajne zasilanie bez wielkiego obciążenia.

To jednak przyszłość. Ogniwa fotowoltaiczne z zastosowaniem nanotechnologii są bardziej wydajne na razie głównie teoretycznie. W praktyce czeka je jeszcze sporo badań, testów i praktycznych sprawdzianów. ■

Mirosław Usidus



Świat kosmicznych dziwów nie przestaje zadziwiać

Matuzalem, wielkie bum i zagadkowa pustka

Im dokładniej przyglądamy się kosmosowi, tym dziwniejszy i bardziej zaskakujący się wydaje. Kosmiczne anomalie i zagadkowe zjawiska rodzą coraz więcej pytań tam, gdzie wolelibyśmy odpowiedzi.

Wszechświat ma, według przyjętych powszechnie obliczeń, 13,8 miliarda lat. Z naukowej wiedzy na temat ewolucji gwiazd wynika wszelako, że gwiazda HD 140283 nazywana „Matuzalemem”, znajdującej się w Drodze Mlecznej, ma 14,5 miliarda lat(!). Czy może istnieć gwiazda starsza niż sam Wszechświat? Chyba nie. Gdzie w takim razie tkwi błąd, w liczeniu wieku gwiazdy czy w szacowaniu wieku Wszechświata? Odpowiedź nie jest zbyt jasna. HD 140283 znajduje się w odległości zaledwie 190 lat świetlnych od nas. Obserwacje tej stosunkowo bliskiej gwiazdy pozwalają nam oszacować dość precyzyjnie jej wiek na 14,46 miliarda lat. Jednak zawartość żelaza wynosząca 0,4 proc. zawartości żelaza w Słońcu, sugerują, że jest ona bardzo stara, za stara wręcz. Szacunki te są jednak obarczone dużą niepewnością wynoszącą prawie miliard lat, co oznacza, że najprostszym sposobem pogodzenia sprzecznych danych jest rozważenie możliwości, że gwiazda Matuzalem wydaje się starsza, niż w rzeczywistości jest, z powodu wydarzeń z przeszłości, których ślady już nie istnieją.

Eksplodujące prawa fizyki

W maju 2023 r. astronomowie donieśli, że udało im się zidentyfikować „największą” kosmiczną eksplozję, jaką kiedykolwiek zaobserwowano, kulę ognia sto razy większą od naszego Układu Słonecznego, która zapłonęła w odległym rejonie Wszechświata ponad trzy lata temu. Wybuch nazwano AT2021lwx. Precyzyjnie rzecz ujmując, nie jest to najjaśniejszy błysk, jaki kiedykolwiek zaobserwowano we Wszechświecie. Ten rekord wciąż należy do wybuchu promieniowania gamma, któremu nadano przydomek BOAT, czyli Brightest Of All Time (z ang.

najjaśniejszy w historii). Philip Wiseman, astrofizyk z brytyjskiego uniwersytetu w Southampton, wyjaśnia, że AT2021lwx został uznany za „największą” eksplozję, ponieważ uwolnił znacznie więcej energii w ciągu ostatnich trzech lat niż krótki błysk BOAT. Astronomowie przeanalizowali kilka możliwych wyjaśnień eksplozji. Jedno z nich traktuje AT2021lwx jako eksplodującą gwiazdę. Jednak błysk jest dziesięć razy jaśniejszy niż jakakolwiek wcześniej widziana supernowa. Zatem rozważa się wyjaśnienie nazywane zaburzeniem pływowym polegające na rozerwaniu gwiazdy na strzępy w trakcie opadania do supermasywnej czarnej dziury. Kojarzy się z modelem kwazara, czyli zjawiska połykania przez supermasywną czarną dziurę ogromnych ilości gazu w centrum galaktyki. Problem w tym, że supermasywne czarne dziury znajdują się w centrum galaktyk, a w przypadku eksplozji tej wielkości należałoby oczekiwać, że galaktyka będzie tak rozległa jak Droga Mleczna. Jednak nie ma śladów galaktyki w pobliżu AT2021lwx.

Inny wielki i dziwny obiekt, który świeci dziesięć milionów razy jaśniej niż Słońce, jest równie tajemniczy i zdaje się łamać prawa fizyki. Tak przynajmniej wynika z publikacji NASA. Należy do kategorii zjawisk znanych astrofizykom jako „ultraluminous X-ray sources” (ULX). Wydzielana przez ten obiekt ilość energii łamie tzw. limit Eddingtona, określający maksymalną jasność obiektu o określonym rozmiarze. Jeśli wartość ta jest przekraczana, to zdaniem naukowców, obiekt powinien się rozpaść. Jednak ULX-y przekraczają ten limit nawet pięćsetkrotnie i nie rozpadają się. Obserwacje opisane w „The Astrophysical Journal”, przeprowadzone przez instrument Nuclear Spectroscopic Telescope Array (NuSTAR), pracujący



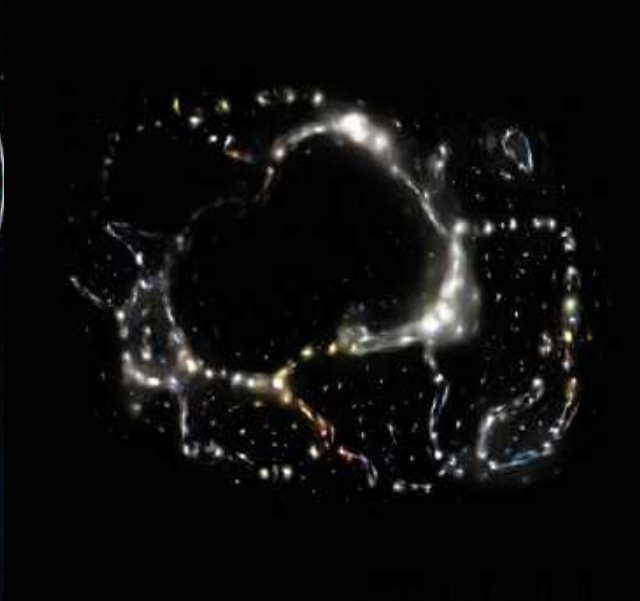
1. Obiekt M82 X-2

w zakresie promieniowania rentgenowskiego, dotyczą ULX o nazwie M82 X-2 (1), zdecydowanie za jasnego w stosunku do rozmiaru. Wcześniej sugerowano, że ekstremalna jasność mogła być rodzajem złudzenia optycznego, ale nowe badania pokazują, że tak nie jest. Jeden z nowszych pomysłów na wyjaśnienie polega na tym, że intensywne pole magnetyczne gwiazdy neutronowej zmienia kształt jej atomów, pozwalając gwiazdzie utrzymać jednorodność, w miarę jak staje się coraz jaśniejsza.

Kosmos jest pusty, ale sąsiedztwo Drogi Mlecznej jest bardziej pusty

Nasza Galaktyka, Droga Mleczna, wydaje się jedną z wielu takich we Wszechświecie. Jednak jej kosmiczne otoczenie znacznie różni się od innych. Chodzi o to, że z obserwacji wynika, iż znajdujemy się w samym środku gigantycznej kosmicznej pustki, największej, jaką kiedykolwiek zaobserwowano, co oznacza, że tak wielkiej pustki nie ma nigdzie, wokół żadnej znanej nam galaktyki. Oczywiście, nie oznacza to, że nie ma w ogóle. Po prostu z tego, co obecnie wiemy, ta, która nas otacza, jest wyjątkowa. Astronomowie po raz pierwszy zaproponowali istnienie tej masywnej pustki, znanej jako pustka KBC (2), w 2013 roku. Od tego czasu dowodów potwierdzających jej istnienie jest coraz więcej.

Zgodnie z teoriami kosmologicznymi materia we Wszechświecie powinna być równomiernie rozłożona w dużych skalach. Zasada ta pozwala naukowcom stosować te same prawa fizyczne do obiektów znajdujących się blisko nas i tych na obrzeżach obserwowalnego Wszechświata. Zakładamy, że Wszechświat jest jednorodny i izotropowy, co oznacza, że wygląda tak



2. Wizualizacja pustki KBC

samo w każdym kierunku i z każdego punktu. Jednak obserwacje przeprowadzone w ciągu ostatniej dekady sugerują, że materia może nie być tak równomiernie rozłożona, jak sugeruje zasada kosmologiczna. Wydaje się ona zbijać w regiony o wysokiej i niskiej gęstości, co powoduje znaczne różnice w rozkładzie materii.

Indranil Banik z Uniwersytetu St. Andrews w artykule opublikowanym w „Monthly Notices of the Royal Astronomical Society” sugeruje, że pustka, w której środkiem jesteśmy, rozciąga się na jakieś dwa miliardy lat świetlnych i jest wystarczająco duża, aby pomieścić dwadzieścia tysięcy galaktyk rozmiarów Drogi Mlecznej ustawionych w rzędzie, od jednego końca do drugiego.

Dla ścisłości warto dodać, że pustka KBC nie jest całkowicie pusta – cechuje się około 20 proc. mniejszą gęstością materii niż przestrzeń poza jej granicami. Różnica ta może nie wydawać się znacząca, ale i tak stanowi dla kosmologów zagwozdkę. Jednym z efektów pustki KBC jest zachowanie pobliskich gwiazd i galaktyk, które oddalają się od nas szybciej, niż oczekiwano by zgodnie z teoriami i stałą Hubble’a. Galaktyki wewnątrz pustki doświadczają z kolei przyciągania grawitacyjnego z zewnątrz, co daje większą lokalną wartość stałej Hubble’a (przesunięcie ku czerwieni). Stała Hubble’a to wartość używana przez kosmologów do opisanego tempa rozszerzania się Wszechświata. Zgodnie z naszym obecnym rozumieniem, stała Hubble’a powinna być taka sama wszędzie we Wszechświecie. Jednak obserwowane tempo ekspansji w naszej kosmicznej okolicy jest wyższe, niż przewidywano, co sugeruje szybszą, niż oczekiwano, ekspansję. ■

Mirosław Usidus



Hakowanie nowoczesnych samochodów dla funkcji premium i nie tylko

Włamać się do auta, by posiedzieć na podgrzewanym fotelu

O możliwości włamywania się, hakowania i przejmowania kontroli nad nowoczesnymi połączonymi z siecią autami pisaliśmy nie raz. Okazuje się, że wyodrębnił się w tej cyberprzestępczości wyspecjalizowany, nastawiony na darmowe korzystanie z luksusów, nurt.

Kilkanaście miesięcy temu grupa niemieckich hakerów włamała się do tesli. Nie chodziło im o zdalne przejęcie kontroli nad samochodem. Nie próbowali uzyskać dostępu do haseł Wi-Fi właściciela ani nie chcieli wykraść numerów kart kredytowych z lokalnej sieci ładowania pojazdów elektrycznych. Ich celem były podgrzewane fotele, funkcja aktywowana dopiero po wyłożeniu przez kierowcę ponad 300 dolarów. Trzech doktorantów z Technische Universität Berlin wraz z niezależnym badaczem (i właścicielem Tesli)

twierdzi, że udało im się złamać dostęp do podgrzewanych siedzeń przez manipulację przy napięciu zasilającym system informacyjno-rozrywkowy samochodu. Udało im się też uzyskać dostęp do wielu jego wewnętrznych systemów i prywatnych danych użytkownika. Zastrzegali, że nie mieli złych intencji. Wręcz przeciwnie.

Gdy słyszymy o hakowaniu samochodów, wyobrażenia podsuwa obrazy i sceny rodem z filmów grozy, w których miliony tesli lub innych aut są zdalnie

przejmowane przez grupy terrorystyczne i zmuszane do ataków np. na szpitale. To na szczęście mało prawdopodobne. Większe ryzyko dotyczy danych osobowych i finansowych związanych z różnymi cyfrowymi dodatkami i funkcjami, które są płatne.

Wzbudza to coraz częściej kontrowersje, bo o ile podgrzewane siedzenia można jeszcze jakoś zrozumieć, to już np. dodatkowe konie mechaniczne czy wspomaganie kierownicy już mniej. A Mercedes-Benz za kilkadziesiąt dolarów miesięcznie oferuje ekstra moc w swoich pojazdach. W BMW usługą premium są nagrania z kamer bezpieczeństwa, zaś funkcje wspomaganie w Fordzie BlueCruise. To coraz częstsze podejście. Koncern General Motors planuje zaoferować ponad pół setki takich dodatkowo płatnych funkcji do 2026 roku.

Zrozumiałe jest, że te posunięcia nie spotkały się z przychylnością kupujących samochody. Plan BMW polegający na pobieraniu opłaty miesięcznej za podgrzewane fotele w Wielkiej Brytanii i Korei okazał się tak niepopularny, że BMW w końcu ogłosiło, że całkowicie porzuci ten pomysł. Firma jednak wciąż oferuje niektóre funkcje, np. cyfrowego asystenta parkowania, w subskrypcji płatnej.

Chociaż funkcje subskrypcji nie są wyłączne dla samochodów elektrycznych, są one nierozzerwalnie związane z rewolucją pojazdów elektrycznych. Opracowywanie i akumulatory do pojazdów elektrycznych to wielkie koszty dla branży motoryzacyjnej, idące w setki miliardów dolarów. A ponieważ pojazdy elektryczne mają generalnie znacznie mniej elementów mechanicznych niż samochody spalinalowe, wymagają niewiele napraw i serwisu, co oznacza,

że producenci samochodów, dostawcy i dealerzy mogą stracić znaczną część przychodów ze sprzedaży części i autoryzowanych napraw. Jeden z dyrektorów Hyundai powiedział niedawno, że jego firma chce, aby 30 proc. przyszłych zysków pochodziło z dystrybucji oprogramowania, funkcji do pobrania, rozrywki w samochodzie i innych funkcji subskrypcyjnych.

Jeśli jednak stawia się na oprogramowanie jako źródło zysku, trzeba się liczyć z ludźmi, którzy potrafią to oprogramowanie zhakować. Przy czym o klientach, którzy chcą zaoszczędzić, chodzi mniej, bardziej o tych, których interesują dane i pieniądze klientów. Za każdym razem, gdy samochód otrzymuje nową aplikację na ekranie dotykowym lub funkcję subskrypcji, stanowi to potencjalną drogę dla cyberwłamywaczy, którzy szukają informacji o karcie kredytowej, danych osobowych i nie tylko. Załóżmy, że ktoś płaci firmie samochodowej 100 złotych miesięcznie za coś takiego, jak podgrzewane fotele. Haker może użyć różnych narzędzi lub technik, aby znaleźć lukę w zabezpieczeniach tej aplikacji i zdalnie się zalogować, a stamtąd uzyskać dostęp do karty kredytowej.

Zagrożenie cyberatakami nie jest niczym nowym dla firm z branży technologicznej. Jednak teraz dołączyć do ich grona chce sektor, który znał się od stu kilkudziesięciu lat głównie na budowaniu silników i karoserii. Nie tylko nie wydaje się do tego dobrze przygotowany, ale może nie rozumieć konieczności natychmiastowego reagowania na takie sytuacje. Jeśli ktoś chce „smartfonów na kółkach”, powinien dobrze zrozumieć, co się z tym wiąże. ■

Mirosław Usidus

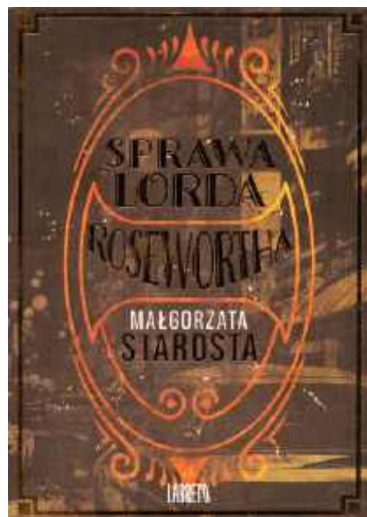
Sprawa lorda Rosewortha

Małgorzata Starosta

Wydawnictwo: Labreto, liczba stron: 288, cena: 54,99 zł

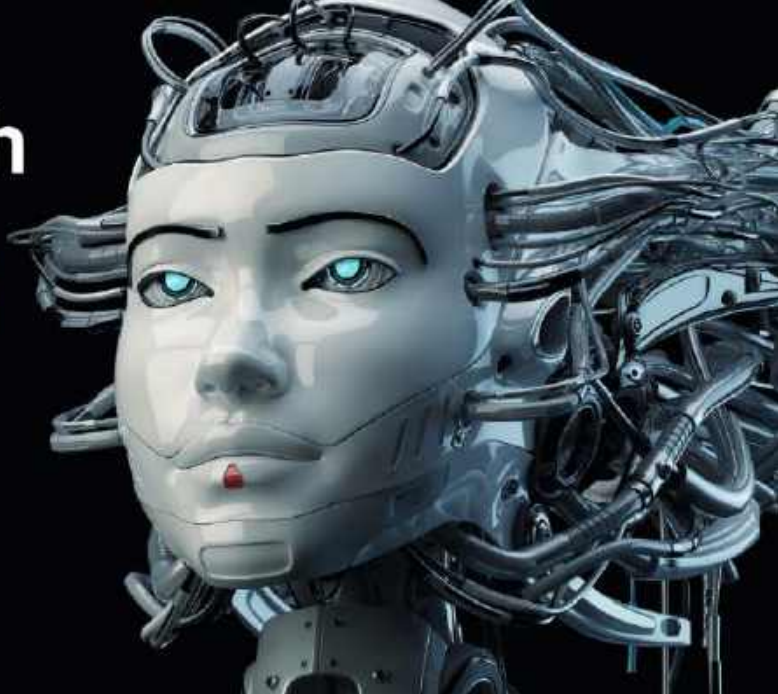
Sierpień 1959 roku. Jonathan Harper, prywatny detektyw borykający się z osobistą stratą, na prośbę dawnej przyjaciółki podejmuje się rozwiązania sprawy zabójstwa Thadeusa Rosewortha, którego śmierć wzburzyła opinię publiczną i przysporzyła nie lada kłopotów Scotland Yardowi. Doskonale znając zwyczaje panujące w świecie arystokracji, Harper spodziewa się kłamstw, intryg, wzajemnych oskarżeń i animozji, jednak ani bystry umysł, ani doświadczenie nie uchronią go przed tym, czego naprawdę dotyczy „sprawa lorda”.

Nieznana przeszłość ofiary, sieć kłamstw i zbrodnia niemal doskonała doprowadzą do zaskakujących odkryć, wystawiając na próbę zaufanie Harpera nie tylko do Roseworthów, ale także do najbliższej mu osoby. Sprawa lorda to brawurowa powieść detektywistyczna osadzona w czasach powojennej transformacji Wielkiej Brytanii, czerpiąca z najlepszych brytyjskich wzorców gatunkowych i flirtująca z czarnym kryminalem.



FunSearch

AI



1. Graficzne przedstawienie FunSearch AI

Umiesz liczyć, licz na siebie, człowieku – głosi stare powiedzenie. Nadzieje pokładane w sztucznej inteligencji, od której oczekuje się, że zmierzy się, i poradzi, z problemami, z którymi my biedzimy się od dawna, mogą okazać się płonne. Nie oznacza to, że AI do niczego naukowcom i inżynierom się nie przyda. Może być pomocna, ale nie oczekujmy od niej noblowskich odkryć.

Czy AI rozwiąże nierozwiązywalne?
I czy ta odpowiedź nam cokolwiek da?

ALGORYTMY
CORAZ BARDZIEJ
POMOCNE,
ALE Z NOBLAMI
JESZCZE SIĘ
WSTRZYMAJMY

W matematyce istnieje wiele problemów, które pozostają nierozwiązane od dziesięcioleci, a czasem nawet stuleci. Problemy te są często tak złożone, że nawet

najbystrzejsze umysły mają trudności ze znalezieniem rozwiązań.

Próbę zmierzenia się z nimi za pomocą generatywnej sztucznej inteligencji podjęli badacze z projektu Google DeepMind. Ich wielki model językowy (LLM) był szkolony na ogromnej ilości literatury matematycznej, prac badawczych i przykładów rozwiązywania problemów. Trening ten obejmował nie tylko podawanie liczb i równań, ale także tekstów wyjaśniających pojęcia matematyczne, teorie i kontekst bardziej ogólny. Powstało coś, co twórcy nazywają procedurą opartą na wstępnym szkoleniu modelu o nazwie FunSearch (1). Po przeszkoleniu modelowi AI przedstawiono nierozwiązany problem matematyczny. Zamiast korzystać z tradycyjnych metod obliczeń matematycznych, sztuczna inteligencja wykorzystała wiedzę o pojęciach matematycznych, zdobytą podczas szkolenia, znajdując, co warto podkreślić, rozwiązanie problemu, które zostało następnie zweryfikowane przez matematyków jako poprawne. Uznano

to za pierwszy przypadek, w którym model sztucznej inteligencji oparty na języku przyczynił się do rozwiązania znanego od dawna problemu. „Używamy LLM jako silnika kreatywności”, wyjaśniał w wypowiedzi dla mediów Bernardino Romera-Paredes, informatyk zaangażowany w projekt DeepMind FunSearch.

FunSearch został zaprezentowany przez Google DeepMind w grudniu 2023 r. Jak czytamy w publikacji na łamach „Nature”, po raz pierwszy duży model językowy (LLM) podjął się odkrycia rozwiązania znanego od dawna problemu naukowego, przedstawiając weryfikowalne i cenne nowe informacje, które nie były wcześniej znane. FunSearch to połączenie dużego modelu językowego o nazwie Codey, który jest wersją PaLM 2 firmy Google, specjalnie dostosowaną do rozumienia i generowania kodu komputerowego, z innymi systemami, które odrzucają nieprawidłowe odpowiedzi i wprowadzają poprawne odpowiedzi do procesu wnioskowania. Zespół badawczy, kierowany przez Pushmeeta Kohla, zastosował metodologię prób i błędów, umożliwiając FunSearch sugerowanie rozwiązań dla problemu początkowo nakreślonego w języku programowania Python.

Proces polega na tym, że Codey proponuje kod do ukończenia programu, a drugi algorytm sprawdza i ocenia sugestie. Powstaje pętla niestannego doskonalenia. Po wielu powtórzeniach FunSearch z powodzeniem wygenerował kod zapewniający poprawne i nieznane wcześniej rozwiązanie złożonego problemu matematycznego ze zbiorem odwracanych kart, który wyewoluował z gry Set, wynalezioną w latach 70. XX wieku przez genetyka Marsha Falco. Talia Set zawiera 81 kart. Każda karta przedstawia jeden, dwa lub trzy symbole, które są identyczne pod względem koloru, kształtu i odcieni, a dla każdej z tych cech istnieją trzy możliwe opcje. Łącznie możliwości te wynoszą $3 \times 3 \times 3 \times 3 = 81$. Gracze muszą odwrócić karty i znaleźć specjalne kombinacje trzech kart, zwane zestawami. Matematycy wykazali, że gracze mają gwarancję znalezienia zestawu, jeśli liczba odwróconych kart wynosi co najmniej 21. Znaleźli również rozwiązania dla bardziej złożonych wersji gry, w których abstrakcyjne wersje kart mają pięć lub więcej właściwości. Pewne nierozwiązane kwestie jednak pozostają. Na przykład, jeśli istnieje n właściwości, gdzie n jest dowolną liczbą całkowitą, to istnieje $3n$ możliwych kart – ale minimalna liczba kart, które muszą zostać ujawnione, aby zagwarantować rozwiązanie, jest nieznana. Problem ten można wyrazić w kategoriach geometrii dyskretnej. Jest to równoważne znalezieniu pewnych układów trzech punktów w n -wymiarowej przestrzeni. Matematycy

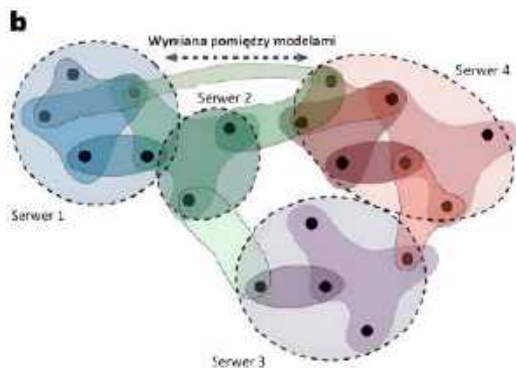
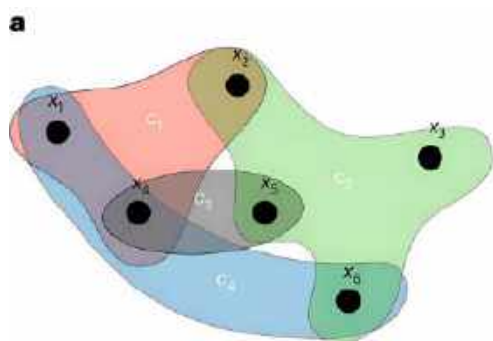
określili granice możliwego ogólnego rozwiązania – biorąc pod uwagę n , odkryli, że wymagana liczba „kart na stole” musi być większa niż podana przez pewien wzór, ale mniejsza niż podana przez inny. Problem polegający na określeniu maksymalnego rozmiaru określonego typu zbioru, co jest zadaniem skomplikowanym, został z sukcesem rozwiązany przez maszynę. Wszechstronność FunSearch została dodatkowo zademonstrowana poprzez zastosowanie go do problemu pakowania pojemników, kolejnego trudnego zadania matematycznego z zastosowaniami w informatyce. FunSearch nie tylko znalazł rozwiązanie, ale także przewyższył metody opracowane przez człowieka.

Zdolność FunSearch do tworzenia zrozumiałego i interpretowalnego przez człowieka kodu wyróżnia go na tle innych algorytmów. W przeciwieństwie do poprzednich narzędzi, które traktowały problemy matematyczne jako łamigłówki podobne do gier takich jak Go czy szachy, FunSearch wykorzystuje duży model językowy działający w tandemie z innymi systemami zaprojektowanymi w celu odfiltrowania nieprawidłowych lub bezsensownych odpowiedzi. Metodologia ta, choć eksperymentalna, okazała się skuteczna. Generując kod, przepis na rozwiązanie, a nie samo rozwiązanie, FunSearch oferuje ciekawsze metody rozwiązywania problemów.

Eksperci nieprzekonani

Sztuczna inteligencja otrzymała niedawno nagrodę w najbardziej elitarnej Międzynarodowej Olimpiadzie Matematycznej (IMO) dla szkół średnich. Każdego roku od 1959 r. uczniowie szkół średnich z ponad stu krajów świata rywalizują w rozwiązywaniu różnorodnych problemów matematycznych obejmujących algebrę, geometrię i teorię liczb. Wielu zwycięzców IMO dostało później, w dorosłym życiu, prestiżowe nagrody matematyczne. W artykule opublikowanym w styczniu tego roku w „Nature”, zespół naukowców z Google DeepMind przedstawił inny model sztucznej inteligencji o nazwie AlphaGeometry, który radził sobie z zadaniami z sekcji geometrii Międzynarodowej Olimpiady Matematycznej.

„Poczyniliśmy duże postępy z modelami takimi jak ChatGPT... ale jeśli chodzi o problemy matematyczne, te [duże modele językowe] zasadniczo osiągają zerowy wynik”, komentował te sukcesy w rozmowie z „Popular Mechanics” Thang Luong z Google DeepMind i autor artykułu na temat AlphaGeometry. „Kiedy zadajesz pytania [matematyczne], model da ci coś, co wygląda jak odpowiedź, ale [w rzeczywistości] nie ma sensu”. Na przykład, sprawy komplikują się, gdy



2. Model hipergraficzny HypOp

sztuczna inteligencja próbuje rozwiązać algebraiczny problem słowny lub problem kombinatoryczny, który wymaga znalezienia liczby permutacji (lub wersji) sekwencji liczb. Aby odpowiedzieć na pytania matematyczne tego kalibru, AlphaGeometry opiera się na kombinacji symbolicznej sztucznej inteligencji, którą Luong opisuje jako precyzyjną, ale powolną, oraz sieci neuronowej bardziej podobnej do dużych modeli językowych, która jest odpowiedzialna za szybkość i kreatywną stronę rozwiązywania problemów.

Tak czy inaczej, eksperci nie są przekonani, że sztuczna inteligencja stworzona do rozwiązywania problemów matematycznych na poziomie szkoły średniej jest gotowa do podjęcia trudniejszych tematów, np. z zaawansowanej teorii liczb lub kombinatoryki, nie mówiąc już o badaniach matematycznych. AlphaGeometry wyróżnia się w tej mierze. Nie oznacza to jednak, że jest ona gotowa do radzenia sobie z matematyką wyższego poziomu.

Marijin Heule z Uniwersytetu Carnegie Mellon pracuje nad zautomatyzowanym weryfikatorem twierdzeń zwanym solverem SAT (testy dla uczniów w amerykańskim systemie edukacyjnym), choć nie ma to bezpośredniego związku z egzaminami w amerykańskich szkołach znanymi pod tą nazwą. „Jeśli chodzi o rozwiązywanie problemów matematycznych lub problemów w ogóle, wyzwaniem dla sztucznej inteligencji jest to, że nie może ona wymyślać nowych koncepcji”, tłumaczy Heule w „Popular Mechanics”. „Przynajmniej w dającej się przewidzieć przyszłości [sztuczna inteligencja] będzie głównie asystować”. Jak wyjaśnia, maszyny naprawdę dobrze radzą sobie z określaniem, czy istnieje niepoprawny argument lub oferowaniem kontrprzykładu. Te oparte na sztucznej inteligencji wskazówki mogą pomóc naukowcom odróżnić ślepe uliczki badawcze od obiecujących ścieżek, ale na razie nic poza tym. Weźmy na przykład ostatnie twierdzenie Fermata. Rozwiązanie tego problemu z teorii liczb

zajął ponad trzy stulecia. Zdaniem Heule, niezwykle trudno byłoby wyjaśnić jego rozwiązanie sztucznej inteligencji, nie mówiąc już o poproszeniu sztucznej inteligencji o jego rozwiązanie. Sztuczna inteligencja może być przydatna do rozpoznawania wzorców lub tak zwanych problemów typu „igła w stogu siana”, w których matematycy szukają czegoś o bardzo konkretnej właściwości w bardzo dużym zbiorze.

Z drugiej strony, według badań przeprowadzonych przez inżynierów z Uniwersytetu Kalifornijskiego w San Diego, narzędzie programistyczne oparte na zaawansowanych technikach sztucznej inteligencji może rozwiązywać złożone, wymagające obliczeniowo problemy szybciej i w bardziej skalowalny sposób niż najnowocześniejsze metody. W artykule, który został opublikowany w maju 2024 r. w „Nature Machine Intelligence”, naukowcy opisują HypOp, framework, który wykorzystuje nienadzorowane uczenie się i hipergraficzne sieci neuronowe (2). Struktura ta potrafi, według nich, rozwiązywać problemy optymalizacji kombinatorycznej znacznie szybciej niż istniejące metody. HypOp potrafi rozwiązać pewne problemy kombinatoryczne, które nie były rozwiązywane tak skutecznie przez wcześniejsze metody.

Jednym z przykładów stosunkowo prostego problemu kombinatorycznego jest ustalenie, ile i jakiego rodzaju towarów należy przechowywać w określonych magazynach tak, by zużywać jak najmniej paliwa podczas dostarczania tych towarów. HypOp może być, jak twierdzą autorzy, stosowany do szerokiego spektrum trudnych problemów w świecie rzeczywistym, z zastosowaniami w odkrywaniu leków, projektowaniu układów scalonych, weryfikacji logicznej, logistyce i wielu innych. Wynika to z faktu, że w przypadku tych problemów rozmiar przestrzeni poszukiwań potencjalnych rozwiązań rośnie wykładniczo, a nie liniowo w stosunku do rozmiaru problemu. HypOp może rozwiązywać te złożone problemy, wykorzystując nowy

algorytm rozproszony, który pozwala wielu jednostkom obliczeniowym na hipergrafie rozwiązywać problem jednocześnie, równolegle, czyli bardziej efektywnie. HypOp wykorzystuje sieci neuronowe hipergraficzne, które mają połączenia wyższego rzędu niż tradycyjne sieci neuronowe w postaci grafu. HypOp może również przenosić uczenie się z jednego problemu do rozwiązywania innych, pozornie odmiennych problemów.

Maszyny liczą, ale nie dowodzą

Od prawie 60 lat istnieje, z czego nie każdy zdaje sobie sprawę, dobrze rozwinięta poddziedzina automatycznej dedukcji, zwykle nazywana Automated Theorem Proving (ATP), która od około 60 lat próbuje budować programy, które mogą automatycznie udowadniać nietrywialne twierdzenia matematyczne. Organizowany jest również coroczny konkurs ATP o nazwie CASC (3), w którym systemy rywalizują w różnych kategoriach.

Jednak za prawie wszystkimi osiągnięciami w automatyzacji dowodzenia wciąż stoi kierujący i nadzorujący człowiek, a komputer wykonuje zwykle niskopoziomą matematyczną „brudną robotę” (tj. rutynowe przetwarzanie symboli, sprawdzanie różnych specjalnych przypadków, które muszą zostać usunięte itp.). Należy jednak pamiętać, że komputery są fantastycznymi narzędziami do takiej pracy.

3. Nagroda w konkursie Automated Theorem Proving CASC



Obecnie wiele osób pracuje od dziesięcioleci nad próbami automatycznego udowadniania poprawności programów. Wielokrotnie więcej wysiłku badawczego wkłada się właśnie w to, a nie w automatyczne udowadnianie twierdzeń matematycznych, ze względu na ich oczywistą wartość komercyjną. Udowadniać znacznie bardziej interesujące rzeczy, gdy kierował nimi inteligentny człowiek.

Nie ma prawdziwych ekspertów, którzy powiedzieliby dziś, że sztuczna inteligencja mogłaby się zmierzyć z wielkimi problemami, takimi jak np. hipoteza Riemanna, jedno z głównych wyzwań matematycznych umieszczone w 1900 roku wśród dwudziestu trzech innych słynnych problemów Davida Hilberta. Hipoteza ta pozostawała nierozwiązana przez cały XX wiek i jest obecnie najstarszym spośród problemów milenijnych, za których rozwiązanie przysługuje milionowa Millennium Prize. Wszystkie problemy milenijne wymagają dowodów, a nie obliczeń. W chwili obecnej nie dysponujemy algorytmami znajdowania i przeprowadzania dowodów. Nawet gdybyśmy je mieli, dowody prawdopodobnie wymagałyby jeszcze potężniejszych maszyn, jak się często mówi, komputerów kwantowych o dziesiątkach tysięcy kubitów mocy obliczeniowej.

Maszyna sama nie wie, czy może rozwiązać każdy problem

Problemem, co do którego były nadzieje, że AI go rozwiąże, jest znane pytanie – „czy $P=NP$?”. To wielkie wyzwanie teoretyczne, którego pełny opis jest niełatwy i nie ma tu na to miejsca. Sformułowany w latach 70. XX wieku niezależnie przez informatyków Stephena Cooka i Leonida Levina, problem to, jak często się określa, pytanie o to, jak łatwo jest rozwiązać i czy da się rozwiązać dany problem za pomocą komputera. Litera P reprezentuje problemy, które okazały się wykonalne do rozwiązania, co oznacza, że czas na obliczenie rozwiązania nie jest poza zasięgiem, i których rozwiązanie jest również łatwe do zweryfikowania, co oznacza, że można sprawdzić, czy odpowiedź jest poprawna. Litery NP oznaczają problemy, których odpowiedź jest również stosunkowo łatwa do zweryfikowania, podobnie jak P, ale dla których nie jest znany prosty sposób na obliczenie rozwiązania. Powszechnie, jako przykład NP, podaje się grę sudoku. Każdy wypełniony diagram sudoku może być dość łatwo sprawdzony pod kątem dokładności, ale zadanie znalezienia rozwiązania rośnie wykładniczo pod względem wymaganego czasu, gdy siatka gry staje się coraz większa. Problem $P=NP$? pyta zatem, czy problemy, które uważamy za trudne do rozwiązania,

NP, ale o których wiemy, że są łatwe do zweryfikowania, mogą w rzeczywistości okazać się zarówno łatwe do zweryfikowania, jak i łatwe do rozwiązania, tak jak problemy P. Negatywna odpowiedź, że P nie równa się NP, oznaczałaby, że niektóre problemy są poza możliwościami obliczeniowymi komputerów, nawet przy ogromnych nakładach – innymi słowy, istnieje górna granica obliczeń. Wyzwania takie jak złamanie niektórych szyfrów wydawałyby się wtedy trudniejsze, poza zasięgiem komputerów.

Problem ten nie doczekał się przekonującego rozwiązania pomimo dziesięcioleci intensywnych badań. Teraz sięgnięto w zmaganiach z nim po generatywną sztuczną inteligencję. Naukowcy przyjęli podejście polegające na podsuwaniu pomysłów modelowi GPT-4 w celu znalezienia odpowiedzi. Pytanie – $P=NP$? jest teraz traktowane jako sesja wielu zapytań z modelem językowym GPT-4. W artykule zatytułowanym „Large Language Model for Science: A Study on P vs. NP”, Qingxiu Dong i współpracownicy programują duży model językowy GPT-4 OpenAI przy użyciu tego, co nazywają metodą sokratejską, czyli kilku tur czatu za pośrednictwem monitu z GPT-4. Metoda zespołu polega na pobieraniu argumentów z wcześniejszego artykułu i podawaniu do GPT-4 w celu uzyskania przydatnych odpowiedzi. Dong i zespół zauważają, że GPT-4 demonstruje argumenty pozwalające stwierdzić, że P w rzeczywistości nie jest równe NP. Twierdzą oni również, że praca ta pokazuje, że duże modele językowe mogą „odkrywać nowe spostrzeżenia”, a te mogą prowadzić do „odkryć naukowych”.

Według prac Takeshi Kojimy i zespołu z Uniwersytetu Tokijskiego i Google Research z 2022 roku, możliwe jest poprawienie zdolności dużych modeli językowych w niektórych zadaniach po prostu poprzez dodanie frazy „pomyślmy krok po kroku”, na początku prompta, wraz z przykładową odpowiedzią. Okazało się, że ta fraza była wystarczająca do wywołania „łańcucha przemyśleń” ze strony modelu językowego. To ten sam typ procedury, który Dong i jego zespół stosują w swojej „metodzie sokratejskiej”. Przez prawie sto rund podpowiedzi autorzy zasypują GPT-4 różnymi prośbami, które zagłębiają się w matematykę $P=NP$, poprzedzając każdą z podpowiedzi stwierdzeniem warunkującym GPT-4, które brzmi tak: „jesteś mądrym filozofem”, „jesteś matematykiem biegłym w teorii prawdopodobieństwa”. Innymi słowy, to rodzaj gry polegającej na naklonieniu GPT-4 do odgrywania roli. Strategia polegała na naklonieniu GPT-4 do udowodnienia, że P w rzeczywistości nie jest równe NP, najpierw zakładając, że tak jest na przykładzie, a następnie znajdując sposób rozpadu twierdzenia,

czyli podejście znane jako dowód przez zaprzeczenie. Trudno powiedzieć, czy wyniki osiągnięte przez Donga i zespół za pomocą GPT-4 faktycznie dowodzą, że P nie równa się NP, ponieważ artykuł Xu i Zhou nie jest jeszcze wystarczająco zweryfikowany naukowo.

AI zadowolona w świecie badań naukowych

Mimo to nie można powiedzieć, że sztuczna inteligencja jest w nauce nieprzydatna. Jeśli chodzi o innowacje, w których, pomimo decydującej w nich roli ludzi, sztuczna inteligencja także ma duży wkład, to można podać sporo przykładów z konkretnych dziedzin. Należą do nich badania materiałów o wysokiej wytrzymałości, przewodników elektrycznych i termicznych, nadprzewodników, katalizatorów, izolatorów i wielu innych nowych materiałów. AI może być używana do automatycznego projektowania i optymalizacji struktur materiałów, które cechują się nowymi, poszukiwanymi właściwościami. Ponieważ osiągnięcia w tej dziedzinie rzadko są przypisywane pojedynczym badaczom, to algorytmy, jako niejako „członkowie zespołów”, również powinny być nagradzane miejscem na liście autorów i wynalazców (4).

Szczególnie dużo sztucznej inteligencji zawdzięcza branża farmaceutyczna. Firma ExclincPharma opracowała w ostatnich latach platformę AI specjalnego przeznaczenia, która służy do skrócenia czasu potrzebnego na prace nad nowymi lekami. Inna firma, Atomwise, opracowała algorytm służący do przyspieszenia procesu opracowywania nowych leków poprzez automatyczne identyfikowanie potencjalnych kandydatów na leki. Z kolei Numenta przygotowała platformę AI do identyfikacji potencjalnych terapii chorób neurodegeneracyjnych, takich jak choroba Alzheimera. Po sztuczną inteligencję sięgnęła również Insilico, opracowując spersonalizowane terapie i leki na bazie genomu.



4. ChatGPT jako współautor publikacji naukowej

Istnieje także wiele różnych algorytmów uczenia maszynowego opracowanych przez samą sztuczną inteligencję. Oto kilka najbardziej popularnych:

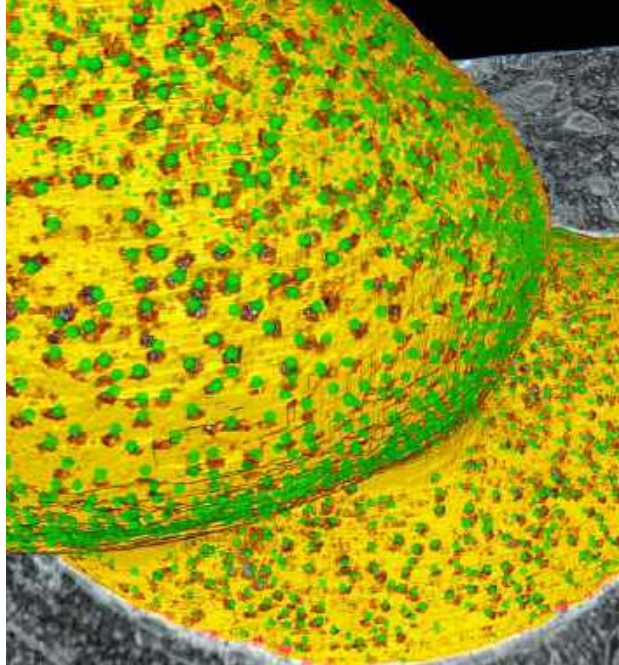
- regresja liniowa – mechanizm służący do prognozowania wartości numerycznych na podstawie danych historycznych;
- drzewa decyzyjne – to z kolei algorytm opracowany do rozwiązywania problemów klasyfikacji i regresji, który opiera się na podejmowaniu decyzji na podstawie serii pytań.

Algorytm k najbliższych sąsiadów (lub algorytm k-nn z ang. k nearest neighbours, KNN) – do klasyfikacji obiektów na podstawie danych o ich cechach i danych o innych, podobnych obiektach;

Także sieci neuronowe służące do nauki modeli AI oraz algorytmy głębokiej nauki powstają przy asyście AI, a niekiedy można mówić, że to sztuczna inteligencja jest ich główną autorką.

W ostatnich latach, po raz pierwszy w historii, naukowcy ze szwajcarskiej politechniki ETH w Zurychu z powodzeniem zautomatyzowali proces modelowania turbulencji w płynach przez połączenie mechaniki płynów ze sztuczną inteligencją. Ich podejście opiera się na połączeniu algorytmów wzmocnionego uczenia maszynowego z symulacjami przepływów turbulentnych przeprowadzonymi w superkomputerze Piz Daint ze Szwajcarskiego Narodowego Centrum Superkomputerowego. Według opisu badań, który ukazał się w periodyku „Nature Machine Intelligence”, naukowcy opracowali nowe algorytmy wzmocnionego uczenia maszynowego (RL) i połączyli je z fizycznym podejściem do modelowania turbulencji.

Naukowcy z Europejskiego Laboratorium Biologii Molekularnej (EBML) połączyli algorytmy sztucznej inteligencji z dwiema najnowocześniejszymi obecnie technikami mikroskopowymi. Dzięki temu czas przetwarzania obrazów skrócił się z dni do zaledwie sekund, a uzyskane obrazy są ostre i dokładne. Wyniki badań zostały opublikowane w „Nature Methods”. Mikroskopia pola świetlnego rejestruje duże obrazy 3D, które pozwalają badaczom śledzić i mierzyć bardzo drobne ruchy, np. bijące serce larwy ryby, przy bardzo dużych prędkościach. Jednak technika ta generuje ogromne ilości danych, których przetwarzanie może trwać wiele dni, a końcowe obrazy zazwyczaj nie mają odpowiedniej rozdzielczości. Mikroskopia arkusza świetlnych skupia się na pojedynczej płaszczyźnie 2D danej próbki w jednym czasie, dzięki czemu badacze mogą obrazować próbki w wyższej rozdzielczości. W porównaniu do mikroskopii pola świetlnego mikroskopia arkusza świetlnego daje obrazy,



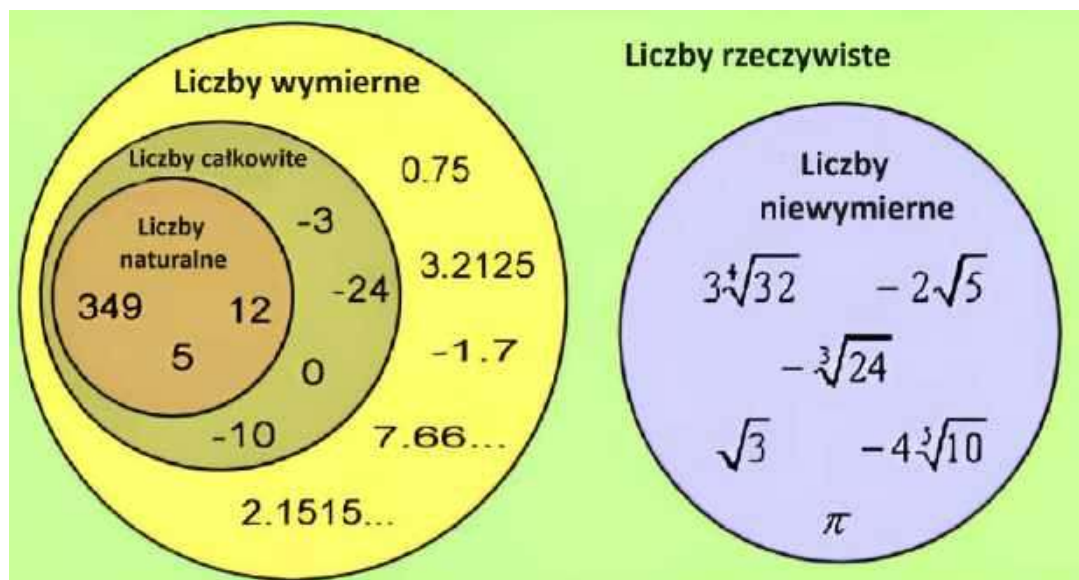
5. Obrazy uzyskane z wykorzystaniem analizy danych mikroskopowych przez algorytmy AI

które są szybsze w obróbce, ale dane nie są tak wszechstronne, ponieważ przechwytyują informacje tylko z jednej płaszczyzny 2D w danym momencie. Aby wykorzystać zalety każdej z technik, naukowcy z EMBL opracowali podejście, które wykorzystuje mikroskopię pola światła do obrazowania dużych próbek 3D i mikroskopię arkusza do szkolenia AI, która następnie tworzy dokładny obraz 3D próbki (5).

Rosyjscy naukowcy z Moskiewskiego Instytutu Fizyki i Technologii, Instytutu Fizyki i Technologii im. Waliewa oraz Uniwersytetu ITMO stworzyli w 2020 r. sieć neuronową, która nauczyła się przewidywać zachowanie systemu kwantowego przez „przyglądanie się” jego strukturze podczas tzw. „spacerów kwantowych”, które można zobrazować jako podróż cząstki w pewnej sieci, na której opiera się obwód kwantowy. Wyniki ich badań zostały opublikowane w „New Journal of Physics”. „Udało nam się wytrenować komputer do samodzielnego przewidywania, czy złożona sieć ma przewagę kwantową”, pisał w publikacji Leonid Fedyczkin z Wydziału Fizyki Teoretycznej MIFT.

Są to oczywiście tylko wrywki z rozlicznych przykładów. Nauka bardzo szeroko korzysta z technik sztucznej inteligencji, od algorytmów pomagających przetwarzać ogromne ilości danych, np. z obserwacji astronomicznych, przez tworzenie nowych metod rozwiązywania problemów matematycznych, po odkrywanie zjawisk i podejść, na które nauka dotychczas nie wpadła w fizyce teoretycznej (temat poruszany w MT). ■

Mirosław Usidus



1. Liczby wymierne i niewymierne

Jaka jest najdziwniejsza liczba rzeczywista, jaką można sobie wyobrazić? Prawdopodobnie wiele osób myśli o liczbach niewymiernych, takich jak pi (π) lub liczba Eulera. Jednak nawet liczby, których nigdy nie można zapisać do końca, wraz z liczbami wymiernymi stanowią jedynie część liczb rzeczywistych (1). Co więcej, według matematyków, większości liczb nie możemy obliczyć.

Większość liczb rzeczywistych jest nieznana – nawet matematykom

CZY COŚ, CZEGO NIE DA SIĘ POLICZYĆ, ISTNIEJE?

Liczby wymierne można przedstawić za pomocą ułamka zwykłego. Pozostałe liczby na osi liczbowej to liczby niewymierne. W ich zbiorze również wyodrębnia się różne kategorie, a większości z nich nie potrafimy sobie nawet wyobrazić. Są trudne lub

niemożliwe do dokładnego zapisania liczby nieskończenie wielkie, nieskończenie małe, urojone lub inne nietypowe przestrzenie liczbowe.

Te, które da się skonstruować i te, których się nie da

Przykładową liczbą niewymierną jest np. pierwiastek kwadratowy z 2, którego reprezentacja dziesiętna jest nieskończona i nie wchodzi w żaden znany okres. Przy czym $\sqrt{2}$ można skonstruować, czyli można wygenerować geometrycznie za pomocą cyrkla i linijki, rysując trójkąt prostokątny o dwóch bokach o długości jednej jednostki. Przeciwprostokątna tego trójkąta ma wtedy długość $\sqrt{2}$. W podobny sposób geometrycznie da się skonstruować także złotą proporcję Φ i wiele innych wartości niewymiernych.

Jednak nawet w starożytności ludzie znali liczby, których nie można było przedstawić w tak prosty, geometryczny sposób. Znanym przykładem jest podwojenie sześcianu. W jaki sposób sześcian o boku długości 1 można przekształcić w sześcian o dwukrotnie

$$L = \sum_{n=1}^{\infty} \frac{1}{10^{n!}} = 0,110001000$$

2. Liczba Liouville'a

większej objętości? Jak orzekł w 1837 roku matematyk Pierre Wantzel, długości krawędzi $\sqrt[3]{2}$ wymaganej dla takiego podwojonego sześciangu nie można skonstruować za pomocą linijki. Jednak $\sqrt[3]{2}$ należy wciąż do liczb algebraicznych, które można zapisać jako rozwiązanie równania wielomianowego. Dla $\sqrt[3]{2}$ odpowiednim równaniem jest $x^3=2$.

Nie ma prostego wzoru, za pomocą którego można obliczyć takie „transcendentalne” lub przestępne liczby. Podaje się, że pierwszą znaną liczbą tego typu była liczba Liouville'a (2), choć do tej kategorii zaliczana jest znana znacznie dawniej liczba π (3). Jeszcze grecki matematyk Archimedes znalazł regułę obliczania π w przybliżeniu. Dziś istnieje wiele algorytmów, które potrafią obliczyć ją do prawie sześciuset milionów miejsc po przecinku. Dysponując wystarczającą mocą obliczeniową i czasem, liczbę tę można wyznaczyć z dowolną dokładnością, przynajmniej w teorii. Jednak wciąż nie będzie obliczona „do końca”. To samo dotyczy liczby Eulera (e), czyli $2\sqrt{2}$.

Chociaż istnieją metody pozwalające stwierdzić, czy liczba jest konstruowalna, trudno jest udowodnić, czy dana wartość jest przestępna, czyli nie należy do liczb algebraicznych. Na przykład w 1934 roku radziecki matematyk Aleksander Gelfond udowodnił, że liczba zespolona e^{π} jest przestępna. Jednak to, czy wartości π^e , $\pi \times e$ lub $\pi - e$ są algebraiczne czy przestępne, do dziś pozostaje niejasne.



3. Symbol liczby pi skonstruowany z cyfr rozwinięcia ułamka

Turing i wkroczenie do świata nieobliczalności

Do XX wieku ludzie zakładali, że liczby przestępne są najbardziej egzotyczną rzeczą, na jaką możemy natrafić w zbiorze liczb rzeczywistych. Jednak w 1937 roku brytyjski matematyk Alan Turing (4) opublikował artykuł na temat czegoś jeszcze „dzikszego”, mianowicie zagadnienia obliczalności liczb. Użył tego terminu do opisanie wszystkich wartości, dla których istnieje reguła obliczeniowa (czyli algorytm), którą komputer może wykonać, aby obliczyć wartość liczbową z dowolnym stopniem dokładności.

Prawie wszystkie znane liczby przestępne, np. π i e , należą z tego punktu widzenia do kategorii obliczalnych. W końcu znamy sposoby (algorytmy) obliczania ich przynajmniej przybliżonej wartości liczbowej. Jak wykazał Turing w swojej pracy,



4. Alan Turing i model tzw. maszyny Turinga



istnieją jednak też liczby nieobliczalne, czyli takie, których wartości nie mogą być przybliżone z dowolną precyzją. To znaczy, że nie mamy pojęcia, jaką mają wartość, nawet w przybliżeniu.

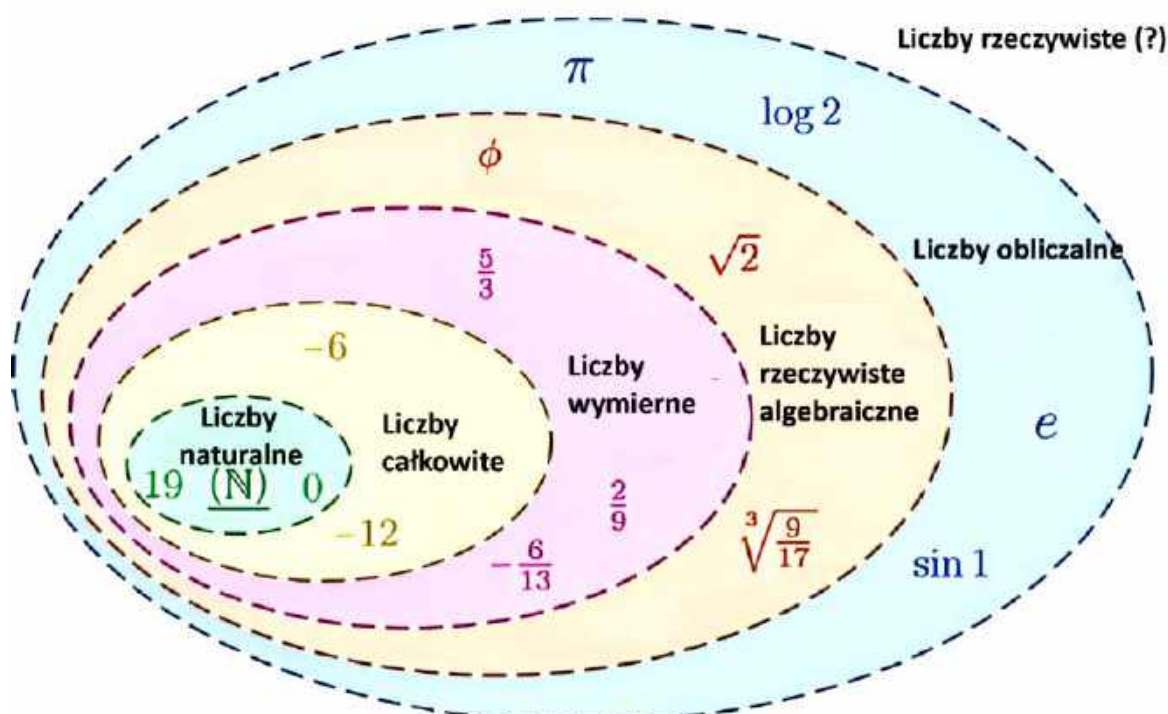
Można to spróbować wyjaśnić w możliwie prosty sposób (w tym artykule nie wchodzimy w gęszcze równań), nawiązując do nieskończonych rozmiarów różnych zbiorów liczbowych, którymi pierwszy zajął się Georg Cantor w XIX wieku, wykazując na przykład, że zbiór liczb naturalnych, całkowitych i wymiernych ma tę samą kardynalność (moc zbioru). Jak wykazał Cantor, kardynalność liczb naturalnych jest najmniejszą możliwą nieskończonością. Nazwał ją „policzalnie nieskończoną”. Liczb rzeczywistych nie można policzyć. Cantor dowodził, że kardynalność liczb rzeczywistych jest z konieczności większa niż liczb naturalnych. Liczby rzeczywiste tworzą zatem zbiór niepoliczalny.

Jak jednak stwierdził Turing, wszystkie liczby obliczalne muszą być policzalne w rozumieniu Cantora. Dla każdej z nich można opracować maszynę, która po prostu oblicza jej wartość. Ponieważ można ponumerować maszyny liczące te liczby, liczby obliczalne są z konieczności policzalne. Nieformalnie możemy myśleć o liczbie obliczalnej tak, jak zrobił to Marvin Minsky, jeden

z ojców AI, a mianowicie jako o „liczbie, dla której istnieje maszyna Turinga, która, biorąc pod uwagę n na jej początkowej taśmie, kończy się n -tą cyfrą tej liczby (zakodowaną na jej taśmie)”.

Z drugiej strony prowadzi to do sytuacji, w której liczby niepoliczalne, dla których nie można zbudować maszyny obliczającej, stanowią większość liczb rzeczywistych – jest ich niepoliczalna liczba (5). Liczby nieobliczalne znamy niewiele. Niektóre znane przykłady liczb nieobliczalnych zostały zdefiniowane przez znany z informatyki problem zatrzymania (haltingu), opracowany przez Turinga. Opisuje się je przez sytuację, w której komputer wykonujący określony zestaw instrukcji w celu rozwiązania problemu (innymi słowy, komputer używający algorytmu) zatrzymuje się w pewnym momencie, zamiast kontynuować obliczenia w nieskończoność. Jak udowodnił Turing, chociaż taka maszyna mogłaby ocenić, czy niektóre algorytmy mogą być wykonywane w skończonym czasie, nie istnieje żadna metoda, która mogłaby to zrobić dla wszystkich możliwych kodów programowania.

Inaczej to ujmując, według Turinga liczba rzeczywista, dla której istnieje algorytm, mechaniczna procedura obliczania jej wartości z dowolną dokładnością, z coraz większą precyzją, jest



5. Zbiory liczb aż do obliczalnych i hipotetyczny świat liczb nieobliczalnych

obliczalna. Zatem na przykład π jest obliczalną liczbą rzeczywistą, $\sqrt{2}$ jest obliczalną liczbą rzeczywistą, a także e jest obliczalną liczbą rzeczywistą. Zbiór obliczalnych liczb rzeczywistych jest tak liczny jak metody obliczania, algorytmy, programy komputerowe. A programów komputerowych jest najwyżej tyle, ile liczb całkowitych 1, 2, 3, 4, 5, ... – można myśleć o programie komputerowym jak o bardzo dużej liczbie całkowitej.

Problem haltingu jest bezpośrednim zastosowaniem twierdzenia o niekompletności innego wybitnego matematyka, Kurta Gödla, które, najogólniej to wyjaśniając, mówi, że nie wszystkie twierdzenia matematyczne można udowodnić. Twierdzenie o niekompletności Gödla z 1931 roku mówi, że jeśli wziąć pod uwagę dowolny skończony zestaw aksjomatów, będą istniały prawdziwe twierdzenia matematyczne, które wymykają się z tego systemu, ale nie można ich udowodnić na podstawie tych aksjomatów.

Argentyńsko-amerykański matematyk Gregory Chaitin w 1975 roku zaproponował liczbę rzeczywistą, która nie jest obliczalna, obecnie nazywaną stałą Chaitina, omega (Ω). Reprezentuje ona prawdopodobieństwo, że losowo skonstruowany program zatrzyma się/zakończy. Załóżmy, że uruchamiamy uniwersalną maszynę Turinga na losowym programie binarnym. W szczególności, za każdym razem, gdy wymagany jest kolejny bit programu, rzucamy monetą i podajemy binarne wyjście do maszyny. Jakie jest prawdopodobieństwo, że maszyna Turinga się zatrzyma? Eksperyment myślowy Chaitina jest prototypowym przykładem problemu zatrzymania. Problem haltingu to problem polegający na określeniu, na podstawie opisu dowolnego programu komputerowego i danych wejściowych, czy program zakończy działanie (zatrzyma się), czy będzie działał w nieskończoność. Aby dokładnie obliczyć stałą Chaitina, trzeba by wiedzieć, które programy zatrzymują się, a które nie, co nie jest możliwe, zgodnie z problemem haltingu. W 2000 roku matematykowi Cristianowi S. Calude i jego kolegom udało się obliczyć pierwsze cyfry stałej Chaitina: 0,0157499939956247687... Oznacza to, że jeśli losowo wygenerujemy program w języku używanym przez Caludea i jego współpracowników, będzie on działał (zatrzymywał się) z prawdopodobieństwem około 1,58 proc. w skończonym czasie działania. Nawet jeśli wynik ma wysoką dokładność, stała Chaitina nie może być obliczona z dowolną precyzją.

Istnieją inne skomplikowane metody definiowania liczb nieobliczalnych. Mimo to, biorąc pod uwagę

obfitość liczb obliczalnych, które znamy, zawsze zaskakujące jest to, że wartości nieobliczalne dominują w zbiorze liczb rzeczywistych.

Zniszczone marzenia o pewnikach

Współczesne badania nad obliczalnością rozpoczęły się około 1900 roku wraz z ogłoszeniem programu Hilberta, mającego na celu aksjomatyzację podstaw matematyki. Stało się to po tym, jak sam Hilbert pokazał w 1899 roku, że spójność geometrii sprowadza się do spójności liczb rzeczywistych. Program Hilberta (1904) został zatem ustanowiony, aby spróbować podobnie wykorzystać skończoność dowodów matematycznych, aby pokazać, że nie można wyprowadzić sprzeczności. Gdyby taki fundament mógł zostać ustanowiony, dowody matematyczne byłyby całkowicie wyrażalne w kategoriach symboli logicznych. Gdyby udało się to osiągnąć, wszystkie dowody, w tym m.in. dowody ostatniego twierdzenia Fermata, hipotezy Riemanna i przypuszczenia Poincarégo, byłyby nieuniknionymi konsekwencjami systemu matematycznego, w którym są wyrażone. I w końcu, jeśli taki „program” zostałby zrealizowany, to można by zbudować wystarczająco sprawną maszynę, pozwolić jej działać wystarczająco długo, aby ostatecznie wszystkie twierdzenia matematyczne zostały ujawnione.

Niecałe 30 lat później wspomniany Kurt Gödel (1906–1978) opublikował pracę „Über formal unentscheidbare Sätze der Principia Mathematica und verwandter Systeme I”, w której zawarł „twierdzenie VI” niszczące marzenie Hilberta o w pełni zaksymatyzowanym, niepodważanym fundamencie matematyki. Gödel pokazał, że jakikolwiek system będziemy konstruować, nieważne jak solidny, zawsze będzie z natury niekompletny. Zawsze będą istniały stwierdzenia wyrażalne w systemie, których nie da się udowodnić, biorąc pod uwagę jego aksjomaty (pewniki).

Czy nieobliczalne/nieopisywalne liczby rzeczywiste „istnieją”? To pytanie o charakterze filozoficznym. Np. zwolennik szkoły platońskiej powie, że istnieją, nawet jeśli nie mamy możliwości, by je skonstruować. Finitysta (jest taki kierunek w „filozofii matematyki”) może powiedzieć, że nie istnieją, dlatego że nie mamy algorytmu do obliczenia lub nawet rozpoznania takiej liczby. Zatem odpowiedź na pytanie – co wynika z tego, że zdecydowanej większości liczb rzeczywistych nie da się obliczyć – zależy od filozofii. I taka może być niespodziewana pointa rozważań o obliczalności liczb. ■

Mirosław Usidus

Jak widać, aby uświadomić ludziom, że fizyka to ciężkie zmagania z problemami, których od wieków nie da się rozgryźć, mimo wielkich postępów, jakie poczyniono, potrzeba popularnego filmu lub serialu, takiego jak „Problem trzech ciał” Netfixa na podstawie bestsellera Cixin Liu (1).

Lista nieznanących od dawna rozwiązania zagadek fizyki jest bardzo długa

NIE TYKO PROBLEM TRZECH CIAŁ

Zagadka ruchów trzech ciał, które wywierają na siebie siły grawitacyjne, nurtuje naukowców od czasów Izaaka Newtona. Choć zrazu może się to wydawać proste zadanie, przewidywanie ich orbit pozostaje dla nas zadaniem niewykonanym. Było to „pierwsze prawdziwe zmartwienie Newtona”, jak to ujął w rozmowie z serwisem „LiveScience” Billy Quarles z Uniwersytetu Valdosta w stanie Georgia w USA.

W układzie składającym się tylko z dwóch ciał, np. planety i gwiazdy, obliczenie ich wzajemnych torów ruchu jest dość proste. Te dwa obiekty będą orbitować wokół swojego środka masy i za każdym razem będą wracać do miejsca, w którym zaczęły cykl. Jeśli jednak dodamy trzecie ciało, takie jak druga gwiazda w układzie podwójnym, sprawa staje się o wiele bardziej skomplikowana. Trzecie ciało przyciąga pozostałe obiekty, wyrывая je z przewidywalnych ścieżek. Ruch ciał zależy od ich stanu początkowego – pozycji, prędkości i mas. Jeśli choćby jeden z tych parametrów ulegnie zmianie pod wpływem dodatkowych sił, ruch wynikowy może się całkowicie zmienić. Fizycy i matematycy pracują nad rozwiązaniem problemu trzech ciał od wieków. Nie ma jednej poprawnej odpowiedzi – jest wiele orbit, które są zgodne z prawami fizyki dla trzech orbitujących obiektów.

Znaleziono rozwiązania tego problemu w niektórych, szczególnych przypadkach. Na przykład, jeśli warunki początkowe są odpowiednie, trzy ciała o równej



1. Okładka książki „Problem trzech ciał”

masie mogą „ścigać się” nawzajem w układzie ósemkowym. Szczególnym hipotetycznym przypadkiem jest też fikcyjny układ Tatooine z *Gwiezdnych Wojen*, w którym pojedyncza planeta krąży wokół dwóch słońc. To układ trzech ciał, ale jeśli planeta znajduje się wystarczająco daleko i krąży wokół obu gwiazd, możliwe jest uproszczenie problemu. Planeta nie wywiera dużej siły na gwiazdy, ponieważ jest o wiele mniej masywna, więc system staje się podobny do układu dwóch ciał. Jak dotąd naukowcy znaleźli ponad tuzin egzoplanet w układach podobnych do wymyślnego Tatooine.

Zdarza się jednak, że orbity trzech ciał nigdy się nie stabilizują, a problem trzech ciał zostaje „rozwiązany” w katastrofie kosmicznej. Grawitacyjnie może spowodować zderzenie dwóch z trzech ciał lub wyrzucenie jednego z ciał z układu na zawsze. W takich wydarzeniach upatruje się pochodzenia „bezzaśpisznych” planet, niekrążących wokół żadnej gwiazdy, których tak wiele znaleźliśmy w kosmosie.

Na planecie ukazanej w filmie i w książce chińskiego autora rozwija się cywilizacja niszczona co pewien czas przez kosmiczny chaos układu trzech ciał. Jednak odradza się. To, w ocenie naukowców, czysta fantazja, gdyż wątpliwe, by w tak niestabilnym układzie mogło powstać i przetrwać życie, a co dopiero cywilizacja.

Zespół matematyków z Uniwersytetu Sofijskiego w Bułgarii ogłosił we wrześniu 2023 r., że znalazł dwanaście tysięcy nowych metod rozwiązania problemu układu trzech ciał. Ich praca została opublikowana w naukowej bazie danych arXiv. Ivan Hristov, matematyk i jego współpracownicy obliczyli te wszystkie teoretyczne orbity w takim układzie za pomocą superkomputera i są przekonani, że przy jeszcze lepszej technologii mógłby on znaleźć ich „pięć razy więcej”. Czy to jest ścieżka do rozwiązania „problemu trzech ciał”? Zobaczmy.

Za mało „energii próżni”?

Młodsza znacznie niż problem trzech ciał mechanika kwantowa to prawdziwe kłębownisko zagadek i pytań bez odpowiedzi. Nie dopuszcza np. nicości, czyli czegoś, co w świecie naszych pojęć funkcjonuje i wydaje się nam oczywistym składnikiem Wszechświata (2). W dowolnym momencie w czasie i przestrzeni energia nigdy nie może być idealnie zerowa – przekonują fizycy, negujący nicość. W „całkowicie pustym” miejscu miejsca mogą powstawać „wirtualne” cząstki – w szczególności para składająca się z cząstki i jej antycząstki, które anihilują się nawzajem i znikają tak szybko, jak się pojawiły. Eksperymentatorzy zaobserwowali rzeczywiste efekty wirtualnych cząstek. Kiedy akceleratory cząstek po raz pierwszy zmierzyły masę bozonu Z, była ona nieco inna niż jego czysta masa, ponieważ czasami zamieniał się on w wirtualny górny kwark. Była to jedna z wielu obserwacji dowodzących istnienia cząstek wirtualnych.

Efektom tych cząstek, które w mgnieniu okaz, powstają i znikają, jest pulsująca w oscylacjach „energia próżni”, która wypełnia kosmos i „wywiera nacisk” na przestrzeń, co uznaje się za jedno z bardziej prawdopodobnych wyjaśnień ciemnej energii, która z kolei ma być główną siłą sprawczą przyspieszającej ekspansji Wszechświata. Problem z energią próżni polega na tym, że jest jej za mało, by wyjaśnić tę ekspansję. Kiedy naukowcy po raz pierwszy zaczęli analizować tę koncepcję, wyliczyli, że energia ta powinna być ogromna, powinna rozpychać Wszechświat tak silnie i szybko, że nigdy nie uformują się żadne gwiazdy i galaktyki. Jednak z wyliczeń i obserwacji wynika, iż energia próżni we Wszechświecie musi być około

120 rzędów wielkości mniejsza, niż przewiduje teoria kwantowa. Rozbieżność ta skłoniła niektórych naukowców do nazwania energii próżni „najgorszym teoretycznym przewidywaniem w historii fizyki”.

Jednak pojęcie to nie zostało porzucone. Wciąż uważa się, że energia próżni jest głównym składnikiem „stałej kosmologicznej”, matematycznej wartości obecnej w równaniach ogólnej teorii względności. Wartość stałej kosmologicznej podana przez ogólną teorię względności jest bardzo mała. Wartość podawana przez kwantową teorię pola jest znacznie większa. To niespójność między fundamentalnymi teoriami. Ogromna rozbieżność między przewidywaną ilością energii próżni a zmierzoną ilością jest zresztą nazywana problemem stałej kosmologicznej. Nie wszyscy zgadzają się, że problem znany od czasów, gdy Einstein umieścił stałą kosmologiczną w równaniach, wymaga naprawy. Stała kosmologiczna jest technicznie tylko liczbą w równaniu, która może przyjąć dowolną wartość, przekonuje znana niemiecka fizyk Sabine Hossenfelder. Jak uważa, fakt, że ma ona taką wartość, jaką ma, jest tylko liczbowym zbiegiem okoliczności.

Wielu fizyków nazywa problem stałej kosmologicznej a także zagadkę energii próżni „czymś zawstydzającym”, jednak alternatywne teorie budzą również mnóstwo wątpliwości. Poza tym naukowcy są bardzo przywiązani do mechaniki kwantowej, która stoi za tym pojęciem i nie chcą szukać „innej fizyki”, bo to grząski teren. Większość proponowanych rozwiązań problemu stałej kosmologicznej dzieli się na trzy kategorie: zmiana równań ogólnej teorii względności, które opisują ekspansję Wszechświata, modyfikacja równań kwantowej teorii pola, które przewidują ilość energii próżni lub wprowadzenie

2. Jak AI wyobraża sobie energię próżni



3. Wizja kosmicznej piany

czegoś zupełnie nowego. Modyfikacja ogólnej teorii względności mogłaby zmienić matematyczną rolę, jaką odgrywa stała kosmologiczna lub całkowicie ją wyeliminować, np. zmieniając sposób, w jaki obliczenia ogólnej teorii względności powinny być stosowane do rozszerzającego się kosmosu. Według tej koncepcji sama materia i fotony mogą być wystarczające, bez dodawania żadnego nowego składnika do Wszechświata, jednak teoria ta wymaga dodatkowych wymiarów, poza trzema wymiarami przestrzeni i jednym czasem. A te są poza naszym zasięgiem poznawczym. Innym podejściem do aktualizacji ogólnej teorii względności jest tzw. sekwestracja. Modyfikuje teorię Einsteina w taki sposób, by odciąć grawitację od wpływu energii próżni. Nie jest to wszystko zbyt jasne, ale poszukiwania alternatyw trwają.

Pojawiają się zresztą nowe dane, które być może są śladem nowego wyjaśnienia problemu energii próżni. Badania fal grawitacyjnych ujawniły np. dziwną aktywność wewnątrz gwiazd neutronowych. Rozbitec atomów w ich poddanych ogromnym siłom jądrach prowadzi do dziwnych zjawisk, nowych faz materii i wzrostów energii próżni tam, gdzie nie powinno być ich w ogóle.

Piana czasoprzestrzenna obok kontrowersyjnej „zasady antropicznej”

Jeśli jednak ogólna teoria względności nie jest problemem, to być może jest nim mechanika kwantowa. Niektórzy teoretycy sugerują, że metoda obliczania energii próżni w kwantowej teorii pola jest nieprawidłowa. Stefan Hollands z uniwersytetu w Lipsku w Niemczech i jego koledzy nie zgadzają się z zastosowaniem regularnych równań

kwantowych do zakrzywionej czasoprzestrzeni, twierdząc, że zostały one zaprojektowane z myślą o płaskiej przestrzeni. Argumentują, że gdyby fizycy byli w stanie poprawnie zmodyfikować je dla zakrzywionej przestrzeni, problem stałej kosmologicznej (i braku wystarczającej energii próżni) zniknąłby. Z kolei Steve Carlipa z Uniwersytetu Kalifornijskiego w Davis proponuje, by traktować czasoprzestrzeń zasadniczo jako „pianę” (3). W takim ujęciu krzywizna przestrzeni podlegałaby ciągłym fluktuacjom w niezwykle małych skalach, znacznie wykraczających poza nasze możliwości pomiarowe. Cała ta skomplikowana struktura zniwelowałaby znaczną część wpływu stałej kosmologicznej, czyniąc ją bardzo małą na poziomie lokalnym.

Jednym z najbardziej znanych, a przez niektórych najbardziej zniechęconych, rozwiązań problemu stałej kosmologicznej jest tzw. zasada antropiczna. Ten kierunek myślenia nie przeczy, że stała kosmologiczna w naszym Wszechświecie ma nieprawdopodobną wartość, ale wyjaśnia to kontekstem czegoś większego, w czym istniejemy, mianowicie multiwszechświata. Jeśli nasz Wszechświat jest tylko jedną bańką w kosmicznym morzu wszechświatów, z różnymi prawami fizycznymi i stałymi w każdym, to musiał istnieć jeden, w którym jest tak akurat jak w naszym. Większość innych nie prowadzi do powstania galaktyk, gwiazd, planet i życia. Ale w tej mnogości musiał powstać też taki jak nasz, prowadzący w konsekwencji do powstania życia i istot, które zadają takie właśnie pytania. Ta teoria oburza wielu fizyków jako „skandaliczna rezygnacja z poszukiwania wyjaśnienia”. Jeden z nich, Lee Smolin, nazywa teorie multiwersum nienaukowymi, grząc, że nie należy wymyślać niesprawdzalnych hipotez tylko dlatego, że dostrojenie naszego Wszechświata nas denerwuje.

Być może w ciągu najbliższych kilku dekad nowe eksperymenty i obserwacje dadzą naukowcom więcej wiedzy o tym, czy stała kosmologiczna (i stojąca za nią energia próżni) jest źródłem ciemnej energii. Projekt Vera C. Rubin Observatory Legacy Survey of Space and Time, który ma rozpocząć się w 2025 roku na teleskopie budowanym obecnie w Chile, powinien znacznie poprawić precyzję pomiarów historii kosmicznej ekspansji.

Wielki Wybuch? Na pewno?

Jeśli ktoś uważa Wielki Wybuch za pewnik, to może zdziwi go, że w fizyce jest to bardziej zagadka i problem niż dogmat. Teoria mówiąca o początkowej wielkiej erupcji nie została wcale zadowalająco udowodniona. Ale nawet jeśli okaże się, że jest poprawna, pozostaje pytanie – co wydarzyło się przed Wielkim Wybuchem? Jak został on zainicjowany? Czy Wszechświat istniał od zawsze? Czy Wszechświat jest częścią multiwersum i tam trzeba szukać źródła Wielkiego Wybuchu? Można zadać mniej ogólne pytanie – dlaczego pozostałość po Wielkim Wybuchu, CMB (lub kosmiczne mikrofalowe tło), ma takie same właściwości (np. temperaturę) we wszystkich kierunkach, zwłaszcza jeśli odległe, rozłączone regiony nigdy nie miały czasu na wymianę informacji między sobą? Albo – dlaczego nie ma żadnych pozostałości wysokoenergetycznych z tej rzekomo gorącej fazy, którą Wszechświat osiągnąłby na wczesnym etapie?

Jedną z możliwości, zaproponowaną przez Alana Gutha (4) na przełomie lat 70. i 80. XX wieku, jest hipoteza, że Wielki Wybuch nie był początkiem, ale był poprzedzony stanem wykładniczo rozszerzającej się pustej przestrzeni, która poprzedza i przygotowuje grunt pod to, co nazywamy gorącym Wielkim Wybuchem. Te rozważania doprowadziły do stworzenia teorii inflacyjnego Wszechświata, w którym przestrzeń była wypełniona pewnym rodzajem energii, być może podobnej do dzisiejszej ciemnej energii, co powodowało, że rozszerzał się nie tylko szybko, ale także nieustannie i bez ograniczeń. Kiedy inflacja dobiegała końca, cała (a przynajmniej jej większość) tej energii miała być przekształcona w cząstki i antycząstki, inicjując fazę Wszechświata, którą utożsamiamy z gorącym Wielkim Wybuchem. Obecnie różne regiony mają tę samą temperaturę i gęstość, ponieważ wszystkie wyłoniły się z tego samego stanu inflacji. Wszechświat, który się wyłonił, wydaje się przestrzennie płaski, ponieważ proces inflacji rozciągnął go tak, że byłby nie do odróżnienia od prawdziwie idealnej płaszczyzny. Nie ma pozostałości wysokoenergetycznych relikwów, ponieważ wszelkie istniejące wcześniej

zostały rozdmuchane, a maksymalna temperatura, jaką Wszechświat osiąga, gdy rozpoczynał się gorący Wielki Wybuch, jest obecnie niewystarczająca do ich ponownego wytworzenia. Innymi słowy, inflacja nie tylko tłumaczy okoliczności Wielkiego Wybuchu, ale rozwiązuje wszystkie trzy główne fizyczne zagadki, które nekły teorię Wielkiego Wybuchu wcześniej.

Uczonym udało się potwierdzić szereg przewidywań hipotezy inflacji, wykluczając właściwie gorący Wielki Wybuch bez inflacji, a także wykluczając szereg modeli inflacyjnych, które nie pasują do danych. Mimo to stawia się szereg pytań dotyczących kosmicznej inflacji, na które nie udzielono jeszcze odpowiedzi, np. czy pierwotne fale grawitacyjne są obecne w naszym Wszechświecie? Chociaż inflacja rozciąga tkaninę Wszechświata tak, że nie można jej odróżnić od płaskiej, fluktuacje kwantowe odcisnęły podczas inflacji mogą również odcisnąć na niej niezerową krzywiznę przestrzenną. Pomiar krzywizny dokładnie w przewidywanym zakresie byłby spektakularnym potwierdzeniem inflacji, jeśli jednak wartość krzywizny będzie się różnić, może to oznaczać kłopoty dla teorii inflacji.

Worek problemów bez dna

Lista problemów wciąż nierozwiązanych przez fizykę jest oczywiście znacznie dłuższa, wręcz ogromna. Nie można jej nawet hasłowo wymienić w tak krótkim artykule. Wiele z tych zagadek, ponieważ dotyczy obiektów, których nie możemy badać z dostatecznie bliska ani obserwować, zapewne pozostanie w katalogu tajemnic jeszcze bardzo długo. Typowym przykładem tej klasy problemów jest pytanie – co dzieje się wewnątrz horyzontu zdarzeń czarnej dziury? Obecne teorie przewidują, że cała materia w czarnej dziurze gromadzi się w jednym punkcie w centrum, ale nie rozumiemy, jak działa ta centralna osobliwość.

4. Alan Guth



Odwieczną zagadką, nie tylko w dziedzinie fizyki, jest czas. Wszystkie teorie fizyczne są formułowane w postaci równań różniczkowych, które wyrażają, jak pewna wielkość zmienia się w odniesieniu do upływu czasu. Jednak teoria względności wskazuje, że czas nie powinien być wyróżniany jako odgrywający wyjątkową rolę w zestawieniu z przestrzenią. Kwantowe równanie pola dla ogólnej teorii względności (równanie Wheelera-Dewitta) traktuje czas symetrycznie z przestrzenią. Sugeruje



5. Próbką opartego na związkach miedzi nadprzewodnika wysokotemperaturowego

to, że czas musi być jak przestrzeń, czyli nie „staje się”, jak nam się to wydaje. Nie jest jednak wcale jasne, co to tak naprawdę oznacza. Być może czas, jakiego doświadczamy, jest zjawiskiem emergentnym (wyłaniającym się), przybliżającym głębszą strukturę. Takie jest założenie pętlowej grawitacji kwantowej. Oznaczałoby to, że żyjemy w całkowitej iluzji i że przemieszczamy się w czasie w taki sam sposób, w jaki przemieszczamy się przez przestrzeń. Przyjęcie takiego rozumienia prowadzi do radykalnej zmiany w sposobie myślenia o geometrii i fundamentalnej strukturze Wszechświata. Jednak kłóci się z tym jak postrzegamy, rozumiemy i „odczuwamy” czas. Nie będzie łatwo człowiekowi z tym sobie poradzić.

Czas to problem, choć fizyczny, to zarazem psychologiczno-filozoficzny. Jest wiele bardziej konkretnych pytań bezodpowiedzi. Dlaczego na przykład kwarki nie mogą istnieć swobodnie? Lub – dlaczego w obserwowalnym

Wszechświecie jest znacznie więcej materii niż antymaterii? Dlaczego jest ogromna rozbieżność między skalą grawitacyjną a typową skalą masy cząstek elementarnych? Jaki mechanizm powoduje, że niektóre materiały wykazują nadprzewodnictwo w temperaturach znacznie wyższych niż około 25 kelwinów? (5) Czy możliwe jest stworzenie teoretycznego modelu opisującego statystykę przepływu turbulentnego (w szczególności jego wewnętrznej struktury)? Ponadto, w jakich warunkach istnieją pro-

ste rozwiązania równań Naviera-Stokesa, opisujących zasadę zachowania pędu dla poruszającego się płynu? Dlaczego emisja elektronów przy braku światła wzrasta wraz z obniżaniem temperatury fotonowielacza? Co powoduje emisję krótkich impulsów światła z implodujących pęcherzyków w cieczy, gdy są one wzbudzone dźwiękiem? Dlaczego siły grawitacyjne są o wiele słabsze niż siły jądrowe, elektryczne i magnetyczne? Dlaczego nie potrafimy w pełni wyjaśnić fizyki prostego przyciągania i odpychania magnetycznego? Czy superpozycja, funkcja falowa i załamanie funkcji falowej są rzeczywistymi procesami fizycznymi, czy tylko konstrukcjami matematycznymi?

Sypać pytaniami z tego przepastnego wora nierozwiązanych problemów fizyki można długo, w pewnym sensie bez końca, gdyż doświadczenie uczy, że odpowiedzi na jedne pytania rodzą kolejne i kolejne. ■

Mirosław Usidus

Science Fiction. Oszustwa, uprzedzenia, zaniedbania i szum informacyjny w nauce

Stuart Ritchie
Wydawnictwo: Wydawnictwo Naukowe PWN, liczba stron: 380, cena: 69,00 zł

Książka, która ujawnia niepokojące i szokujące sygnały świadczące o kryzysie, jaki przechodzi współczesna nauka. Nauka to sposób, w jaki rozumiemy świat. Jednak niewłaściwa interpretacja i koloryzowane statystyki sprawiły, że ogromna liczba badań naukowych stała się bezużyteczna lub, co gorsza, wprowadzająca w błąd. Obecny system finansowania badań oraz publikacji ich wyników nie tylko nie chroni nas przed pomyłkami i słabościami naukowców, ale wręcz aktywnie zachęca badaczy do ingerowania w otrzymywane rezultaty, co może mieć katastrofalne konsekwencje. Stuart Ritchie, psycholog i popularyzator nauki, demaskuje przekłamania, manipulacje oraz szkodliwe nawyki naukowców, podając w wątpliwość prawdziwość i obiektywizm wielu badań naukowych. Na stronach swojej książki obnaża prawdziwe oblicze współczesnej nauki – okazuje się bowiem, że wiele powszechnie akceptowanych i bardzo wpływowych teorii bazuje na nierzetelnych badaniach.



Źródło i natura życia pozostają jednymi z największych niezgłębionych tajemnic nauki. Wciąż nie mamy pełnego wyjaśnienia mechanizmu sprawiającego, że coś jest „żywe”, a coś innego – materią nieożywioną.

Wyłoniło się i wyłania nadal, a my nie mamy pojęcia – jak i dlaczego?

TAJEMNICA ŻYCIA

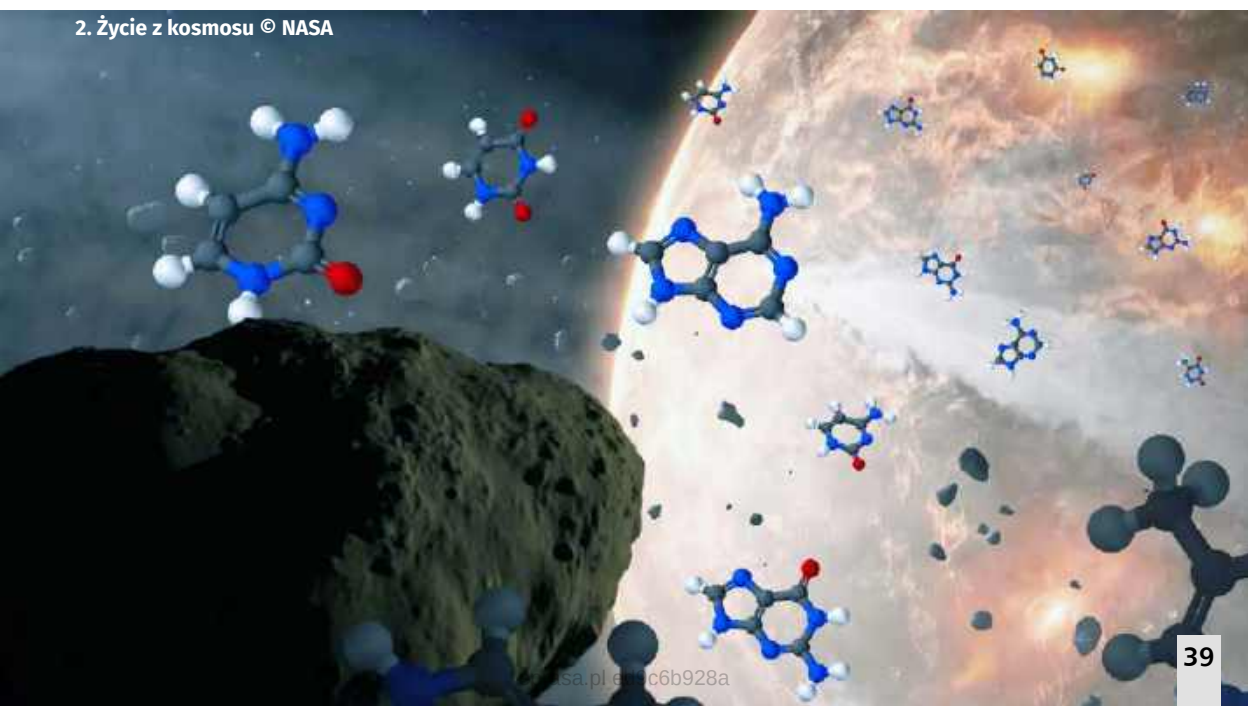
Z tego co wiemy, Ziemia jest jedyną planetą, na której istnieje życie (1). Jeśli życie jest tak wyjątkowe, to dlaczego? Jak powstało życie na Ziemi? Poszukiwanie początków życia na Ziemi to zadawanie innych ważnych pytań, np. o to, jak rzadkie jest to zjawisko (bo co do tej wyjątkowości, też mamy wątpliwości) i jakie jest prawdopodobieństwo jego powstania gdziekolwiek indziej poza Ziemią.

Dążenie do zrozumienia, dlaczego na Ziemi istnieje życie, musi rozpocząć się od wody, ponieważ woda jest jedyną rzeczą na Ziemi, której potrzebuje całe życie.



1. Ziemia i życie

2. Życie z kosmosu © NASA



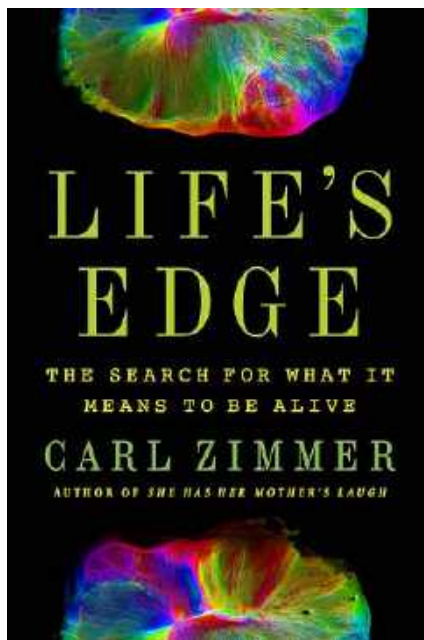
Różne formy życia mogą przetrwać na bardzo różnych źródłach pożywienia, bez tlenu i światła, ale nic nie żyje bez wody. Sama woda też jest pewną zagadką. Naukowcy nie do końca wiedzą, w jaki sposób woda pokryła dwie trzecie powierzchni naszego globu. Czy została dostarczona przez komety uderzające w naszą planetę? A może była od początku istnienia Ziemi w jakiś sposób głęboko ukryta?

A gdy pojawiła się woda, to w jaki sposób pojawiło się życie? Niektórzy dziś przypuszczają, że życie w ogóle nie zaczęło się na Ziemi (2). Wciąż jednak, jak się zdaje, większość naukowców sądzi, że życie powstało tu, na naszej planecie. Od dziesięcioleci próbują odtworzyć w laboratoriach warunki panujące na wczesnej Ziemi, z nadzieją że w końcu będą w stanie stworzyć coś podobnego do pierwszych prostych komórek, które miały samorzutnie powstać tutaj miliardy lat temu. W latach pięćdziesiątych XX wieku naukowcy Harold Urey i Stanley Miller wykazali, że możliwe jest zsyntetyzowanie aminokwasu glicyny, jednego z najbardziej podstawowych budulców życia, przez traktowanie gazów, które prawdopodobnie wypełniały atmosferę miliardy lat temu, ciepłem i symulowanymi wyładowaniami atmosferycznymi. Od tego czasu naukowcom udało się stworzyć lipidowe struktury, które wyglądały bardzo podobnie do błon komórkowych. Udało im się również stworzyć cząsteczki RNA, które przypominają uproszczone DNA. Ale uzyskanie wszystkich tych składników życia w laboratorium i złożenie ich w prostą komórkę – jak na razie okazuje się niewykonalne.

Jeśli znasz definicję, to umiesz odróżnić życie od jego braku

Czym w ogóle jest życie? „Nikt nie był w stanie zdefiniować życia, a niektórzy powiedzą ci, że nie jest to możliwe”, uważa Carl Zimmer, dziennikarz naukowy i autor książki „Life's Edge: The Search for What it Means to be Alive” (3). Oczywiście, istnieją setki definicji życia, ale jego zdaniem nie ma jednoznacznej odpowiedzi na pytanie początkowe. „Nasze mózgi w sensie biologicznym są dostrojone do rozpoznawania takich rzeczy jak ruch. Jesteśmy niejako zaprogramowani do rozpoznawania życia. Nie oznacza to jednak, że wiemy, czym ono jest”, kwituje autor.

Problem polega na tym, że dla każdej definicji życia naukowcy mogą znaleźć wyjątek. Weźmy na przykład definicję życia NASA: „Życie to samowystarczalny system chemiczny zdolny do darwinowskiej ewolucji”. Definicja ta wykluczałaby jednak wirusy, które nie są „samowystarczalne” i mogą przetrwać i replikować się tylko w innych organizmach.



3. Okładka książki „Life's Edge: The Search for What it Means to be Alive” Carla Zimmera

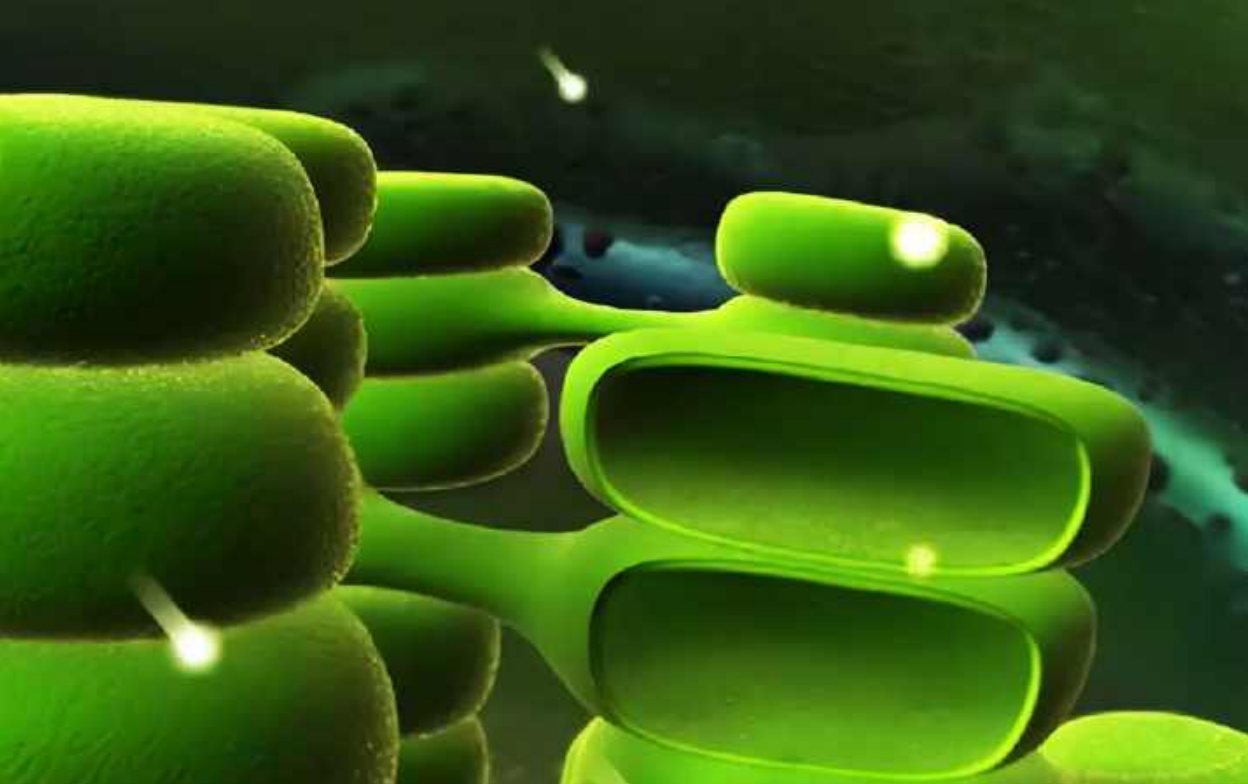
Definicja życia wydaje się bardzo przydatna, gdy szukamy jego oznak poza Ziemią, na innych planetach i różnego rodzaju ciałach niebieskich. Potrzebne jest narzędzie do niezawodnego rozpoznawania życia i odróżniania go od „nieżycia”. Zimmer powątpiewa, czy przy naszym obecnym stanie wiedzy jest to możliwe.

Zresztą, zaczynając pytania od samych siebie, nie rozumiemy jako ludzie w pełni relacji między fizycznymi procesami mózgu a świadomością, emocjami i subiektywnym doznaniem bycia „żywym”. Pomimo nagromadzonej wiedzy, nadal nie pojmujemy w całości złożonych szlaków biochemicznych i interakcji w żywych komórkach. Nie wyjaśniono w zadowalający sposób, skąd wzięła się ogromna ilość informacji genetycznej zakodowanej w DNA. Jest wiele dalszych pytań bez odpowiedzi.

Proces montażu a może mechanika kwantowa?

Gdyby ktoś pytał, to fizyka również nie potrafi wyjaśnić fenomenowi życia. Po prostu nie ma wzoru, który wyjaśniałby różnicę między żywą a martwą formą materii. Z tej perspektywy życie wydaje się po prostu w dość tajemniczy sposób „wyłaniać” z nieożywionych części, takich jak cząstki elementarne i większe skupiska masy.

Nie znaczy to, że fizyka nie próbuje. Jednym z jej najnowszych podejść do zagadnienia życia jest tzw. teoria



4. Kwantowe podłoże fotosyntezy

montażu, opisana niedawno w „Nature”. Koncentruje się na złożoności i informacjach (takich jak DNA). U podstaw teorii montażu leży idea, że obiekty mogą być definiowane nie jako niezmiennie byty, ale przez historię ich powstawania. Przenosi to uwagę na procesy, dzięki którym złożone konfiguracje są konstruowane z prostszych bloków konstrukcyjnych. Bloki konstrukcyjne można składać podobnie jak klocki Lego, tworząc cząsteczki składowe życia. Teoria proponuje „indeks montażu”, który określa ilościowo minimalne kroki lub najkrótszą ścieżkę wymaganą do zbudowania obiektu. Miara ta śledzi stopień „selekcji” niezbędnej do uzyskania zespołu obiektów – odnosząc się do pamięci, takiej jak DNA, wymaganej do stworzenia żywych istot. W końcu żywe istoty nie powstają spontanicznie. Wymagają DNA jako wzorca do tworzenia nowych wersji.

Zdaniem autorów, eksperymentalna weryfikacja w laboratorium wykazuje, że cząsteczki związane z życiem, takie jak hormony i metabolity (produkty reakcji metabolicznych), są rzeczywiście bardziej złożone i wymagają więcej informacji do złożenia niż cząsteczki, które nie są jednoznacznie związane z życiem, takie jak dwutlenek węgla. Teoria montażu ma też oferować wgląd w pochodzenie życia. Dzieje się tak, ponieważ mówi ona, że istnieje punkt, w którym cząsteczki stają się tak złożone, że zaczynają wykorzystywać informacje do tworzenia swoich kopii,

co wymaga pamięci i informacji. To rodzaj progu, w którym życie powstaje z nieżycia. Nie da się uzyskać żywych organizmów ani technologii, którą tworzą bez wysokiego poziomu pamięci i selekcji.

Naukowcy testujący teorię montażu planują dokładniej zbadać pochodzenie życia, tworząc w laboratorium rodzaj chemicznej zupy. W tej zupie z czasem mogłyby powstawać zupełnie nowe cząsteczki, albo poprzez dodawanie różnych reagentów, albo przez przypadek. Dostrajając szybkość reakcji i warunki, uczeni mogliby zbadać punkt przejścia od nieżycia do życia – i dowiedzieć się, czy jest on zgodny z przewidywaniami teorii montażu.

Od lat trzydziestych XX wieku naukowcy podejrzewali, że to zjawiska kwantowe stoją za fotosyntezą, procesem kluczowym dla organizmów żywych. W 2007 r. zespół naukowców z Lawrence Berkeley National Laboratory (Berkeley Lab) z Departamentu Energii USA przedstawił pierwsze dowody na to, że tak właśnie jest.

W procesach fotosyntezy rośliny absorbują fotony światła za pomocą komórek zwanych chromoforami (4). Te z kolei, wzbudzone, uwalniają quasi-cząstki zwane ekscytonami, które transportują energię do ośrodka reakcji. Tutaj może ona zostać przekształcona w energię chemiczną, którą roślina jest w stanie metabolizować. Cały ten proces zachodzi w jednej miliardowej części sekundy, z blisko stuprocentową sprawnością.

Prędkość jest konieczna, aby uniknąć strat energii i rozpraszania się jej w postaci ciepła. Greg Engel i jego koledzy z Berkeley Lab wykazali, że zamiast podróżować tą czy inną drogą, energetyczne wzbudzenie wykorzystuje superpozycję. Naukowcy wykorzystali do swoich eksperymentów siarkową bakterię *Chlorobium tepidum*. Jest to jeden z pierwszych organizmów, które zaczęły dokonywać fotosyntezy i występuje na Ziemi już od ponad miliarda lat. Obniżyli temperaturę bakterii do 77 kelwinów (-196°C). Następnie, posłali krótkie impulsy światła laserowego przez ciało bakterii. Śledzili je za pomocą dwuwymiarowej spektroskopii elektronicznej. Chcieli dokładnie poznać, jak przepływała tam energia. Zauważyli, że impuls przemieszcza się nie w linii prostej, lecz w postaci fal. Ze względu na zjawisko kwantowej koherencji ekscyton może, jako fala, „wykryć” wszystkie możliwe drogi, znaleźć najbardziej efektywną z nich i wybrać ją. Wyniki tego badania zostały opublikowane w czasopiśmie „Nature”. W kolejnych latach przeprowadzono szereg innych badań wykazujących to samo, mianowicie, fotosyntezę działającą na zasadzie koherencji kwantowej.

Poznamy tajemnicę życia, patrząc na jego lustrzane odbicie?

Cząsteczka składowa kodu genetycznego życia, DNA, spiralna helisa, jest prawoskrętna. Cząsteczki kodowane przez DNA, białka, są lewoskrętne. Nawet skromne cukry, takie jak glukoza, mają skręcony kształt. Dlaczego ta skrętność, zwana chiralnością (5), występuje w cząsteczkach (skupiskach atomów),



5. Obrazowe przedstawienie dwóch rodzajów chiralności

z których zbudowane jest całe życie na Ziemi? Czy ta skrętność nie jest przypadkiem tym, czego szukamy, zadając sobie pytanie, co odróżnia życie od materii nieożywionej? Dasze pytania to: jak i kiedy zdecydowano o lewo- lub prawoskrętnym chiralnym skręceniu składników życia?

Wiemy, że jest to dla organizmów żywych bardzo ważne. Jeśli podamy bakteriom lewoskrętne aminokwasy, to włączą je do swoich białek. Jeśli podamy im prawoskrętne aminokwasy, prawdopodobnie je zignorują, nawet jeśli są to te same cząsteczki, ale odwrócone w lustrzany sposób. Niektórzy badacze uważają, że po drugiej stronie lustra również mogą istnieć lustrzane formy życia, które być może kiedyś nauczymy się tworzyć. Powstają już lustrzane wersje białek i kwasów nukleinowych DNA. Cząsteczki te mogłyby być np. lekami, tym bardziej skutecznymi, że ich lustrzana struktura powinna umożliwiać im działanie poza zasięgiem wzroku zwykłych mechanizmów obronnych organizmu służących do rozbijania obcych cząsteczek.

Biochemik Ting Zhu z uniwersytetu w Hangzhou w Chinach dąży od lat do zbudowania lustrzanych wersji kluczowych molekularnych składników życia. W ciągu ostatnich kilku lat Zhu i jego koledzy stworzyli szereg lustrzanych wersji kluczowych biomolekuł w żywych komórkach: DNA, RNA oraz enzymów służących do ich replikacji i tłumaczenia ich sekwencji na białka. Nie ma jeszcze pełnego zestawu, ale sądzą, że są bardzo blisko. W teorii może być możliwe złożenie tych składników w syntetyczne jednostki podobne do komórek, które mogą się replikować i metabolizować: rodzaj prymitywnej formy życia, ale odwróconej w stosunku do każdego znanego organizmu, a zatem pierwszej prawdziwie nienaturalnej formy życia. Nikt nie wie, jak takie lustrzane komórki dogadywałyby się ze zwykłymi. Może ignorowałyby się nawzajem lub rywalizowały? Jest też pytanie, czy nasza planeta przypadkiem nie przeprowadziła już raz tego eksperymentu miliardy lat temu, kiedy powstawało życie?

Życie w lustrzanym odbiciu (6) zostało po raz pierwszy wyobrażone przez Louisa Pasteura. W 1848 roku wydedukował on, że produkt uboczny produkcji wina zwany kwasem winowym może krystalizować w dwóch lustrzanych formach krystalicznych, ponieważ ich cząsteczki składowe mają przeciwną chiralność. Pasteur uważał, że chiralność molekularna jest fundamentalną cechą istot żywych, tym właśnie, co odróżnia je od świata nieożywionego. Zastanawiał się, czy chiralność jest indukowana przez siły takie jak magnetyzm, elektryczność lub światło. Przeprowadził serię



6. Żywe organizmy po dwóch stronach lustra

eksperymentów z krystalizacją związków w polach magnetycznych, uprawą roślin w świetle słonecznym z odwróconą polaryzacją itd. Eksperymenty te wiele nie dały. Ostatnio naukowcy z Uniwersytetu Harvarda wysunęli przypuszczenie, że magnetyzm w minerałach mógł rzeczywiście odegrać rolę w nadaniu najwcześniejszemu życiu jego określonej chiralności, ale to wymaga weryfikacji. Na początku lat dziewięćdziesiątych XX w. Stephen Kent z uniwersytetu w Chicago i jego współpracownicy chemicznie zsyntetyzowali z cząsteczek aminokwasów w lustrzanym odbiciu (z odwrotną chiralnością) enzym proteazę HIV-1, który wirus HIV wykorzystuje do degradacji innych białek. Pomogło to w opracowaniu leków zwalczających AIDS poprzez blokowanie działania enzymu.

Ting Zhu chce teraz stworzyć całą nową formę życia poprzez odwrócenie skrętności cząsteczek składowych, np. całą bakterię w lustrzanym odbiciu, z odwróconymi wszystkimi kwasami nukleinowymi, białkami i cukrami. Jeśli taki organizm zostanie kiedykolwiek stworzony, to zdaniem badaczy, nie powinien wyglądać i działać inaczej niż normalna bakteria. Dopiero po powiększeniu poszczególnych cząsteczek różnica stałaby się widoczna. Na razie jednak wszystko pozostaje w sferze fantazji. Życie w lustrzanym odbiciu może pomóc w zgłębieniu „jednej z największych tajemnic związanych z powstaniem życia na Ziemi” – mówi Zhu – jego „homochiralności” (wszystkie wykrywalne cząsteczki mają taką samą chiralność), a inaczej rzecz ujmując – dlaczego życie wykorzystuje tylko kwasy D-nukleinowe i białka L. Niektórzy uważają, że wybór był całkowicie przypadkowy. Być może wśród prebiotycznych bloków budulcowych istniały różne typy cząsteczek składowych z różną skrętnością, a tylko przypadkowe wahania stężeń jednego typu,

wzmocnione przez procesy sprzężenia zwrotnego, doprowadziły do przewagi określonego typu. Zhu wskazuje inną możliwość. „Czy życie na Ziemi jest naprawdę homochiralne, czy po prostu nie mamy odpowiednich narzędzi do poszukiwania drugiej wersji chiralnej?” pyta. Co, jeśli „lustrzana” wersja życia istnieje w niektórych zakątkach środowiska naturalnego, których nie jesteśmy w stanie eksplorować? Może istnieje „ciemna biosfera”, cały ekosystem nieznanych lustrzanych form życia.

Zagadki od góry do dołu i w drugą stronę

Nauka z powodzeniem zastosowała już technikę redukcji, aby określić elementy składowe życia. Redukcja to poszukiwanie wyjaśnień lub skutecznych opisów działania systemu przez skupienie się na jego elementach składowych w mniejszej skali. Jeśli interesuje nas ludzkie ciało, to redukcję prowadzą od narządów do komórek, przez DNA do genów, biomolekuł i tak dalej.

Podejście to odniosło sukces. Jednak nie wystarczy do zrozumienia fenomenu życia. Według innego nurtu badawczego, kluczem jest zrozumienie życia jako złożonego systemu adaptacyjnego, czyli takiego, w którym organizacja i przyczyna występują na wielu poziomach. Znaczenie mają nie tylko cząsteczkowe i atomowe bloki konstrukcyjne, lecz także relacje elementów składowych w górę i w dół skali, systemy połączonych sieci, od genów do ekosystemu i w drugą stronę. To „głębsze rozumienie” wydaje się wciąż czymś relatywnie mglistym. Na razie mamy problemy z podstawami i fundamentami. Ich skomplikowane sieci powiązań to kolejny etap zagłębiania się w tajemnicę życia. ■

Mirosław Usidus

Czekając na aerostrady przyszłości

Rozmowa z Tomaszem Patanem, założycielem, szefem ds. technologii, i wynalazcą Jetson ONE.

„Młody Technik”: Kiedy Jetson ONE trafi do zwykłej sprzedaży? Czytaliśmy w mediach pochodzące z jesieni 2023 r. zapowiedzi, że stanie się to w bieżącym roku. Czy rzeczywiście jest na to szansa?

Tomasz Patan: Pierwsze egzemplarze Jetson ONE trafią w ręce naszych klientów w trzecim kwartale tego roku. Postawiliśmy sobie to za cel, a cały zespół pracuje ciężko, żeby zrealizować marzenia naszych przyszłych pilotów.

MT: Sprzedaż to jedno, a zezwolenie na loty w przestrzeniach nad terenami zurbanizowanymi, gdzie mieszkają ludzie i jest gęsty transport naziemny, to drugie. Amazon, który zapowiadał wprowadzenie dronów dostawczych na takich terenach do 2020 r., obecnie nie znajduje się dużo bliżej realizacji tego projektu niż dekadę temu, gdy o tym zaczęto mówić. Jak Pan widzi perspektywy uzyskiwania zezwoleń na latanie Jetsona ONE w tej przestrzeni w USA i w Europie?

TP: Od samego początku projektowałem nasz latający pojazd w oparciu o istniejące w USA od wielu lat regulacje pojazdów typu „Ultralight”. Jetson ONE wpisuje się perfekcyjnie w ramy tych regulacji, które dają wspaniałą możliwość lotów o charakterze rekreacyjnym bez posiadania licencji pilota. Działą to na podobnej zasadzie jak na przykład paralotniarstwo, gdzie można latać swobodnie poza gęsto zabudowanym terenem oraz w określonej odległości od lotnisk (przestrzeń klasy G).

Współpracujemy i prowadzimy rozmowy z cywilnymi urzędami lotnictwa kilku państw Unii Europejskiej, czego rezultatem będą podobne możliwości lotu na Starym Kontynencie. We Włoszech uzyskaliśmy certyfikację „Ultralight” już w zeszłym roku, a niedługo Jetson ONE otrzyma tu własną, oficjalną certyfikację pojazdu „eVTOL”, która da możliwość latania każdemu po zaliczeniu krótkiego szkolenia



1. Tomasz Patan
(© Tomasz Patan)

Tomasz Patan

Ma 39 lat, urodził się w Gdańsku, gdzie po 3 latach zrezygnował ze studiowania architektury na Politechnice Gdańskiej, aby poświęcić się pasji projektowania maszyn latających.

Fascynuje go najnowsza technologia, dzięki której może realizować futurystyczne projekty mobilności.

Większość życia spędził, jak to określa, z „głową w chmurach”, poświęcając się rozwijaniu kreatywnego myślenia, umiejętności projektowania, prototypowania i budowy pojazdów latających. Stworzył od podstaw firmę budującą innowacyjne drony do filmowania z powietrza. Obecnie buduje rewolucyjne rozwiązania transportu w Jetson oraz najnowszym projekcie Volonaut.

Następny zaplanowany rozdział i wyzwanie to rakiety i eksploracja kosmosu.



2. Dwóch pilotów i ich powietrzne bolidy na moment przed wspólnym lotem w formacji (© Tomasz Patan)

w naszej szkole pilotów eVTOL. Takich możliwości nie brakuje również w Polsce, gdzie także, dzięki kategorii UL115, możliwe będzie latanie Jetsonem ONE po odbyciu krótkiego treningu w naszej szkole.

MT: W źródłach czytamy, że człowiek latający Waszym eVTOL-em nie będzie musiał mieć licencji pilota. W kontekście tego, o czym była mowa w poprzednim pytaniu, czyli lotów na terenie zamieszkanym, brzmi to nieco kontrowersyjnie. Dlaczego Pańskim zdaniem mamy się nie bać, że ta maszyna spadnie na kogoś np. idącego po ulicy albo nie wleci komuś przez szybę do domu? Przecież nawet operowanie

dźwostkiem też wymaga umiejętności, a gdy w grę wchodzi bezpieczeństwo i życie ludzi oraz prędkości ok. 100 km/h, umiejętności te stają się bardzo ważne.

TP: Jak wspominałem wyżej, loty nad gęsto zaludnionym terenem nie są legalne. Komputer pokładowy Jetsona ONE utrzymuje go stabilnie w powietrzu nawet bez udziału pilota. Ruchy dźwostka nie są bezpośrednio przekładane na ruch pojazdu w powietrzu, są one niejako instrukcją, w którą stronę, z jakim przyspieszeniem czy jaką prędkością pilot chciałby się poruszać. Komputer pokładowy w oparciu o dane z licznych sensorów dokonuje setek kalkulacji na sekundę i dba



3. Prezentacja najnowszej odstony prototypu przedprodukcyjnego Jetson ONE na targach elektrycznej mobilności w Rzymie jesienią 2023 r. (© Tomasz Patan)



4. Tak bajkowo prezentuje się Jetson ONE na tle zachodzącego słońca na słynnej plaży Santa Monica przy Los Angeles (© Tomasz Patan)

o bezpieczeństwo oraz płynność lotu, nie dopuszczając do niebezpiecznych sytuacji. Dołożyliśmy naprawdę dużo starań, aby latanie stało się de facto prostsze niż prowadzenie samochodu. Taka zresztą jest jedna z misji Jetsona – to „demokratyzacja” latania. O bezpieczeństwo lotu dba również pełna redundancja systemów napędu, komputera pokładowego, awioniki. Jest wbudowany spadochron balistyczny, radar utrzymujący stały pułap lotu i pełen asortyment czujników, które utrzymują optymalne parametry lotu.

MT: A co, jeśli wydarzy się wypadek? Będą ofiary w ludziach? Przykład Tesli czy śmiertelnego wypadku z udziałem testowego samochodu autonomicznego Ubera pokazuje, że choć na drogach jest mnóstwo wypadków, w których giną co dzień ludzie, uwagę opinii publicznej i mediów skupiają konstrukcje nowe, które obiecują zmianę modelu transportu. Chociaż osobiste pojazdy latające, które proponujecie, mogą w sensie ogólnym być transportem lepszym, szybszym i nawet bezpieczniejszym, to pojedynczy wypadek może szybko zakończyć tę przygodę.

TP: Dokładamy ogromnych starań, aby zminimalizować taką możliwość i mamy pełną świadomość konsekwencji takich potencjalnych wydarzeń. Z drugiej strony rewolucja w transporcie i mobilności, którą

5. Trudno nie wpaść w zachwyt sam na sam z kapsułą Dragon SpaceX świeżo po powrocie z orbity okołoziemskiej. Takie momenty ogromnie motywują do pracy nad ambitnymi projektami (© Tomasz Patan)



wdrażamy, jest zbyt ważna i wręcz nieunikniona, aby popadać w negatywne nastawienie czy poddać się po pojedynczych przypadkach takich niefortunnych zdarzeń. Analogicznie do wdrożenia i ewolucji samochodu czy samolotu ponad sto lat temu.

MT: Projekt powietrznej aeromobilności w Dubaju z wykorzystaniem eVTOL-i EHang zakłada autonomię pojazdów, czyli osoby na pokładzie będą jedynie pasażerami z niewielkim bagażem, podróżującymi po określonych, kontrolowanych trasach. Czy nie sądzi Pan, że władze, jeśli w ogóle zgodzą się na tego typu aeromobilność miejską, to jedynie w takim modelu, w którym bezpieczeństwo nie będzie zależało od ludzi, których umiejętności prowadzenia nie są w żaden sposób weryfikowane? Czy wyobraża Pan sobie Jetson ONE jako maszynę autonomiczną, sterowaną bez udziału człowieka?

TP: Wyobrażam sobie Jetson Three albo Four jako częściowo lub w pełni autonomiczny „latający samochód”, dzięki któremu niemal każdy mieszkaniec aglomeracji podróż z jednego krańca miasta na drugi odbędzie w kilka minut, nie w godzinę czy dwie. W rzeczywistości opracowałem już kilka rozwiązań i koncepcji, dzięki którym stanie się to faktycznie możliwe w ciągu kilku lat. Naszą misją jest przeniesienie transportu z ziemi w powietrze i urzeczywistnienie tego, co dobrze znamy ze świata science fiction.

MT: A może Jetson ONE szuka niszy wyczynowo-sportowej? Wtedy problemy, o których mowa w poprzednich wątkach, nie mają takiego znaczenia, a rośnie znaczenie osiągnięć, szybkości, pułapu, manewrowości. Lotnictwo samolotowe po ok. dwóch dekadach od pierwszego lotu braci Wright wkroczyło w erę ogromnej popularności wyścigów i innych zawodów lotniczych. Czy sądzi Pan, że branża eVTOL-i, latających samochodów czy motocykli, jakkolwiek to definiować, również czeka złota era sportów wyczynowych, tak jak to miało miejsce z samolotami?

TP: Jedną z pierwszych myśli po uniesieniu się w powietrze było – „to kiedy przyjdzie mi się pościgać z drugim pilotem Jetsona?”. Wykonaliśmy już wspólne przeloty z dwoma pojazdami na raz w powietrzu i jest to niesamowita frajda. Dodatkowo niedługo ustawimy tor ułożony z nadmuchiwanymi kilkumetrowych pylonów, aby zacząć promować wyścigi jetsonów. Zapowiada się atrakcyjne widowisko i zacięta rywalizacja, w której biorące udział zespoły z całego świata będą miały możliwość udoskonalania technologii. Sport i rywalizacja są najlepszym inkubatorem innowacji i sprytnych rozwiązań „out of the box”.

MT: Zasięg Jetson ONE wygląda na wystarczający do krótkich miejskich lotów, jednak korzystanie



6. Silniki turboodrzutowe, niespełnione marzenia z dzieciństwa nabierają formy w postaci najnowszego prototypu pojazdu napędzanego w ten właśnie sposób (© Tomasz Patan)

z napędu elektrycznego wiąże się ze wszystkimi problemami akumulatorowymi z długim czasem ładowania na czele. Jakże Pan i Pańscy współpracownicy macie pomysły na doskonalenie i przewyżnianie ograniczeń akumulatorów w przyszłości?

TP: Od samego początku stawiałem na pięknie proste i nieskomplikowane rozwiązania, to dotyczy napędu, co do którego chcemy, aby zawierał jak najmniej części i wymagał minimalnej ilości serwisowania. Obecnie tylko napęd elektryczny bazujący na bateriach spełnia te wymagania. Na szczęście teraz na całym świecie nad doskonalszymi typami baterii pracują rzesze mądrych ludzi i będzie miało to przełożenie na zasięg i czas lotu kolejnych wersji Jetsona. Oczywiście, mając na uwadze obecne ograniczenia technologiczne, nie rozłożyliśmy rąk w bezsilności i uzbroiliśmy nasz pojazd w wymienne baterie, wystarczy minuta, aby znów wzbić się w powietrze. Obecnie czas ładowania to zaledwie godzina, więc tu też źle nie jest.

MT: Rozszerzając nieco problem z poprzedniego pytania – co Pan sądzi o perspektywach elektryfikacji lotnictwa? Jak wiadomo, odtworzenie mocy, wydajności i przede wszystkim gęstości energii paliwa lotniczego, za pomocą akumulatorów, to wielkie wyzwanie, któremu na razie nie udało się sprostać, mimo wielu wysiłków, wykorzystania doświadczeń samochodowej Formuły E, zaawansowanych prac Rolls-Royce'a itd. Czy widzi Pan szansę na przełom w elektryfikacji lotnictwa na dużą skalę w perspektywie pięciu-dziesięciu lat?



7. Latanie to jedno z największych marzeń ludzkości, banalnie prosty w pilotażu Jetson ONE pozwala spełnić je niemal każdemu, dając przy tym ogromną radość i dawkę adrenaliny (© Tomasz Patan)

TP: Zdecydowanie widzę taką szansę, elektryfikacja mobilności dopiero się rozpędza, a efekty licznych badań i udoskonaleń już niedługo będą widoczne. Kluczem będzie zwiększenie wydajności i gęstości ogniw baterii, ale nie takie rzeczy ludzkość już przerabiała. Jestem z natury wielkim optymistą i wierzę, że taki przełom nastąpi w ciągu dwóch najbliższych dekad.

MT: Rozmowie naszej partneruje Polska Fundacja Fantastyki Naukowej. Przekazała własne pytania. Jedno z nich nawiązuje do tego, że w fantastyce naukowej latające samochody (lub podobnego typu pojazdy służące do przewozu ludzi) często są charakterystycznym elementem futurystycznych wizji świata. Wizja ich powszechnego wykorzystania utrwaliła się chociażby dzięki takim filmom jak „Powrót do przyszłości”, „Piąty element”, czy też „Blade Runner”. Dlaczego Pańskim zdaniem ludzie marzą o tego rodzaju pojazdach?

TP: Ludzie od zawsze marzyli o lataniu, w szczególności w formie łatwo dostępnej, a najlepiej jako urządzenie czy pojazd, który można mieć na własność. Wizerunek „latającego samochodu” niesamowicie stymuluje wyobraźnię, bo taki pojazd pozwala dotrzeć w niemalże każde miejsce bardzo szybko, omijając korki i wyzwalając emocje towarzyszące oderwaniu się od ziemi i problemów codzienności. Jest też oczywiście kwestia widowiskowego sposobu, w jaki przedstawiane są takie pojazdy, gdzie zazwyczaj bohaterowie używają go do pościgu czy w kaskaderskich manewrach, każdy ma w sobie trochę z dziecka, a taki obraz dobrze z tymi emocjami rezonuje.

MT: Kolejne pytanie PFFN zawiera spostrzeżenie, że kultura popularna przyzwyczaiła nas do konkretnych wyobrażeń powietrznych samochodów, które jednak mocno różnią się od tego, co obecnie obserwujemy, jeśli chodzi o ich projektowanie, stylistykę oraz napęd. Sądzi Pan, że za jakiś czas takie pojazdy zaczną przypominać swoje odpowiedniki z filmów science fiction, czy raczej nie powinniśmy na to liczyć?

TP: Myślę, że to w dużej mierze zależy od odwagi twórców i strategii, jaką przyjmują. Osobiście nie wyobrażam sobie, aby „latający samochód” okazał się ostatecznie tylko zachowawczą modyfikacją klasycznego samolotu i zapewniam, że tą ścieżką Jetson nie pójdzie. Chcemy budować fantastyczne pojazdy razem z filmów science fiction, będzie futurystycznie, kompaktowo z cechami supersamochodów – tak, aby każdy chciał się stać ich posiadaczem. Oczywiście realizacja bardzo ambitnych planów oznacza pewne kompromisy, jest większe ryzyko, więcej włożonej pracy oraz poświęceń, ale dzięki temu projekt jest większym wyzwaniem wartym wstawiania codziennie o 5 rano i zasypiania z głową pełną rozwiązań.

MT: Fundacja pyta też o uczenie maszynowe, w której to dziedzinie zachodzi ostatnio wielki postęp. Dostrzega Pan jakieś obszary, na których sztuczna inteligencja mogłaby wesprzeć wykorzystanie latających samochodów przez przeciętnego użytkownika, np. w zakresie pilotażu czy autonomii lotu?

TP: AI i autonomia staną się nieodłącznym elementem tej branży. Aby ziszczała się wizja prawdziwych



*Rozmowa z Tomaszem Patanem
z Jetson ONE, przeprowadzona została
w partnerstwie z Polską Fundacją
Fantastyki Naukowej, od której pochodzi
część pytań zadanych rozmówcy „Młodego
Technika”.*

„latających samochodów”, będzie trzeba zaprząć automatyzację i kontrolę nad parametrami lotu oddać w ręce autopilota. My będziemy mogli cieszyć się relaksującą podróżą w powietrzu, docierając do celu bez stresu i dużo szybciej, niż jest to możliwe dziś. AI wyznacza optymalną trasę, modyfikując ją na bieżąco w trakcie lotu z uwzględnieniem natężenia ruchu czy warunków atmosferycznych, włączy bezpiecznie nasz pojazd do ruchu na wyznaczonej w powietrzu, trójwymiarowej „aerostradzie”.

MT: Kto oglądał futurologiczne ryciny z przełomu XIX i XX wieku, ten wie, że nasi prapradziadkowie wyobrażali sobie nasze czasy z niebem wypełnionym dużymi i małymi maszynami latającymi. Do pewnego stopnia tak się stało, ale o demokratyzacji latania, a precyzyjniej mówiąc, sterowania maszynami przez każdego, kto chce, raczej nie ma mowy. Czy

to się kiedyś stanie? Czego potrzeba, by te stare wizje się ziściły?

TP: Takie same obawy budziły pojawiające się na drogach ponad sto lat temu samochody, dziś czynność prowadzenia auto wydaje się banalnie prosta i często zapominamy, jak bardzo skomplikowana to interakcja. Oczywiście, jak już wspomniałem, nie będzie innej drogi do kolejnej rewolucji w przemieszczaniu się niż automatyzacja i to sporo uproszczy wdrożenie jetsonów w niedalekiej przyszłości. Zawsze będzie też możliwość pilotowania manualnie takich pojazdów w kontrolowanych warunkach, na przykład rekreacyjnie czy po prostu dla weekendowej frajdy po odbyciu odpowiedniego szkolenia.

MT: Dziękuję za rozmowę.

Rozmawiał **Mirosław Usidus**

Latem, o tej samej porze

Annabel Monaghan

Wydawnictwo: Insignis, liczba stron: 307, cena: 44,99 zł

Idealna lektura na nostalgię za latem, opowiadająca o zaręczonej kobiecie, która po czternastu latach od ostatniego spotkania staje twarzą w twarz ze swoją pierwszą miłością. Spędzała z nim każde wakacje od piątego do siedemnastego roku życia – do momentu, gdy złamał jej serce, podając w wątpliwość wszystko, co dotychczas wyobrażała sobie na temat ich miłosnej historii i samej siebie. Zasady plażowania: spaceruj długo po piasku; do każdego koktajlu wkładaj parasolkę; NIE spotykaj swojej pierwszej miłości. Życie Sam zmierzają we właściwym kierunku. Ma idealnego narzeczonego – lekarza Jacka, świetną pracę na Manhattanie (chyba że ją zwolni) i właśnie zamierza odwiedzić miejsce, gdzie odbędzie się jej wesele, w pobliżu rodzinnego domu na plaży na Long Island. Wszystko powinno przebiegać zgodnie z planem, lecz w chwili przyjazdu Sam zauważa, że coś nie gra. Wyatt jest tutaj. Jej Wyatt. Ale nie ma powodu, by trzydziestoletnia zaręczona kobieta wpadała w panikę przy faciecie, który złamał jej serce, gdy miała siedemnaście lat, prawda? A jednak, będąc z powrotem na tej plaży i słysząc znajome dźwięki gitary Wyatta rozchodzące się w nocnym powietrzu z sąsiedniego domu – jakby zatrzymał się czas – czuje, że wspomnienia powracają: dotyk jego skóry na jej skórze, wspólne noce w domku na drzewie i prawda o ich rozstaniu. Sam przypomina sobie, kim kiedyś była, a gdy Wyatt ponownie wkracza w jej życie, ich więź jest tak samo niezaprzeczalna, jak zawsze. Sam będzie musiała dokonać wyboru.





Górnictwo pozaziemskie

Kopiąc wśród gwiazd

Astrogórnictwo niejednokrotnie rozpałało wyobraźnię ludzi. Kiedyś było domeną fantastyki i futurologii, dziś co roku odbywają się konferencje w Luksemburgu, Kolorado czy w Krakowie o tematyce surowców kosmicznych. Dyskusja o legalności, tworzeniu polityki, prawa oraz standardów toczy się na płaszczyźnie Organizacji Narodów Zjednoczonych oraz innych organizacji. Dlatego zacznijmy od tego, czym astrogórnictwo jest, jakie są jego korzenie oraz jak to się ma dzisiaj do eksploracji kosmosu.

O górnictwie kosmicznym mówimy najczęściej przy omawianiu metod i procesów wydobywania, oczyszczania, magazynowania i transportu surowców kosmicznych. Definicja surowca kosmicznego jest zależna natomiast od kontekstu – pewne definicje znajdziemy w ustawodawstwie, zaś nierzadko odmienne w opracowaniach stricte technicznych lub koncepcyjnych. Najprościej jest uznać, że surowcem kosmicznym jest każda naturalnie występująca poza Ziemią materia, którą jesteśmy w stanie wydobyć lub pozyskać oraz wykorzystać. Jest to definicja obszerna, chociaż nieprecyzyjna, co pokażemy w dalszych rozważaniach.

Najczęściej termin górnictwo kosmiczne albo precyzyjnie pozaziemskie (uszanujmy Kopernika, Ziemia też znajduje się w kosmosie), asteroidalne lub księżycowe dotyczy wielkoskalowego, przemysłowego wydobywania takich potencjalnych surowców jak minerały ziem rzadkich, platynowce, woda, izotopy o znaczeniu energetycznym (hel-3) czy substancje niezbędne systemom podtrzymywania życia lub do budowy pojazdów i części zamiennych dla urządzeń (metale takie jak żelazo, nikiel, cyna czy tytan). Małoskalowe, doraźne i wykonywane na potrzeby pozyskania surowca będzie stanowiło element ISRU – In-situ Space Resource Utilization. W odróżnieniu od koncepcji górnictwa kosmicznego ISRU przyjmuje szersze spektrum surowców i zasobów takich jak temperatury otoczenia lub promieniowanie gwiazdowe czy mikrogravitacja. ISRU towarzyszy często element produkcji – wytwarzania elementów lub struktur z surowców pozyskanych z otoczenia (naturalnych lub ze złomu kosmicznego), takich jak części zamienne, paliwo lub regolito-pochodna osłona stacji załogowych i habitatów mieszkalnych.

Dla pewnej przejrzystości należy wspomnieć, że produkcja kosmiczna jest elementem działalności związanej z surowcami kosmicznymi i ważnym dla górnictwa kosmicznego, gdyż utylizuje surowce na potrzeby masowej lub wyspecjalizowanej produkcji. Jednak różnica między produkcją zasilaną przez surowce pozyskane z górnictwa kosmicznego a tą zasilaną przez małoskalowe wydobywanie to właśnie kwestia skali oraz celu produkcji. ISRU przypomina bardziej rzemiosło oparte na dostępnych zasobach, np. regolicie i surowcach znajdujących się w otoczeniu, niż przemysłowe pozyskanie tego samego surowca. Jednak ta różnica w pewnym momencie może stać się niesamowicie płynna, gdyż wytwory ISRU mogą stanowić podstawę drobnych usług lub oferowanych produktów (np. paliwo lub gotowe pyłochrony dla lądowisk), natomiast część surowców pozyskiwanych przemysłowo będzie stanowić podstawę zaopatrzenia i konserwacji maszyn oraz infrastruktury systemu wydobywania.



Różnica między tymi podejściami do działalności związanej z surowcami kosmicznymi ma także kontekst polityczny. Wydobywanie na własne potrzeby, szczególnie drobne, lub wykorzystanie na potrzeby rozbudowania własnych placówek powierzchniowych i orbitalnych niesie ze sobą mniejsze kontrowersje w środowisku międzynarodowym niż typowe górnictwo nastawione na zysk i skierowane ku odbiorcom znajdującym się na naszej planecie macierzystej.

Jest to powód sporu w środowisku prawa międzynarodowego, gdyż odnosi się do koncepcji bogacenia się części państw w wyniku eksploatacji bogactw znajdujących się poza obszarami suwerennych państw (dno morskie, przestrzeń kosmiczna). Dodatkowo sam akt prowadzenia prac wydobywczych na ciele niebieskim budzi obawy o możliwość naruszenia zasady niezawłaszczalności,

będącej jedną z podstaw międzynarodowego prawa kosmicznego. Przyczyn umieszczenia takiej zasady w prawie kosmicznym nie należy upatrywać w popularności „Star Treka” w latach 60., a w zimnej wojnie oraz zakończeniu drugiej wojny światowej.

Przepisy prawa kosmicznego były tworzone z myślą o aktualnych problemach i zagrożeniach, wśród których była groźba uczynienia z przestrzeni kosmicznej teatru wojennego, poligonu doświadczalnego dla broni atomowej lub zaognianie sytuacji na Ziemi przez spory terytorialne na Księżycu. Dlatego też problemem obecnego stanu prawa kosmicznego jest dostosowanie regulacji, mających swoje oparcie w traktatach z czasów zimnej wojny (które wciąż obowiązują) do nowych przedsięwzięć i sytuacji. Nie ulega wątpliwości, że gdyby te traktaty spisywano 40...60 lat wcześniej, zawierałyby przepisy dotyczące wytyczania obszarów suwerennych, czasu prowadzenia prac górniczych i rozwiązywania konfliktów, a z kolei pisane dzisiaj dawałyby większą swobodę i decyzyjność w materii kosmicznej podmiotom pozarządowym (związkom zawodowym, organizacjom NGO czy korporacjom).

Wracając jednak do naszej rzeczywistości, spór o legalność astrogórnictwa wynika z odmiennego pojmowania operacji astrogórnich oraz partykularnych interesów krajów działających (oraz niedziałających) w przestrzeni kosmicznej.

Sama idea astrogórnictwa może wydawać się dość nowa, chociaż o astrogórnikach przeczytamy tak

u Garreta S. Servissa, jak i Konstantego Ciołkowskiego. Nawet starsi czytelnicy mogą pamiętać prace Andrzeja Marksa czy Janusza Thora, gdzie występuje zagadnienie kopalnictwa księżycowego. Co jednak jest ważne w dzisiejszych czasach, to zmiana podejścia do surowców.

Początkowo surowcem wartym sprowadzania nie były platynowce czy kobalt, lecz złoto, rubiny i szmaragdy. Jest to dość ciekawe, zważywszy że nawet w fantastyce przełomu XIX i XX wieku pojawiały się pomysły na metalurgię kosmiczną (z wykorzystaniem nieważkości lub energii słońca). Niemniej wziąć pod uwagę należy, że w fantastyce astrogórnicy często byli tylko elementem tła lub zwyczajnie traktowanym po macoszemu pretekstem dla poprowadzenia fabuły. Zdarzały się też utwory traktujące poważniej zagadnienie kopalni kosmicznych, nierzadko w kontekście relacji pracy człowieka i robota (na przykład u Asimova), jednak surowce pozyskiwane zaczynały zmieniać swój charakter. Już w latach 30. XX wieku można zauważyć, że wydobywane są metale konstrukcyjne na potrzeby fabryk w kosmosie lub ziemskich. Nierzadko było wydobywane jakieś „unobtanium”, fikcyjna cudowna substancja lub pierwiastek niezbędny dla jakiejś cywilizacji. Były to na przykład antymateria, monobieguny magnetyczne, pierwiastek zero, kryształy dilithium, beskar czy „egzotyczna materia”. Przyprawa z Diuny mogłaby zostać potraktowana jako surowiec kosmiczny pochodzenia biologicznego, jeśli poluzujemy stosowane przez nas definicje. Niemniej to do autorów książek, filmów i gier zależy wyjaśnienie metod pozyskania lub zastosowania surowców, jakie opisują. Zupełnie inaczej traktowane są surowce w książkach Bena Bovy z serii *Układ Słoneczny* (ang. *Grand Tour*), a odmiennie w utworach Bohdana Peteckiego lub Dennisa Taylora. Tam gdzie widziane są jako podstawa ekonomii, spojrzenie będzie odmienne od wizji, gdzie służą one bezpośrednio na prawom, budowie, uzupełnieniu zapasów czy innym potrzebom np. samoreplikujących się sond.

Techniki wydobywania surowców kosmicznych będą różnić się znacznie od tych, jakie stosuje się powszechnie wobec analogicznych surowców na Ziemi. Powodem tego są warunki pozaziemskie i stan, w jakim występuje taki surowiec (jak np. woda). Pierwsze misje demonstracji technologii wydobywania (poza technologiami rozpoznania złoża, magazynowania, transportu czy właściwego przemieszczania maszyny zaopatrzonej w sprzęt o przeznaczeniu wydobywczym) będą musiały być dostosowane do konkretnego ciała niebieskiego oraz jego obszaru, na którym mają funkcjonować. Nie jest to odmienne od wysyłania sond i lądowników. Różnica polega na stałej przemysłowej pracy i współdziałaniu





panujące na pozaziemskim ciele niebieskim oraz brak dostępności łatwej do wykorzystania cieczy w sporych ilościach. Jest to problem związany między innymi z ograniczeniem tarcia, oczyszczaniem terenu, przemieszczaniem zwiercin oraz chłodzenia urządzeń. Dlatego możliwe jest, że przy górnictwie, które będzie wydobywać surowce z różnych skał, potrzebne będą nowe rozwiązania techniczne i podejście do możliwości użycia np. materiałów wybuchowych w zmniejszonej (w stosunku do ziemskiej) lub mikro-grawitacji. Dlatego zanim powstaną maszyny zdolne wydobywać minerały ziem rzadkich, tytan lub osm, będzie konieczne przetestowanie podstawowych systemów przemysłowej pracy na Księżycu, nie mówiąc o systemach telekomunikacji LunaNet czy systemach pozycjonowania satelitarnego.

Czy w astrogórnictwie jest przyszłość? Owszem, głównie dla infrastruktury pozaziemskiej, chociaż nie jest wykluczone, że surowce kosmiczne będą stanowić część produktów sprowadzanych z kosmosu na rynki ziemskie. Ma jednak długą drogę technologiczną przed sobą, a ponadto nie sprzyja niezakłóconemu rozwojowi astrogórnictwa klimat polityki kosmicznej (hamowanie inicjatyw przeciwnika jest ważniejsze od obiektywnego rozwoju). Większość ludzi jednak nie zauważy infrastruktury górniczej i technologii przetwórstwa, jaka będzie w nim stosowana. Zobaczmy produkty, które im spowszednieją, tak jak nawigacja i mapy w telefonie. ■

Kamil Muzyka

wielu systemów (mobilnych i stałych). Ponadto pierwsze prace astrogórnictwa będą prowadzone na zasobach lodu wodnego na Księżycu nie ze względu na wartość ekonomiczną, a ze względu na prostotę wykorzystanych systemów. System pozyskania, magazynowania i transportu wody będzie dla inżynierii kosmicznej oraz astrogórnictwa ważny ze względu na testowanie urządzeń pracujących w warunkach księżycowych i ich zużycia oraz rozpoznania np. nieznanych dziś zjawisk, które mogą zajść na którymś etapie ich pracy. Systemy wodne mają bardzo małą ilość części ruchomych, co zmniejsza też podatność na awarie. Natomiast bliskość Księżyca pozwala na lepszą komunikację i pracę np. w tandemie lub zespole robotów. Podobnie rzecz będzie się miała z użyciem regolitu na potrzeby budowy lub osłony struktur księżycowych. Regolit nie jest gotowym produktem, który ma jednolitą topliwość i właściwości materiałowe czy kształt, by dało się bezpośrednio z niego tworzyć filament do drukarki lub przędze z włókna krzemianowego.

Przy omawianiu systemów wodnych musimy pamiętać, że w tym kontekście system bezwodny nie jest jego przeciwieństwem. System wodny wymieniony w tekście jest systemem pozyskania wody z lodu znajdującego się na ciele niebieskim, najczęściej przez podgrzewanie lodu pod osłoną w celu jego rozpuszczenia lub parowania. Natomiast o systemach bezwodnych mówić będziemy w kontekście systemów górniczych, które nie stosują płuczki, ze względu na warunki



**O tych, co przekuli innowacyjne wizje w biznesowy sukces**

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

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

To oni kształtują cywilizację technologiczną.

To bohaterowie naszych czasów.

Z Tajwanu na podbój świata – **Jen-Hsun „Jensen” Huang**

Jako współzałożyciel i wieloletni szef firmy NVIDIA jest jednym z najważniejszych „sprawców” rewolucji w przetwarzaniu grafiki komputerowej, a potem, gdy okazało się, że procesory graficzne jego firmy to klucz do rozwoju sztucznej inteligencji, także jedna z głównych postaci trwającej rewolucji AI.



1. Jensen Huang

CV: Jen-Hsun „Jensen” Huang

Data i miejsce urodzenia: 17.02.1963, Tainan, Tajwan

Adres zamieszkania: Santa Clara, Kalifornia, USA

Obywatelstwo: amerykańskie, tajwańskie

Stan cywilny: żonaty, dwoje dzieci

Majątek: ok. 12 mld dolarów

Kontakt: jensenuhuang@nvidia.com

Edukacja: Szkoła średnia w Aloha, stan Oregon, USA, Uniwersytet stanowy Oregonu – 1984, Uniwersytet Stanforda – 1992

Doświadczenie zawodowe: praca w LSI Logic i Advanced Micro Devices (AMD), od 1993 – założyciel i prezes NVIDIA

Zainteresowania: tenis stołowy, gotowanie, gry komputerowe



2. Jensen Huang (w środku) z Chrisem Malachowskim i Curtisem Priemem i logo NVIDIA w tle

Jen-Hsun „Jensen” Huang (1) urodził się siedemnastego lutego 1963 roku w Tainan na Tajwanie. Gdy miał pięć lat, jego rodzina, czyli również oczywiście on, przeniosła się do Tajlandii. Gdy miał dziewięć lat, wraz z bratem został wysłany do Stanów Zjednoczonych, by zamieszkać z wujem w Tacoma w stanie Waszyngton. Po kilku latach jego rodzice także przenieśli się do USA i osiedlili się w Oregonie, gdzie Huang ukończył szkołę średnią Aloha High School. Jego ulubionym sportem był wówczas tenis stołowy – był w nim tak dobry, że trafił do krajowego rankingu graczy. Warto dodać też, że przeskoczył dwa lata edukacji i ukończył szkołę już w wieku szesnastu lat. Wcześniej też zaczął pracować – był kierowcą i kelnerem w restauracji.

Jednocześnie Huang kontynuował naukę. Zdobył tytuł licencjata inżynierii elektrycznej na stanowym uniwersytecie Oregonu w 1984 roku, a następnie tytuł magistra w tej samej dziedzinie na prestiżowym Uniwersytecie Stanforda w 1992 roku.

Od gier ku sztucznej inteligencji

Zanim wraz z Chrisem Malachowskim i Curtisem Priemem założył firmę NVIDIA (2), pracował jako dyrektor w LSI Logic i projektant mikroprocesorów w Advanced Micro Devices (AMD). Jego kariera zawodowa w innych firmach nie trwała jednak długo. W 1993 roku, gdy Huang miał 30 lat, powstała NVIDIA, a on objął stanowisko jej prezesa.

Trójka założycieli skupiła się na projektowaniu i wytwarzaniu procesorów graficznych (GPU) dla branży gier wideo. Pierwszym produktem firmy NVIDIA

był NV1, procesor graficzny (GPU), który umożliwiał tworzenie grafiki 3D dla komputerów osobistych. Opracowali także m.in. układ RIVA 128, jeden z pierwszych procesorów graficznych integrujących transformację, przycinanie i oświetlenie w jednym chipie. Ta innowacja sprawiła, że NVIDIA znalazła się w czołówce branży przetwarzania grafiki. Firma szybko zyskała uznanie za swoją innowacyjną technologię graficzną, a jej procesory graficzne stały się popularne w branży gier.

Początkowo dysponowali skromnym jak na Dolinę Krzemową kapitałem w wysokości 40 tys. dolarów. Wkrótce jednak udało im się pozyskać inwestycje w wysokości 20 milionów od firm venture capital (inwestycje wysokiego ryzyka). Dzięki temu wsparciu, NVIDIA przetrwała trudne pierwsze lata i ostatecznie zdobyła solidną, trwałą pozycję na rynku półprzewodników. Weszła na giełdę w 1999 roku. Gdy cena jej akcji osiągnęła 100 dolarów za sztukę, Huang zrobił sobie tatuaż z logo firmy na lewym ramieniu.

W XXI wieku NVIDIA stopniowo wychodziła poza branżę gier w kierunku takich dziedzin jak sztuczna inteligencja, obliczenia mobilne, technologie pojazdów autonomicznych i sieci społecznościowe. Układy NVIDIA z architekturą CUDA stały się podstawą infrastruktury umożliwiającej przeprowadzanie skomplikowanych obliczeń na ogromną skalę i zrewolucjonizowały procesy uczenia głębokich sieci neuronowych.

Ważnym momentem w rozwoju potęgi firmy było stworzenie procesora graficznego GeForce 256, który był przełomem w dziedzinie grafiki komputerowej

i stał się podstawą dla wielu późniejszych rozwiązań firmy NVIDIA w dziedzinie uczenia maszynowego. W 2006 roku Huang wprowadził na rynek pierwszą kartę graficzną przeznaczoną specjalnie do zastosowań naukowych i obliczeń związanych z AI – Tesla GPU. Jej następczyni, Tesla P100, była procesorem graficznym zaprojektowanym specjalnie do zadań związanych ze sztuczną inteligencją. Była pierwszym układem GPU obsługującym operacje tensorowe, które są niezbędne dla algorytmów głębokiej nauki maszynowej. Od tego czasu NVIDIA wydała kilka produktów skoncentrowanych na sztucznej inteligencji, w tym Tesla V100, DGX-1 i Jetson AGX Xavier. Produkty te były wykorzystywane w różnych zastosowaniach, w tym w autonomicznych samochodach, obrazowaniu medycznym i w robotyce.

Wartość rynkowa NVIDIA przekroczyła ostatnio trzy biliony dolarów, co ustawia ją w gronie najwyższych wycenianych spółek świata. Huang, który ma 3,5 proc. akcji firmy (co czyni z niego obecnie multimiliardera), został obsypany tytułami „Biznesmena Roku” magazynu „Fortune” w 2016 roku i „Najlepszego Prezesa” Harvard Business Review w 2020 roku, Fellow IEEE, co jest najwyższym wyróżnieniem przyznawanym przez Instytut Inżynierii Elektrycznej i Elektronicznej USA. W 2021 roku magazyn „Time” umieścił go na liście 100 najbardziej wpływowych osób na świecie.

Huang jest znany ze swojego charakterystycznego stylu, często pojawia się w skórzanej kurtce. Od dzieciństwa jest zapalonym graczem w gry komputerowe, co było nie bez wpływu na powstanie NVIDIA. Szczególnie lubi „Cywilizację”. Jako człowiek nadzwyczajnie bogaty stał się znany ze swojej działalności filantropijnej. W 2022 roku przekazał 50 milionów dolarów na rzecz swojej Alma Mater, uniwersytetu w Oregonie, na utworzenie instytutu superkomputerów. Ponadto przekazał 30 milionów dolarów Stanfordowi na budowę szkoły inżynierskiej im. Jen-Hsuna Huang.

Prawo Huang

Szef NVIDIA znany jest z tzw. prawa Huang, opartego na obserwacjach zmian w dziedzinie informatyki i inżynierii. Opiera się na tezie, że postęp



3. GeForce 256 firmy NVIDIA

w dziedzinie procesorów graficznych (GPU) rośnie w tempie znacznie szybszym niż w przypadku tradycyjnych jednostek przetwarzających w komputerach (CPU). Obserwacja ta koresponduje z prawem Moore’a, które przewidywało, że liczba tranzystorów w układzie scalonym podwaja się mniej więcej co dwa lata. Według prawa Huang, wydajność procesorów graficznych będzie podwajać się co dwa lata. Obserwacja stojąca za owym prawem została poczyniona przez Jensena podczas konferencji GPU Technology Conference (GTC) 2018, która odbyła się w San Jose w Kalifornii. Zauważył on wtedy, że procesory graficzne NVIDIA były „25 razy szybsze niż pięć lat temu”, chociaż prawo Moore’a przewidywało jedynie dziesięciokrotny wzrost.

Oprócz zainteresowania nowym oryginalnym spojrzeniem szefa NVIDIA pojawiły się głosy krytyki. Dziennikarz Joel Hruska, pisząc w serwisie „ExtremeTech” w 2020 roku, stwierdził, że „nie ma czegoś takiego jak prawo Huang”, nazywając je „iluzją”, która opiera się na zyskach możliwych dzięki prawu Moore’a; i że jest zbyt wcześnie, aby stwierdzić, że prawo istnieje. Organizacja badawcza Epoch odkryła, że w latach 2006...2021 wydajność cenowa GPU (pod względem FLOPS/\$) podwajała się mniej więcej co 2,5 roku, czyli znacznie wolniej, niż przewidywało prawo Huang.

Niezależnie od tego, jak ostatecznie zostanie ocenione prawo Huang, jego dorobek trudno kwestionować. Jest właścicielem ponad stu patentów w dziedzinie grafiki komputerowej, architektury GPU i sztucznej inteligencji. I zapewne nie powiedział jeszcze w tej dziedzinie ostatniego słowa. ■

Mirosław Usidus



POLSKA FUNDACJA
FANTASTYKI NAUKOWEJ

młody
m.technik

POWRÓT DO PRZYSZŁOŚCI

Fantastyka naukowa znów w „Młodym Techniku”

Nowy dom



Pierwsza wojna wyniszczyła półtora kontynentu. Druga odcisnęła swe piętno na dwóch. Trzecia wgryzła się w ziemię pozostałych, wsiąkła w nią i wszyscy po raz kolejny zaręczali, że nigdy więcej, a potem nastąpiła czwarta. Ta, która kazała stracić wszelką nadzieję.

KMR Mrok, siódmy okręt z floty trzydziestu sarkofagów cywilizacji, po czterech latach od przekroczenia magnetycznych baniek płaszcza systemowego w końcu wszedł do właściwego układu planetarnego 4-Tarczy – piątej co do jasności gwiazdy tej konstelacji. Kapitan Vizvet stała przy bulaju i patrzyła na tutejsze słońce. Kulę wodoru i helu nawet z tak daleka jaśniejszą od innych.

– Będziemy musieli się nauczyć, że już piątą nie jest... Mrok, planety?

Nim sarkofagi zabrały genetyczne archiwum i kapsuły hibernacyjne, sprawdzili wszystkie trzydzieści układów, ale i tak co jakiś czas Vizvet pytała inteligencję statku o najnowsze obserwacje.

– Potwierdzam istnienie pięciu planet.

Miało być siedem! Zaciśnięła powieki.

– Ekosfera?

– Zero.

Musiały kryć się za gwiazdą centralną...

Podchodzili w płaszczyźnie ekliptyki, Vizvet wołałaby od bieguna, nie żeby to cokolwiek zmieniło. Jeśli naukowcy popełnili błąd, typując układ 4-Tarczy, to podróż Mroku miała zakończyć się porażką. Nie mieli dość paliwa, by lecieć dalej. A inni? Jedenaście lat temu dogonił ich agonálny krzyk KMR Ciska. Czy kolejne sygnały były w drodze? Czy został tylko Mrok?

Siedem minut światła dalej statek potwierdził istnienie trzech kolejnych planet. Dwie znajdowały się w ekosferze. To ich szukali. Nowa, ósma, palona żarem 4-Tarczy, nie miała znaczenia.

Vizvet odetchnęła pierwszy raz od dawna.

W hibernatorach Mroku czekało trzy tysiące dwanaście osób gotowych budować nowy dom od zaraz. W chłodniach leżały zarodki – klony tych, co pragnęli przede wszystkim żyć – kolejnych dziesięciu tysięcy, ale to ci pierwsi byli niepowtarzalni. Dla nich KMR Mrok stanowił jedyną szansę.

– Atmosfery?

– Tak.

– Skład? – głos jej zadrżał.

N₂: 2,99%, O₂: 17,55%, CO₂: 77,43%

N₂: 78,08%, O₂: 20,86%, CO₂: 0,09%

Zobaczyła pod powiekami. Pierwsza fatalna. Ale druga? Może nie idealna, ale lepsza, bliższa optimum. Byli w stanie ją dostosować.

– Mrok.

Statek nie odpowiedział, ale słuchał.

– Ślady istnienia cywilizacji na powierzchni? – wyszeptała.

Wstrzymała oddech. W imię domu Vizvet gotowa była uruchomić artylerię Mroku, ale wolałaby nie budować nowego początku na trupach słabszych.

– Źródła światła, brak. Źródła ciepła, brak. Sygnatury geograficzne, brak.

Odetchnęła.

Marcelina zaklęła. Siedziała w HQ HS-39-Triton-A głęboko pod lodową powłoką neptunowego księżyca, przecinaną elastycznymi rurociągami kryjącymi kable sygnałowe. To właśnie przez jeden z nich spłynęła informacja, że pasywne sensory na powierzchni coś zobaczyły. Jedenaście ziemskich dni wcześniej autonomiczne obserwatorium na Eris zarejestrowało obiekt poruszający się po nienaturalnej trajektorii. Aarne zaklasyfikował go jako potencjalny błąd pomiaru, ale teraz Tryton też go widział.

– Statek – szepnęła.

SI habitatu potwierdziła.

Długi na dwieście osiem metrów, szeroki na trzydzieści trzy z hakiem i wysoki na niemal pięćdziesiąt. Mógł tu być w dowolnym celu, ale rozproszony uniwersytet od dawna twierdził, że nawet jeśli życie w Słonecznym powstało pierwsze, to młodsze cywilizacje w końcu ich dogonią, a potem próbują prześcignąć. I gdy ludzkość postanowiła przeczekać regenerację Ziemi po jej nieuchronnym wyniszczeniu ukryta pod lodem księżyców gazowych olbrzymów, to inni mogą wybrać inaczej. Nie będą czekać. Zamiast patrzeć, jak ogień, powodzie oraz huragany pożerają ślady ich cywilizacji, obserwować, jak w atmosferze płoną kolejne porzucone satelity oraz stacje, oni polecą ku innym gwiazdom, by znaleźć nowy dom.

– Błagam. Szukajcie Marsa. Zaczęliśmy go grzać, nim pieprznęło. Oddamy go wam bez jednego wystrzału...

Kapitan Vizvet patrzyła na planetę przez niewielkie okno. Zachwyt ją paraliżował. Nigdy nie widziała żadnej planety na własne oczy. Urodziła się na Mroku jako klon pierwszej kapitan Vizvet, ale uczyła się o planetach od matki i od inteligencji statku. Wiedziała, że takie jak ta były darem. Stary dom miał pięć kontynentów, na planecie 4-Tarczy Mrok znalazł sześć. Tu masa lądowa skupiała się głównie na jednej półkuli, w domu była rozłożona bardziej równomiernie; szczegóły.

– Witajcie w nowym domu – szepnęła kapitan Vizvet ze ściśniętym z radości gardłem.

– Obiekt minął Marsa. Prawdopodobieństwo, że celem jest Ziemia 98,851% – zameldowała SI.

Marcelina zaklęła głośno.

Aarne spojrzał na nią pytająco.

– Nienawidzę rozproszonego.

– Dlaczego?

– Bo mieli rację. – Drżała. Dłonie miała pokryte potem i śmierdziała strachem. Marzyła o porozumieniu, poznaniu, ale ryzyko było zbyt wielkie. – Aktywacja pola na Mare Nubium. – Odetchnęła głęboko. – Przepraszam...

Kapitan Vizvet wisiała zmrożona strachem ponad pokładem mostka. Inteligencja Mroku zameldowała o nagłe wzbierającym źródle energii na satelicie Nowego Domu.

– Wstępna analiza: działo jonowe. Cel: sarkofag cywilizacji numer siedem.

– Uni... ■

Marta Magdalena Lasik

Ataki bez plików i inne nowe cyberzagrożenia

Wiedza chroni najlepiej

Dwa lata temu firma Kaspersky Lab poinformowała o odkryciu nietypowego, pierwszego podobno w swoim rodzaju, ataku złośliwego oprogramowania „bezplikowego”. Nieużywające plików jako nośników ataki nie są wprawdzie nowością w ogóle, ale tym razem złośliwy kod „bez pliku” schował się w systemie plików, czyli wyjątkowo perfidnie.

Mówiąc językiem specjalistów, w lutym 2022 r. po raz pierwszy zaobserwowano technikę umieszczania kodu powłoki w dziennikach zdarzeń systemu Windows w ramach ataku złośliwego oprogramowania. Pozwoliło to na ukrycie w systemie plików „bezplikowego” trojana aktywowanego na ostatnim etapie ataku.

Czerwony kod i kolejne „atrakcje”

Złośliwe oprogramowanie (1) to ogólny termin odnoszący się do każdego software’u zaprojektowanego w celu uszkodzenia lub wykorzystania do przestępczych celów dowolnego programowalnego urządzenia, usługi lub sieci. Cyberprzestępcy zazwyczaj używają takich narzędzi do wydobywania danych, które można dalej wykorzystać w celu uzyskania korzyści finansowych lub innych. Zaś bezplikowa odmiana złośliwego oprogramowania to taki rodzaj cyberataku, w którym „malware” nie umieszcza pliku wykonywalnego na dysku, korzystając z plików już w systemie istniejących i wykonywanych. Działa w pamięci RAM. Oznacza to, że tradycyjne oprogramowanie antywirusowe ma trudności z jego wykryciem – nie ma bowiem określonego pliku wskazującego na infekcję.

Uważa się, że pierwszy przykład bezplikowego malware’u pojawił się w 2001 roku wraz z wirusem Code Red. Wykorzystywał lukę przepełnienia bufora w serwerach internetowych Microsoft IIS. W efekcie ponad 350 tysięcy serwerów zostało dotkniętych awarią, a na stronach hostowanych wyświetlany był komunikat „Witamy w <http://www.worm.com>! Zhakowany przez Chińczyków!” (2). Kod za to odpowiadający nie pozostawił żadnych plików ani trwałych śladów na dysku twardym, ponieważ działał jedynie w pamięci zainfekowanych maszyn. W kolejnych latach liczba



1. Komputerowe malware – wizualizacja

podobnych infekcji rosła. Tylko w 2018 roku tego typu ataki stanowiły około 35 proc. wszystkich kampanii cyberprzestępczo-hackerskich.

Atak bezplikowy może mieć postać ciągów tekstowych pobranych z serwera WWW, a następnie przekazanych jako parametry do programu konfigurującego takiego jak PowerShell, który interpretuje skrypty. Złośliwy kod zostanie następnie wykonany bezpośrednio w pamięci z PowerShell, bez śladu na dysku. W wykrytym przez Kaspersky Lab ataku bezplikowym wykorzystano wiele technik,

Welcome to <http://www.worm.com> !

Hacked By Chinese!

2. Komunikat wyświetlany w ramach ataku Code Red

m.in. komercyjne pakiety do testów penetracyjnych i pakiety antywyciekowe, w tym skompilowane w języku programowania Go, a także kilka programów trojańskich infekujących atakowaną maszynę. Aby uzyskać dalszy dostęp do systemu, hakerzy wykorzystali dwie różne metody – za pośrednictwem komunikacji sieciowej HTTP i przez korzystanie z nazwanych wcześniej ciągów. Złośliwe oprogramowanie było ukryte w duplikacie istniejącego pliku, dodanym do ciągu.

Rosnące zagrożenie atakami „bezplikowymi” jest na fali rosnącej, razem z innymi nękającymi użytkowników komputerów i sieci zagrożeniami, takimi przede wszystkim jak ransomware, w których złośliwe oprogramowanie szyfruje pliki na komputerze ofiary i żąda okupu za ich odblokowanie. Stało się to jednym z największych zagrożeń dla firm i osób indywidualnych. Popularność zdobywają obecnie także opisywane przez MT niejednokrotnie zagrożenia wobec IoT – wraz ze wzrostem popularności Internetu Rzeczy rośnie liczba tworzonych z tych urządzeń botnetów, służących głównie do ataków DDoS. Nowością ostatnich lat jest cryptojacking, nielegalne wykorzystywanie mocy obliczeniowej ofiary do wydobywania kryptowalut. Wraz ze wzrostem popularności chmur obliczeniowych hakerzy zaczęli wykorzystywać luki w zabezpieczeniach takich usług do kradzieży danych lub przejęcia kontroli nad urządzeniami.

Podobnie ataki z wykorzystaniem uczenia maszynowego do tworzenia bardziej zaawansowanych i przekonujących ataków, np. deepfake. Potwierdzają to nowe badania przeprowadzone przez firmę Menlo Security. „Pracownicy wykorzystują sztuczną inteligencję w swojej codziennej pracy. Kontrole nie mogą jej po prostu zablokować, ale nie możemy też pozwolić jej działać dziko”, mówi Andrew Harding, wiceprezes Menlo Security, w wywiadzie dla VentureBeat. „Organizacje wzmacniają środki bezpieczeństwa, ale jest pewien haczyk. Większość z nich stosuje te zasady tylko w odniesieniu do domeny, co już nie wystarcza”. Ta fragmentaryczna taktyka nie nadąza za stale pojawiającymi się nowymi platformami generatywnej sztucznej inteligencji. Raport ujawnił, że liczba prób przesyłania plików w związku z działalnością usług AI wzrosła o 80 proc. w ciągu pół roku. Ryzyko

w tym przypadku wykracza daleko poza potencjalną utratę danych w wyniku przesyłania plików.

Badacze ostrzegają, że generatywna sztuczna inteligencja może również poważnie wspomagać oszustów stosujących tzw. phishing, czyli podszywanie się w celu wyłudzenia danych. Niedawno oszuści ukradli 25,6 miliona dolarów przy użyciu sztucznej inteligencji, podszywając się pod dyrektora finansowego firmy podczas sfigowanej wideokonferencji, w której uczestniczyły fałszywe wideoawatary pracowników firmy. Wykorzystali fałszywe głosy i zmanipulowane filmy, aby nakłonić pracownika międzynarodowej korporacji w Hongkongu do przelewu.

Rosyjska wojna cyberpsychologiczna

Nowością jest też nasilenie cyberwojny, w której uczestniczą rosyjskie i chińskie wyspecjalizowane jednostki hakerskie, niekoniecznie stosując środki techniczne, czyli wirusy, malware i łamanie zabezpieczeń. Bardzo często są to działania propagandowo-socjotechniczne oparte na wiedzy społeczno-psychologicznej i manipulacjach dużymi grupami ludzi. Wraz z pojawieniem się nowych narzędzi AI arsenał stosowanych przez nich środków znacznie się poszerzył.

Opublikowany pod koniec ubiegłego roku raport centrum analiz zagrożeń firmy Microsoft pt. „Russian Threat Actors Dig In, Prepare to Seize on War Fatigue” (z ang. – „Rosyjscy agenci i twórcy zagrożeń kopią, przygotowują się do wykorzystania zmęczenia wojną”) pokazuje te nowe techniki działania w kontekście wojny na Ukrainie, np. wyłudzone nagrania wideo amerykańskich celebrytów, które wkomponowane w odpowiedni kontekst zyskują antyukraiński wydźwięk. Początkowo uważano, że śmierć Jewgienija Prigożyna, który był właścicielem Wagner Group i farmy trolli Internet Research Agency, osłabi aktywność rosyjskiej propagandy w Internecie. Jednak Microsoft w swoim raporcie opisuje szeroko zakrojone operacje wpływu prowadzone przez rosyjskie podmioty, które nie są powiązane z Prigożynem.

Latem, w okresie wzrostu plonów i żniw, Rosja penetrowała przedsiębiorstwa rolne, kradła dane, wdrażała złośliwe oprogramowanie i wykorzystywała ataki wojskowe do niszczenia zboża. Raport Microsoftu pokazuje silne powiązanie między działaniami wojskowymi, propagandowymi i cyberatakami. Na przykład w ciągu czterech dni pod koniec lipca 2023 r., po wycofaniu się Moskwy z Czarnomorskiej Inicjatywy Zbożowej, Rosja zaatakowała obiekty infrastruktury zbożowej w Odessie za pomocą pocisków manewrujących i równocześnie przeprowadziła cyberatak na ukraińską organizację zajmującą się sprzętem



3. Ilustracja rosyjskich operacji cyberwojennych pochodząca z raportu Microsoftu ©Microsoft

rolniczym. Rozpowszechniała w prorosyjskich mediach zarzuty, że Ukraina, USA i NATO nadużywały korytarza zbożowego do celów terrorystycznych, a nie pomocy humanitarnej. We wrześniu 2023 r. Rządowy Zespół Reagowania na Incydenty Komputerowe Ukrainy (CERT-UA) ogłosił, że ukraińskie sieci energetyczne są na celowniku, a Microsoft Threat Intelligence obserwował ślady aktywności rosyjskiego wywiadu wojskowego (GRU) w ukraińskiej infrastrukturze energetycznej od sierpnia do października 2023 r.

Były inne operacje. Wiązany z Rosją podmiot Storm-1099, oskarżany przez europejską organizację DisinfoLab o masowe fałszowanie stron internetowych od wiosny 2022 r., podjął działania propagandowe skierowane w zwolenników Ukrainy na świecie. Grupa tworzy serwisy informacyjne takie jak Reliable News Network (RNN) w celu rozsiewania antyukraińskiej propagandy, łącząc świat cyfrowy i i działania organizacyjne, wspierając np. demonstracje. W przeszłości celem Storm-1099 były

kraje Europy Zachodniej, zwłaszcza Niemcy, ale obecnie skupił się na Izraelu i Stanach Zjednoczonych, w związku z wojną i wyborami prezydenckimi w USA. Serwisy Storm-1099 propagowały np. fake newsy, że Hamas nabył ukraińską broń. Jednocześnie prowadzono działania typowo hakerskie. Grupa Forest Blizzard próbowała uzyskać dostęp do instytucji militarnych na Zachodzie za pośrednictwem phishingu.

Amerykańskie struktury wywiadowcze uznały te zagrożenia za na tyle poważne, że postanowiły głębiej zbadać techniki wojny „cyberpsychologicznej”. Organizacja Intelligence Advanced Research Projects Activity (IARPA) ogłosiła niedawno, że „poszukuje informacji na temat metod, badań, ustaleń, podejść i odpowiednich wskaźników do scharakteryzowania efektów poznawczych w operacjach cybernetycznych”. „Techniki stosowane obecnie w reklamie internetowej, kampaniach politycznych, handlu elektronicznym i grach online z powodzeniem wykorzystują słabe punkty ludzkiej psychologii”, czytamy w RFI wydanym przez IARPA. „Cyberprzestępcy często wykorzystują podobne ludzkie ograniczenia, stosując inżynierię społeczną”. „Jakie wewnętrzne i zewnętrzne czynniki kierują zachowaniem cyberoperatorów?” – to jedno z ośmiu pytań zadanych przez IARPA ekspertom ze środowisk akademickich i przemysłu prywatnego.

IARPA i nie tylko ona wierzy, że szersza świadomość i znajomość cyberpsychologicznych technik manipulacji już jest niezłym sposobem na walkę z nimi. Z eksperymentów wynika, że ataki na osoby mające nawet ogólną wiedzę o tych technikach były znacznie mniej skuteczne. ■

Mirosław Usidus

Zaginiona księgarnia

Evie Woods

Wydawnictwo: StoryLight, liczba stron: 456, cena: 44,99 zł

Przy cichej ulicy w Dublinie zaginiona księgarnia czeka, by ją odnaleźć...

Opalina, Martha i Henry zbyt długo godzili się na rolę drugoplanowych bohaterów.

Kiedy w końcu postanawiają przejąć stery własnego życia, los prowadzi ich pod pewien adres, gdzie zapomniana księgarnia czeka, by ktoś ją odnalazł. Trójka niczego niepodważających nieznajomych przekonuje się, że ich własne historie są tak samo niezwykle, jak te, które opisano na stronach ich ukochanych książek.

Odkrywając tajemnice księgarnianych pótek, przenoszą się do cudownego świata... gdzie nic nie jest takie, jakie się wydaje.





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

Dawid Przepiórka był jednym z czołowych polskich szachistów pierwszej połowy XX wieku i znanym na świecie autorem kompozycji szachowych. Z jego inicjatywy powstał Polski Związek Szachowy, którego był wieloletnim wiceprezesem i mecenasem. Zwyciężył w pierwszych mistrzostwach Polski i z polską drużyną wygrał olimpiadę szachową w Hamburgu w 1930 roku. Jako pierwszy w historii Polak, otrzymał tytuł Honorowego Członka Międzynarodowej Federacji Szachowej.

Dawid Przepiórka

Dawid Przepiórka to wybitny polski szachista i kompozytor szachowy (1). Urodził się 22 grudnia 1880 roku w Warszawie w rodzinie zamożnych właścicieli ziemskich i przedsiębiorców pochodzenia żydowskiego. Syn Izraela Przepiórki, bogatego warszawskiego inwestora budowlanego i Dwojry (Debory) Kon, córki lubelskiego rabina Szneura Zalmana Ladiera. Zginął w pierwszych miesiącach 1940 r., rozstrzelany przez hitlerowców w Palmirach.

Był miłośnikiem szachów, uznawanym za genialne szachowe dziecko, a potem znanym na całym świecie głównie z kompozycji szachowych i wytrawnym znawcą gry końcowej. Dawid nauczył się grać w szachy w wieku siedmiu lat, chociaż nikt z jego rodziny nie umiał grać w tę grę. Jako dwunastolatek pokonał znanego francusko-polskiego mistrza szachowego Jana (Jeana) Taubenhauza (jednego z prekursorów gry symulta-

nicznej na ziemiach polskich) podczas rozgrywanej przez niego symultany. Prasa warszawska donosiła wówczas, że „Przepiórka wygrał partię w błyskotliwym stylu”, ale przebieg tej sensacyjnej partii nie został zachowany. Pierwsza kompozycja szachowa autorstwa Przepiórki została zamieszczona w 1896 r., kiedy miał 16 lat.



**2. Dawid Przepiórka, Bohe-
mia, 25 stycznia 1903. Mat w 2
posunięciach**

W latach 1903–1905 Dawid Przepiórka studiował na wydziale matematyki Uniwersytetu Warszawskiego. W latach 1905–1914 przebywał w Niemczech, kontynuując studia na uniwersytecie w Getyndze, a następnie w Monachium. Podczas studiów w Monachium poznał niemiecką szkołę problemową i stworzył wiele zadań, głównie trzy- i czterochoďówek. Jego wielochoďówki były na najwyższym światowym poziomie i wiele z nich, również w dzisiejszych czasach, zadziwia swoją pomysłowością i wzorową realizacją. Lata pierwszej wojny światowej (1914–1918) spędził w Genewie (Szwajcaria). Po wojnie wrócił do Warszawy i zasilili szeregi Warszawskiego Towarzystwa Zwolenników Gry Szachowej. W 1932 r. ukazała się książka H. Weeninka „David Przepiórka a Master of Strategy”, zawierająca zbiór 130 jego zadań. W sumie opublikował 160 zadań szachowych, z czego 21 znalazło się w Albumach FIDE. Swoje zadania publikował na łamach fachowych



**1. Dawid Przepiórka – polski szachista, kompozytor
zadań szachowych i matematyk,
źródło: <https://tiny.pl/dn441>**



Wyniki indywidualne olimpiady szachowej w Hamburgu w 1930 roku

Sz.	Imię i nazwisko	1	2	3	4	5	6	7	8	9	0	1	2	3	4	5	6	7	Pkt.	Gier	%
1	A. Rubinstein	1	1	1	=	1	1	=	1	=	1	1	1	=	1	1	1	1	15.0	17	88
2	K. Tartakower	1	1	0	=	=	1	=	*	=	1	=	1	1	1	=	1	1	12.0	16	75
3	D. Przepiórka	1	*	=	=	*	1	*	1	=	0	*	=	0	1	1	1	1	9.0	13	69
4	K. Makarczyk	=	0	1	*	=	1	=	=	*	0	0	*	*	1	1	1	=	7.5	13	58
5	P. Frydman	*	=	*	=	1	*	=	1	0	*	1	=	0	*	*	*	*	5.0	9	56

czasopism zagranicznych (głównie niemieckich) oraz w dziennikach i tygodnikach krajowych: „Tygodniku Ilustrowanym”, „Kurierze Warszawskim”, „Wędrowcu”, „Kłosach”, „Ziarnie”, „Sporcie”, „Tygodniku Szachowym”, „Ilustracji Polskiej” i „Neue Lodzer Zeitung”.

A to jedna z kompozycji szachowych Dawida Przepiórki opublikowana w czasopiśmie „Bohemia” 25 stycznia 1903 roku. Rozwiązanie zadania znaleźć można na końcu tego artykułu.

Uzyskany na turnieju w Coburgu (Niemcy) w 1904 r. tytuł szachowego mistrza pozwolił Dawidowi Przepiórce uczestniczyć w turniejach międzynarodowych. W szachach klasycznych największe indywidualne sukcesy odniósł w 1926 r., kiedy wygrał turniej w Monachium (przed Bogolubowem i Spielmannem) oraz pierwsze mistrzostwa Polski w Warszawie (3).

W 1928 r. w Hadze zdobył tytuł wicemistrza świata amatorów (za Maxem Euwe). Reprezentował Polskę na dwóch olimpiadach szachowych (1930 i 1931). Wraz z drużyną w 1930 r. w Hamburgu zdobył złoty medal olimpijski (indywidualnie uzyskał wynik 9 z 13 rozegranych partii), a rok później w Pradze – srebrny (indywidualnie 10/17).

W olimpiadzie w Hamburgu startowało 18 drużyn, zwyciężyła reprezentacja Polski (4) przed Węgrami i Niemcami. Lider polskiej drużyny – Akiba Rubinstein, grający w tej olimpiadzie na pierwszej szachownicy, osiągnął wtedy niezwykle wynik – wygrał trzynaście partii i zremisował w czterech!

A oto jedna z partii rozegranych przez Dawida Przepiórkę, który podczas olimpiady w Hamburgu w 1930 roku pokonał słynnego Marcela Duchampa. Marcel Duchamp (1887–1968) był jednym

3. Mistrz Polski Dawid Przepiórka w symultanie szachowej z 27 członkami Towarzystwa Miłośników Gry Szachowej im. Józefa Dominika w Krakowie, źródło: Narodowe Archiwum Cyfrowe, <https://tiny.pl/dn44j>





4. Zdobywcy złotego medalu na Olimpiadzie Szachowej w roku 1930 w Hamburgu, od lewej stoją: Paulin Frydman, Ksawery Tartakower, Stefan Rotmil, Akiba Rubinstein, Kazimierz Makarczyk, Dawid Przepiórka, Marian Wróbel, źródło: <https://tiny.pl/dn44l>

z najważniejszych artystów awangardowych XX wieku, prekursorem większości nowoczesnych prądów: dadaizmu, a następnie pop-artu, happeningu i konceptualizmu. Uznawany jest za najlepszego szachistę wśród malarzy. Pisałem o nim w artykule Marcel Duchamp – szachy i sztuka w nr 8/2023 „Młodego Technika”.

Marcel Duchamp – Dawid Przepiórka, Hamburg 25.07.1930, runda 15

1. d4 Sf6 2. c4 e6 3. Sc3 d5 4. Gg5 Sbd7 5. e3 c6 6. c:d5 e:d5 7. a3 Gd6 8. Gd3 O-O 9. Sf3 We8 10. O-O Sf8 11. Hc2 Gg4 12. Sh4 h6 13. G:f6 H:f6 14. Sf5 Wad8 15. S:d6 W:d6 16. h3 Gc8 17. Se2 Hg5 (diagram 5) 18. f4 (błędna decyzja osłabiająca centralnego piona e3, którego nie można będzie obronić, białe powinny grać 18. Kh2 lub 18. Sg3) 18. He7 19. Wf3 We6 20. Waf1 W:e3 21. Sg3 W:f3 22. W:f3 He1+ 23. Sf1 Sd7 24. g4 Hh4 25. Se3 Sf8 26. f5 (Osłabia pozycję białego króla, lepsze było 26. Hf2 lub 26. Sg2) 26...Sh7 27. Sg2 Hf6 28. Wf4 Sg5 29. Kh2 Hd6 (diagram 6) 30. Kg3 (Bezpieczniejsze było 30. Hf2) 31. h4 a5 32. Hd1 b6 33. Ha4 Gd7 34. Hc2 c5 35. d:c5 b:c5 36. Hd1 Gc6 37. Hf1 Kf8 38. Kh2 Sf7 39. Gb5 G:b5 40. H:b5 We4 41. Kg3 d4 42. H:a5 d3 43. Hd2 We2 44. Hd1 W:g2+ 45. K:g2 H:f4 46. H:d3 H:g4+ 0-1 (diagram 7).



5. Marcel Duchamp – Dawid Przepiórka, pozycja po 17...Hg5



6. Marcel Duchamp – Dawid Przepiórka, pozycja po 29...Hd6



8. Restauracja Gastronomia na rogu Nowego Świata i Alei Jerozolimskich, źródło: NAC, <https://tiny.pl/dn444>

Od 1920 (z niewielkimi przerwami) redagował dział szachowy w „Tygodniku Ilustrowanym”, jak również w „Kurierze Polskim” i tygodniku „Dookoła Świata”. W latach 1926–1933 był redaktorem naczelnym i współwydawcą miesięcznika „Świat Szachowy”. Z inicjatywy Dawida Przepiórki powstał Polski Związek Szachowy, którego był wiceprezesem w latach 1926–1939. Od 1933 r. kierował pracami Warszawskiego Koła Problemistów. Jedną z jego pasji było gromadzenie literatury szachowej, miał najbogatsze w kraju zbiory, liczące ponad tysiąc tomów.

Był wieloletnim wiceprezesem



7. Marcel Duchamp – Dawid Przepiórka, pozycja końcowa, w której Marcel Duchamp poddał partię



Polskiego Związku Szachowego i mecenasem, który wspierał finansowo PZSzach. Dawid Przepiórka był właścicielem dużego domu na skrzyżowaniu warszawskich Alei Jerozolimskich i Nowego Świata (8). Zarówno kamienica, jak i znajdująca się tam restauracja Gastronomia (jedna z najpopularniejszych restauracji w międzywojennej Warszawie) zapewniały mu stały dochód.



9. Linie prowadzenia pionów – linie Przepiórki

Miłość Przepiórki do szachów i jego ogromny majątek znalazły szczęśliwe połączenie. Bez jego hojnego wkładu finansowego Polacy nie mogliby zorganizować w Warszawie Olimpiady Szachowej w 1935 roku. Hojność Przepiórki została w Warszawie bardzo doceniona. Za zasługi w organizacji szachowej olimpiady w Warszawie w 1935, której był inicjatorem i finansowym mecenasem, został odznaczony Złotym Krzyżem Zasługi (12 października 1937). W 1936 roku jako pierwszy w historii Polak otrzymał tytuł Honorowego Członka Międzynarodowej Federacji Szachowej FIDE (Fédération Internationale des Échecs).

Linia prowadzenia piona – linia Przepiórki

Jeżeli oba króle są w pobliżu piona, to król strony silniejszej powinien starać się wyjść przed swojego piona i doprowadzić go do pola przemiany.

Strona silniejsza wygrywa, jeżeli jej król zdoła zająć **linię prowadzenia piona**, zwaną **linią Przepiórki**. Linie prowadzenia piona tworzą trzy pola (nazywane też polami kluczowymi piona), znajdujące się o dwa

10. Dawid Przepiórka (po prawej) podczas partii szachowej z Marianem Wróblem, źródło: <https://tiny.pl/dn44n>



11. Dawid Przepiórka, źródło: Narodowe Archiwum Cyfrowe, <https://tiny.pl/dn44k>

pola przed pionem, gdy pion nie dotarł jeszcze do linii środkowej i bezpośrednio przed pionem, gdy pion przekroczył linię środkową.

Na diagramie 9 dla piona stojącego na polu b4 linie Przepiórki tworzą pola a6, b6 i c6, a dla piona na polu g5, będącego już w obozie przeciwnika, linie Przepiórki tworzą pola f6, g6 i h6.

W przypadku pionów w kolumnach skrajnych a i h zasada linii Przepiórki nie ma zastosowania.

Po ataku Niemiec na Polskę we wrześniu 1939 r. Przepiórka najpierw stracił mieszkanie, a potem kamienicę na rogu Nowego Świata i Alei Jerozolimskich, a on sam przeprowadził się do mieszkania swojego najbliższego współpracownika i przyjaciela Mariana Wróbla, jednego z najlepszych kompozytorów szachowych na świecie (10).

W styczniu 1940 roku Dawid Przepiórka został aresztowany przez gestapo wraz z grupą ok. 30 szachistów w kawiarni szachowej mieszczącej się w prywatnym mieszkaniu Franciszka Kwiecińskiego przy ulicy Marszałkowskiej 76 (róg Hożej). W Areszcie Śledczym przy ul. Daniłowiczowskiej w Warszawie Przepiórka wygłosił dla współwięźniów swój ostatni wykład poświęcony końcówkom: król i pion przeciwko królowi, w którym przedstawił stworzoną przez siebie regułę szachową. Aresztowany wraz z Przepiórką Zygmunt



12. Certyfikat tytułu mistrz FIDE kompozycji szachowej, źródło: <https://tiny.pl/dn442>

Szulce kilka lat po wojnie opublikował wspomnienie o tych wydarzeniach, opisał regułę i nazwał ją „linią Przepiórki”. Część aresztowanych szachistów została po kilku dniach wypuszczona. Osoby pochodzenia żydowskiego zostały przewiezione do więzienia Pawiak i w okresie styczeń–kwiecień 1940 rozstrzelane przez Niemców podczas masowych egzekucji w Palmirach.

Polski Związek Szachowy zorganizował szereg międzynarodowych turniejów szachowych (memoriałów) pamięci Przepiórki, m.in. w 1950 (Szczawno-Zdrój) – zwyciężył Paul Keres, w 1957 (Szczawno-Zdrój) – zwyciężył Jefim Geller, w 1983 (Warszawa) – zwyciężył Josif Dorfman i w 2019 (Warszawa) – zwyciężył Alexey Kislinsky.

W 2023 Międzynarodowa Federacja Szachowa pośmiertnie odznaczyła Dawida Przepiórkę (11) tytułem mistrza FIDE kompozycji szachowej (12). ■

Zadania do samodzielnego rozwiązania



Zadanie 1
13. Vaccaroni vs. Mazocchi,
1891
Mat w 3 posunięciach



Zadanie 2
14. Jerzy Konikowski, 1967
Mat w 2 posunięciach

Kompozycja szachowa z tego artykułu:

Dawid Przepiórka, Bohemia, 25 stycznia 1903. Mat w 2 posunięciach

Rozwiązanie: 1.Sd7! 1...Ka6 2.Sc5# 1...Kc6 2.Sd6# 1...Kc8 2.Sd6# 1...Ka8 2.He4#

Rozwiązanie zadań z MT 7/2024

Zadanie 1

G. Heathcote, „American Chess Bulletin”, 1911

Mat w 2 posunięciach

Rozwiązanie: 1...Sd4!

1...H:d4+ 2.H:h7# lub 1...G:d4 2.Hb1#

Zadanie 2

G. Rinder, „Schach”, 1961

Mat w 2 posunięciach

Rozwiązanie: 1.Hd5+!! K:d5 2.Wf5#

Archiwalne odcinki o tematyce szachów
<http://bit.ly/2VohMA1>





Szczegółowa odpowiedź na to pytanie zależy od tego, czy mówimy o tradycyjnej żarówce z żarnikiem, świetlówce kompaktowej czy lampie diodowej LED. Jakkolwiek u podstaw świecenia każdego z tych urządzeń stoją odmienne procesy fizyczne, to łączy je jedno: prowadzą one do wzbudzenia elektronów w materiale na wyższy poziom energetyczny.

Światło elektryczne

Dlaczego żarówka świeci?

W większości przypadków wyższy poziom energetyczny są mniej stabilne (niemniej istnieją wyjątki od tej reguły), zatem elektrony wracają na poziom niższy, emitując nadmiar energii jako promieniowanie elektromagnetyczne.

Odrobina historii

Człowiek od zamierzających czasów korzystał z różnych źródeł światła, aby oświetlać swoje domostwa. Do takich źródeł zaliczały się na przykład świece, lampy spalające oleje roślinne albo naftę. Miały jednak kilka wad. Najważniejszą z nich była łatwość powodowania pożaru przy nieostrożnym obchodzeniu się z ogniem. Dość istotny był również efekt nierównomiernego świecenia w miarę wypalania się paliwa.

Pierwsze prototypy tradycyjnych żarówek powstały na przełomie lat trzydziestych i czterdziestych XIX wieku. Początkowo urządzenia te nie miały praktycznego zastosowania, były jednak interesujące jako potencjalne źródło oświetlenia, bezpieczne w użyciu i niezanieczyszczające otoczenia. W ciągu kolejnych trzydziestu lat uzyskano prototyp żarówki nadającej się do praktycznego wykorzystania. Wkrótce potem w USA rozpoczęła się masowa produkcja żarówek, których konstrukcja nie zmieniła się znacząco aż do obecnych czasów.

Ze względu na niską wydajność żarówek (tylko 5% pobranej energii zostaje zamienione w światło) od połowy XX wieku coraz częściej zastępowano je urządzeniami wykonywanymi w innych technologiach. Jednym z takich urządzeń była lampa wyładowcza (świetlówka), która początkowo używana była jedynie w różnego rodzaju instytucjach

(szkoły, biura, fabryki). Dopiero na przełomie XX i XXI wieku jej konstrukcja ewoluowała w kierunku kompaktowego urządzenia, które można zamontować w miejscu tradycyjnej żarówki.

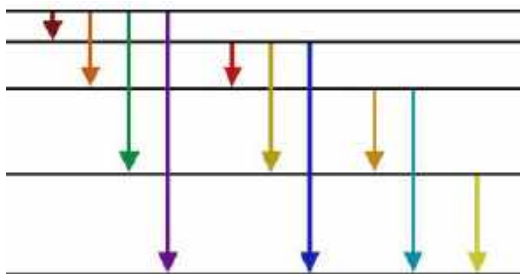
Wadą świetlówki okazało się przede wszystkim nie-naturalne światło, wyraźnie inne od światła słonecznego. Nie bez znaczenia jest również fakt, iż lampy gorszej jakości mogą emitować szkodliwe promieniowanie UV. Urządzenia te zawierają ponadto toksyczną rtęć, dlatego konieczny jest ich recykling.

W ostatnich latach lampy wyładowcze zostały praktycznie całkowicie wyparte przez lampy diodowe LED. W przypadku tych ostatnich urządzeń łatwiej uzyskać światło białe o widmie zbliżonym do widma światła słonecznego.

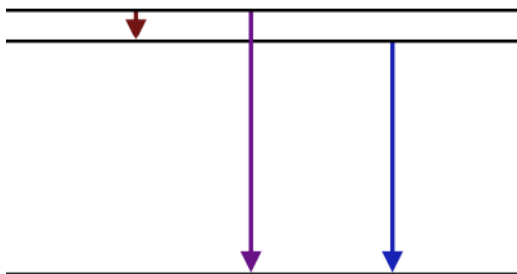
Kilka słów o mechanizmach świecenia

W tradycyjnej żarówce prąd płynący przez żarnik powoduje rozgrzanie go do temperatury około 3000°C. Dzieje się tak dlatego, że rozpędzone nośniki prądu (w tym przypadku elektrony), zderzając się z materiałem, z którego wykonany jest żarnik, oddają mu swoją energię. To z kolei powoduje wzbudzenie elektronów znajdujących się w żarniku na wyższe poziomy energetyczne.

W przypadku ciała stałego poziomy te są bardzo liczne, tworząc gęste pasma energetyczne. Jeśli zostaną one obsadzone przez wzbudzone elektrony, to istnieje ogromna ilość możliwości powrotu na poziom niższy. Powracające elektrony wytracają nadmiar energii poprzez emisję kwantów promieniowania. Energia każdego kwantu, a zatem także widziany przez człowieka kolor światła, zależy wyłącznie od różnicy energii



1. Rozgrzane ciało stałe emituje równocześnie wiele kwantów światła o różnych energiach. Powstaje widmo ciągłe, odbierane przez oko ludzkie jako światło jednobarwne



2. W przypadku pobudzonych do świecenia gazów istnieje zazwyczaj tylko kilka przejść w zakresie światła widzialnego. Powstaje widmo dyskretne składające się z wyraźnie oddzielonych linii, charakterystyczne dla każdego pierwiastka

pomiędzy poziomem, z którego nastąpiło przejście a poziomem docelowym.

Sytuacja ta została w dużym uproszczeniu zilustrowana na **rysunku 1**. Jednoczesna emisja kwantów o różnych energiach prowadzi do powstania widma ciągłego, będącego mieszaniną wszystkich kolorów. Przy odpowiednio wysokiej temperaturze żarnika widmo takie jest przez oko ludzkie odbierane jako światło zbliżone do światła słonecznego.

Nieco prościej wygląda sytuacja w lampach wydławowych, w których do świecenia pobudzane są rozrzedzone gazy. Struktura poziomów energetycznych gazów jest znacznie uboższa niż w przypadku ciał stałych i do tego jest charakterystyczna dla każdego pierwiastka. Można zatem dobrać taki gaz, który będzie emitował kwanty promieniowania pobudzające do świecenia konkretny luminofor. Podobne rozwiązanie stosuje się w lampach diodowych, skonstruowanych w oparciu na półprzewodnikach.

Sprawdź swoją wiedzę

W półprzewodnikach zazwyczaj istnieje możliwość przeskoku elektronu pomiędzy dwoma ściśle określonymi poziomami. W rzeczywistości jeden z tych poziomów jest pasmem przewodzenia, co dla uproszczenia

zostało pominięte w tym artykule. Wyobraź sobie, że na **rysunku 2** środkowy poziom energetyczny przesuwamy bardzo blisko poziomu górnego, tworząc pasmo energetyczne. Jakie światło będzie dominować w widmie diody LED o takiej strukturze poziomów?

Wskaż prawidłową odpowiedź:

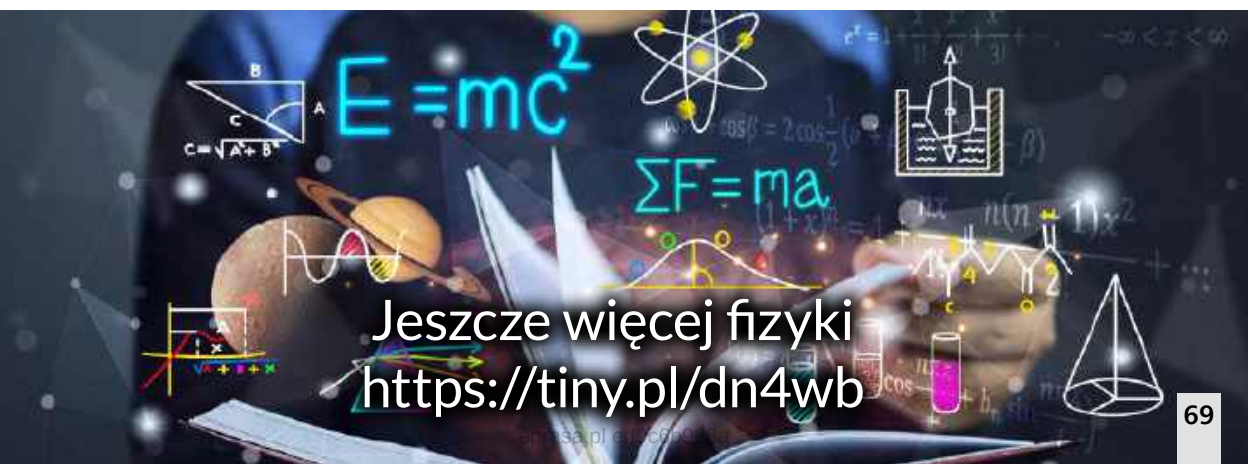
- A. Czerwone lub podczerwone, odpowiadające strzałce z lewej strony rysunku.
- B. Fioletowe lub ultrafioletowe, odpowiadające strzałce środkowej.
- C. Niebieskie, odpowiadające strzałce z prawej strony rysunku.
- D. Nie będzie emitowane żadne światło.

Dla nauczyciela

Niniejszy materiał może zostać wykorzystany na lekcji fizyki w szkole ponadpodstawowej zarówno w zakresie podstawowym, jak i rozszerzonym do realizacji punktów podstawy programowej związanych z powstawaniem widm emisyjnych ciał, w szczególności punktu X.4 (zakres podstawowy) oraz XI.4 (zakres rozszerzony). ■

Joanna Borgensztajn

Odpowiedź do zadania: B



Jeszcze więcej fizyki
<https://tiny.pl/dn4wb>



Cyberbezpieczeństwo

Cieńko określić, kiedy doszło do pierwszego ataku hakera, ale za prekursora w tej dziedzinie uchodzi trzynastoletni David Dennis, który w 1974 roku dokonał pierwszego zarejestrowanego ataku typu DoS (odmowa dostępu do zasobów). Od tego czasu wiele się zmieniło, niestety na niekorzyść użytkowników sieci internetowej, a w związku z tym ludzkość musiała wymyślić sposoby radzenia sobie z zagrożeniem cybernetycznym. Potrzeba rozwinęła naukę, pozwoliła na opracowanie nowych technologii i wciąż kształtuje nowe rozwiązania, których autorami są specjaliści od cyberbezpieczeństwa. Zapraszamy na kierunek, dzięki któremu możemy czuć się dużo bardziej bezpieczni w cyberświecie.

Planując przyszłość zawodową w cyberbezpieczeństwie, można skorzystać z wielu ścieżek rozwoju, z których każda ma swoje zalety i specyfikę. Polskie uczelnie oferujące specjalizacje z zakresu cyberbezpieczeństwa, takie jak Uniwersytet Warszawski czy Politechnika Warszawska, zapewniają solidne programy studiów, które kładą nacisk na praktyczną wiedzę z obszaru bezpieczeństwa systemów komputerowych i sieci. Akademia Górniczo-Hutnicza w Krakowie oraz Politechnika Wrocławska mają w swojej ofercie kierunki informatyczne, obejmujące istotne zagadnienia z zakresu cyberbezpieczeństwa. Wojskowa Akademia Techniczna w Warszawie oferuje specjalizacje bezpośrednio związane z ochroną informacji, co jest okazją dla tych, którzy szukają specjalistycznej ścieżki kariery. Studia na renomowanych uczelniach dają możliwość zdobycia szerokiej wiedzy teoretycznej i praktycznej, co jest kluczowe w branży coraz bardziej narażonej na cyberzagrożenia. Jednak studia nie są jedyną drogą do zdobycia wiedzy i kwalifikacji. Szkolenia online na platformach takich jak Coursera,

edX, Udemy czy Pluralsight oferują szeroki wybór kursów, które można realizować zdalnie, elastycznie dostosowując je do własnych potrzeb i tempa nauki. Uczestnictwo w tego typu szkoleniach może być znaczącym uzupełnieniem wiedzy zdobytej na uczelni oraz umożliwia uzyskanie cennych certyfikatów, takich jak CISSP, CEH czy CompTIA Security+. Wybór pomiędzy studiami a szkoleniami zależy od indywidualnych preferencji i celów zawodowych. Studia oferują solidne przygotowanie teoretyczne i praktyczne, podczas gdy szkolenia online mogą zapewnić dodatkową specjalizację i certyfikaty, zwiększając atrakcyjność na rynku pracy.

Decydując się na studia, należy wziąć pod uwagę, że jest to kierunek bardzo popularny, a w związku z tym można spodziewać się licznej konkurencji. W osiągnięciu sukcesu i uzyskaniu upragnionego indeksu pomóc może dobre zdanie egzaminu maturalnego. Uczelnie często wymagają poziomu rozszerzonego z matematyki i informatyki. Choć może wydawać się to mniej oczywiste, fizyka na tym kierunku jest ważna ze względu na jej

zastosowania w zrozumieniu podstawowych zasad działania technologii komputerowych, a także w niektórych specjalistycznych obszarach, dlatego zdawanie jej na maturze jest jak najbardziej wskazane. I oczywiście nie można zapominać o znajomości języka obcego, szczególnie angielskiego. Jest on kluczowy w środowisku międzynarodowym branży IT, gdzie wiele zasobów, dokumentacji i rozwiązań jest dostępnych właśnie w tym języku. Planując swoją karierę w cyberbezpieczeństwie, należy przygotować się na okres studiów trwający przynajmniej 5 lat. Jednakże, aby stać się pełnoprawnym specjalistą w tej dziedzinie, może być konieczne zdobycie dodatkowych certyfikatów branżowych oraz odbycie staży czy praktyk, co może wydłużyć czas edukacji o kolejne 2 lata. Cyberbezpieczeństwo to dziedzina, która stale się rozwija i ewoluuje w odpowiedzi na nowe zagrożenia. Dlatego ważne jest ciągle doskonalenie swoich umiejętności, uczestnictwo w konferencjach branżowych oraz śledzenie najnowszych trendów i technologii.

W trakcie nauki studenci mają możliwość profilowania kierunku rozwoju swojej wiedzy. Specjalizowanie się w określonym obszarze pozwala zdobyć bardziej zaawansowane umiejętności w danej dziedzinie. Oto kilka przykładów specjalizacji, które mogą być dostępne na różnych uczelniach: bezpieczeństwo sieci komputerowych, a więc ochrona i zabezpieczanie sieci, wykrywanie i neutralizacja zagrożeń oraz zarządzanie, kryptografia i ochrona danych, w ramach których omawiane są techniki szyfrowania, protokoły kryptograficzne, bezpieczeństwo danych, ochrona prywatności, analiza zagrożeń i incident response, które obejmują wykrywanie, analizę i reagowanie na incydenty bezpieczeństwa, zarządzanie ryzykiem, forensykę cyfrową, bezpieczeństwo aplikacji oraz zarządzanie bezpieczeństwem informacji.

Wybór odpowiedniej uczelni i dostanie się na nią jest już połową sukcesu. Kolejne pół zależeć będzie od systematyczności w nauce oraz indywidualnych preferencji. Studiowanie na tym kierunku nie jest łatwym zadaniem. Oprócz „królowej nauk”, studenci muszą zmierzyć się z kilkoma innymi wymagającymi przedmiotami, z których jednym z najbardziej zaawansowanych jest kryptografia. Jest to dziedzina nauki, która nie tylko wymaga solidnych podstaw matematycznych, ale także umiejętności analitycznych i logicznego myślenia. Do głównych wyzwań, jakie studenci muszą pokonać, należą zagadnienia matematyczne. Przedmiot ten opiera się na zaawansowanej matematyce, szczególnie w zakresie teorii liczb oraz algebry. Studenci muszą zrozumieć i stosować skomplikowane koncepty matematyczne, takie jak

teoria grup, krzywe eliptyczne czy funkcje skrótów. Dobrze rozwinięte umiejętności matematyczne są kluczowe do efektywnego zrozumienia i tworzenia algorytmów kryptograficznych. Kolejnym wyzwaniem jest praktyczne wdrożenie oraz analiza algorytmów. Studenci muszą być w stanie nie tylko zaimplementować różne metody szyfrowania i podpisów cyfrowych, ale również przeprowadzać ich analizę pod kątem bezpieczeństwa i efektywności. Zrozumienie, jak działają algorytmy na poziomie kodu oraz umiejętność ich oceny pod względem odporności na ataki, jest kluczowe dla przyszłych specjalistów ds. cyberbezpieczeństwa. Podsumowując, w trakcie studiów znajdzie się czas na życie studenckie, ale dopiero po opanowaniu matematyki i zdanii kryptografii.

Nie ulega wątpliwości, że rosnące zagrożenia cybernetyczne są głównym czynnikiem napędzającym popyt na specjalistów w tej dziedzinie. Ataki takie jak ransomware czy phishing stają się coraz bardziej złożone i wymagają zaawansowanej wiedzy oraz umiejętności, aby je skutecznie neutralizować. Firmy i instytucje na całym świecie poszukują wykwalifikowanych profesjonalistów, którzy mogą zapewnić bezpieczeństwo ich systemów informatycznych i danych. Jednym z kluczowych atutów kariery w cyberbezpieczeństwie są atrakcyjne zarobki. Niedobór specjalistów w tej dziedzinie prowadzi do tego, że nawet początkujący mogą liczyć na konkurencyjne wynagrodzenia. Ponadto specjaliści mogą pracować w różnych sektorach gospodarki, od IT i finansów po zdrowie, administrację publiczną czy przemysł, co dodatkowo zwiększa ich możliwości zawodowe. Aby zmaksymalizować swoje szanse na rynku pracy po ukończeniu studiów, warto inwestować w dodatkowe kwalifikacje i doświadczenie praktyczne. Udział w stażach, praktykach oraz konkursach związanych z bezpieczeństwem cybernetycznym nie tylko zwiększa wiedzę, ale także zapewnia cenne referencje. Certyfikaty od renomowanych organizacji są również cenione przez pracodawców i mogą otworzyć drzwi do bardziej zaawansowanych stanowisk.

Studia na kierunku cyberbezpieczeństwo nie tylko otwierają drzwi do dynamicznie rozwijającej się branży, ale również zapewniają stabilność zawodową i satysfakcjonujące wynagrodzenia. Dzięki ciągłemu rozwojowi technologicznemu i rosnącemu zagrożeniu cybernetycznemu zapotrzebowanie na ekspertów ds. bezpieczeństwa będzie się utrzymywać na wysokim poziomie. Dla ambitnych absolwentów szkół średnich, którzy podejmą wyzwanie studiów w tej dziedzinie, przyszłość zawodowa wydaje się niezwykle obiecująca. ■

Michał Pacholski

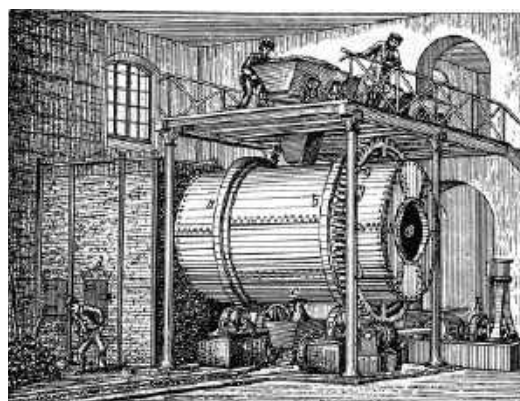
Błędne ścieżki prowadzą do celu (2)

Czy „spacer z flogistonem” był udany? Przechadzki sprzyjają rozmyślaniom, o czym wiedzieli już starożytni. Choć ze współczesnego punktu widzenia teoria flogistonowa jest błędna – przecież spalanie to energiczne łączenie się z tlenem – to jednak spalana substancja coś traci (i nie mówimy tu o produktach gazowych). Domyślasz się, co to jest?

Sukcesy...

Teoria flogistonowa powstała na bazie obserwacji spalania drewna, węgla drzewnego i kamiennego. Pozostały popiół waży mniej niż użyte paliwo, wywnioskowano więc, że z substancji ulatuje hipotetyczny fluid ognia, czyli flogiston (produktów gazowych nie brano pod uwagę). Teoria, rozszerzona także na substancje nieorganiczne, a zwłaszcza metale, była pierwszym w historii chemii nowoczesnym wytłumaczeniem szerokiego zakresu obserwowanych przemian. Dokonano także wyraźnego rozróżnienia pomiędzy chemią nieorganiczną i organiczną. W tej pierwszej można było poddać substancję szeregowi reakcji i odzyskać ją w wyjściowej postaci w takiej samej ilości, np. spalić metal, działaniem kwasu zmienić powstałą ziemię (tlenek) w sól, z soli wytrącić wodorotlenek, przez prażenie wodorotlenku uzyskać tlenek, który pod wpływem ogrzewania z węglem znowu przekształcał się w metal. Z substancjami organicznymi taki ciąg przemian był niewykonalny, zresztą dobrze wiesz, że np. zwęglonego ziemniaka nie da się przywrócić do stanu początkowego. Za przekształcenia substancji organicznych miała odpowiadać tajemnicza *vis vitalis*, czyli „siła życiowa” obecna tylko w świecie przyrody ożywionej (jej kariera potrwała do roku 1828).

Na podstawie teorii flogistonowej opracowano technologię dwóch ważnych procesów. Pierwszym z nich była synteza kwasu siarkowego, do dziś zaliczanego do najważniejszych produktów chemii. Do tej pory otrzymywano jego niewielkie ilości w wyniku prażenia naturalnie występujących siarczaków. Gdy kwas siarkowy stał się szeroko dostępny, można było za jego pomocą produkować inne ważne odczynniki (np. kwas azotowy i solny), wiele użytecznych substancji oraz oczyszczać metale. Kwas siarkowy wykorzystano także do drugiego ze wspomnianych procesów – produkcji sody, czyli węgla



1. Piec obrotowy do suszenia sody otrzymywanej w metodzie Leblanca (ilustracja z książki Hermanna Osta „Lehrbuch der Technischen Chemie”, 1890)

sodu. Tę bardzo ważną substancję (surowiec do produkcji szkła i mydła) uzyskiwano z nielicznych złóż naturalnych oraz ze spalania wodorostów. Do metody opracowanej przez Nicolasa Leblanca potrzebne były sól kamienna i kwas siarkowy, a w rezultacie połączenia w jednej wytwórni obu procesów powstawały pierwsze kombinaty chemiczne (1). Współcześnie stosuje się już inne technologie, ale nadal wielkość produkcji kwasu siarkowego („krwi przemysłu chemicznego”) świadczy o zaawansowaniu rozwoju przemysłowego kraju, a sody używanej do wytwarzania mydła – o jego poziomie cywilizacyjnym.

...i problemy

Z alchemii wyodrębniła się farmacja, której adepci zajmowali się wytwarzaniem medykamentów, do czego niezbędna była waga. Powszechne zastosowanie tego przyrządu postawiło przed teorią flogistonową problem nie do rozwiązania. O ile popiół

rzeczywiście ważył mniej niż wzięte do spalania drewno, o tyle ziemie metali były cięższe niż metale, z których powstały. Jak zatem mogło dojść do utraty flogistonu, skoro zwiększyła się masa substancji?

Teorii nie odrzuca się od razu, ale modyfikuje jej założenia, tak też postąpiono z flogistonem. Ponadto nie we wszystkich przypadkach obserwowano wzrost masy utlenianych metali, fakt ten spowodowany był jednak użyciem do ich podgrzewania soczewek skupiających promienie słoneczne. Sposób ten pozwalał uzyskiwać tak wysokie temperatury, że część powstających tlenków odparowywała i rzeczywiście stwierdzano ubytek masy. W XVIII wieku, po przełomowych pracach Newtona, nikt nie wierzył w ujemną masę flogistonu, ale zakładano, że jest on bardzo lekki, wszak ulatywał ze spalanej substancji. Ponieważ miał wypełniać wolne przestrzenie w każdym ciele palnym, zatem, na mocy prawa wyporu Archimidesa, po spaleniu ciało było cięższe niż przed tym procesem. Oczywiście następowało tu pomieszanie pojęć masy i ciężaru (a także gęstości), które – choć w potocznym języku stosowane wymiennie – nie są tym samym. Dlaczego zatem spalane substancje organiczne zachowywały się inaczej i faktycznie ich ubywało? Teoria nie potrafiła udzielić racjonalnej odpowiedzi i coraz więcej chemików podchodziło do niej nieufnie.

Jak schwycić ducha?

Do początków wieku XVIII chemia była „dwuwymiarowa”, czyli w bilansie masy uwzględniała tylko ciała stałe i ciecze. Oczywiście gazy znano od dawna, wielokrotnie stwierdzano wydzielanie „ducha” lub „powietrza” podczas przemian, ale nie zwracano uwagi na te ulotne produkty. Chemia gazów miała dopiero nadejść, a jej początkiem stało się wynalezienie **wanny pneumatycznej** przez angielskiego botanika **Stephena Halesa** (uprzednio zbierano gazy do zwierzęcych pęcherzy lub worków skórzanых – w domu możesz użyć baloników – ale nie były to wygodne sposoby). Przykładem takiego urządzenia jest probówka wypełniona wodą, odwrócona i zanurzona w płaskim krystalizatorze również wypełnionym wodą, do środka probówki wprowadzony jest wężyk z kolby, w której wydziela się gaz. Zestaw ten jest używany w szkole, a i w domowym laboratorium okazuje się przydatny (2).

Kilka praktycznych uwag dotyczących używania wanny pneumatycznej. Probówkę napełnij wodą, zakryj palcem i odwróconą umieść w większym naczyniu z wodą. Możesz również na wierzch wypełnionej „z czubkiem” probówki położyć kawałek bibuły i odwrócić ją – woda się nie wyleje (dla pewności całą operację wykonaj w zlewku). Do probówki wkładasz



2. Stephen Hales (1677–1761) i jego wanna pneumatyczna

wężyk i czekasz, aż wydzielający się gaz wypchnie z niej wodę, potem probówkę zakorkuj pod wodą i przenieś w inne miejsce. Wężyk koniecznie wyjmij z wody, zanim zakończysz ogrzewanie naczynia reakcyjnego, inaczej podciśnienie wciągnie wodę do środka i może ono pęknąć.

W opisywany sposób zbierasz gazy słabo rozpuszczalne w wodzie, np. wodór czy tlen. Te, które rozpuszczają się lepiej, zbierasz nad gorącą wodą (rozpuszczalność gazów maleje ze wzrostem temperatury) lub wręcz do naczynia umieszczonego w powietrzu (np. amoniak).

Powietrze to nie pierwiastek

Alchemicy wielokrotnie stwierdzali wydzielanie się gazów, zwykle jednak nie interesowali się tymi „duchami”. W połowie XVII wieku flamandzki lekarz i alchemik **Johann van Helmont** opisał szczegółowo jeden z nich, który powstawał w wyniku spalania drewna i oddychania. Nadał mu nazwę **gazu leśnego** (łac. *gas silvestris*, Helmont jest również autorem pojęcia „gaz”), jak się domyślasz, był to dwutlenek węgla. Po 100 latach szkołki chemik **Joseph Black** ponownie odkrył CO₂ podczas badań nad substancjami alkalicznymi. Gaz ten wydzielał się podczas działania kwasów na węglany, skały wapienne, muszle i skorupki jaj. Black opracował też znaną do dziś procedurę wiązania dwutlenku węgla za pomocą wody wapiennej (roztwór wodorotlenku wapnia pochłania CO₂ z wydzielaniem słabo rozpuszczalnego węglanu wapnia, co powoduje zmętnienie roztworu), skąd można go odzyskać działaniem kwasów. Z powodu łatwego wiązania gaz ten nazwano **związanym powietrzem** (ang. *fixed air*), a chemicy zaczęli uważać, że i podczas spalania metali również wiązany jest jakiś rodzaj powietrza (3).

Roztwarzanie metali w kwasach z wydzielaniem wodoru to także reakcja znana już alchemikom. Angielski arystokrata, a przy tym chemik i fizyk (jako pierwszy wyznaczył wartość stałej grawitacji), **Henry Cavendish**



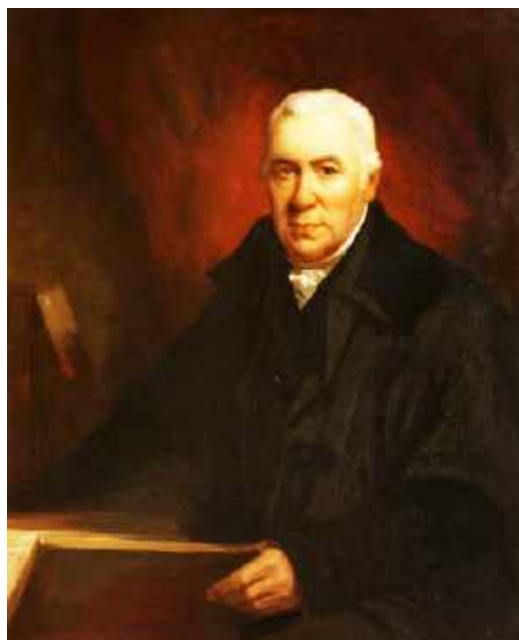
3. Z lewej: Joseph Black (1728–99), jego prace rozpoczęły badania nad składem powietrza, odkrywca dwutlenku węgla i magnezu. Z prawej: pod wpływem kwasu ze skorupki jaj wydzielą się „związane powietrze” (dwutlenek węgla) – to charakterystyczna reakcja, jakiej ulegają węglany

określił ten gaz jako odrębną substancję w roku 1766. Miał to być ni mniej, ni więcej tylko sam ...flogiston! Cavendish rozumował następująco: metale tracą flogiston w wyniku spalania, ale także rozтворяjąc się w kwasach (sole są niepalne), zatem to, co ulatuje po wrzuceniu np. żelaza do kwasu siarkowego, musi być flogistonem. Dodatkowo wodór spalał się wybuchowo (nazwano go **powietrzem palnym**, ang. *inflammable air*), co potwierdzało, że faktycznie był to od dawna poszukiwany fluid ognia (4).

Daniel Rutherford, uczeń Josepha Blacka, chciał się dowiedzieć, ile flogistonu może zostać pochłonięte przez powietrze atmosferyczne. W tym celu



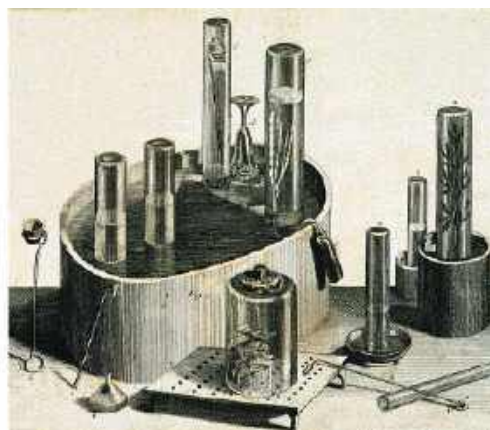
4. Henry Cavendish (1731–1810), odkrywca wodoru



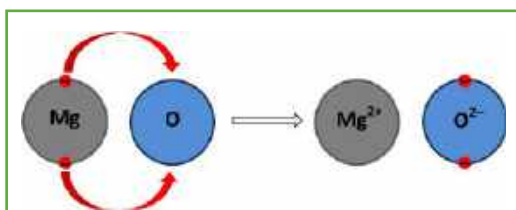
5. Daniel Rutherford (1749–1819), odkrywca azotu

spalał w nim węgiel, a powstający gaz był absorbowany w roztworze zasadowym. Gdy w próbce powietrza już nic nie chciało się palić, Rutherford zbadał pozostałość po reakcji. Okazało się, że gaz ten ma mniejszą gęstość od zwykłego powietrza atmosferycznego i nie podtrzymuje palenia ani nie nadaje się do oddychania. W ten sposób w roku 1772 pojawiło się **powietrze flogistonowane** (czyli nasycone flogistonem), jak nazwał je flogistyk Rutherford, obecnie znany jest jako azot (5).

Istnienia w powietrzu czynnika odpowiedzialnego za spalanie i oddychanie domyślano się już w XVII wieku, ale oficjalnie pojawił się on w świecie nauki w roku 1774. Wtedy to angielski chemik i duchowny **Joseph Priestley** dokładnie zbadał powietrze, które wydzielało się z ogrzewanego tlenku rtęci. Gaz ten znakomicie podtrzymywał palenie oraz ułatwiał oddychanie, zatem – zgodnie z poglądami flogistyków, do których należał i Priestley – był on całkowicie pozbawiony flogistonu, skoro tak chętnie go pochłaniał. Nie powinna być zatem zaskoczeniem nadana mu nazwa – **powietrze deflogistonowane**. Dwa lata wcześniej tego samego odkrycia dokonał **Carl Wilhelm Scheele**, skromny szwedzki aptekarz, a przy tym doskonały chemik (zidentyfikował aż 6 nowych pierwiastków), który jednak spóźnił się z publikacją swojej pracy. Scheele także był flogistyką i tlen (bo o tym gazie mowa) dla niego również stanowił powietrze pozbawione flogistonu. W każdym razie powietrze atmosferyczne



6. Odkrywcy tlenu Carl Wilhelm Scheele (1742–1786) i Joseph Priestley (1733–1804) oraz ówczesna aparatura stosowana do doświadczeń z gazami (ilustracja z książki J. Priestleya „Experiments and Observations on Different Kinds of Air”, 1775)



7. Podczas utleniania elektrony pełnią funkcję „flogistonu”

Co traci spalana substancja?

Spalanie (gwałtowne łączenie z tlenem z towarzyszącymi efektami cieplnymi i świetlnymi) to szczególny rodzaj szerszego zjawiska, czyli utleniania. Dodatkowo zjawisko to rozumiane jest nie tylko jako reakcja z tlenem, ale zwiększenie stopnia utlenienia polegające na utracie elektronów (pamiętaj, że elektrony mają ładunek ujemny). Oczywiście elektrony nie znikają, lecz są przytaczane przez substancję utleniającą, np. tlen (on z kolei ulega redukcji, czyli zmniejszeniu stopnia utlenienia). Jeżeli tak interpretować spalanie, elektrony stanowią swoisty „flogiston” (7).

okazało się mieszaniną co najmniej dwóch składników, z których jeden był całkowicie nasycony flogistonem (azot), a drugi zupełnie go pozbawiony (tlen). Nastąpił koniec arystotelesowskiego żywiołu (6).

W tym samym roku Priestley spotkał w Paryżu młodego francuskiego chemika i przedstawił mu wyniki swoich prac nad deflogistonowanym powietrzem. Tym uczonym był **Antoine Lavoisier**, a co wynikło z tego spotkania, dowiesz się w następnym odcinku. ■

Krzysztof Orliński



więcej chemii na stronie
<https://tiny.pl/dtp7>

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

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



Twierdzenie Halla o małżeństwach

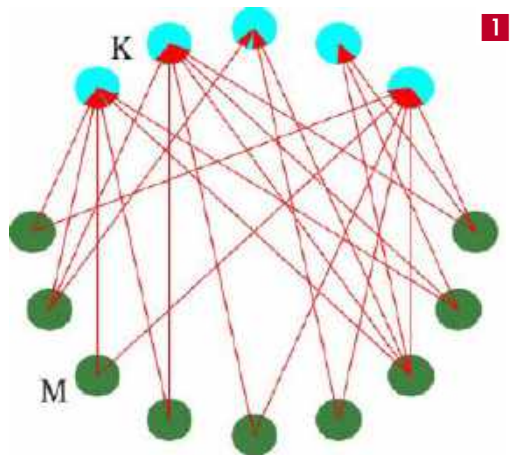
Wakacje już zbliżają się do końca; ale jeszcze trochę potrwać. Temat kącika matematycznego jest tylko pozornie wakacyjny i mało poważny: twierdzenie Halla o małżeństwach. Bo czy matematyka może się przydać w biurze matrymonialnym? Hm, ciekawe, jak by sobie poradziły z problemem wyboru programu oparte na sztucznej inteligencji.

Matematycy lubią nadawać twierdzeniom żartobliwe formy. Mamy na przykład „twierdzenie o kanapkach”: każdą kanapkę z masłem i szynką można przekroić jednym cięciem tak, by przepołowić bułkę, masło i szynkę. W analizie matematycznej jest „twierdzenie o dwóch policjantach”: jeżeli idziesz trzymany z dwóch stron przez dwóch policjantów, to trafisz do komisariatu. Chodzi tu o prosty fakt, znany każdemu, kto uczył się podstaw analizy matematycznej: mamy trzy ciągi a, b, c takie, że $a \leq b \leq c$. Jeżeli dwa skrajne dążą do tej samej granicy g , to trzeci też ma granicę g . W analizie funkcjonalnej (polskiej specjalności matematycznej lat międzywojennych) mówi się też o twierdzeniu o ziemniakach: skończony zbiór ziemniaków można zapakować do jednego worka – ale tłumaczenie, o co „naprawdę” chodzi, wykracza poza ramy tego kącika.

Trzeba szczerze powiedzieć, że takie żartobliwe podejście spotyka się często z niezrozumieniem i niekiedy wręcz niechęcią do matematyki. „Dlaczego moje pieniądze, prostego podatnika, mają iść na takie bzdury”? Jediną odpowiedzią jest chyba odwołanie się do poczucia humoru ludzi – oraz przekonanie ich, że to właśnie nie tylko żart, że to jednak czemuś służy, choćby refleksji intelektualnej. A co jest poważne, a co nie, jest bardzo często umowne. Skończyła się olimpiada. Czy

rzucanie kijem albo kopanie skórzanej (no, od wielu lat już nie skórzanej) piłki jest czymś poważnym?

Po tym przydługim wstępie przejdę do... spraw męsko-damskich w majestacie matematyki. Otóż na Marsa leci wieloosobowa załoga. Podróż jest długa, trzeba wytrzymać. W załodze są kobiety i mężczyźni – tych jest trochę więcej. Po pewnym czasie zaczynają się tworzyć pary. W regulaminie nie jest to zakazane. To przecież pozwala lepiej znieść kosmiczną nudę i frustrację.





W załodze, jak to w każdej grupie ludzi, jedni lubią się bardziej, inni mniej. I oto wykres (1). Mamy dwie grupy, K i M – domyślności Czytelników pozostawiam, skąd taki wybór oznaczeń. Każda osoba z grona K rozważa, z kim ewentualnie mogłaby się związać. Pokazują to strzałki. W wersji „małżeńskiej”: za tego faceta ewentualnie mogłabym się wydać. Inni niech spadają! Co prawda w stanie nieważkości (lećmy przecież na Marsa) to jest niewykonalne.

I teraz wreszcie zagadnienie matematyczne: czy takie przyporządkowanie jest możliwe? Czy każda osoba z grona K znajdzie takiego M, który jest akceptowalny? Oczywiście nie jest to symetryczna sytuacja: jeżeli $M > K$, to nie każdy M znajdzie swoją K. Cóż, życie (2).

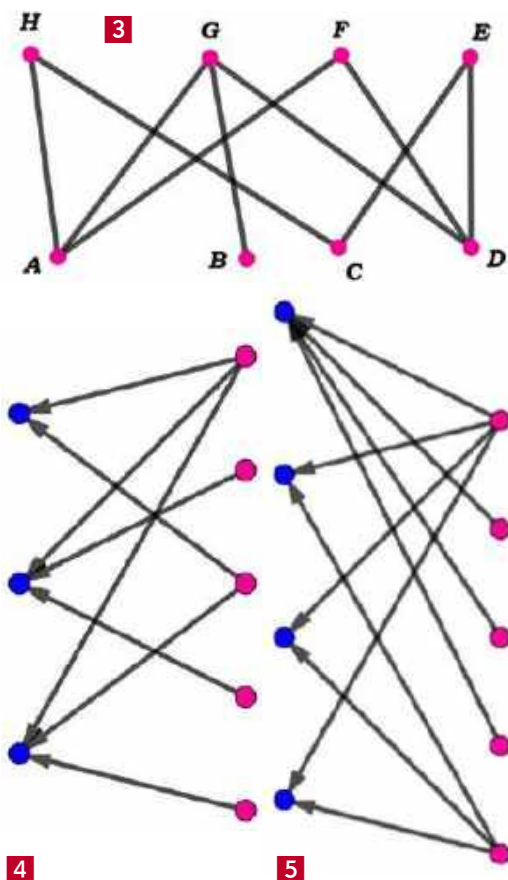
O to właśnie chodzi w twierdzeniu, udowodnionym w 1935 roku przez Philipa Halla: takie przyporządkowanie jest możliwe, jeżeli tylko spełniony jest dość oczywisty warunek (zwany teraz warunkiem Halla). Choć oczywisty, jest trochę skomplikowany. Łatwiej go sformułować właśnie w terminach matrymonialnych. Zamiast długiego „tego mężczyznę mogłabym poślubić”, będę pisał „znam go”.

Po pierwsze: każda kobieta zna co najmniej jednego mężczyznę.

Po drugie: każde dwie kobiety znają co najmniej dwóch mężczyzn.

Po trzecie: każde trzy kobiety znają (łącznie) co najmniej trzech mężczyzn.

Po czwarte: każde cztery kobiety znają (łącznie) co najmniej czterech mężczyzn.



Po piąte... i tak dalej.

Jest dość jasne, że jeżeli którykolwiek z tych warunków nie jest spełniony, to nie każda pani znajdzie pana, który jej choć trochę odpowiada. Twierdzenie, o którym opowiadam, mówi, że jest także przeciwnie: jeżeli warunek Halla jest spełniony, to odpowiednie przyporządkowanie istnieje. W locie na Marsa każda astronautka znajdzie swojego astronautę. Gdyby liczba K była równa liczbie M (czyli parytet byłby 50–50), mogłoby być i odwrotnie.

Spójrzmy na następne rysunki. Układ na **rysunku 3** nazywa się w matematyce grafem dwudzielnym: są w nim dwa odrębne zbiory, a połączenia są tylko między jednym a drugim – nie ma natomiast połączeń w obrębie jednego zbioru. Na **rysunku 4** warunek Halla jest spełniony, na **rysunku 5** nie: trzy pierwsze niebieskie kropki są połączone tylko z dwiema czerwonymi.

Dowód twierdzenia Halla jest ciekawy, ale może tylko dla matematyków. Natomiast ważniejsze jest co innego, co zilustruję dygresją. Jest dość oczywiste, że w Warszawie (Krakowie, Poznaniu, Gdańsku itd.) znajdują się osoby, które mają tyle samo włosów na głowie

(i nie jest to liczba zero, bo łysych panów znajdzie się wielu). To bardzo prosto zrozumieć: w każdym z tych miast jest więcej niż powiedzmy, pół miliona ludzi. Mamy mniej włosów na głowie, zatem muszą być osoby równe pod względem liczby włosów. To typowy dowód matematyczny: wykazaliśmy, że takie osoby są. Informatyk powie, że ta informacja nic mu nie daje. Potrzebny jest algorytm, który takie osoby znajdzie.

Podobnie jest z twierdzeniem Halla. Wiemy już, że taki *matching* istnieje. Jak go znaleźć? Algorytm nie jest trudny – w odróżnieniu od hipotetycznego programu na szukanie osób tak samo owłosionych. Nie podam szczegółów, ale idea jest taka: zaczynamy przyporządkowywać, kolejne K do kolejnych M . Jeżeli w którymś momencie się zatnie, cofamy się o krok i zmieniamy. Jeżeli wyczerpiemy wszystkie możliwości i nie posuniemy się naprzód, cofamy się o następny krok i szukamy, szukamy. Twierdzenie Halla gwarantuje, że za którymś razem musi się udać. Jest to algorytm powolny, ale jak najbardziej wykonalny.

Od mniej więcej dwóch lat panoszy się sztuczna inteligencja. Ja nazywam ją sobie Apolonią Inteligentną. Niektórzy wnioskuje, że zabije ona naszą cywilizację, a potem nas wszystkich. Zobaczmy, ale podobnie było z wieloma wynalazkami. A ja poprosiłem Apolonię o napisanie programu na szukanie przyporządkowania w twierdzeniu Halla. Uff, najtrudniej było „jej” wytłumaczyć, jaki graf ma wziąć. Trwało to długo. Nie jestem zręcznym programistą, średnio zaawansowany ktoś napisałby szybciej.

Jak wspomniałem na wstępie, wątpię, żeby to wszystko przydało się w biurze matrymonialnym. Siła twierdzenie Halla – jak większości zagadnień matematycznych – tkwi w przykładach i zastosowaniach. Jest taka maksyma, że lot w abstrakcję musi zaczynać się i kończyć na Ziemi. Astronauci powinni też powrócić z Marsa, choć podobno zgłaszają się już ochotnicy do lotu w jedną stronę. Ostatecznie w dawnych czasach wielu ludzi emigrowało do Ameryki ze świadomością, że zostaną tam już na zawsze.

Wróćmy do diagramu z rysunku 1. Odczytamy go inaczej. W pewnej dużej firmie są prace do wykonania: $K_1, K_2, K_3, \dots$. Zbiór prac oznaczmy przez K . Zbiór pracowników niech nazywa się M . Nie każdy pracownik potrafi wykonać każdą pracę, ale do każdej nadaje się kilka osób. Jeżeli spełniony jest warunek Halla, to da się do każdej pracy przydzielić fachowca. Więcej: algorytm (który omówiłem w zarysie) powie, kogo przydzielić do jakiej pracy. Jeżeli firma jest naprawdę duża, to być może ręcznie nie da się tego zrobić. Podobnie jest z układaniem planu lekcji

w dużej szkole: do każdego przedmiotu jest kilkoro nauczycieli. Za tak zwanych „moich czasów” oddelegowana do tej pracy wprawna nauczycielka poświęcała na to kilka dni, zużywając kilka ołówków i gumek.

Inne zastosowanie twierdzenie Halla jest może mniej przydatne, chociaż matematyka w sporcie to oddzielny temat i nie chodzi tylko o liczenie punktów. Często po meczach siatkarskich, hokejowych itp. wybierany jest najlepszy zawodnik na danej pozycji. Jak zauważyłem, jest to trochę sztuczne; chodzi o dodatkowe emocje i może nagrody. Oto więc.

Zadanie. W lidze siatkarskiej bierze udział $2n$ zespołów. Rozgrywki są systemem „każdy z każdym”, po jednym meczu (bez rewanżu). Nie ma remisów. Po każdej kolejce jednemu zespołowi, który wygrał swój mecz w tej kolejce, wręczony zostaje pamiątkowy medal. Wykaż, że można tak to zrobić, żeby w każdej kolejce medal dostał inny zespół.

Rozwiązanie zrozumiemy na przykładzie $n=3$ (sześć zespołów), niech to będą zespoły a, b, c, d, e, f . Mamy pięć kolejek, niech $K=\{1,2,3,4,5\}$, $M=\{a,b,c,d,e,f\}$. Utwórzmy graf: do każdego dochodzą trzy strzałki do trzech elementów zbioru M . Po chwili zastanowienia się dochodzimy do wniosku, że warunek Halla jest spełniony, bo:

W każdej kolejce jest zwycięzca (a nawet trzech).

W każdych dwóch kolejkach jest dwóch zwycięzców (a nawet trzech).

W każdych trzech kolejkach jest na pewno trzech zwycięzców.

W każdych czterech kolejkach jest na pewno co najmniej czterech zwycięzców. Meczów jest 12. Przypuśćmy, że jest tylko trzech zwycięzców. Muszą wygrywać wszystkie mecze, bo $3 \cdot 4 = 12$. Ale w którejś kolejce muszą grać ze sobą, nie mogą więc wygrywać wszystkich.

W pięciu kolejkach jest na pewno co najmniej pięciu zwycięzców, bo nie może być dwóch drużyn, które przegrały wszystkie mecze, a zatem pomysł wręczania medalu za każdym razem innemu zespołowi jest wykonalny. Diagram na **rysunku 6** pokazuje rozgrywki z udziałem klubów z Łodzi, Sierpca, Łap, Żywca, Pruszkowa i Szczecina. Nazwy oczywiście są wymyślone przeze mnie. Kropki oznaczają zwycięzców, a czerwone kropki te zespoły, które zostały wybrane jako „drużyny kolejki”. Zauważmy, że choć Ołówek Pruszków wygrał wszystkie mecze, ani razu nie został wyróżniony. Oczywiście tu nie stosowałem algorytmu, tylko wybrałem „na oko”, a co by zrobił algorytm, nie sprawdziłem.

Oto inne zastosowanie. Stworzę smutną fabułę. Pewien obszar leśny podzielono na 100 działek

Kto z kim?

1 NG, MO, ZP

2 MG, NZ, PO

3 ZG, MP, NO

4 PG, OZ, MN

5 OG, NP, MZ

budowlanych. Wycięto drzewa, drewno sprzedano do Finlandii, zabetonowano podjazdy. Każda działka ma tyle samo metrów kwadratowych. Działki sprzedano patodeveloperowi, a niezależnie od tego strażacy – przestrzegając stosownych przepisów – podzielili na mapie teren na 100 obszarów o tej samej powierzchni, ale zupełnie inaczej. Wykaż, że można wykonać 100 przyłączeń wody, żeby na każdej działce budowlanej i na każdej strażackiej znalazło się przyłącze.

Rozwiązanie. Niech K to będą działki budowlane, M – strażackie. Rysujemy graf od K do M i od M do K, łącząc te działki, które chociaż częściowo się pokrywają. Warunek Halla jest spełniony, można kopać.

Twierdzenia Halla ma efektowne zastosowanie w... pewnej sztuczce karcianej. Jest oto dwóch magików. Jeden z nich podaje komuś z publiczności talię 124 różnych kart i prosi o wybranie pięciu. Następnie pokazuje (publicznie!) drugiemu kolejno tylko cztery. Ten zgaduje piątą!

Wydaje się to niemożliwe – kart jest za dużo, żeby odtworzyć jedną z brakujących. Ale od czego matematyka? Zauważmy po pierwsze, że pierwszy sztukmistrz decyduje, którą kartę schować. Ile jest możliwych wyborów pięciu kart ze 124? Czytelnicy, którzy pamiętają coś ze szkoły, wiedzą, że jest to

$$124 \cdot 123 \cdot 122 \cdot 121 \cdot 120$$

$$1 \cdot 2 \cdot 3 \cdot 4 \cdot 5$$

Nie musimy tego obliczać. Zauważmy tylko, że mianownik jest równy 120 i zostaje iloczyn $124 \cdot 123 \cdot 122 \cdot 121$. Teraz obliczmy, ile jest ciągów czterech kart? Pierwszą

Kolejka 1 2 3 4 5

Napraw Łódź	•	•			
Młot Sierpc		•		•	
Zabierz Lapy	•				•
Piwo Żywiec			•	•	•
Ołówek Pruszków	•	•	•	•	•
Grzebień Szczecin			•		

kartę ciągniemy na 124 sposoby, drugą na 123, trzecią na 122, czwartą na 121. Mamy więc tyle samo możliwości: $124 \cdot 123 \cdot 122 \cdot 121$. Jest zatem szansa, aby każdy zbiór zakodować innym ciągiem.

Niech zbiorami będą M (mężczyźni w twierdzeniu Halla), ciągami K (kobiety). Warunek Halla jest spełniony, więc można ich pożenić. Innymi słowy, można każdy zbiór pięciu kart jednoznacznie zakodować ciągiem czterech z nich. Osobną sprawą jest algorytm, jak to zrobić, pewien efektowny, autorstwa M. Klebera, znalazłem w czasopiśmie „Mathematical Intelligencer”, tom 24, numer 1 (rok 2002). Zachęcam do zajrzenia tam, a na koniec jeszcze jedno zadanie.

Rozkładamy 52 karty w prostokącie o 13 rzędach poziomych i 4 pionowych (kolumnach). Wykaż, że da się wybrać po jednej karcie z każdej kolumny tak, by nie powtarzały się wielkości (to znaczy, by mieć jednego asa, jedną dwójkę, jedną trójkę i tak dalej).

Jak to zrobić? Zastosować twierdzenie Halla. Jak? Mniej więcej tak, jak przy wyborze najlepszej drużyny siatkarskiej i przyłączy wody do działek.

Twierdzenie Halla znajduje zastosowanie i w matematyce z najwyższej półki. Ma też wiele uogólnień, związanych z trudniejszymi zagadnieniami praktycznymi, na przykład uwzględniających stopień preferencji poszczególnych wyborów, na przykład małżeńskich, ale i stopień kwalifikacji pracowników do poszczególnych zadań (jak w przykładzie z firmą), ale to już inna historia. Zastosowanie „matrymonialne” to tylko przykrywką, trochę z obliczeniem na efekt „wow”. ■

Archiwalne artykuły z matematyki:
<https://tiny.pl/c9cgz>



Szkoła Wynalazców

dozwolone do lat 15

Mieliście zadanie co się zowie twórcze, wymagające ekstremalnej fantazji. Chodziło o wynalezienie nowej idei dla dronów, a więc należało wymyślić nowe zastosowania dronów.

Wydaje się, że drony już są ostatecznie wyeksploatowane jako bezzałogowy środek transportu „wszystkiego”. Rzeczywiście, drony przenoszą: ładunki wybuchowe, środki ratunkowe dla osób zablakanych w dżungli, lekarstwa, organy do przeszczepów, krew do transfuzji, materiały informacyjne, szpiegowanie z lotu ptaka i mnóstwo jeszcze innych rzeczy. Oczywiście istnieje szereg klas wielkościowych dronów: od dronów – owadów, do dronów o wielkości małego samolotu. Jednakże wspólną ich cechą jest, jak w szkolnym zadaniu: „z miejsca A do miejsca B coś przenieść”. Czy więc drony mogą tylko „coś przetranszować” i tylko podglądać? To pytanie było właśnie do was skierowane. Co mogłyby jeszcze ciekawego drony wykonywać?

Damian Borycki: Jeżeli drony mogą coś przetranszować i dostarczyć pod wskazany adres, to mogą na pewno coś zabierać. Przed wielu laty – gdy drony były jeszcze w „powijkach”, były już jednak dość proste drony śmigłowce. Różni dowcipnicy korzystali z nich, polując na panie opalające się w stroju Ewy na specjalnie ogrodzonym kawałku plaży w Sopocie i jednak przysyłające piersi ręcznikami, dla ochrony przed zbyt silnym słońcem. Nadlatywał wtedy dron, pazurkami chwycił ręcznik i zmykał.

No cóż, człowiek to podobno homo sapiens, ale przede wszystkim homo habens fun (człowiek bawiący się). Podobną funkcję dziś spełniają mewy, które bardzo często porywają plażowiczom trepki, ręczniki itp.

Zdzisław Kowalski: Można by wykorzystać drony do usuwania wronich gniazd z kominów. Dron nadlatywałby nad otwór wylotowy komina, opuszczał coś w rodzaju harpuna i wznosił się, zabierając gniazdo.

Normalnie robią to kominarze, ale dron mógłby pomóc. Problem w tym, że jego harpun mógłby zahaczyć o jakiś element stały komina i wtedy dron zostałby uwięziony!

Roman Godowski uważa, że drony – odpowiednio wyposażone – mogłyby być użyte do mycia okien budynków wysokościowych.

Pewnie by mogły, ale ich wydajność byłaby chyba niezbyt wielka, bo nie mogą się zbyt blisko zbliżyć do okien.

Wymienionym kolegom gratuluje i zachęcam do udziału w dalszych konkursach.

Nowe zadanie

Zadanie „na czasie”. Liczymy tu na młodzieńczą fantazję. Chodzi o najazd nielegalnych imigrantów na naszą wschodnią granicę. Obecna sytuacja jest nie do utrzymania: żołnierze mają broń i ostrą amunicję, ale nie wolno im jej używać, są rozliczani z liczby wystrzelonych pocisków i karani za ich – zdaniem władzy – nieuzasadnione użycie. No to co? Czy zatrudnić pedagogów, którzy będą na granicy łagodnie tłumaczyli imigrantom niestosowność ich zachowania? Należałoby wymyślić jakiś sposób, niemiły i odstręczający tych „biednych ludzi”, ale nieczyniący im krzywdy, o strzelaniu do nich nie wspominając. Jednakże obecnie nasze władze zapewniają im pełny komfort i bezpieczeństwo, tolerując najdziksze zachowania, ze śmiertelnym zranieniem jednego z naszych żołnierzy włącznie i co? No więc spróbujcie zaproponować jakiś sposób na zniechęcenie imigrantów do atakowania naszej granicy.

Zaproponować humanitarny sposób na zlikwidowanie naporu imigrantów na naszą wschodnią granicę (rysunek 1).

Zadanie na pewno nie jest łatwe, ale liczymy na młodzieńczą fantazję! Propozycje prosimy nadsyłać do końca września br.



Klub Wynalazców

bez ograniczeń wieku

Mieścieście zadanie typowe dla technologii kształtowego wycinania laserem lub tlenem różnych kształtek, w tym także w rurach. Chodziło więc o to, żeby zaproponować sposób usuwania kropelek metalu, przyklejonych do wnętrza rury po wycinaniu tlenem lub laserem otworów kształtowych w jej ściankach.

Zjawisko jest znane i „jakoś tam” pokonywane. Dla rur prostych, to jeszcze pół biedy. Problem jest, gdy rura jest powyginana. Cisnące się na usta metody przedmuchiwania strumieniem piasku, w przypadku rur powyginanych, odpadają. Pamiętamy, że rury mają już powycinane różnokształtne otwory, przez które piasek oczywiście wyleci w siną dal, zasypując halę. Oczyma duszy widzimy coś w rodzaju głowicy frezarskiej, ale znów ta powyginana rura! Zadanie wymagało więc sprytu i wyobraźni. Zobaczmy, jak sobie z tym radzili nasi koledzy:

Zbigniew Krajewski do czyszczenia zakrzywionych rur proponuje głowicę typu frezarskiego, osadzoną na wale giętkim, który co kilkadziesiąt cm musiałby być stabilizowany przez okrągłe tarcze podtrzymujące wał mniej więcej w centrum rury. Średnica freza powinna być dobrana z małym luzem w stosunku do średnicy wewnętrznej rury.

Prawdopodobnie najprostszy i skuteczny sposób. Być może elastyczne osadzenie płytek skrawających głowicy dałoby lepszy efekt, bo głowica nie zacinaby się w przypadku nieco grubszego narostu kropek.

Zygmunt Fijałkowski proponuje wykonanie narzędzia w formie „miotły”, tzn. z rozchyłającymi się sprężystymi ramionami, zaopatrzonymi w kształtki wieloostrowkowe, skrawające. Ta „miotła” powinna być osadzona na wale giętkim lub przegubowym i obracając się z dużą szybkością, powinna poradzić sobie z naklejonymi grudkami metalu.

Interesującą i prostą koncepcją. Prawdopodobnie mogłaby obsłużyć pewien zakres średnic. Zasadniczą zaletą „miotły” jest to, że obracana z możliwie dużą prędkością obrotową korzystałaby z docisku do ścianek rury siłami odśrodkowymi.

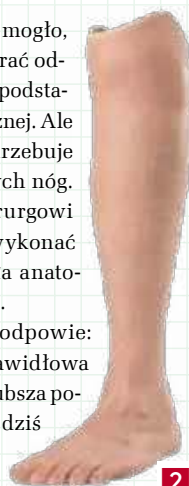
Obu kolegom gratuluję i zapraszam do następnych zadań.

Nowe zadanie

Tym razem zadanie z pogranicza medycyny i mechaniki: w przypadku stosowania protez nóg bardzo ważne jest, aby sztuczna noga była dokładnie taka sama

jak druga, żywa. Wydawać by się mogło, że nie jest to trudne – należy pobrać odlew gipsowy z żywej nogi i na jego podstawie wykonać formę dla nogi sztucznej. Ale to tak nie działa, bo nikt nie potrzebuje dwóch lewych ani dwóch prawych nóg. Co robić? A więc: doradźcie chirurgowi ortopedzie i protetykowi, jak wykonać protezę np. prawej nogi, aby była anatomicznie prawidłowa (**rysunek 2**).

Część naszych czytelników odpowie: a po co komu anatomicznie prawidłowa proteza; wystarczy, żeby była z grubsza podobna do nogi zdrowej, a poza tym dziś widzi się protezy w ogóle niepodobne do żywej nogi, za to funkcjonalne (**rysunek 3**). W zasadzie prawda, ale spróbujmy cofnąć się o jakieś 70–80 lat. Nie było skanerów 3D ani programów komputerowych 3D, w ogóle patrząc na to z dzisiejszego punktu widzenia, to było „średniowiecze”! No właśnie. Mimo to w tamtych czasach ludzie potrafili robić piękne protezy. Wszystkim życzymy takiego podejścia do tego zadania i przypominamy termin nadsyłania odpowiedzi: koniec sierpnia br.



2



3



Vademecum Młodego Wynalazcy

Spotykamy się często z wypowiedziami typu: „eee, jakby to była prawda, to dawno by to wymyślili Amerykanie”. Jest to echo naszych polskich kompleksów w stosunku do narodów, które miały mniej dramatyczną historię i mogły więcej uwagi poświęcić wynalazczości i w ogóle twórczości. Oznacza też naszą niewiarę w możliwości Polaków. Przy wszystkich naszych narodowych wadach akurat w wynalazczości mamy się czym pochwalić. Z kilkoma osiągnięciami „naszych” warto się zapoznać chociaż w pigułce.

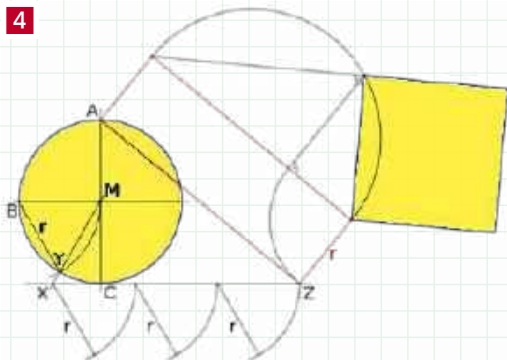
Nieco zapomnianą, ale wielką postacią był, żyjący w latach 1631–1700 **Adam Adamandy Kochański** herbu Lubicz – polski uczonec i duchowny katolicki; matematyk, fizyk, astronom, inżynier-mechanik, filozof i bibliotekarz; członek zakonu jezuitów. Najbardziej znanym jego osiągnięciem wynalazczym była „kwadratura koła”, inaczej mówiąc konstrukcja geometryczna, pozwalająca z dużą dokładnością narysować kwadrat o powierzchni równej powierzchni danego koła, wyłącznie przy użyciu cyrkla i linijki. Konstrukcja genialnie prosta, matematycznie piękna (**rysunek 4**), pozwoliła wyznaczyć wartość liczby π z... Błąd względny wartości π Kochańskiego w stosunku do najlepszych obecnych oznaczeń dokładnością do czterech miejsc po przecinku. Wartość π wg Kochańskiego wynosiła 3,14153..., wobec współcześnie obliczonej (z bardzo wielką liczbą cyfr po przecinku) wartości: 3,14159 wynosi 0,000059. Na tle trwających tysiące lat prób dokładnego wyrażenia stosunku obwodu koła do jego średnicy osiągnięcie Kochańskiego to sukces światowy. Setki prób, od starożytnych przybliżeń $\pi=22/7$, czyli 3,142857, do późniejszych, włączając w to wartość π podaną w podręczniku wydanym w 1566 r. przez Stanisława Grzebskiego: „Cirkumferentia (czyli obwód koła) jest tak wielka, jako trzy diametry i siódma część diametru

bez małego kęska”. Rozczulające jest to: „bez małego kęska” – symbol bezradności wielu pokoleń matematyków wobec upartej liczby π . Czy gonitwa za jak najdokładniejszym wyznaczeniem wartości π (nawet do bilionów miejsc po przecinku!) ma jakikolwiek praktyczny sens?

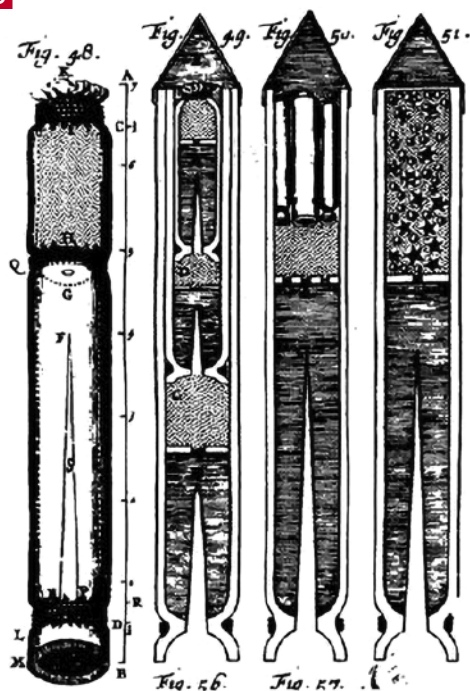
W bardzo wyjątkowych przypadkach sięgamy po wartość z pięcioma miejscami po przecinku. Gdy kielecki Chemar wykonywał zbiorniki kuliste o średnicy 16 m, to gdyby obwód obliczyć, korzystając z popularnego przybliżenia $\pi=3,14$, wyszłoby 50240 mm, a licząc dokładniej – 50265 mm. Różnica 25 mm nie do przyjęcia przy trasowaniu segmentów kuli. Warto dodać, że każdy przeciętny kalkulator ma zapisaną, oznaczoną symbolem EXP liczbę π , z dokładnością do 9 cyfr po przecinku.

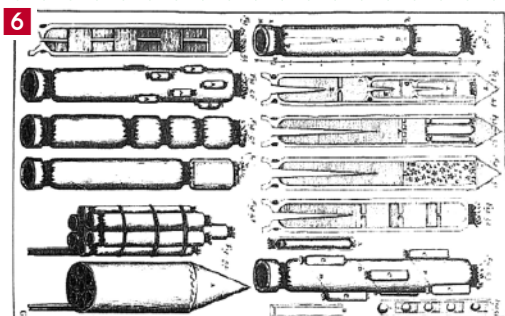
Ważną postacią, godną przypomnienia jest **Kazimierz Siemienowicz**, z wykształcenia artylerzysta, żyjący i działający w czasach króla Władysława IV. Mamy pełne prawo nazwać go ojcem techniki raketowej na kontynencie europejskim, Rakiety były jednak znane dużo wcześniej w Indiach i później w Chinach i Mongolii. Siemienowicz wsławił się jako wybitny teoretyk techniki raketowej. Swoje pomysły i dociekania

4



5





zawarł w wydanej w 1650 roku Księdze, której łaciński tytuł brzmi: „*Artis magnae artilleriae pars prima, studio et opera Casimiri Siemienowicz*”, czyli: „Sztuki wielkiej artylerii część pierwsza”. Księga ta stała się szybko bestsellerem na europejskim rynku, doczekała się wielu tłumaczeń. Tłumaczenie na język polski ukazało się w 1963 roku. Nie można się temu dziwić, bo w wiekach wcześniejszych popularnym językiem teoretyków sztuki wojennej była łacina i polscy specjaliści nie potrzebowali tłumaczenia.

Księga Siemienowicza jeszcze dziś budzi podziw dla wszechstronności i kreatywności autora. Szkice rakiet zawierają rysunki rakiet wielostopniowych, rakiet ze skrzydełkami typu „delta”, rakiet z dodatkowymi silnikami raketowymi, wspomagającymi, itp. (rysunki 5 i 6). Można powiedzieć, że Siemienowicz wyprzedził swoją epokę o ok. 300 lat!

Kolejną postacią godną przypomnienia jest **Ernest Malinowski**, właściwie Ernest Adam Stanisław Malinowski herbu Ślepowron (ur. 5 styczeń 1818 na Wołyniu, zm. 2 marca 1899 w Limie) – polski inżynier drogowy i kolejowy, bohater obrony Callao w 1866 r., budowniczy kolei w Peru i Ekwadorze, projektant i budowniczy Centralnej Kolei Transandyjskiej. Kolej tę zbudowano dla połączenia Limy z bogatym w minerały regionem Cerro de Pasco i żyzną doliną Jauja, celem transportu bogactw z górskich rejonów Peru do portu w Limie. Niemożliwe stało się możliwe. E. Malinowski podjął się zaprojektowania konstrukcji całej linii kolejowej i nadzorowania jej budowy. Ogółem na całej trasie z Limy do La Oroya wykuto 63 tunele o łącznej długości ponad 6000 m, wybudowano 61 mostów i wiaduktów o łącznej długości ok. 2 km i wykonano 10 nawrotów (zmian czoła pociągu).

Najbardziej znany jest wybudowany w 1872 r. wiadukt Verrugas (najdłuższy na trasie, ok. 200 m) – **rysunek 7**, łączący brzegi wąwozu o tej samej nazwie. Opiera się na trzech filarach ze stalowych kratownic, z których najwyższy ma prawie 77 m.

Gdy patrzysz na ten słynny wiadukt, wręcz niewiarygodne wydaje się, że jeździły po nim pociągi. Lekka kratownicowa konstrukcja filarów ujawnia fakt, że w czasie gdy była projektowana, nie było komputerów i całość konstrukcji powstawała na deskach kreślarskich, obliczana z pomocą suwaka logarytmicznego (kto dziś umie się nim posłużyć?). Być może dlatego, patrząc na to z dzisiejszego punktu widzenia, układ prętów kratownic między filarami uważamy za nieco dziwny. Nie umniejsza to osiągnięcia Malinowskiego, który robił to wszystko w latach 70. XIX wieku. Przy projektowaniu i budowie tej linii Ernest Malinowski musiał rozwiązać cały szereg nietypowych problemów, czyli dokonać wielu wynalazków. Jednym z takich nowatorskich rozwiązań było pokonanie sporego wzniesienia metodą 10 nawrotów, czyli zmian czoła pociągu, który jechał do przodu, a następnie do tyłu, na wyższy poziom i tak 10 razy. Było to znacznie szybsze niż jazda pętlami, okalającymi górską przeszkodę. Podobnych problemów było oczywiście o wiele więcej. Malinowski doczekał się uznania w Peru, a także w... Polsce.

Znacznie bliższy naszym czasom jest kolejny wybitny wynalazca polski – **Jan Czochralski**. Był to wyjątkowo uzdolniony młody i ambitny człowiek. Po otrzymaniu świadectwa dojrzałości, gdy był już pewien, że nie grozi mu za to kara, podarł je, mówiąc: „Proszę przyjąć do wiadomości, że nigdy nie wydano bardziej krzywdzących ocen”. Brak tego dokumentu zamknął młodemu Czochralskiemu drogę do dalszej kariery nauczycielskiej i naukowej. Nie powstrzymało go to jednak od oświadczenia, że wróci w rodzinne strony dopiero wtedy, gdy zdobędzie sławę. Późniejszy etap to studia w Berlinie i dalsza praca naukowa. W latach 1911–1914 był asystentem Wicharda von Moellendorffa, z którym opublikował swoją pierwszą pracę poświęconą krystalografii metali, a dokładniej podwalinom późniejszej teorii dyslokacji. Od tamtej pory właśnie metalografia stała się dziedziną, której Czochralski poświęcił



8



resztę życia. Pierwsze publikacje były poświęcone zastosowaniu aluminium w elektronice. Największy rozgłos przyniosła mu odkryta w 1916 r. metoda pomiaru szybkości krystalizacji metali, wykorzystywana obecnie do produkcji monokrystalów krzemu. Według popularnej anegdoty metodę tę odkrył przypadkowo, zanurzając pióro w tyglu z gorącą cyną zamiast w kałamarzu. Ta metoda dziś jest niesłychanie ważna w dziedzinie produkcji półprzewodników wszelkich typów. Ma zasięg ogólnosiwiatowy (**rysunek 8**). Dzieje Czochralskiego obfitują jednak też w momenty dramatyczne, zgodnie z porzekadłem o „polskim oddziale w piekle, w którym nie ma straży piekielnej”. Niesłusznie oskarżany, bez dowodów, został jednak ostatecznie zrehabilitowany i 29 czerwca 2011 r. Senat Politechniki Warszawskiej ogłosił rehabilitację Czochralskiego. No cóż, lepiej późno niż wcale...

Ostatnią postacią z tej mikroserii wynalazców polskich jest **Olga Malinkiewicz**, doktor fizyki i wynalazca (**rysunek 9**).

Na oko sympatyczna, młoda kobieta, z którą chętnie poszlibyśmy na dyskotekę. Tymczasem pani Olga to poważna naukowiec, fizyczka i wynalazca technologii produkcji perowskitu, sztucznego minerału o właściwościach fotowoltaicznych. Jej droga życiowa jest typowa dla osób wybitnych, które nie mieszczą się w standardowych ramach. Doceniona w... przedszkolu, ujemnie oceniona w szkole podstawowej, trafiła po maturze (uzyskanej w innej szkole) na Uniwersytet Warszawski, który opuściła po uzyskaniu stopnia licencjata. Na dalsze studia przeniosła się do Barcelony, gdzie uzyskała stopień magistra na Politechnice Katalońskiej. Jeszcze w czasie studiów,

w 2009 roku, rozpoczęła pracę w instytucie ICFO (The Institute of Photonic Studies), a następnie na Uniwersytecie w Walencji. W tej uczelni w 2017 obroniła doktorat (w The Institute for Molecular Science). Po powrocie do Polski, mimo różnych biurokratycznych trudności, z kolegami założyła firmę: Saule Technologies, której zadaniem było uruchomienie doświadczalnej produkcji perowskitu w odmianach. Jedną z najciekawszych właściwości perowskitu produkowanego w Saule Technologies jest możliwość nanoszenia go na elastyczne folie (**rysunek 10**), tekstylia, a także na ściany budynków i dachów.

Najbliższe problemy to: dalsze podnoszenie sprawności elementów fotowoltaicznych opartych na technologii perowskitu, zapewnienie odporności na czynniki atmosferyczne i studia nad efektywnym wykorzystaniem tego wynalazku. Pani Olga zyskała sławę międzynarodową, została wielokrotnie uhonorowana odznaczeniami, dyplomami itp. Jest przykładem wynalazcy nowoczesnego, którego osiągnięcia oparte są na solidnej wiedzy i oczywiście pracy.

Te kilka sylwetek polskich wynalazców wybrałem nieprzypadkowo. Pytani o te właśnie postaci uczniowie jednej ze szkół średnich nie wiedzieli nic! A przecież poza tą bardzo krótką galerią polskich wynalazców mamy tysiące postaci wielkich twórców techniki. Niestety duża ich część w okresie rozbiorów pracowała w obcych firmach i uczelniach. Pozostają nam dylematy w rodzaju: czy Maria Skłodowska-Curie to francuska uczona polskiego pochodzenia, czy polska uczona, wykształcona we Francji i pracująca na francuskiej uczelni...

No więc warto sobie przypomnieć Reja piszącego o tym: że „Polacy nie gęsi i swój język mają”. Muszę tu jednak przypomnieć, że Rej miał na myśli „gęsi język”, którego Polacy nie używają! ■

Prezes Klubu Wynalazców
Champion TRIZ
Jan Boratyński

9



10



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 8/2024

Uznajemy za ciekawą ideę przechowywania przyborów rysujących i malujących w kontrolowanej atmosferze ze stałą wilgotnością. Niezależnie od szczegółów, tego rodzaju pomysł wydaje się praktyczny i daje się zrealizować.

Autorem pomysłu jest Antoni Bukowski

1 Witold Karwacki – jako student miał w programie obozy poligonowe w ramach studium wojskowego. Zmora zajęć było czyszczenie broni – wszelkiego rodzaju, a więc od „kałasza” do armaty kal. 85, wz. 1944. Witold uważa, że klasyczne czyszczenie lufy „pakulami” za pomocą wyciorów to anachronizm niegodny XXI wieku. Proponuje opracowanie naboju czyszczących, którymi ładowałoby się broń, naciskało spust i gotowe! W przypadku szczególnie zanieczyszczonej lufy można by zabieg powtórzyć. Byłoby to błogosławieństwo dla żołnierzy, którzy obecnie muszą pracowicie bawić się czyszczeniem broni pod okiem bezwzględniego pana kaprała.

Myśl jest dobra i ma już swojego „przodka”. Do dokładnego oczyszczenia lufy armatniej stosowano „przebijanie” przewodu lufy za pomocą okrągłego, lipowego klocka. Kłoczek przebiegał się z pomocą wycioru i młota. Efekt był niezły, ale robota długa i czasami niemiła, gdy kłoczek miał za dużą średnicę i grzęzł gdzieś w połowie lufy. Propozycję kolegi na pewno da się zrealizować i chyba trzeba pomysł podrzucić wojskowemu. Gdyby ideę „naboju czyszczącego” zaakceptowano, to można by w taśmach do karabinów maszynowych np. co setny nabój ostry wstawiać nabój czyszczący!

2 Jakub Gregorczyk – uważa, że ostatnie pomysły przedstawicieli naszych władz, zwłaszcza władz regulujących sprawy wojskowe, zasługują na szczególną uwagę. Decyzja o wydaniu żołnierzom pilnującym wschodniej granicy broni automatycznej i ostrych naboju, z jednoczesnym zakazem ich używania, nawet do ostrzegawczego strzelania w powietrze, karanie żołnierzy za zużycie ok. 48 naboju – wyczerpuje definicję paranoidealnej głupoty. Oznacza to, że kandydaci do wysokich stanowisk, zwłaszcza tych, które dotyczą naszej obronności, powinni oprócz testu na IQ przejść specjalistyczne badania psychologiczne, a może psychiatryczne. Jest nie do przyjęcia sytuacja, gdy bezpieczeństwo 38 milionów ludzi zależy od stanu psychobiologicznego kilku osób.

Sprawa jest bardzo poważna i ma wymiar ogólnowiatowy. Zbyt często potworne skutki rodziło przekazanie władzy w ręce jednostki lub niewielkiej grupy osób. Aż nadto dobrze znane przykłady to Hitler, Stalin, ostatnio Putin i... następni już czekają. W Polsce aż takich ekstremalnych jednostek nie mamy, jednakże obronność to temat najwyższej wagi. Nieszczęściem demokracji jest to, że taka propozycja, jaką kolega przedstawił, w Polsce nigdy nie przejdzie.

3 Marek Strzeński – proponuje opracowanie aplikacji do smartfonów, dotyczącej jadłospisów przy stosowaniu diety ketogenicznej. Idziemy do sklepu z wykazem „papierowym” artykułów, które mamy zamiar kupić, niestety nie zawsze są one dostępne. Aplikacja na smartfonie powinna natychmiast zaproponować artykuły ekwiwalentne w sensie składu białek, tłuszczów i węglowodanów. Dieta ketogeniczna jest najwartościowszą dietą leczniczą, pozwalającą wyleczyć nawet stwardnienie rozsiane, uważane za „nieuleczalne”. Warto zadbać o oprzyrządowanie tej diety na smartfonie.

Pojawia się ostatnio coraz więcej metod medycyny zintegrowanej i być może kiedyś będzie można wyleczyć się u „normalnego” lekarza. Komerccjalizacja medycyny, farmacji, kształcenia studentów medycyny, korupcja i „sponsoringowanie” lekarzy, doprowadziły do sytuacji przykrej: ludzie boją się iść do szpitala!

4 Marek Chorabik ma gospodarstwo i pracuje maszynami. Po powrocie z pola w zależności od stanu gleby i pogody maszyny są często niemilosiernie zabłocone. Splukiwanie ich wodą to koszty, a poza tym błoto na terenie zabudowań gospodarskich. Marek rozważa możliwość zdmuchiwania błota strumieniem sprężonego powietrza. Słyszał o niemieckich „wzmocniaczach energii” takiego strumienia, pozwalających podnieść nawet dwudziestokrotnie skuteczność wydmuchu.

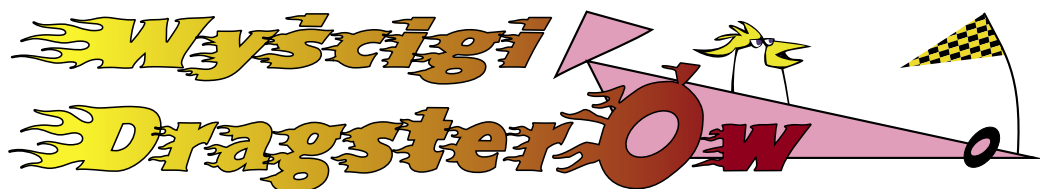
Takie wzmocniacze rzeczywiście są, nawet dostępne w Polsce. Są stosunkowo niedrogie, jednakże należałoby przeprowadzić próby, czy można adaptować wzmocniacz do zadania zdmuchiwania błota? W każdym razie pomysł wart przemyślenia.

5 Antoni Bukowski jest architektem i w swojej pracy niezależnie od grafiki komputerowej używa kolorowych flamastrów do szkicowania różnych wariantów projektu. Pisaki, mimo zamykania ich wieczkami, po jakimś czasie wysychają i nie nadają się do niczego. Antoni proponuje opracowanie pojemnika na pisaki, zapewniającego wilgotność taką, żeby pisaki mogły tam przetrwać chociaż parę tygodni.

Wysychanie pisaków dotyczy także długopisów, piórotek „automatycznych”, itp. Pojemnik, o jakim pisze Antoni, rzeczywiście mógłby pomóc. Projekt i wykonanie nie powinny być trudne, a dla kogoś, kto używa tych rzeczy, byłaby to duża pomoc.



Budujemy wyścigowe modele „dragsterów” z napędem sprężynowym



Wyścigi dragsterów to rodzaj sportu motorowego, w którym dwa pojazdy muszą od startu zatrzymanego w jak najkrótszym czasie pokonać odcinek drogi lub toru o długości 1/4 mili, często osiągając zawrotne prędkości i przyspieszenia. Wyścigi wywodzą się ze Stanów Zjednoczonych i są tam bardzo popularne. Takie zawody rozgrywane są też w Polsce – są to Mistrzostwa Polski w Wyścigach Równoległych im. Jana Wdowczyka.

Nasz model, choć nie będzie demonem prędkości ani przyspieszenia, może sprawić wiele frajdy, szczególnie emocjonujące będą wyścigi modeli przeprowadzone na wzór prawdziwych zawodów, można je przeprowadzić na szkolnym korytarzu czy w sali gimnastycznej. W moim przypadku wyścigi odbywały się w modelarni

na dystansie ośmiu metrów. Model można też wykorzystać jako pomoc naukową na lekcji fizyki, np. do zademonstrowania zamiany energii potencjalnej w kinetyczną czy zobrazowania ruchu jednostajnie przyspieszonego.

Model został zaprojektowany przy użyciu darmowego programu Free CAD i jest możliwy



Wydrukowane części modelu



Narzędzia oraz wkręty do montażu ramy



Opony kół z uszczelek hydraulicznych



Montaż ramy



Zakładamy opony



Szpulka napędowa

do wydrukowania na najprostszej, niskobudżetowej drukarce 3D.

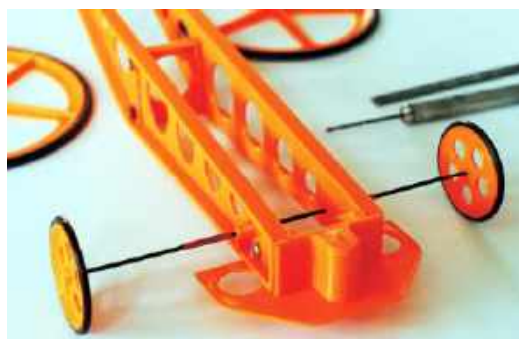
Budujemy model

Budowę modelu rozpoczynamy od wydrukowania głównych części, które znajdują się w dołączonych plikach (<https://mlodytechnik.pl/dragster>).

Do druku można użyć dowolnego filamentu, np. PLA lub PETG, projekt przygotowany jest pod dyszę o średnicy 0,4 mm, trzeba też zwrócić uwagę na precyzyjne

wypoziomowanie stołu i grubość warstw, dla modeli, które drukowałem, grubość pierwszej warstwy była ustawiona na 0,3 mm, a pozostałe warstwy na 0,2 mm.

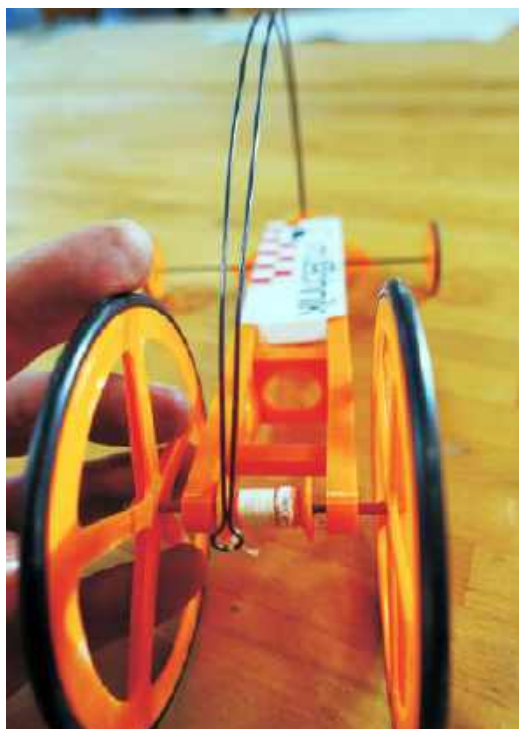
Otwory w ramie na osie kół oraz otwory w kołach należy rozwiertić do wymaganej średnicy, w wydrukach są tylko otwory pilotujące, do wiercenia nie należy używać wiertarek szybkoobrotowych, najlepiej rozwiercać ręcznie, używając wiertła zamocowanego w uchwycie. Ramę skręcamy za pomocą wkrętów do tworzyw sztucznych, trudno je będzie

**Koła i osie****Przednie zawieszenie****Tylna oś z rolką napędową****Sprężyna napędu****Sprężynę należy dobrze wkleić klejem cyjanoakrylowym**

zakupić, ale można je odzyskać ze zużytych zabawek czy plastikowych obudów elektroniki, ostatecznie można zastosować małe wkręty do drewna lub śrubki 3 mm. Przednią oś należy wykonać z drutu stalowego o średnicy 0,8...1,5 mm i długości 14...16 cm. Tylną oś wykonamy z drutu 1,5...2 mm, jej długość to 7 cm. Otwory w kołach i szpulce napędowej rozwiernamy tak, żeby pasowały na wcisk, natomiast otwory w ramie muszą zapewnić swobodny obrót osi kół. Przednią oś zabezpieczmy przed przesuwaniem dwoma odcinkami rurki termokurczliwej, natomiast koła i szpulkę przyklejamy kropelkami kleju cyjanoakrylowego.

**Naciągamy sprężynę**

Tak jak w prawdziwym dragsterze, żeby zapewnić wymaganą przyczepność kół, potrzebne są dobre opony, które w naszym modelu wykonane są



Model gotowy do jazdy

z gumowych uszczeltek (oringów), w przypadku przednich kół są to uszczelki śrubunku 1 cal, natomiast opony tylnych kół to uszczelki filtra wody. Takie uszczelki są dostępne w sklepach hydraulicznych lub marketach budowlanych na stoiskach z armaturą hydrauliczną.

Jako napęd modelu zastosowano element sprężysty, analogicznie jak w łuku lub kuszy, z tym że w naszym przypadku naprężona cięciwa nie wystrzeliwuje strzały, a rozwija się ze szpulki, napędzając koła modelu. Od zastosowanej sprężyny będą zależały osiągi modelu, sprężyna napędowa w moim przypadku wykonana jest z dwu drutów stalowych („fortepianowych”) o średnicy 1,2 mm i długości 30...35 cm.



Numer startowy – grafika pomniejszona 50% (plik do pobrania ze strony <https://mlodytechnik.pl/dragster>)

Można ją wykonać też z pojedynczego drutu o średnicy 1,5 mm lub pręcika z włókna szklanego o średnicy 2 mm, wykonanie tego elementu pozostawiam Waszej pomysłowości i kreatywności. Sprężynę należy zamontować w otworze znajdującym się z przodu modelu, wcześniej go rozwiercając tak, żeby uzyskać pasowanie na wcisk i zabezpieczyć klejem, do górnego końca sprężyny przywiązujemy mocną nitkę zakończoną pętelką, którą założymy na zaczep szpulki napędowej, następnie naciągamy sprężynę, kręcąc tylnymi kołami. Model stawiamy na podłodze i puszczamy. Napęd poprawnie wykonanego modelu powinien wystarczyć na przejechanie około czterech metrów, dalszą część dystansu model pokonuje, jadąc siłą bezwładności.

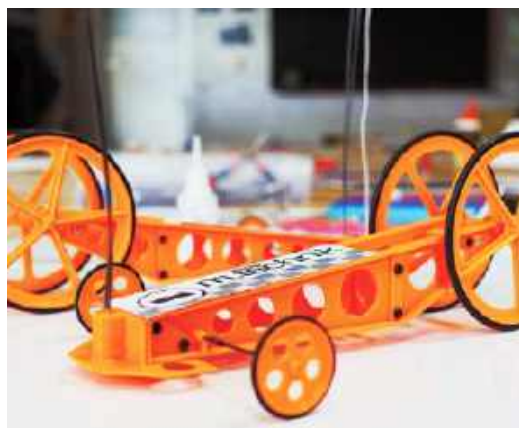
Gotowy model można ozdobić banerkiem z numerem startowym – propozycje numerów startowych z logo MT i AVT lub do projektu według własnego pomysłu dołączone są do plików, należy je wydrukować na karcie z bloku technicznego, wyciąć i nakleić na model.



Start



Pół sekundy po starcie



Gotowe modele



Wyścigi najlepiej przeprowadzić na wzór prawdziwych, wyznaczając dystans 5 do 10 metrów, natomiast jeśli mamy jeden model, można mierzyć czas przejazdu na określonym dystansie znając długość dystansu i czas przejazdu, możemy obliczyć średnią prędkość modelu.

Moim podopiecznym w modelarni na dystansie ośmiu metrów udało się uzyskać czas 2,8 sekundy, ale sądzę, że bez problemu uda się go Wam poprawić. ■

Mariusz Wrona

PRZEGLĄD TECHNICZNY Przyszło 5-lecie lotnictwa

W referacie, wygłoszonym na tegorocznym kongresie Instytutu Przewozów (Institute of Transport) w Anglii, płk. Briston, mówiąc o „Przyszłym 5-leciu lotnictwa”, wskazał na utrudnienie jakie powstaje dla organizacji komunikacji lotniczej skutkiem przerw zimowych. Wydatki zimą są dość znaczne, zaś dochodów nie ma, nadto po przerwaniu lotów są pewne trudności z ponownem ich podejmowaniem. Ponieważ nie można narazie rozwinąć większej komunikacji osobowej zimą, że względu na nieodpowiednią pogodę oraz krótki dzień, należy zwiększyć w tym okresie ruch towarowy. Ten ostatni wspomaga już komunikację osobową w stopniu zadowalającym na linii London-Bruksella-Kolonja. Wówczas gdy ruch osobowy w styczniu r. b. spadł na tej linii, w porównaniu z lipcem 1923, w stosunku 1: 13, ruch towarowy łączony z osobowym (w odpowiednim przeliczeniu) zmienił się w tym samym okresie tylko jak 1: 4. Linia przewoziła po 50 t przesyłek dziennie i była zmuszona często odmawiać ich przyjmowania. Ciężar opłacający swój przewóz stanowi około 25% całej wagi samolotu z ładunkiem. Gdyby waga samolotu z silnikiem mogła być zmniejszona o 1/2 kg na 1 KM, wówczas ładunek użyteczny, opłacający się, mógłby być znacznie powiększony. Silniki chłodzone wodą uważa autor za odpowiedniejsze niż chłodzone powietrzem i wroży im większe zastosowanie w przyszłości. Spodziewa się on, iż wkrótce uda się uzyskać konstrukcję płatowców o wadze do 9 kg na 1 KM, z czego ok. 3 kg będzie stanowił ładunek użyteczny. Omawiając oszczędności na paliwie do silników, wskazuje autor system Ricardo wtryskiwania zdwojonego,

z zastosowaniem jako paliwa nafty, oraz system 2-paliwowy (nafta + spirytus), przy czym ten ostatni jest wtryskiwany dodatkowo samoczynnie.

19 sierpnia 1924

Skrzyżowanie Wisły przewodami elektrycznymi w Chełmnie

Dnia 3 sierpnia r. b., odbyła się w Chełmnie, na Pomorzu, uroczystość pierwszego w Polsce skrzyżowania przewodami elektrycznymi rzeki o dużej rozpiętości, bo 612 m. Jest to w naszych warunkach fakt który należy podnieść z uznaniem, jako jeden z ważnych kroków na drodze elektryfikacji i wyzyskania energii wodnej, zwłaszcza że dokonano go w tak trudnych warunkach gospodarczych. Budowę tę przeprowadził Związek Elektryfikacyjny 3 powiatów Pomorza: Chełmińskiego, Toruńskiego i Świeckiego, mający na celu wybudowanie sieci elektrycznych o napięciu 15000 woltów, czerpiącej energię z elektrowni w Gródku. Energiczna praca Związku, na którego czele stoi p. starosta dr Bobke, zaś kierownictwo techniczne spoczywa w rękach p. inż. A. Hoffmanna, posunęła znacznie naprzód elektryfikację Pomorza.

26 sierpnia 1924

PRZEGLĄD ELEKTROTECHNICZNY Ze Zjazdu Elektrotechników Czesko-Słowackich

Zanim złożę dokładne sprawozdanie ze zjazdu, w którym uczestniczyłem jako delegat Stowarzyszenia Elektrotechników Polskich i Politechniki Warszawskiej, chciałbym podzielić się z czytelnikami Przeglądu kilku swymi wrażeniami z pierwszego zetknięcia się z elektrotechnikami czeskiemi. Przemysł elektrotechniczny w Czechach i elektryfikacja kraju znacznie wyprzedza Polskę, to też Związek

Elektrotechników Czesko-Słowackich jest liczniejszy, bogatrzy i sprawniejsi od naszego Stowarzyszenia. Zjazd wypadł imponujący. Odczyty dotyczyły się postępów wiedzy i elektryfikacji kraju, wreszcie dotyczyły przepisów i norm. Jednocześnie była otwarta wystawa nowości elektrotechn., wytwarzanych w Czechach (kilkadziesiąt firm uczestniczyło w tej wystawie!). Zwiedzanie fabryk i wycieczki do sąsiednich elektrowni okręgowych dopełniały całości. W dziedzinie przepisów i norm Czesi zrobili wiele. Przepisy opracowują samodzielnie i wydają je w postaci książek. Większość przepisów uzyskała sankcję państwową. Na Zjeździe uczestniczyło sześciu gości zagranicznych; trzech Francuzów, dwóch Rosjan z Petersburga i niżej podpisany. Przyjmowani byliśmy z prawdziwie słowiańską gościnnością i wylewnością. Mowy nasze były przerywane i przyjmowane hucznymi oklaskami, a co ważniejsze poświęcono wiele czasu na konferowanie z nami w sprawach fachowych. Przemysłowcy informowali się o stanie naszego rynku, kierownicy Związku Elektrotechnicznego zaproponowali nam współpracę na polu przepisów i norm elektrotechnicznych, a wszyscy zdadzali żywe zainteresowanie się Polską. Jestem głęboko przekonany, że po tym Zjeździe stosunki między polskim a czeskim światem elektrotechnicznym znacznie się ożywią zarówno w dziedzinie nauki, jak i przemysłu.

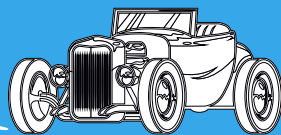
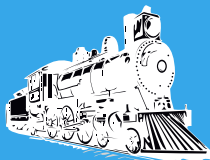
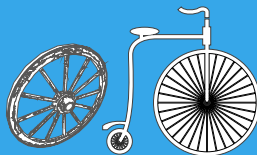
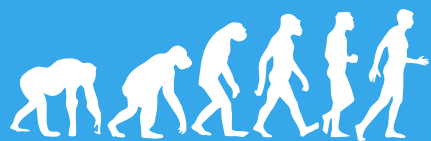
15 sierpnia 1924

PRZEGLĄD PRZEMYSŁOWO-HANDLOWY Z przemysłu papierniczego

Ogólny zastój w życiu gospodarczym Polski do pewnego stopnia dotknął i przemysł papierniczy. Jednak zastój w przemysle papierniczym nie jednakowo odbił się na wszystkich papierniach.

Papiernie produkujące papiery drukowe – gazetowe mają i obecnie całkowity zbyt zapewniony; to też pracują one pełną parą. Nadal w gorszych warunkach znajdują się papiernie, produkujące lepsze i wykwintne gatunki papieru. Hurtownicy posiadają jeszcze pewne zapasy tych gatunków papieru, powstrzymują się przeto w obecnym czasie ogólnego braku gotowizny chwilowo od dalszych znaczniejszych zakupów. To też i papiernie te częściowo, lecz nie w zbyt szerokich ramach, zmniejszyły swą zdolność wytwórczą. Jednak zapasy papieru lepszych gatunków nie są zbyt wielkie u hurtowników. Wobec tego spodziewana jest jesienią poprawa sytuacji. Pomimo chwilowej niezbyt pomyślnej konjunktury w przemyśle papierniczym, – czynione są usilne starania w kierunku czy to budowy nowych papierni, czy to rozszerzenia starych przez postawienie nowych maszyn papierniczych. Chwilowo mamy do zanotowania realizujące się już projekty budowy 2 nowych papierni i projekty rozszerzenia dwóch starych papierni; – tak, iż Polskę przybywa ogółem 5 maszyn papierniczych na ogólną ilość 17 obecnie egzystujących. Przyczem wydajność tych nowych maszyn będzie znacznie większa, niż wydajność obecnie czynnych maszyn. Do obecnej produkcji papieru koło 360–400 wagonów miesięcznie Polskę przybywa produkcja koło 200 wagonów miesięcznie, a więc obecna produkcja wzmagą się o 50%. Na tem jednak jeszcze nie kończą się projekty dalszych budowli: sfery finansowo-przemysłowe w dobrze rozumiałym własnym interesie pragną i nadal lokować swe kapitały w tej gałęzi przemysłu, wierząc w jego pomyślną przyszłość rozwojową.

23 sierpnia 1924



starożytność

średniowiecze

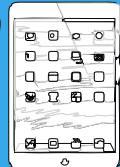
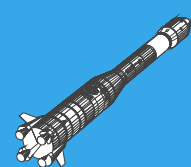
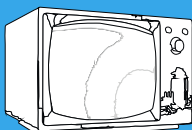
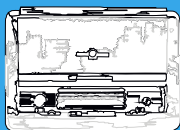
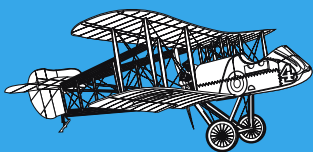
XVI w.

Lody

Mrożone napoje i desery są znane od co najmniej 4000 lat p.n.e., kiedy w Mezopotamii zaczęto budować pierwsze pomieszczenia do przechowywania lodu. Substancja o konsystencji śniegu, prawdopodobnie używana do chłodzenia wina, była sprzedawana na ulicach Aten w V wieku p.n.e., zaś rzymski cesarz Neron pijał, według przekazów, mrożone napoje z dodatkiem miodu. Rzymska książka kucharska z I wieku zawiera przepisy na słodkie desery posypane śniegiem. Źródła z czasów dynastii Tang w Chinach opisują słodki napój z mrożonego mleka bawolego z dodatkiem kamfory. Fresk w grobowcu chińskiego księcia Zhanghuai (1) z siódmego wieku n.e. przedstawia szlachetne damy trzymające „Su Shan”, deser wykonany przez stopienie Su (kremowego masła wprowadzonego do środkowych Chin przez koczownicze ludy z północy), ściekającego na talerz, dekorowanego kwiatami i wreszcie zamrażanego. Niektóre źródła opisują produkty podobne do lodów pochodzące z Persji już w 550 r. p.n.e. Persowie według tych przekazów produkowali i podawali mrożony deser faloodeh i sorbety przez cały rok. Początki kakigōri, japońskiego deseru z lodem i syropem, sięgają okresu Heian.

Przełomem w produkcji lodów było odkrycie, że lód zmieszany z solą uruchamia egzotermiczną reakcję chemiczną, która tworzy schłodzoną zawiesinę o znacznie niższej temperaturze zamarzania niż zwykła woda. Zanurzone w kąpeli z takiej solanki kryształki lodu łatwo tworzyły się w różnych płynnych miksturach. Regularne mieszanie zapobiegające formowaniu się dużych kryształów lodu pozwoliło uzyskać mrożoną pianę, którą można było nabierać łyżką. Znałe są przekazy, że ta technika znana była w świecie arabskim już w XIII wieku. Najwcześniejszy znany pisemny proces sztucznego wytwarzania lodu znany jest nie z tekstów kulinarnych, ale z XIII-wiecznych pism syryjskiego historyka Ibn Abu Usaybia – w jego książce „Kitab Uyun al-anba fi tabaqat-al-atibba” (Księga źródeł informacji o klasach lekarzy), dotyczącej medycyny, w której Ibn Usaybi’a (2) przypisuje ten proces jeszcze starszemu autorowi, Ibn Bakhtawayhi, o którym nic nie wiadomo.

Imperium Mogolów w Indiach wykorzystuje sztafety jeźdźców do sprowadzania lodu z Hindukuszu do stolicy na terenie współczesnego Delhi, który używano do produkcji kulfi (3), popularnego do dziś na subkontynencie indyjskim mrożonego deseru mlecznego, często opisywanego jako „tradycyjne indyjskie lody”.



XVI-XVII w.

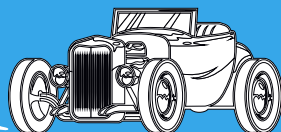
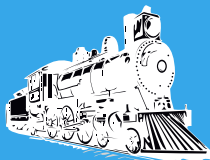
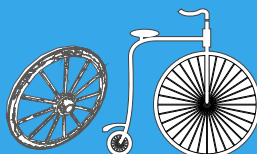
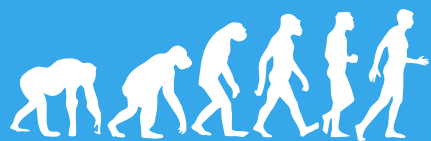
Pomimo wiedzy o efekcie chłodzenia występującym, gdy dodaje się sól do lodu, w Europie dopiero w drugiej połowie XVII wieku zaczęto wytwarzać sorbety i lody przy wykorzystaniu tego procesu. O wprowadzeniu lodów do nowożytnej Europy krąży wiele niepotwierdzonych faktograficznie przekazów. Zastugi są czasami przypisywane Marco Polo podróżującemu do Chin, chociaż nie ma o tym wzmianki w żadnym z jego pism. Według innej legendy włoska księżna Katarzyna Medycejska wprowadziła aromatyzowane lody sorbetowe do Francji, kiedy przywozła ze sobą włoskich kucharzy do Francji po poślubieniu księcia Orleanu (Henryka II Francuskiego) w 1533 roku. W rzeczywistości jednak we Francji w okresie Medyceuszy nie było żadnych włoskich kucharzy, a ponadto wiadomo, że lody znane były we Francji jeszcze przed narodzinami Katarzyny Medycejskiej. Angielski król Karol I był podobno pod takim wrażeniem „zamarznętego śniegu”, że zaoferował swojemu własnemu producentowi lodów dożywotnią emeryturę w zamian za utrzymanie receptury w tajemnicy, by lody mogły być królewskim przywilejem. Według „L'Isle des Hermaphrodites” praktyka chłodzenia napojów lodem i śniegiem pojawiła się w Paryżu, zwłaszcza na dworze, już w XVI wieku. Narrator zauważa, że przechowywano lód i śnieg, które później dodawano do wina. Praktyka ta rozwijała się za panowania Ludwika XIII i była prawdopodobnie krokiem w kierunku stworzenia lodów, jakie znamy.

1660–92

We Francji w 1660 r. Ludwik XIV zezwolił na otwarcie pierwszej w Paryżu lodziarni pod nazwą „Cafe Procope”. Od tego czasu lody były serwowane przez cały rok, a nie tylko w sezonie letnim. We Francji w „Catalogue des Marchandises rares”, wydany w Montpellier przez Jeana Fargeona, wymieniono jakiś rodzaj mrożonego sorbetu, chociaż jego skład nie został podany, Fargeon określił, że był on spożywany w stanie zamrożonym przy użyciu pojemnika, który został zanurzony w mieszaninie lodu i saletry. W 1682 roku publikacja pt. „Le Nouveau confiturier françois” podaje przepis na rodzaj lodów, zwany „neige de fleur d'orange” (tłum. „śnieg z kwiatów pomarańczy”). Pierwszy przepis na lody smakowe w języku francuskim pojawia się w 1674 roku w „Recueil de curiositez rares et nouvelles de plus admirables effets de la nature” Nicholasa Lemery'ego. Przepisy na sorbety zostały zamieszczone w publikacji „Lo Scalco alla Moderna” Antonio Latiniego z 1694 roku. Przepisy na aromatyzowane lody zaczynają pojawiać się w „Nouvelle Instruction pour les Confitures, les Liqueurs, et les Fruits” François Massialota (4), poczynawszy od wydania z 1692 roku. W 1686 r. Włoch Francesco dei Coltelli otwiera kawiarnię z lodami w Paryżu, a produkt staje się tak popularny, że w ciągu następnych 50 lat w Paryżu otwarto kolejne ćwierć tysiąca lokali tego typu.



1. Fresk z grobowca księcia Zhanghuai, 2. Stara rycina mająca przedstawiać Ibn Abu Usaybia (po prawej), 3. Współczesna forma kulki, 4. Okładka książki Massialota, 5. Ilustracja z książki L'Art de Bien Faire les Glaces d'Office Emy'ego



1768

Ukazuje się książka kucharska „L'Art de Bien Faire les Glaces d'Office” autorstwa M. Emy'ego, pierwsza poświęcona w całości przepisom na aromatyzowane lody i desery (5).

XIX w.

Lody stają się popularne i niedrogie w Anglii (6) dzięki działalności szwajcarskiego emigranta Carlo Gattiego, który założył w 1851 roku pierwsze stoisko z lodami przed stacją Charing Cross w Londynie, sprzedając lody w gałkach umieszczanych w wafelkach w cenie pensa za taką porcję. Wcześniej lody były drogim przysmakiem wyższych i bogatszych sfer. Gatti zbudował specjalny pojemnik do przechowywania lodu. W 1860 r. Gatti rozszerzył działalność i zaczął importować lód do wytwarzania lodów na dużą skalę z Norwegii. Sprowadzanie i przechowywanie lodu do produkcji tych przysmaków przestawało stopniowo być konieczne po wynalezieniu chłodziarek i lodówek. Pierwszą chłodziarkę opracował australijski wynalazca James Harrison. Jednak w większości źródeł jako wynalazca chłodziarki podawany jest bawarski inżynier Carl von Linde, który w 1871 roku zastosował system chłodzenia w browarze Spaten w Monachium, aby umożliwić produkcję piwa latem. Środkiem schładzającym był eter dimetylowy lub amoniak (eter metylowy zastosował także Harrison).

1843

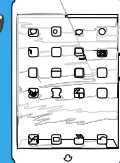
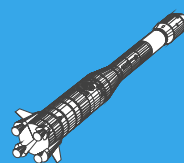
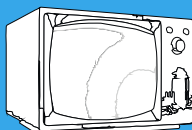
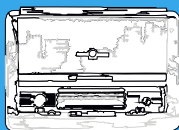
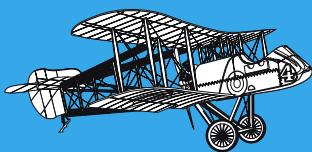
W Stanach Zjednoczonych Nancy M. Johnson wynajduje ręczną maszynę do produkcji lodu. Produkcja lodów była wcześniej bardzo pracochłonna i zajmowała często wiele godzin. Nowe urządzenie znacznie ułatwiało i przyspieszało pracę. Numer patentu na sztuczną zamrażarkę to US3254A (7). Sztuczna zamrażarka zawierała ręczną korbę, która po przekręceniu obracała i obracała dwie sąsiednie szerokie, płaskie płyty zawierające szereg otworów, które pomagały w ubijaniu lodów, czyniąc je bardziej jednorodną masą, a jednocześnie ułatwiając usuwanie kryształków lodu z wewnętrznych ścian cylindrycznego pojemnika. Drewniana wanna zawierała mieszaninę soli i pokruszonego lodu, topiąc w ten sposób pokruszony lód, ale obniżając temperaturę roztworu poniżej punktu zamarzania w wyniku obniżenia temperatury topnienia cieczy przez sól. To, w połączeniu z roztworem lodów, pobierało energię cieplną z lodów, zamrażając je.

1874

Wynalezienie napoju lodowego przypisuje się Amerykaninowi Robertowi Greenowi, chociaż nie ma jednoznacznych dowodów, że rzeczywiście był pierwszy. Tradycyjny przekaz głosi, że w bardzo upalny dzień Greenowi zabrakło lodu do napojów smakowych, które sprzedawał, i zamiast tego użył lodów waniliowych kupionych od sprzedawcy z sąsiedniego kramu, wymyślając nowy napój.

1888–1904

Ważną postacią w historii lodów jest Brytyjka Agnes Marshall, zwana „królową lodów”, autorka i popularyzatorka tego przysmaku. To w jej książce kucharskiej znajduje się pierwsza wzmianka o wypiekanych rożkach jako naczyniu na porcję lodów. Opisuje „rożki robione z migdałów i pieczone w piekarniku” (8). Rożek do lodów został spopularyzowany w Stanach Zjednoczonych na Światowych Targach w St. Louis w stanie Missouri w 1904 roku.



1920–26

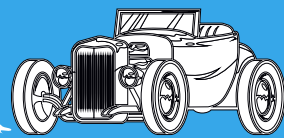
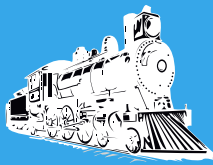
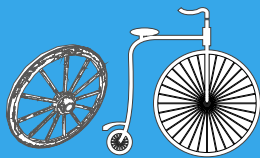
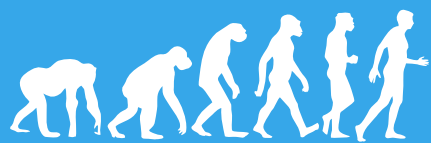
Pierwsze lody na patyku powstały podobno przez przypadek – w wyniku zamarznięcia szklanki z oranżadą z wetkniętym w nią mieszadłem, pozostawionej za oknem, na mrozie. Lody na patyku opatentował Amerykanin Harry Burt w 1924 r. W 1920 roku, kiedy Burt prowadził lodziarnię i firmę cukierniczą w centrum amerykańskiego miasta Youngstown, opracował polewę czekoladową, która była „kompatybilna” z lodami. Według zeznań w procesie o prawa do patentu, złożonych przez wdowę po Burcie ponad dekadę po jego śmierci, Burt wpadł na pomysł włożenia drewnianego patyczka do pokrytego czekoladą batonika z lodami w 1920 lub 1921 roku. Wkrótce pojawił się produkt jego firmy Good Humor (9) o nazwie Eskimo Pie. 30 stycznia 1922 r. Burt złożył wniosek patentowy, który obejmował proces i aparaturę produkcyjną, a także sam produkt. W październiku 1923 r. przyznano mu patenty zarówno na proces, jak i sprzęt produkcyjny, ale nie na produkt. To wydarzenie otworzyło drogę do bitew prawnych w nadchodzących latach. W 1925 roku Good Humor złożył pozew przeciwko Citrus Products Company i Popsicle Corporation, które sprzedawały własne wersje lodów na patyku. Negocjacje z Citrus Products Company w celu uzyskania licencji od Good Humor zakończyły się niepowodzeniem, a Good Humor dobrowolnie wycofał pozew w 1926 roku.

1926

Ważnym wydarzeniem w historii lodów było wprowadzenie lodów miękkich, które zawierają więcej powietrza, co obniża koszty. Maszyna do lodów miękkich napędzania łożek lub naczynek z króćca. Urządzenie takie opatentował w 1926 r. Charles Taylor z Buffalo w stanie Nowy Jork. Założona przez niego firma Taylor Company (10) nadal produkuje maszyny do lodów. Lody miękkie mają historycznie więcej ojców. W 1936 roku Tom Carvel opracował własną tajną formułę lodów miękkich, a także opatentował maszyny do lodów o bardzo niskiej temperaturze. Do wynalazku lodów miękkich, zwanych w USA „soft serve”, prawa rości sobie również firma Dairy Queen, która recepturę opracowała w 1938 roku. W latach 60. XX wieku producenci maszyn do lodów wprowadzili do automatów zmechanizowane pompy powietrza, zapewniając lepsze napowietrzanie.



6. Uliczna sprzedaż lodów w Londynie w XIX wieku, 7. Nancy M. Johnson i ilustracja jej maszyny do lodu z wniosku patentowego, 8. Rozdział książki Agnes Marshall poświęcony wypiekaniu wafelków w kształcie różka, 9. Stara reklama lodów na patyku firmy Good Humor, 10. Współczesne miękkie lody firmy założonej przez Charlesa Taylora



WYNALEZKÓW Typy lodów i desrów lodowych

Lody mleczne

Wytwarzane z mleka, żółtek, jaj, cukru lub innych dodatków smakowych. Jako substancje smakowo-zapachowe, oprócz wanilii, można stosować: kakao, kawę, orzechy, owoce świeże lub przetwory owocowe. Przy produkcji lodów mlecznych o smaku kakaowym, kawowym, owocowym itp. substancje smakowe i zapachowe, takie, jak kakao, kawa lub owoce, dodaje się również w procesie schładzania.

Parfiany – lody śmietankowe

Od lodów mlecznych różnią się tym, że do ich produkcji używa się nie mleka, ale śmietanki. Aromatyzowane są tak samo jak i mleczne. W końcowej fazie zamrażania kompozycji mlecznej dodaje się śmietankę o zawartości 10...33 proc., po czym kompozycję poddaje się dalszemu napowietrzeniu i zamrożeniu. Podobnie jak lody mleczne, lody śmietankowe można sporządzać z dodatkiem kakao, kawy, orzechów, owoców i ich przetworów.

Lody owocowe

Lody owocowe produkuje się na bazie wodnego roztworu cukrowego. Otrzymany roztwór zagotowuje się, a następnie poddaje się ochładzaniu i zamrażaniu. W czasie zamrażania dodaje się uprzednio umyte, oczyszczone i rozdrobnione owoce lub przetwory owocowe. Po dodaniu owoców masę lodową poddaje się dalszemu procesowi napowietrzania i zamrażania, a następnie hartowaniu.

Lody mleczno-owocowe

To połączenie owocowego orzeźwienia z kremową konsystencją lodów mlecznych. Do masy mlecznej dodaje się zarówno zmiksowane owoce, jak też ich kawałki.

Lody cassate

Lody cassate otrzymuje się przez formowanie kilku rodzajów lodów, kremowej bitej śmietany i owoców. Do formowania służą specjalne formy w kształcie półkuli z otwieranym płaskim dnem. Poszczególne rodzaje mas lodowych nakłada się do formy tak, aby po przekrojeniu jej na połowę widoczne było równomierne, półkuliste rozmieszczenie tych mas. Najczęściej stosuje



się zestaw składający się z trzech rodzajów lodów oraz pozostałych składników. Najlepszy efekt uzyskuje się, stosując jako zewnętrzną warstwę lody mleczne lub śmietankowe o smaku waniliowym i kolejno następne warstwy lodów owocowych, kremu bita śmietana połączonych z owocami i lodów o smaku kakaowym lub kawowym.

Sorbety

Przecier z soku owocowego, gotowanego z cukrem, następnie ucieranego aż do otrzymania konsystencji półpłynnej „pomadki”. Odmianą z dodatkiem nabitka, jednak w znacznie mniejszej proporcji niż lody mleczne, jest sherbet.

Kulfi

Te tradycyjne indyjskie lody wytwarzane są poprzez powolne podgrzewanie słodzonego mleka, aż cukier się skarmelizuje, a mleko skondensuje. Mieszanina jest następnie zamrażana w foremkach w kształcie stożków bez uprzedniego ubijania. W rezultacie kulfi jest gęstsze i bardziej przypomina krem niż znane nam lody.

Dondurma

Lody kremowe i na tyle gęste, że można je jeść widelcem i nożem. Ten turecki przysmak zawiera mastyks – (jadalny) organiczny klej wytwarzany z żywicy drzewa mastyksowego. ■

M.U.

Klipsch JUBILEE (2)

W pierwszej części artykułu poświęconego konstrukcji Jubilee przedstawiliśmy jej genezę, formę i jeden z dwóch przetworników – niezwykle, średnio-wysokotonowy, promieniujący przez dużą tubę. W drugiej części przenosimy się do sekcji niskotonowej, a także objaśniamy zasadę działania całego systemu, „kontrolowanego” i korygowanego przez zwrotnicę aktywną.

Klipsch Jubilee to pod wieloma względami niekonwencjonalny zespół głośnikowy. Łączy różne światy – elementy klasycznej techniki tubowej, tradycyjny styl wzorniczy, ultranowoczesny driver średnio-wysokotonowy i zwrotnicę aktywną z DSP.

Tuba, która bez wspomagania korekcją wzmacniałaby najniższe częstotliwości aż do granicy pasma akustycznego, musiałaby mieć ogromne wymiary. Już Jubilee są egzotycznie duże, a przecież nie jest to konstrukcja, która za pomocą samej obudowy sięga liniowo do 20 Hz. Coś jej w tym pomaga...

Głośników niskotonowych z zewnątrz nie widać, co jest typowe dla największych obudów Klipscha (Klipschorn, La Scala), w których tuba jest uformowana przed głośnikiem.

Promieniowanie od tylnych stron membran pary 30-cm głośników, zainstalowanych jeden nad drugim w wewnętrznej pionowej przegrodzie (znajdującej się niedaleko za centralną częścią frontu), pobudza układ bas-refleks, zajmujący mniej więcej połowę całej objętości obudowy. Jego otwory też są ukryte wewnątrz konstrukcji, trzy tunele sięgając w głąb komory bas-refleks, zostały zainstalowane w tej samej przegrodzie co głośniki, a więc ich promieniowanie, razem z promieniowaniem przedniej strony membran głośników, biegnie przez tubę, która jest symetrycznie rozdzielona na dwie sekcje – lewą i prawą.



Tuba niskotonowa jest „pozaginana”, aby można ją było „spakować” w obudowie o przynajmniej względnie „ustawnej” formie. Mimo to tuba nie jest bardzo długa – od głośników do wylotów jest ok. 1 metra, a łączna powierzchnia obydwu okien też nie jest tak duża, aby zapewnić wzmocnienie najniższych częstotliwości. Tutaj z pomocą przychodzi bas-refleks, dostrojony do ok. 23 Hz, jednak nie będzie on w stanie dać takiego podbicia, jakie daje tuba w zakresie kilkuset herców. Czy tuba działa wzmacniając na promieniowanie z bas-refleksu, skoro znajduje się przed jego „wylotami”? Tak, ale... jak już stwierdziliśmy, nie na częstotliwości najniższe. Będzie więc potrzebny jeszcze jeden element całej układanki – korekcja elektryczna. Według szacunków i pomiarów obudowa tubowa Jubilee działa wzmacniając powyżej 100 Hz, a same głośniki są filtrowane



Po zdjęciu maskownic zobaczymy wyloty tubowych labiryntów i wzmacniające je wieńce, ale głośniki są ukryte głębiej

dolnoprzepustowo powyżej 300 Hz (gdzie przetwarzanie przejmuje średnio-wysokotonowy).

Jubilee to rzadko spotykany system ze zwrotnicą aktywną.

Sygnał ze źródła (przedwzmacniacza albo odtwarzacza z regulowanym poziomem sygnału) dostarczany jest do specjalnego zewnętrznego urządzenia, które znajduje się w komplecie. Tutaj następuje podział i korekcje, i dopiero stąd niezależne sygnały przygotowane dla obydwu sekcji kolumny biegną do końcówek mocy – dlatego potrzebne są dwie „na stronę”. Wzmocniony w nich sygnał dociera już bezpośrednio do przetworników, bez typowego dla kolumn pasywnych filtrowania w zwrotnicy opartej na elementach biernych, bo zadanie to wykonała wcześniej zwrotnica aktywna.

Jubilee, w odróżnieniu od kolumn aktywnych, nie mają własnych końcówek mocy (trzeba się w nie zaopatrzyć samodzielnie, dlatego też zwrotnica aktywna nie mogła zostać zainstalowana wewnątrz kolumny, bo nie wsadzimy tam nawet kupionych wzmacniaczy).

Systemy z aktywną zwrotnicą występują rzadko, ale nie są nowością. Tak jak w kolumnach w pełni aktywnych, aktywna zwrotnica ułatwia przeprowadzenie dokładnego filtrowania i korekcji, ustalenie optymalnych charakterystyk amplitudowych i fazowych poszczególnych sekcji, głównie w celu uzyskania najlepszej charakterystyki wypadkowej całego zespołu.

Filtry aktywne działają na sygnale niskopoziomym i korygując go, nie wprowadzają strat energii,



Strzałki ilustrują bieg ciśnienia w tubowym labiryncie, prowadzącym od głośników i wylotów tuneli bas-refleks (znajdują się poniżej)



Aktywna zwrotnica Jubilee to zewnętrzne eleganckie urządzenie. Dwa regulatory z prawej strony pozwalają niezależnie ustalić poziom sygnału wyjściowego dla końcówek mocy podłączonych do obydwu sekcji, dzięki czemu mogą one być różnego typu i różnej mocy

takich jak filtry biernie, które wyłącznie tłumią sygnały już wzmacnione.

Nowoczesne zwrotnice aktywne coraz częściej opierają się na działaniu procesorów DSP, a więc operują na sygnale cyfrowym (i tak jest też w Jubilee), co jeszcze zwiększa elastyczność działania, chociaż wymaga przekonwertowania analogowego sygnału wejściowego na cyfrowy... a potem z powrotem na analogowy.

Konstrukcje w pełni aktywne stały się tym bardziej opłacalne pod względem korzyści akustycznych i funkcjonalnych w relacji do kosztów, gdy powiązano je z transmisją bezprzewodową i przyjmowaniem sygnałów cyfrowych ze źródeł mobilnych. Zastosowanie samej aktywnej cyfrowej zwrotnicy, bez pełnej integracji systemu, zainstalowania wzmacniaczy, uruchomienia strumieniowania, nawet bez udostępnienia wejść cyfrowych, wciąż ma swoje zalety, nie jest to jednak rozwiązanie ultranowoczesne z punktu widzenia współczesnego użytkownika, zmusza do większego rozbudowania systemu (dwa razy więcej końcówek mocy, więcej kabli). Dlaczego więc Klipsch zdecydował się na takie rozwiązanie? Czy nie mógł opanować charakterystyk ogólnie prostego układu dwudrożnego klasyczną zwrotnicą bierną, co przecież ćwiczył z powodzeniem od 75 lat? Wyprowadzeniem zwrotnicy aktywnej na zewnątrz nie wzbogacił systemu o żadną dodatkową funkcjonalność.

Są dwa ważne powody zastosowania zwrotnicy aktywnej w Jubilee.

Już w opisie średnio-wysokotonowego wspomnieliśmy, że niezależnie od jego fenomenalnej konstrukcji i możliwości, jego charakterystyka (nawet wraz

z odpowiednią tubą, którą dodał Klipsch) jest daleka od liniowości i wymaga silnej korekcji.

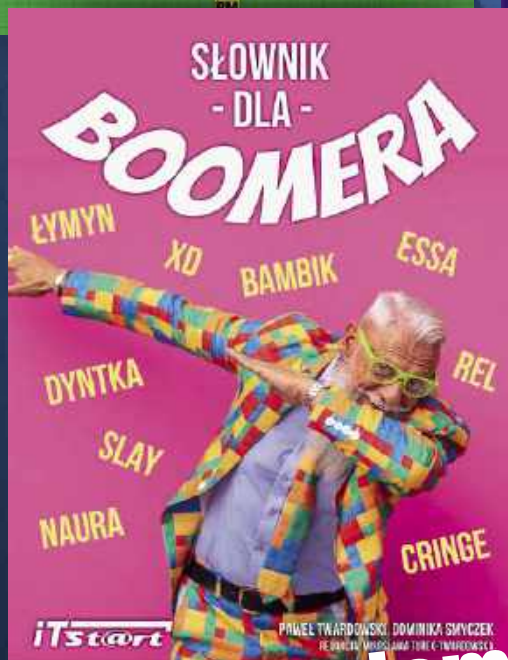
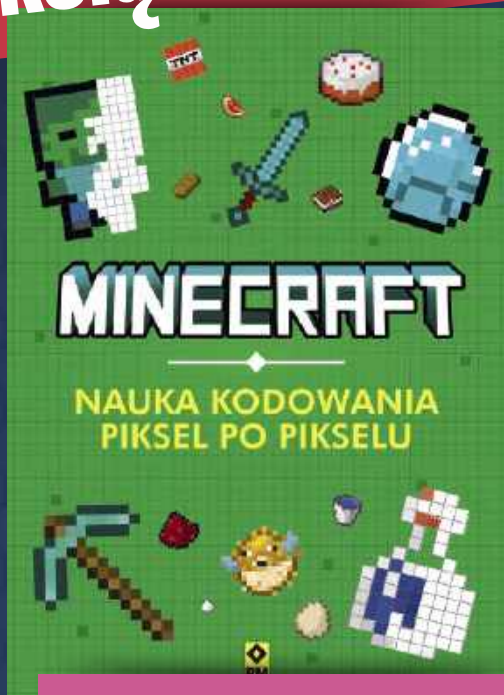
Charakterystyka przed korekcją dość łagodnie, ale konsekwentnie opada powyżej 2 kHz, różnica między średnim poziomem w zakresie 500 Hz...2 kHz a średnim poziomem w zakresie 10...20 kHz wynosi aż 15 dB; nawet taką charakterystykę można by skorygować filtrem biernym, tłumiąc średnie tony do poziomu wysokich, ale lepiej podciągnąć wysokie do średnich, co może przeprowadzić tylko filtr aktywny, korygując przy okazji pomniejsze mankamenty charakterystyki.

Wyda się jednak, że tylko z tego powodu konstruktor nie zdecydowałby się na filtr aktywny; kluczowa jest korekcja w drugiej sekcji, na dolnym skraju pasma. Dzięki niej charakterystykę rozciągnięto do granicy zakresu akustycznego. Bez pomocy aktywnej zwrotnicy opada ona poniżej 100 Hz, mając przy 20 Hz spadek 12 dB; to i tak bardzo dobry wynik, zapewniający w zasadzie pełną słyszalność aż do tej częstotliwości w warunkach normalnego pomieszczenia (choć nie bez wpływu jego rezonansów). Klipsch jednak niemal wyrównał charakterystykę 20 Hz, co wymagało silnej korekcji „dodatniej” poniżej 100 Hz; korygowanie za pomocą filtrów biernych tak, aby poziom przy 100 Hz został stłumiony o 12 dB względem 20 Hz, byłoby trudne (duże elementy filtrów, straty mocy) i prowadziłoby do znacznego obniżenia efektywności.

Wyniki pomiarów przedstawione charakterystykami poszczególnych sekcji, zarówno przed korekcją, jak i po korekcji, na różnych osiach, przedstawiono w pełnym teście AUDIO 4/2024. ■

Andrzej Kisiel

Książki w Ulubionym Kiosku



z rabatem do 30%

Zobacz pełną ofertę - ponad 500 tytułów
na www.UlubionyKiosk.pl

epresa.pl ed9c6b928a