



nr 6. czerwiec 2023

e-suplement www.mt.com.pl



Tu przejrzysz
i kupisz ten numer

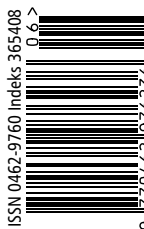
NEWS 24/7
przeglądaj codziennie
na swoim smartfonie

młody **m.technik**

Ciekawi świata są zawsze młodzi



ŁAMANIE GŁOWY **Czy mózg się obroni?**



ISSN 0462-9760 Indeks 365408
9 4770462 197623
cena: **14,90 zł** (w tym 8% VAT)

FIZYKA W SZKOLE
Rozpad promieniotwórczy (2)

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



**do 6* wydań
gratis!**

* Cena prenumeraty rocznej wynosi 163,90 zł.

Przy zamówieniu prenumeraty dwuletniej w cenie 268,20 zł

oszczędność wynosi równowartość sześciu wydań „Młodego Technika”

Wszystkie opcje prenumeraty i e-prenumeraty znajdziesz na stronie

www.UlubionyKiosk.pl

prenumerata@avt.pl

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

konto 18 1050 1012 1000 0024 3173 1013

eprasa.pl 0e17570fd7



Prenumerata
dla szkół i placówek
oświatowych
30% taniej!

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

Roczny dostęp online
kosztuje 100,00 zł

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

Ostatnia reduta w szarych komórkach

Wielka inwazja inwigilacyjna trwa. Jedni tłumaczą zamachy na naszą prywatność, na dane i szczegóły dotyczące najintymniejszych spraw tym, że to dla naszego dobra. Inni po prostu robią na tym pieniądze i utrzymują władzę. W tej sytuacji ostatnią twierdzą, w której możemy się ukryć, jest to, co wydaje nam się bezpiecznie zamknięte w czaszce. Do naszych myśli się nie dobiorą – myślimy (nomen omen).

Otóż jak wynika z rozlicznych informacji o badaniach i eksperymentach na mózgu, których przegląd znaleźć można w tym numerze „Młodego Technika”, także i ta ostatnia reduta prywatności może być, jeśli już nie jest, na celowniku. Istnieją podstawy rozwiązań technicznych pozwalających na wgląd w treść mentalnej aktywności człowieka. Powstały już pierwsze obrazy narysowane na podstawie skanowania fal mózgowych lub pochodzące z chipów podłączonych do układu nerwowego.

Czy sztuczny mózg „wyłoniłby” z siebie świadomość?

Mózg jest od lat intensywnie mapowany w celu rozpoznania funkcji i powiązań każdego składającego się nań biotrybika. Techniki te znajdują od dłuższego czasu zasto-

sowanie w medycynie, choć daleko jeszcze do pełnej i szczegółowej mapy mózgu człowieka. Mamy za to już całościowe obrazy układów nerwowo-mózgowych prostszych stworzeń.

Ów pęd do hakowania i rozpracowania mózgu może dla wielu brzmieć niepokojąco. Jednak gdy dochodzimy do kwestii wyższego poziomu, np. świadomości, czegoś, co w końcu z mózgiem jest tradycyjnie silnie związane, nauka, psychologia i fizyka (bo od niedawna i ta dziedzina ma w tej sprawie wiele do powiedzenia) są nieco bezradne, przerzucając się teoriami i hipotezami, które trudno weryfikować.

Dochodzimy do zagadnienia sztucznego mózgu, czyli prób zbudowania dokładnej kopii tego, co mamy pod czaszką. Takiej kopii, która działałaby tak samo i była tak samo wydajna, a może nawet, po dołączeniu wspomagania AI, wydajniejsza niż nasze półtora kilo pofałdowanej tkanki. W pewnym sensie zbudowanie wiernego „cyfrowego bliźniaka” ludzkiego mózgu byłoby bardzo cenne – dowiedzielibyśmy się, czy świadomość „wyłania” się z materii, czy jednak jest czymś innym. To frapująca kwestia, ale zapewne przyjdzie trochę poczekać na odpowiedzi w tej sprawie.

Mirosław Usidus

Spis treści

Temat numeru:

Lamanie głów. Czy mózg się obroni?

- 24 • Ostatnia rubież prywatności – czy na pewno? Podglądacze myśli już na horyzoncie
- 29 • Pokazać wszystko – jak się wszystko łączy się ze wszystkim – i jak to wszystko działa. Postępy w mapowaniu mózgu
- 34 • Świadomość – czym jest i skąd się bierze. Z materii zrodzona czy wręcz przeciwnie?
- 41 • Sztuczny mózg – idea, która nie chce odejść. Organoid do organoida

Technika

- 8 Info Zoom
- 16 Dodaj do obserwowanych Horyzonty mgłą spowite
- 17 • Supersamoloty, nie tylko hipersoniczne. Czy *Top Gun* to tylko film?
- 20 • Napęd oparty na toroidzie wokół śmigła. Po co tyle hałasu?
- 22 • Roboty wspinające się po pionowych powierzchniach. Magnetyczny wspinacz
- 44 Raport MT: Proton – najbardziej skomplikowana rzecz, jaką znamy. Tym bardziej inny, im bardziej mu się przyglądamy
- 54 Nasi idole – liderzy innowacji: Polski wkład w światowe kuchenne rewolucje – Stefan Poptawski

m.technik

- 58 e-Technologie: Nowa architektura sieci – splinternetowe alternatywy. Fatum wielkiego podziału wraca do internetu

Szkoła

- 62 Chemia inna niż w szkole: Pierwiastkowy high life
- 66 Edukacja przez szachy: Oliwia Kiotbasa wicemistrzynią Europy
- 70 MT studiuje: Chemia budowlana
- 72 Fizyka bez granic: Czy substancja promieniotwórcza straci z czasem swoje właściwości? (42)
- 74 Matematyka z ludzką twarzą: Apolonia Inteligentna. Magiczny kafelek
- Klub i Szkoła Wynalazców
- 78 • Szkoła Wynalazców, dozwolone do lat 15
- 79 • Klub Wynalazców, bez ograniczeń wieku
- 80 • Vademecum Młodego Wynalazcy
- 83 Pomysły genialne, zwariowane i takie sobie
- 84 Na warsztacie: Ekranoplan MT-2023
- Odkryj historię wynalazków
- 88 • Przekładnie mechaniczne
- 92 • Klasyfikacja przekładni mechanicznych
- 93 Konkurs „As Majsterkowania”

Hobby

- 94 Akademia audio: Monachium High-End 2023. TUZIN głośnych POMYSŁÓW
- 2 Prenumerata
- 3 Od wydawcy
- 6 Listy, Facebook
- 99 Sędziwy Technik – 100 lat temu prasa pisała

Nasz mózg to nasza ostatnia reduta prywatności, ostatni bastion, w którym możemy się skryć, pewni, że nikt naszych myśli nie może inwigilować. To jednak może się zmienić, bo techniki „czytania” a nawet hakowania mózgu nie stoją w miejscu. Powstają coraz dokładniejsze jego mapy, a my zastanawiamy się, czym jest świadomość, bo maszyny zdają się do niej zbliżać.

W tym wydaniu MT m.in.:

- **Horyzonty mgłą spowite: Supersamoloty, nie tylko hipersoniczne.** Skunk Works, legendarne biuro projektowe Lockheed Martina, współpracowało z twórcami sequela *Top Gun* przy tworzeniu filmowego hipersonicznego samolotu Darkstar. Jednak ostatnie wydarzenia kładą się zastanowić, czy to była jedynie filmowa fikcja.
- **Nasi Idole: Stefan Poptawski.** Polski wkład w światowe kuchenne rewolucje.
- **e-Technologie:** Fatum wielkiego podziału wraca do internetu.

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

• Adres wydawnictwa:
03-197 Warszawa, ul. Leszczyńska 11,
tel. 22 257 84 98, faks: 22 257 84 00,
e-mail: avt@avt.pl, http://www.avt.pl

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

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

• Redaktor Wydania:
Wojciech Marciniak

• DTP:
MAD Sp. z o.o.
e-mail: dtp@mad.media.pl

• Konsultacja graficzna:
Małgorzata Jabłońska

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

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



ŁAMANIE GŁOWY

Czy mózg się obroni?

23



44

Proton niezmiennie zmienny

List miesiąca

Zestaw Burmestra

W styczniowym numerze „Młodego Technika” z 2023 na str. 69 ukazał się interesujący artykuł Profesora Michała Szurka pod tytułem „Krzywe, krzywki i krzywki”. Autor w początkowej części tego artykułu pisze o powszechnie używanych przed kilkudziesięciu laty krzywkach kreślarskich. Oto cytat ze wspomnianego artykułu: „Wiem, że takie krzywki były w użyciu jeszcze w latach trzydziestych XX wieku, Ciekawe, kto je wymyślił i opatentował. Szukałem wiadomości w Internecie, ale nie znalazłem.

Kierując się bezinteresowną chęcią pomocy Panu Profesorowi i uzupełnienia wiedzy Czytelników, chciałbym odpowiedzieć na to pytanie i dodać kilka krótkich uwag. Wspomniane krzywki, pokazane w cytowanym artykule na fotografii 1, wymyślił oraz opatentował niemiecki uczyony Ludwig Burmester i te przyrządy są znane jako „zestaw Burmestra”, albo „krzywki francuskie”. Burmester urodził się 5 maja 1840 r. Był matematykiem i zajmował się głównie geometrią oraz kinematyką ruchów na powierzchniach płaskich. Badał m.in. tory zakreślane przez wybrane punkty mechanizmów płaskich, składających się z dźwigni połączonych ze sobą za pomocą przegubów. Stosował i rozwinął dział geometrii nazywany geometrią rzutową. Burmester początkowo pracował jako nauczyciel w Łodzi, a następnie przeniósł się do Drezna. Zmarł 20 kwietnia 1927 r.

Krzywki Burmestra były powszechnie używane jeszcze w latach siedemdziesiątych i na początku lat osiemdziesiątych XX w. Najczęściej korzystali z nich kreślarze, uczniowie techników i studenci politechnik, a także plastycy i projektanci mody. Krzywki wykonywano z różnych materiałów, głównie z drewna i tworzyw sztucznych. Wspomniana w artykule „dykta” to stara, potoczna nazwa sklejki, czyli materiału wytworzonego z cienkich warstw drewna, posmarowanych klejem, nałożonych na siebie i sprasowanych. Słoje w sąsiednich warstwach są skierowane do siebie prostopadle, co zwiększa wytrzymałość tego materiału i odporność na odkształcenia pod wpływem wilgoci oraz czyni go bardziej izotropowym. Następnie stosowano tworzywa sztuczne, celuloid i polistyren. Warto wspomnieć, że celuloid był też używany do produkcji pierwszych taśm filmowych, ale zrezygnowano z niego ze względu na łatwopalność i liczne pożary. Bezpieczniejszym materiałem, stosowanym w późniejszych latach, był polistyren. Często brzegi krzywków miały uskok, który ułatwiał ręczne wykonywanie rysunków technicznych tuszem kreślarskim. Używane wówczas do tego celu przyrządy – grafiony, miały w górnej części uchwyt a w dolnej dwa sprężyste ostrza, których odległość była



regulowana nakrętką. Między ostrzami umieszczano kroplę tuszu. Przypadkowe zetknięcie tego tuszu np. z brzegiem linijki lub krzywika powodowało uciążliwe kleksy.

Wspomniana w cytowanym artykule stalówka „rondówka”, bardziej znana pod nazwą „redisówka”, i służyła głównie do pisania pismem technicznym. Jej dolny koniec stanowiła okrągła płytka, zapewniająca stałą grubość linii, ale redisówki były też używane do kreślenia linii na rysunkach technicznych i dawały mniej kleksów. Zalecane użycie krzywików polegało na wybraniu takiego fragmentu brzegu, który przechodził przez co najmniej trzy wcześniej wyznaczone punkty kreślonej krzywej, np. fragmentu wykresu. Takie postępowanie umożliwiało otrzymanie „bardziej gładkiej” krzywej, czyli niemającej widocznych załamań. Linie tworzące brzegi krzywików składały się z fragmentów krzywych o zmiennych promieniach krzywizny – głównie były to krzywe Eulera i klotoidy. Później krzywki o ustalonym kształcie zostały zastąpione przez krzywki bardziej uniwersalne, czyli długie taśmy wykonane z giętkiego tworzywa sztucznego. Takie taśmy można było łatwo ukształtować palcami, tak żeby przechodziły przez wyznaczone punkty. Prowadząc wzduż tak dogiętych taśm dowolny przyrząd piszący, wykreślało się wyznaczoną krzywą.

Ostatecznie krzywki zostały wyparte przez rozwój technologii informacyjnej i edytory graficzne, np. powszechnie dziś stosowane programy typu CAD. Krzywki Burmestra są dziś raczej ciekawostką, używaną przez niektórych plastyków i hobbystów. I jeszcze drobna uwaga. W cytowanym artykule na rysunku 1, najmniejszy, środkowy krzywik to krzywik eliptyczny. Po prawej stronie (patrzac z pozycji czytelnika) jest pokazany krzywik paraboliczny, a dolny krzywik jest krzywikiem hiperbolicznym. Kończąc, przesyłam serdeczne pozdrowienia dla Pana Profesora wraz z życzeniami wielu inspirujących i ciekawych artykułów.

Stanisław Bednarek

Wydział Fizyki i Informatyki Stosowanej Uniwersytetu Łódzkiego

Komentarz autora artykułu Michała Szurka:

Istotnie, historia tak niegdyś niezbędnego narzędzia, jak krzywki, jest bardzo ciekawa. Trochę szkoda, że właściwie podzieliły już one los suwaka logarytmicznego – chociaż jeszcze się trzymają i do rysowania ładnych łuków na pewno będą się przydawać. Najlepszy dowód to to, że wciąż są produkowane.

Krzywe w ogóle są ładne. Ich czas to druga połowa siedemnastego wieku i pierwsza połowa osiemnastego – stulecie krzywych. Badano je wtedy bardzo intensywnie. Znajdowano je u starożytnych Greków i średniowiecznych Arabów, analizowano w nowym ujęciu, szukano dowodów, wynajdowano nowe (i krzywe, i dowody). Włoski uczony Guido Grandi (1671–1742) pisał w 1728 roku o krzywych jako o „kwiatach geometrii”. Niektóre miały istotnie poetyckie nazwy. Na przykład cykloida była nazywana piękną Heleną geometrii. Przypomnę, co to jest klotoida. Na drogach szybkiego ruchu, a w szczególności na torach szybkiej kolei, prosty tor nie może na zakręcie od razu przechodzić w łuk okręgu. W punkcie przejścia krzywizna doznałaby gwałtownej zmiany: z zera na... niezerową. Czy jest krzywa, która umożliwia łagodne przejście z prostej do łuku i odwrotnie? Jest. Nazywa się klotoida i z samych jej równań wynika, że spełnia nasz warunek: jej krzywizna jest wprost proporcjonalna do długości łuku. Rośnie też proporcjonalnie siła odśrodkowa. Gdy jedziemy po takim właśnie łuku, musimy powoli, ale jednostajnie obracać kierownicą (nie jest do końca poprawne, gdyż promień skrętu nie jest tak prosto związany z kątem obrotu kierownicy). Oto przedstawienie parametryczne klotoidy (liczba a jest dowolnym parametrem):

$$x(t) = \int_0^t \cos \frac{s^2}{2a} ds, \quad y(t) = \int_0^t \sin \frac{s^2}{2a} ds$$

W Internecie można znaleźć wiele zdjęć pokazujących, że łuki autostrad są istotnie klotoidalne. Zresztą, w rozporządzeniu Ministerstwa Transportu i Gospodarki Morskiej z 1999 roku właśnie klotoidy są wymieniane jako najczęściej stosowane krzywe przejściowe.



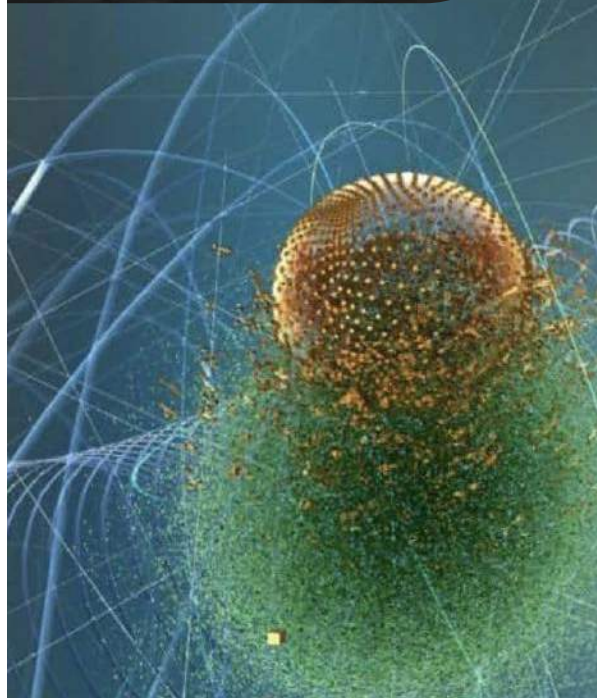
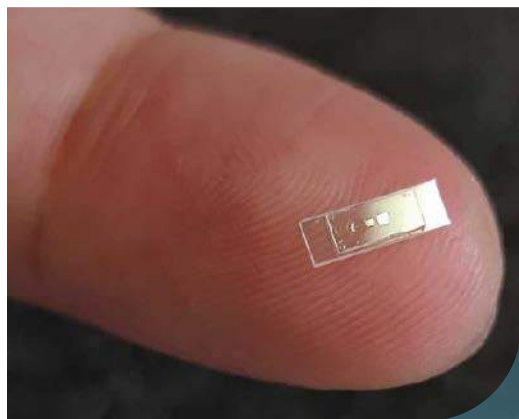
CYBERWOJNA

Chińska aplikacja zakupowa, czyli malware

Jak alarmują specjaliści w dziedzinie cyberbezpieczeństwa, Pinduoduo, chińska aplikacja typu e-commerce, szpieguje użytkowników w taki sposób, że naruszenia prywatności i bezpieczeństwa danych sięgają w niej „zupełnie nowego poziomu”. Ponadto, gdy się ją już zainstaluje, jest trudna do usunięcia z urządzenia mobilnego. W marcu aplikacja firmy PDD z korzeniami w Chinach została zawieszona w sklepie Google’a.

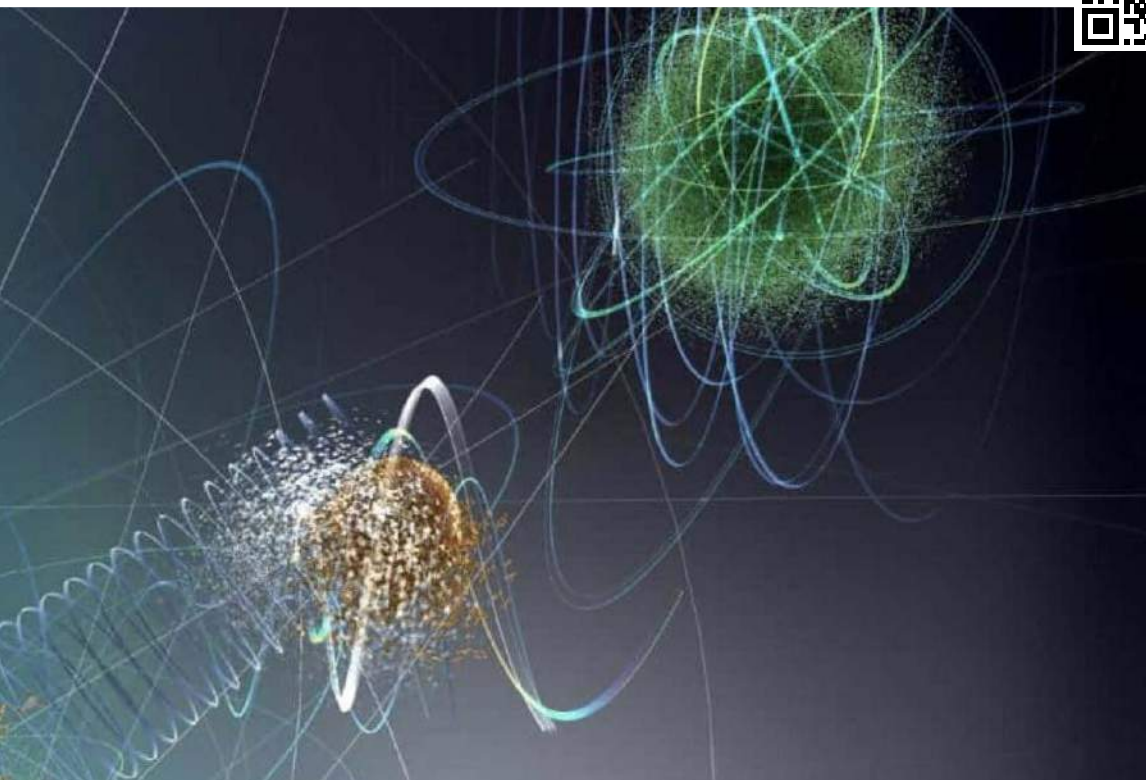
Sledztwo, do którego serwis CNN zaangażował wielu ekspertów z całego świata, potwierdziło istnienie wbudowanego w chińską aplikację oprogramowania typu malware, korzystającego z luk w zabezpieczeniach systemu operacyjnego Android. „Nie widzieliśmy jeszcze takiej aplikacji głównego nurtu, która próbowałaby dokonać eskalacji swoich uprawnień do tego stopnia, by uzyskać dostęp do rzeczy, do których nie powinna uzyskać dostępu”, komentuje w wypowiedzi dla serwisu Mikko Hyppönen, główny badacz w WithSecure, fińskiej firmy zajmującej się cyberbezpieczeństwem.

Informacje na temat szpiegowania przez Pinduoduo pojawiają się w momencie debaty toczonej na temat chińskiego TikToka, którego zablokowania w USA domaga się wielu polityków i innych osób publicznych. Zwraca się także uwagę na inną sprzedażową apkę PDD o nazwie Temu, której popularność na świecie rośnie w szybkim tempie. ■



Za pierwszy znany przypadek, gdy makroskopowa ilość materii została wprowadzona w stan superpozycji kwantowej w sposób zamierzony i precyzyjnie zaplanowany, uznano wyniki badań naukowców ze szwajcarskiej politechniki ETH w Zurychu, którzy wraz z kolegami z Danii i Niemiec umieścili kryształ szafiru zawierający 10^{16} atomów w stan kwantowy nazywany superpozycją. Uczeń pobili tym samym i to w sposób miażdżący poprzedni rekord liczby splatanych atomów, który wynosił dwa tysiące.

Zespół wykorzystał tak zwane rezonatory fal akustycznych. Są to, mówiąc w uproszczeniu, niewielkie płytki szafirowe, które wprawia się w drgania, a następnie te drgania są mierzone. Aby wywołać drgania, które reprezentują kwantowe stany superpozycji,



SPLĄTANIE

Kwantowe efekty w skali makro – to pierwszy taki eksperyment

kryształ jest sprzęgany za pomocą efektu piezoelektrycznego (tworzącego pole elektryczne, gdy materiał jest odkształcony) z obwodem nadprzewodzącym, który działa jak bit kwantowy lub kubit, znany z komputerów kwantowych. Kubit może przyjąć jeden z dwóch możliwych stanów kwantowych lub ich superpozycje. Sprzęgając kubit z kryształem, można przenieść stan superpozycji kubitu na zbiorowe drgania atomów w kryształach. Co więcej, kubit może być następnie użyty do wykrywania stanu wibracji kryształu.

Cele tego niezwykłego eksperymentu, opisanego w naukowym periodyku „Physical Review Letters”, to dokładny pomiar czasu trwania takiego efektu w wibrującej strukturze krystalicznej

i udowodnienie, że weryfikacja eksperymentalna założeń mechaniki kwantowej jest możliwa na obiektach znacznie masywniejszych niż dotychczas używane przez fizyków w tego rodzaju badaniach. Zespół schłodził kryształ, który wibrował prawie sześć miliardów razy na sekundę, do setnej części stopnia powyżej zera absolutnego, tak by zminimalizować fluktuacje termiczne. Po wprowadzeniu kryształu w określony stan kwantowy badacze wykrywali ten stan po upływie zmiennej ilości czasu za pośrednictwem kubita. Dzięki temu mogli określić, czy stan wibracyjny kryształu był rzeczywiście kwantowy, czy też można go opisać mechaniką klasyczną. W swoim eksperymencie wykryli cechy kwantowe w drganiach kryształu trwających do 40 mikrosekund. ■



TECHNIKA MOTORYZACYJNA

Rekordowa moc w małym silniku

Firma LiquidPiston skonstruowała XTS-210, niewielki silnik z rotacyjnym układem podobnym do silnika Wankla, który dostarcza tyle samo mocy, ile tradycyjne tłokowe silniki Diesla o pięciokrotnie większych rozmiarach i masie.

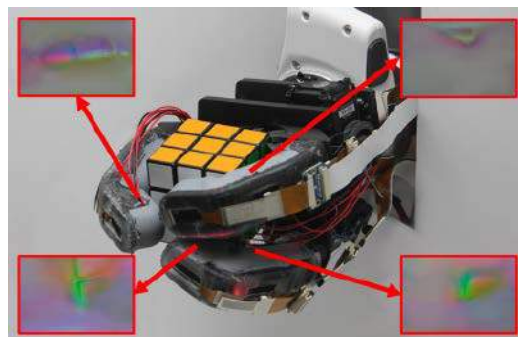
Przeznaczona do zastosowań wojskowych, lotniczych i biznesowych jednostka napędowa ma masę zaledwie 19 kilogramów i gabaryty nie większe od piłki do koszykówki.

Główne części robocze nowego typu silnika to wirnik i wał. Twórcy silnika z LiquidPiston niejako odwrócili typową konstrukcję wankla, w której tłok o kształcie zbliżonym do trójkąta obraca się w obudowie podobnej kształtem do orzeszka ziemnego. W XTS-210 to wirnik jest ukształtowany podobnie do orzeszka, a obudowa ma zbliżony do trójkąta przekrój. Dzięki takiemu układowi udało się wykorzystać dwie największe zalety silnika wysokoprężnego – dużą kompresję i bezpośredni wtrysk.

Jedną z najbardziej interesujących cech XTS-210 jest możliwość stosowania w nim wielu różnych rodzajów paliw, w tym oleju napędowego i nafty/paliwa lotniczego. Firma dąży w tej konstrukcji do uzyskania mocy około 20 kW (26,8 KM) i momentu obrotowego 29,4 Nm, w obu przypadkach przy 6500 obrotów na minutę. ■



Animacja prezentująca zasadę działania silnika XTS-210:
<https://bit.ly/3qbo48N>



ROBOTY

Robotyczna dłoń rozpoznaje od razu to, co złapała

Naukowcy z MIT opracowali robotyczną dłoń, która wykorzystuje czujniki dotykowe o wysokiej rozdzielczości, aby dokładnie zidentyfikować obiekt już po jednokrotnym jego uchwyceniu z 85-procentową trafnością. Palec tego manipulatora składa się ze sztywnego, wykonanego techniką druku 3D endoszkieletu, który jest umieszczony w obudowie i zamknięty w przezroczystej silikonowej „skórze”.

Miękka warstwa zewnętrzna wyposażona została w liczne czujniki wysokiej rozdzielczości wbudowane w miękką warstwę wspomnianej „skóry”. Czujniki, które wykorzystują kamerę i diody LED do zbierania informacji wizualnych o kształcie obiektu, zapewniają nieprzerwaną detekcję na całej długości palca. Każdy palec rejestruje szczegółowe dane dotyczące wielu części obiektu jednocześnie. Po udoskonaleniu projektu badacze zbudowali manipulator z dwoma palcami ułożonymi w kształt litery Y z trzecim palcem jako przeciwnym kciukiem. Robotyczna dłoń rejestruje sześć obrazów, gdy chwyta obiekt (dwa z każdego palca) i wysyła te obrazy do algorytmu uczenia maszynowego, który wykorzystuje je jako dane wejściowe do identyfikacji obiektu.

Szybka identyfikacja obiektu, jeśli sprawdzi się w praktycznych zastosowaniach, będzie przełomem technicznym w tej dziedzinie. Dotychczas konstrukcje tego typu polegały na instalacji wszystkich czujników na opuszkach palców robota, więc obiekt musiał znajdować się w pełnym kontakcie z nimi, aby został zidentyfikowany, co wiązało się często z potrzebą wielokrotnego chwytania. Inne projekty wykorzystują czujniki o niższej rozdzielczości rozmieszczone wzdłuż całego palca, ale nie rejestrują tak wielu szczegółów, więc znów wymagane są tu wielokrotne powtórzenia chwytów. ■



MARS

Ziemijski model marsjańskiej bazy czeka na odważnych

NASA zaprezentowała konstrukcję CHAPEA (pełna nazwa – Crew Health and Performance Exploration Analog) symulującą habitat marsjański, przystosowaną do rocznego pobytu ochotników, którzy chcieliby przetestować warunki zbliżonego do analogicznego pobytu w bazie zbudowanej na powierzchni Czerwonej Planety.

CHAPEA, zainstalowana w Centrum Kosmicznym Johnsona w Houston w Teksasie, składa się z przestrzeni o powierzchni ok. 158 metrów kwadratowych, podzielonej na dziewięć pomieszczeń, wśród których są sypialnie, wspólna łazienka i toaleta oraz przestrzeń wspólna. Obiekt został zbudowany przy użyciu wielkogabarytowych drukarek 3D, co również jest częścią testów mających na celu sprawdzenie, czy podobne metody budowy mogą być zastosowane na Marsie.

Początek eksperymentu zaplanowany jest na miesiące letnie tego roku. Podczas rocznego pobytu czworo ochotników, z których każdy ma wykształcenie naukowe, ale nie jest astronautą po specjalistycznym przeszkoleniu, będzie mieszkało w habitacie i pracowało jako zespół, wykonując zadania podobne do tych, które marsonauci będą wykonywać na Czerwonej Planecie. W trakcie

12-miesięcznego pobytu będą m.in. uprawiać sałatę na żywność, prowadzić badania naukowe, chodzić na „spacery po Marsie” i obsługiwać szereg zrobotyzowanych maszyn. Aby doświadczenie było jak najbardziej realistyczne, wolontariusze będą również zmuszeni do radzenia sobie z wyzwaniami środowiskowymi, takimi jak izolacja, ograniczenia dotyczące zasobów i awarie sprzętu. Monitorowanie stanu fizycznego i psychicznego każdej osoby będzie ważną częścią testu. Uczestnicy będą mogli utrzymywać kontakt z rodziną i przyjaciółmi, ale komunikacja będzie miała 20-minutowe opóźnienia, tak jak ma to w rzeczywistości miejsce między Ziemią a Marsem. ■





EKSPLOKACJA KOSMOSU

Europa rusza do księżyców Jowisza – Starship Muska eksploduje

Europejska Agencja Kosmiczna zainaugurowała nową misję kosmiczną, której celem jest eksploracja lodowych księżyców Jowisza, które uważane są za jedno z najciekawszych miejsc do poszukiwania życia pozaziemskiego w całym naszym Układzie Słonecznym. Sonda Jupiter Icy Moons Explorer, w skrócie JUICE, wystartowała z Gujana Space Center w Gujanie Francuskiej na pokładzie ciężkiej rakiety ESA Ariane 5.

Misja JUICE wykona przeloty wokół trzech najbardziej fascynujących księżyców Jowisza – Kallisto, Europy i Ganimedesa – zbierając przy tym dane za pomocą wielu pokładowych instrumentów naukowych. Podróż potrwa ponad osiem lat. Po drodze zaplanowanych jest kilka grawitacyjnych asyst na orbitach wokół Wenus, Marsa i Ziemi. Po dotarciu w pobliże Jowisza sonda ma w planie przeprowadzenie 35 przelotów wokół ciał w układzie tej wielkiej planety.

Wkrótce po starcie JUICE miała miejsce pierwsza próba orbitalnego lotu ogromnego statku Starship zbudowanego przez SpaceX. Sukces, jak to określają komentatorzy, był połowiczny. Start był udany, jednak nie udało się zakończyć całego lotu maszyny, ponieważ statek eksplodował kilka minut później, jeszcze w atmosferze. Firma Elona Muska będzie teraz badać przyczyny wybuchu. Kolejny start zaplanowany jest ogólnie w kilkumiesięcznej perspektywie czasowej. ■



Relacja wideo
ze startu i lotu Starshipa:
<https://bit.ly/3I0t21P>



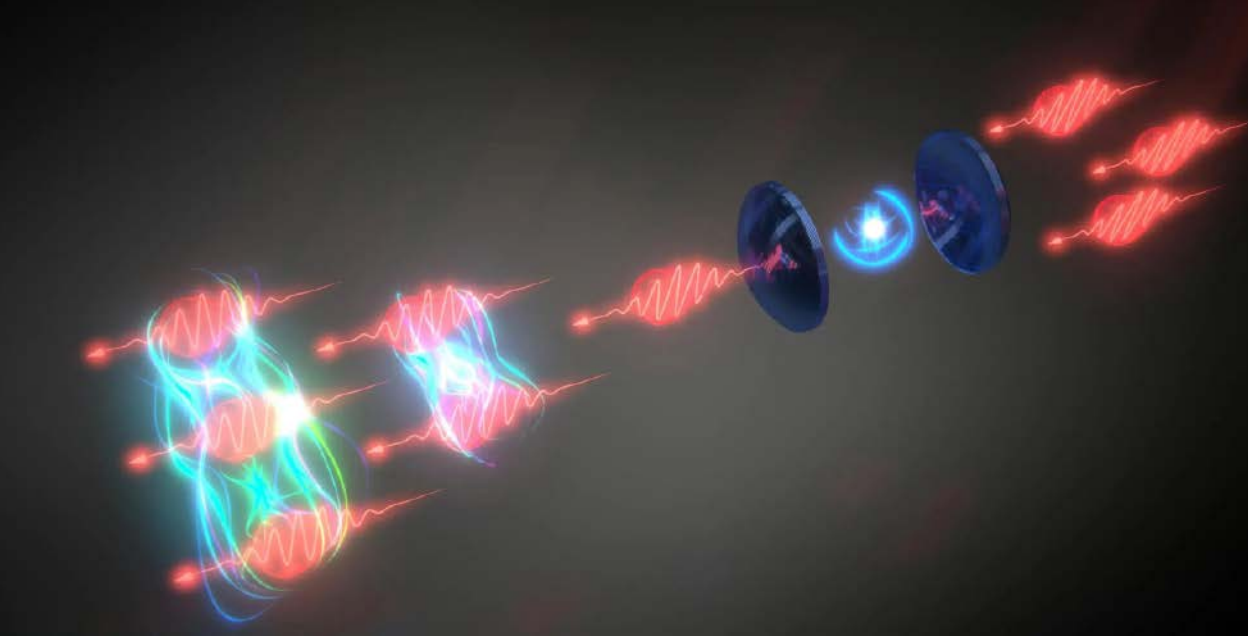
BIODRUK 3D

Drukowanie tkanek wewnątrz pacjenta

Zespół inżynierów biomedycznych z Australii opracował niewielkiego, giętkiego robota podobnego do chirurgicznego endoskopu, który może być wykorzystany do drukowania metodą 3D biokomponentów wewnątrz ludzkiego ciała. Konstruktorzy urządzenia mają nadzieję, że procedura ta pomoże usprawnić w przyszłości proste procedury medyczne, a z czasem nawet bardziej złożone operacje.

System nazwany F3DB wyposażony jest w trzysosiową głowicę drukującą, która może wyginać się i skręcać za pomocą układów hydraulicznych na końcówce miękkiego ramienia robotycznego. Dysza drukująca może wytwarzać wstępnie zaprogramowane kształty lub da się obsługiwać ręcznie, jeśli wymagane jest drukowanie bardziej złożonych lub nieokreślonych wstępnie elementów.

Najmniejszy prototyp ma średnicę przekroju ramienia około 11–13 milimetrów (mm), czyli podobnie jak komercyjne endoskopy, jednak, jak zapewniają twórcy F3DB, w przyszłości będzie można go skalować do jeszcze mniejszych rozmiarów. Obecnie, według przedstawicieli australijskiego zespołu, nie ma komercyjnie dostępnych technik podobnego typu, a tkanki, jeśli mają być wprowadzane do organizmu pacjenta, hodowane są pozaustrojowo i umieszczane operacyjnie. ■



FOTONY

Fizycy na drodze do „kwantowego światła”

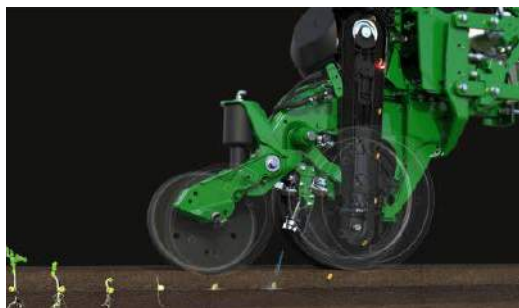
Według publikacji, która ukazała się w „Nature Physics”, międzynarodowy zespół fizyków po raz pierwszy powodzeniem przeprowadził manipulację fotonami, elementarnymi cząstkami światła. Choć być może dla laika nie brzmi to szczególnie sensacyjnie, fizycy opisują to jako wielki kwantowy przełom, który otwiera ogromne możliwości naukowe i technologiczne.

O ile fizyka doskonale radzi już sobie z kontrolowaniem kwantowo splątanych atomów, o tyle osiągnięcie tego samego za pomocą światła okazało się znacznie trudniejsze. W ramach nowego eksperymentu, zespół z uniwersytetu w Sydney i uniwersytetu w Bazylei w Szwajcarii wystrzelił pojedynczy foton i parę związanych fotonów w stronę kropki kwantowej (sztucznie

stworzonego atomu), dzięki czemu można było zmierzyć bezpośrednie opóźnienie czasowe pomiędzy fotonem samym w sobie a tymi, które były związane.

Udało się osiągnąć stymulowaną emisję i detekcję pojedynczych fotonów, a także małych grup fotonów w pojedynczym atomie, co doprowadziło do ich silnej korelacji, innymi słowy uzyskania „kwantowego światła”. „Demonstrując możliwość manipulacji stanami związanymi fotonów, zrobiliśmy ważny pierwszy krok w kierunku zaprzęgnięcia światła kwantowego do praktycznych zastosowań”, wyjaśnia w komunikacie Sahand Mahmoodian z uniwersytetu w Sydney. Kolejne kroki to wykorzystanie tego podejścia w budowie lepszych komputerów kwantowych. ■

120 miliremów wynosi dawka promieniowania otrzymywana przez osoby pracujące na nowojorskiej stacji Grand Central, co jest poziomem, który byłby przekroczeniem norm w elektrowni jądrowej.



AGROTECHNIKA

Ziarnko do ziarnka w ekspresowym tempie

Dzięki opracowanej przez firmę John Deere nowej technice ExactShot rolnicy mogą używać robotów elektrycznych do sadzenia i nawożenia nasion w niezwykle szybkim tempie przy zachowaniu wysokiej precyzji. Według udostępnionej przez firmę informacji system pozwala na wysianie 6600 nasion w ciągu zaledwie trzech sekund.

System ExactShot składa się z 54 modułowych robotów 56v i czujników służących do rejestrowania momentów, w których każde pojedyncze nasiono trafia do gleby. Gdy to następuje, robot rozpyla jedynie taką ilość nawozu, jaka jest potrzebna, około 0,2 ml, bezpośrednio na nasiona precyzyjnie w momencie, gdy zagłębiają się w ziemię.

Operator wprowadza parametry dozowania, mierzy prędkość i warunki pracy. Ciśnienie w systemie jest obliczane na podstawie dawki, prędkości jazdy i danych wejściowych dyszy. Ziarno jest rozpoznawane automatycznie, a czas opóźnienia oprysku jest obliczany na podstawie prędkości jazdy w rzędzie, prędkości taśmy i prędkości oprysku. Przedstawiciele firmy podkreślają znaczne oszczędności (sięgające powyżej 60 proc. ilości nawozu) i poprawę wydajności całego procesu siewu, co oczywiście ma wpływ na późniejsze plony. ■

8 lub więcej litrów
oleju napędowego dziennie
wciąż wycieka z zatopionego
w 1941 roku w Pearl Harbor
na Hawajach amerykańskiego
pancernika USS „Arizona”.



GEOMETRIA

Figura einstein, czyli kapelusz

Matematycy odkryli pierwszy znany przykład „figury einstein”, z której można ułożyć mozaikę, na powierzchni, na podobieństwo np. kafelków terakotowych w łazience, w taki sposób, że nie ma żadnych pustych miejsc we wzorze, zaś ułożenie „kafelków” nie powtarza się na całej powierzchni, co czyni układ „mozaiką aperiodyczną”. Składającą się z ośmiu mniejszych kształtów w kształcie rombów trzynastoboczną figurę nazwano „the hat” (z ang. „kapelusz”).

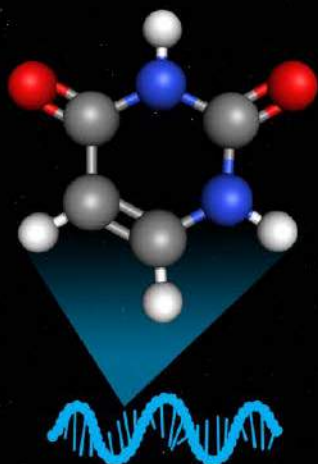
Dotychczas matematycy znali jedynie takie mozaiki aperiodyczne, które nie składały się z pojedynczego kształtu. Ta wyjątkowa figura zidentyfikowana została przez Davida Smitha, niezawodowego matematyka, który sam siebie określa jako „pomysłowy majsterkowicz kształtów”, i opisana w pracy opublikowanej w serwisie arXiv.org. Według opisu kształt ten to „polilatawiec” (ang. „polykite”), czyli układ złożonych w jedno, mniejszych kształtów latawca.

Choć nazwa „einstein” kojarzy się z wielkim fizykiem, to w rzeczywistości wywodzi się z niemieckiego słowa „ein Stein”, co oznacza „jeden kamień”. Niepowtarzające się wzory mają powiązania ze światem rzeczywistym. Materiałoznawca Dan Shechtman otrzymał w 2011 roku Nagrodę Nobla w dziedzinie chemii za odkrycie kwazikryształów, materiałów, w których atomy ułożone są w uporządkowaną strukturę, która nigdy nie powtarza tego samego ułożenia. Nowo odkryty aperiodyczny kafelek może mieć dalsze życie w dziedzinie badań materiałoznawczych. ■

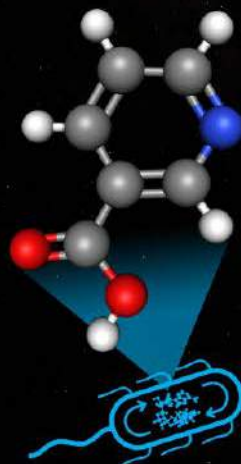


Przekształcenia
kształtów „einstein”:
<https://bit.ly/43zna4i>

URACYL



NIACYNA



UKŁAD SŁONECZNY

Cząsteczki składowe życia na asteroidzie

Jeden z kluczowych elementów składowych RNA odkryli japońscy badacze w próbkach gruntu przetransportowanych na Ziemię przez próbnik Hayabusa 2 z asteroidy Ryugu. Odkrycie wskazuje na to, że podstawy życia mogły zostać przyniesione na Ziemię spoza naszej planety a także że rudymentarne formy życia mogą istnieć gdzie indziej w Układzie Słonecznym.

Japońskie odkrycie dotyczy uracylu, jednej z pięciu nukleobaz, które tworzą nasz kod genetyczny, witaminy B3 (niacyny) i szeregu innych cząsteczek organicznych, które występują na powierzchni kosmicznej skały. Pięć nukleobaz to adenina, guanina, cytozyna, tymina i uracyl. Łączą się z rybozą

i fosforanami, tworząc DNA i RNA, struktury przypominające drabinę, które składają się na kod genetyczny życia na Ziemi.

„Skoro uracyl i inne nukleobazy występują w przestrzeni kosmicznej, oznacza to, że również składniki kwasów nukleinowych [DNA i RNA] są obecne w tym środowisku”, powiedział serwisowi „Live Science” Yasuhiro Oba, astrochemik z Uniwersytetu Hokkaido w Japonii. „Moim zdaniem, trudno jest wykluczyć możliwość, że niektóre formy życia są obecne w środowiskach pozaziemskich”. Kolejną możliwą konsekwencją jest możliwość pochodzenia także życia ziemskiego z elementów składowych pochodzących z kosmosu. ■

6,5 km/h wynosi liniowa prędkość obrotu Wenus wokół własnej osi, co oznacza, że zdrowy człowiek mógłby z łatwością poruszać się szybciej po powierzchni tej planety niż planeta się obraca.



PLANETA ZIEMIA

♦ NASA prowadzi aktywny monitoring niespotykanej dotychczas anomalii w polu magnetycznym Ziemi – gigantycznego regionu o niższym natężeniu pola magnetycznego rozciągającego się nad powierzchnią globu pomiędzy Ameryką Południową a południowo-zachodnią Afryką, nazywanego anomalią południowoatlantycką. ♦ Jak wynika z ostatnio opublikowanych wyników przeprowadzonych przez Narodowy Uniwersytet Australii studiów fal sejsmicznych wygenerowanych w ok. dwustu trzęsieniach Ziemi, w wnętrzu naszej planety znajduje się kula z litego, stałego, a nie ciekłego, metalu, najprawdopodobniej stopu niklu i żelaza, o średnicy ok. 650 km. ♦

ENERGIA

♦ Naukowcy z Wiedeńskiego Instytutu Technologicznego opracowali nowy typ opartego na materiałach ceramicznych ogniwa, nazywany baterią tlenowo-jonową, która, jak zapewniają, jest bardziej ekologiczna, trwalsza i mniej łatwopalna niż akumulatory litowo-jonowe, choć przyznają, że jest wciąż mniej wydajna i nagrzewa się do wysokich temperatur. ♦ Specjalizujący się w systemach wykorzystania energii geotermalnej startup z amerykańskiego Houston o nazwie Fervo Energy zaprezentował na pustyniach stanu Nevada system, który służy zarówno do pozyskiwania energii z gorącej wody wydobywanej spod ziemi, jak też do magazynowania energii pod powierzchnią przez całe godziny, a nawet dni, co w praktyce tworzy wielki geotermalny akumulator energetyczny. ♦

BADANIA KOSMOSU

♦ W publikacji w czasopiśmie „Nature Geoscience”, zespół uczo-

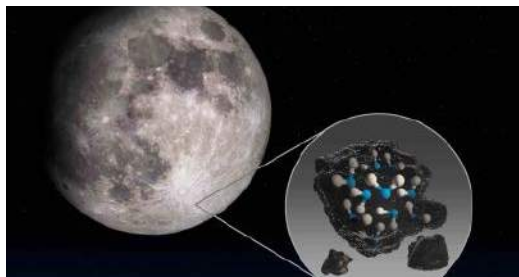
nych oszacował, że drobne szklane granulki, które są wszechobecne w księżycowym gruncie, mogą łącznie zawierać nawet 270 bilionów kilogramów wody, co wystarczyłoby do wypełnienia nią stu milionów basenów olimpijskich, stanowiąc ilość wystarczającą do zaopatrzenia astronautów i baz księżycowych.

♦ Teleskop Kosmiczny Hubble’a odkrył oddalony od Ziemi o 390 milionów lat świetlnych obiekt oznaczony jako Z 229-15, który wymyka się klasyfikacjom, gdyż wykazuje cechy aż kilku kategorii ciał niebieskich na raz, zatem – może być kwazarem, ale także aktywnym jądrem galaktycznym lub galaktyką Seyferta, typem galaktyki spiralnej bądź nieregularnej, charakteryzującej się jądrem o dużej jasności. ♦

NOWE MATERIAŁY

♦ Nadprzewodnictwo w temperaturze bliskiej temperatury pokojowej (20,5°C) wykazuje opisany w „Nature” materiał nazywany n-domieszkowanym wodorkiem lutetu i choć wciąż potrzebne jest do uzyskania tego efektu ciśnienie sięgające dziesięciu tysięcy atmosfer, to jest ono znacznie niższe niż zwykle wymagane w wysokotemperaturowych nadprzewodnikach. ♦ Grupa badawcza z hiszpańskiego uniwersytetu w Alicante opracowała proces otrzymywania rozpuszczalnego w wodzie tworzywa sztucznego na bazie skrobi ziemniaczanej, które wkrótce zostanie wprowadzone na rynek, zaś oprócz rozpuszczalności w wodzie materiał ma być także w pełni kompostowalny i biodegradowalny, nadając się do stosowania jako elastyczna folia, przede wszystkim w torbach i opakowaniach. ♦ Xiaoyu Song i Leslie Schoop z Uniwersytetu Princeton i ich koledzy, za pomocą procesu zwanego chemiczną eksfoliacją, wyprodukowali z dwusiarczku wolframu tusz o właściwościach nadprzewodzących po schłodzeniu do temperatury 7,3 kelwina (-266°C), co może okazać się praktycznym rozwiązaniem przy tworzeniu obwodów w komputerach kwantowych, które mogłyby być drukowane w temperaturze pokojowej i chłodzone w celu osiągnięcia nadprzewodnictwa. ■

M. U.





1. Zdjęcie udostępnione przez Lockheed Martin samolotu Darkstar użytego w filmie *Top Gun. Maverick*

Supersamoloty, nie tylko hipersoniczne

Czy *Top Gun* to tylko film?

Wiadomo było, że legendarny dział zaawansowanych projektów Lockheed Martin, Skunk Works, współpracował z producentami sequela *Top Gun* przy tworzeniu filmowego hipersonicznego samolotu Darkstar, którym latał gwiazdor filmu Pete „Maverick” Mitchell. Jednak dalsze wydarzenia każą się zastanowić, czy to była jedynie filmowa fikcja.

Uważni obserwatorzy zaawansowanych projektów lotniczych zauważyli, że w marcu 2023 r. w serii tweetów gratulujących twórcom *Top Gun* nominacji do Oscarów, Lockheed zdawał się sygnalizować, że być może istnienie samolotu takiego jak Darkstar (1) to nie tylko filmowa fantazja. Lockheed przyznawał już wcześniej, że jego projektanci ze Skunk Works poważnie podeszli do zadania zaprojektowania filmowego supersamolotu, potwierdzając, że zbudowano pełnowymiarową makietę proponowanego odrzutowca zdolnego do osiągnięcia prędkości 10 Ma. Makieta była tak przekonująca, że, jak twierdził producent filmu, Jerry Bruckheimer, Chiny zmieniły trajektorię swojego satelity szpiegowskiego, aby sfotografować samolot podczas kręcenia filmu. „Oni myśleli, że to jest prawdziwe”, mówił w rozmowie z serwisem Sandboxx News w 2022 r.

Lockheed przypomniał przy okazji swoją legendarną konstrukcję SR-71 Blackbird, pisząc m.in., iż jest to „nadal najszybszy załogowy samolot odrzutowy”. Opracowany w latach 60. XX wieku przez Lockheed Skunk Works i znanego inżyniera lotniczego Clarence’a Johnsona, SR-71 był głównym amerykańskim samolotem obserwacyjno-rozpoznawczym na dużych wysokościach aż do jego wycofania pod koniec lat 90. Zdolny do lotu z prędkością 3,2 Ma (3900 km/h), SR-71 nadal

jest rekordzistą świata w kategorii najszybszego oddychającego powietrzem samolotu załogowego.

Potem w zadaniach wywiadowczych Pentagon postawił na swój arsenał satelitów szpiegowskich, bezzałogowych statków powietrznych (UAV) oraz samolotów stealth piątej generacji. Jednakże rozwój broni antysatelitarnej, technologii antydostępowych i rozpraszających oraz technologii stealth spowodował, że planiści obronni zaczęli ponownie rozważać konstrukcje samolotów hipersonicznych, zdolnych do penetracji chronionej przestrzeni powietrznej i unikania obrony powietrznej.

Lockheed Martin’s Skunk Works miał podobno około 2007 r. rozpocząć prace nad hipersonicznym bezzałogowcem. W 2013 roku firma ujawniła, że rzeczywiście pracuje nad konstrukcją samolotu z dwoma silnikami scramjet, zaprojektowanego do prędkości przelotowej 6 Ma (powyżej siedmiu tysięcy km/h), nazwanego SR-72, a potocznie nazywanego „synem Blackbirda”. Rozwój SR-72 był owiany tajemnicą. Jednak od 2018 roku Lockheed Martin zapowiedział, że prototyp SR-72 ma odbyć pierwszy lot do 2025 roku, a samolot może wejść do służby w latach trzydziestych. W przeciwieństwie do swojego poprzednika, hipersoniczny SR-72 jest podobno samolotem bezzałogowym, zdolnym, oprócz



2. X-43 skonstruowany przez NASA

operacji nadzoru i rozpoznania, do wykonywania misji ofensywnych. Krąży pogłoski, iż „syn Blackbirda” ma być pierwszym na świecie samolotem hipersonicznym zdolnym do wyrzeliwania pocisków hipersonicznych.

Obserwatorzy techniki lotniczej zauważyli, że samolot Darkstar pojawiający się w *Top Gun. Maverick* przypomina grafikę koncepcyjną SR-72. Przed premierą filmu w maju 2022 roku dyrektor ds. komunikacji międzynarodowej Lockheeda, John Neilson, przyznawał, że pojawiły się plotki, iż film umożliwi „podgląd możliwego wyglądu Lockheed Martin SR-72”. Kokpit Darkstara nie ma widoczności do przodu. Pilot filmowy musi polegać na systemie monitorowania, aby zobaczyć, co jest przed samolotem. Podobne rozwiązanie widać w eksperymentalnym samolocie naddźwiękowym firmy Lockheed Martin, X-59 QueSST, który ma przejść testy w 2023 roku w ramach projektu NASA Low-Boom Flight Demonstrator. Nazwa „Darkstar” pochodzi ponadto od prototypu drona, RQ-3 opracowanego przez firmy Lockheed i Boeing w połowie lat 90.

Przedstawiciele Lockheeda Martina, poproszeni o komentarz, nie odrzucili pomysłu, że firma może mieć już hipersoniczny samolot zdolny do lotu z prędkością ponad sześciokrotnie większą od prędkości dźwięku. Na pytanie o to, czy Darkstar lub coś podobnego rzeczywiście istnieje, rzecznik firmy powiedział: „Nie mam nic dodatkowego do przekazania związanego z językiem użytym w tweecie”.

NASA już dawno śmiga hipersonicznie

Tymczasem wspomniana NASA już dość dawno temu zbudowała hipersoniczny „samolot” o prędkości 9,6 Ma. Mowa o X-43, maszynie, której zbudowano w sumie trzy egzemplarze, a jej dziewiczy lot odbył się jeszcze w czerwcu 2001 roku. Wersja X-43A, podczas trzeciego i ostatniego lotu w listopadzie 2004 roku, osiągnęła prędkość 9,6 Ma (11 852 km/h) nad Oceanem Spokojnym na zachód od Kalifornii na wysokości około 33 500 m. Zostało to oficjalnie uznane za rekord prędkości dla samolotu z napędem odrzutowym w Księdze Guinnessa.

Dla porównania, poprzedni rekord dla pojazdu odychającego powietrzem należał do rakiety napędzanej silnikiem odrzutowym, która osiągnęła nieco więcej niż 5 Ma. Zaś największą prędkość osiągniętą przez samolot napędzany rakietą X-15 wynosiła 6,7 Ma.

Jednak program X-43C został zawieszony na czas nieokreślony w marcu 2004 roku, a jego następcą został Boeing X-51 Waverider. X-51, który zakończył swój pierwszy hipersoniczny lot w maju 2010 roku, nie osiągnął tej samej prędkości co X-43. Poruszał się z prędkością 5 Ma (5300 km/h) i na wysokości 21000 m. Podobnie jak w przypadku X-43, X-51 używał B-52 jako platformy startowej.

Wysyp projektów

Projektów hipersonicznych samolotów, zarówno komercyjnych, jak i tych o charakterze bardziej militarnym, nie brakuje. Sęk w tym, że istnieją one bardziej na papierze lub na pulpicie grafika tworzącego efektowne wizualizacji niż w rzeczywistości.

Startup z Atlanty o nazwie Hermeus opracowuje hipersoniczny samolot, który przewoziłby pasażerów z prędkością co najmniej 6000 km/h, pięć razy większą niż prędkość dźwięku. Pomysł jest taki, że samolot podczas startu z pasa startowego korzystałby z napędu turbodrzutowego. Potężny silnik strumieniowy włączałby się w środkowej części lotu, a nad lądowaniem czuwałby znów bardziej tradycyjny napęd odrzutowy. Hermeus oczekuje, że jego pierwszy prototyp samolotu – Quarterhorse – będzie gotowy do lotu do 2023 roku. Po tym czasie przejdzie do budowy jednego o nazwie Darkhorse. Wszystko to będzie przygrywką do budowy 20-miejscowego samolotu hipersonicznego firmy nazwanego Halcyon (3). Ponieważ przepisy wciąż czynią nielegalnym podróżowanie z prędkością naddźwiękową nad lądem, Halcyon może być ograniczony do latania na trasach transatlantyckich (jego przewidywany zasięg około 6,5 tys. km nie wystarczy na trasy transpacyficzne).

Inny startup, Venus Aerospace z Teksasu, opracowuje z kolei bolid pasażerski Stargazer, który miałby osiągać nawet 9 Ma, ale to jest raczej samolot kosmiczny, trochę



3. Wizualizacja konstrukcji hipersonicznej Halcyon firmy Hermeus

inna kategoria, bliższa pomysłom lotów pasażerskich za pomocą rakiety Starship firmy SpaceX Elona Muska.

Departament Obrony USA zlecił też niedawno opracowanie kolejnego samolotu testowego docelowo o prędkościach hipersonicznych australijskiej firmie o nazwie Hypersonix. Defense Innovation Unit (DIU) poszukiwał przedsiębiorstw, które mogłyby zapewnić szybkie, niedrogie, wielokrotnego użytku hipersoniczne platformy testowe, najlepiej z alternatywnymi systemami naprowadzania, nawigacji, kontroli i komunikacji. Powinny też mieć funkcje napędowe, takie jak oddychanie powietrzem i funkcje dla cykli kombinowanych. Stworzony przez Hypersonix samolot DART AE służy do testowania szybkich platform, części, czujników, komunikacji i systemów kontroli. Firma podaje, że DART AE ma strumieniowy silnik scramjet zasilany wodorem, który może rozpędzać samolot do 7 Ma. Przewiduje się, że pierwszy lot samolotu odbędzie się na początku 2024 roku.

Także zespół brytyjskiego przemysłu, wojska i rządu ujawnił w ubiegłym roku plany dostarczenia hipersonicznego samolotu. To program Eksperymentalnego Hipersonicznego Pojazdu Powietrznego (HVX). Opracowane przez Reaction Engines rozwiązania silnika o cyklu kombinowanym SABRE stanowią techniczną bazę tego programu. W połączeniu z technologią turbin gazowych Rolls-Royce'a, mają się mierzyć z trudnymi problemami związanymi z lotem hipersonicznym. Ponadto w ramach programu prowadzone są prace projektowe nad eksperymentalnymi koncepcjami pojazdów hipersonicznych. Na pokazach w Farnborough w ub. roku zaprezentowano koncepcyjny pojazd hipersoniczny z jednym silnikiem – Concept V.

Chińczycy mają swoje ambicje

Chiny podobno rozwijają i testują swoją technologię hipersoniczną (4) w bezprecedensowym

tempie. Gdyby wierzyć oficjalnym komunikatom, to do 2035 roku Chiny zamierzają zbudować całą hipersoniczną flotę pasażerską, która wykorzystywać ma orbitę okołozemską, aby udać się do dowolnego miejsca na świecie w mniej niż godzinę. Jest to więc znów koncepcja nieco podobna do lotów pasażerskich Starshipem obiecywanych przez Elona Muska. Mimo że program ten został wyśmiany przez zachodnich obserwatorów, Chiny prawdopodobnie z całą powagą pracują nad tymi projektami.

Niektóre publikacje wspominają o chińskich pracach nad napędem magneto hydrodynamicznym, znanym również jako akcelerator MHD, który w skrócie polega na napędzaniu pojazdów przy użyciu jedynie pól elektrycznych i magnetycznych, bez ruchomych elementów, poprzez zastosowanie magneto hydrodynamiki do przyspieszania elektrycznie przewodzącego materiału pędnego (cieczy lub gazu). Płyn jest kierowany do tyłu, a w reakcji na to pojazd przyspiesza do przodu. Aby osiągnąć maksymalną wydajność, silnik magneto hydrodynamiczny będzie musiał być połączony z innymi rodzajami się dopiero technikami, takimi jak systemy szybkiego chłodzenia i silniki detonacyjne.

Pekińska firma Space Transportation (znana w Chinach jako Lingkong Tianxing) snuje plany opracowania pojazdu latającego przewożącego pasażerów, który może mknąć po niebie z prędkością nawet



4. Ilustracja przedstawiająca prawdopodobnie chińską konstrukcję samolotu hipersonicznego



1600 metrów na sekundę, ponad dwukrotnie większą niż Concorde. Firma opublikowała w 2022 r. animowany film reklamowy, na którym pasażerowie (nie wymagający kasku ani kombinezonu kosmicznego) wsiadają na pokład czegoś, co wydaje się 12-miejscowym samolotem kosmicznym, który mieści się pod aerodynamiczną strukturą w kształcie delty, flankowaną przez dwie rakiety wspomagające. Pojazd wyrzeliwuje pionowo w niebo, a po osiągnięciu wysokości przelotowej oddziela się od rakiet wspomagających i z prędkością 7000 kilometrów na godzinę przemierza granicę kosmosu, lądując pionowo w miejscu przeznaczonym za pomocą podwozia typu trójnóg.

Wcześniej, w 2018 roku Chiny ujawniły, że budują 265-metrowy tunel aerodynamiczny, który może być używany do testowania prototypów samolotów hipersonicznych o prędkości do 25 Ma (30 625 km/h) w Chińskim Państwowym Kluczowym Laboratorium Dynamiki Gazów Wysokotemperaturowych Chińskiej Akademii Nauk.

Space Transportation ma przed sobą, jak się zdaje, wciąż jeszcze sporo pracy nad rozwiązaniem wyzwań technicznych i biznesowych, ale wydaje się, że ma za sobą błogosławieństwo państwa, co w Chinach ma kluczowe znaczenie. ■

Mirosław Usidus

Napęd oparty na toroidzie wokół śmigła

Po co tyle hałasu?

Śmigła i śruby napędowe są projektowane tak, by oddziałując z medium, zazwyczaj powietrzem lub wodą, za pomocą ruchu obrotowego wypychać lub odpychać się od tego medium, napędzając pojazd latający lub pływający.

Przez bardzo długi czas o żadnych dużych zmianach w konstrukcji tych elementów nic nie słyszano. Samoloty napędzane śmigłami wciąż używają głównie skrzydeł, aerodynamicznych łopatek o konstrukcji zbliżonej do bambusowych śmigiełek, jakimi zabawiali się chińskie dzieci 2400 lat temu. Wzrost wydajności w działaniu nowszych konstrukcji był zaskakująco niewielki w stosunku do drewnianych śmigieł, które bracia Wright opracowali w 1903 roku. W łodziach nadal używa się tych samych mniej więcej śrub napędowych, których pierwsze wersje wypróbowywano ok. 1700 roku.

1. Czterowirnikowy dron z toroidalnymi śmigłami opracowanymi na MIT



Obręcz wokół śmigła

Jednym z kluczowych problemów związanych ze śmigłami wielopłatowymi jest hałas, które robią podczas pracy. Ludzie są zwykle najbardziej wrażliwi na dźwięki o częstotliwościach od około 100 Hz do 5 kHz. Ma to sens ewolucyjny, gdyż to właśnie w tej części skali akustycznej słyszymy samogłoski, które są kluczowe dla komunikacji werbalnej. Niestety śmigła samolotów również wydzielają dźwięki w tym zakresie.

Zespół konstruktorów z MIT, pracując nad cichszym samolotem, zasilanym strumieniem jonów, szukał sposobu na zmniejszenie poziomu tego hałasu, wejrzał w zamierzchłą przeszłość aeronautyki, znajdując tam wśród niekiedy dziewiętnastowiecznych jeszcze prototypów maszyn latających konstrukcje wykorzystujące pierścienie, obręcz (toroid) otaczający element napędowy, śmigło lub coś innego o zbliżonym kształcie. Za pomocą drukarki 3D technicy wykonali szereg prototypów śmigieł obwieszonych pierścieniami. Według ich relacji, podczas pracy, np. napędzając drony (1), brzmią przy tradycyjnych śmigłach jak szumiąca bryza, wydając znacznie mniej inwazyjny dźwięk.

Zdaniem twórców nowej konstrukcji, śmigła były cichsze, gdyż wiry generowane są w nich na całej długości śmigła, a nie tylko na samej jego końcówce.

To ma sprawiać, że w efekcie rozpraszają się szybciej w atmosferze, nie rozchodząc się tak intensywnie jak vibracje z tradycyjnych śmigieł.

Toroidy na obwodach śmigieł mogą również działać jako osłony zwiększające bezpieczeństwo ruchomych elementów. Trzeba jednak przyznać, że dodają one do masy konstrukcji, co wpływa negatywnie na długość pracy baterii. Mogą także być podatne na podmuchy wiatru, co może grozić destabilizacją drona lub jeszcze większym zużyciem energii do zasilania systemu stabilizującego. Można wspomnieć o innych potencjalnych wadach takiej konstrukcji. Są to dość skomplikowane kształty, więc znacznie trudniejsze do wyprodukowania niż standardowe śmigła wykorzystujące tanie i łatwe formowanie wtryskowe. Najprawdopodobniej powinny być drukowane w 3D. Ale nawet jeśli podwoją lub potroją cenę śmigieł, są to te tańsze części drona i ogólny wpływ na koszty może nie być aż tak trudny do zniesienia.

Testy terenowe nie wskazywały jednak wybitnie negatywnych efektów ubocznych. Okazało się, że najlepiej działający projekt o symbolu B160 był nie tylko cichszy przy danym poziomie ciągu niż najlepsze standardowe śmigło, które testowano porównawczo, ale również wytwarzał więcej ciągu przy danym poziomie mocy. Oczywiście wyniki te wymagają dodatkowych testów i weryfikacji.

Na tym etapie nie jest jasne, czy projekty tego typu mogą zastąpić tradycyjne śmigła w samolotach o stałej konstrukcji, czy też w elektrycznych taksówkach powietrznych VTOL. Te ostatnie już teraz wydają się być znacznie cichsze od helikopterów, ale jeśli w końcu zaleją przestrzeń miejską szybkim, tanim i ekologicznym transportem lotniczym, to łączny poziom hałasu będzie się liczył.

Nie tylko ciszej, ale również wydajniej

Aerodynamika i hydrodynamika są ze sobą ściśle powiązane. Jak się okazuje, istnieje już produkt bliski komercjalizacji w dziedzinie pojazdów pływających, który przyjmuje bardzo podobne podejście do konstrukcji śruby napędowej. Firma Sharrow Marine uzyskała, jak podaje, znakomite wyniki, jeśli chodzi o nowe typy śrub do łodzi, które wykorzystują pętle toroidalne (2) zamiast standardowych łopatek.

Konstrukcje te nie tworzą wirów krawędziowych – głównego źródła strat energii i zaskakująco dużego składnika ogólnego hałasu silników jednostek pływających. Zysk w sprawności jest znaczący w wodzie przy średnich obrotach, wypełniając wyraźną lukę w krzywej przyspieszenia łodzi i oszczędzając ogromne ilości energii. Znacznie zmniejszając



2. Śruba napędowa do łodzi firmy Sharrow Marine

ilość cieczy, która „ześlizguje się” z powierzchni śruby zamiast być przez nią efektywnie wypychana, toroidalne śruby zasysają większe ilości wody i przesuwają łódź o większy dystans w przeliczeniu na jeden obrót. Regularnie podwajają prędkość, jaką łódź może osiągnąć przy niższych i średnich obrotach, radykalnie poszerzając efektywny zakres obrotów silnika. Zmniejszają też zużycie paliwa o około 20 proc. To bardzo dużo, biorąc pod uwagę ogromne zapotrzebowanie na energię w łodziach napędzanych śrubami i skalę tego przemysłu.

Firma Sharrow twierdzi, że dodatkowym pozytywnym efektem jest znaczne zmniejszenie tendencji łodzi do przechylania się do tyłu podczas przyspieszania. Zamiast tego, cała łódź unosi się z wody. Ponadto następuje znaczna redukcja hałasu. Firma twierdzi, że jej śrubę można zamontować na każdym silniku zaburtowym, a następnie mknąć z prędkością 48 km/h na tyle cicho, że można prowadzić rozmowę na pokładzie bez podnoszenia głosu.

Sharrow sprzedaje już swoje śruby nowego typu. Pasują do szerokiej gamy popularnych silników zaburtowych większości głównych producentów. Haczyk musi być, a w tym przypadku jest nim cena, ok. dziesięciokrotnie wyższa niż tradycyjnych konstrukcji. Jednak, jeśli ktoś intensywnie użytkuje łódź, koszt ten powinien zwrócić się mu w znacznie mniejszym zużyciu paliwa. ■

Mirosław Usidus



Prezentacja napędu
łodzi firmy Sharrow:
<https://bit.ly/3MDgFa1>



Prezentacja robota MARVEL:
<https://bit.ly/3MjmtPl>

Roboty wspinające się po pionowych powierzchniach

Magnetyczny wspinacz

Widzieliśmy to na filmach. Maszyny pełzające po pionowych ścianach a nawet po sufitach. Widzieliśmy nawet ludzi dokonujących tego samego za pomocą instrumentarium przywierającego do ścian. W rzeczywistości skonструowanie sprawnie poruszającego się po pionowej powierzchni robota nie jest takie proste.

Robot wspinający się, który mógłby szybko poruszać się po powierzchniach takich jak ściany i sufity, ma większą operacyjną przestrzeń roboczą w porównaniu z innymi robotami naziemnymi. Jednak umiejętności określane ogólnie jako wspinaczka robotów są obecnie ograniczone do małych prędkości lub prostych zadań.

Specjaliści z Koreańskiego Instytutu Zaawansowanych Nauk i Technologii KAIST, Hae-Won Park, Yong Um i Seungwoo Hong, postanowili spróbować przezwyciężyć ograniczenia, konstruując operującego swobodnie, nie na uwięzi, czworonożnego robota wspinaczkowego o nazwie MARVEL (skrót od angielskiej nazwy – „magnetically adhesive robot for versatile and expeditious locomotion”), zdolnego do zwinnego i wszechstronnego poruszania się na podłożach ferromagnetycznych.

MARVEL ma być wykorzystywany w środowiskach i sytuacjach, które wiążą się z pewnym ryzykiem dla człowieka. Jak opisują go twórcy w publikacjach, m.in. w „Science”, MARVEL przewyższa znane do tej pory roboty wspinaczkowe pod względem szybkości wspinania i zdolności do wykonywania różnych ruchów. Osiąga najwyższą wśród podobnych maszyn prędkość pionowego i odwróconego chodu.

Najważniejsze zastosowane w konstrukcji ośmiokilogramowego MARVEL-a innowacje to konstrukcja stopy wykorzystująca magnesy elektropermanentne i elastomery magnetoreologiczne, dzięki którym robot może zmieniać właściwości magnetyczne każdej ze swoich stóp, oraz modelowe sterowanie predykcyjne dostosowane do stabilnego wspinania.

W eksperymentach robot osiągnął na sufitach i pionowych ścianach prędkości przemieszczania się odpowiednio 0,5 metra (1,51 długości ciała) na sekundę i 0,7 metra (2,12 długości ciała) na sekundę. Co więcej, robot radził sobie ze złożonymi zadaniami, takimi jak pokonywanie szczelin o szerokości dziesięciu centymetrów,



1. Robot MARVEL wspina się po ścianie zbiornika

pokonywanie przeszkód o wysokości do pięciu centymetrów oraz przechodzenie z podłogi na ściany i ze ścian na sufity. Konstruktorzy w ramach pokazów zademonstrowali umiejętność MARVEL-a polegającą na wspinaniu się po zakrzywionej powierzchni zbiornika pokrytego farbą o grubości do 0,3 milimetra oraz warstwą rdzy i kurzu (1). Koreańczycy z powodzeniem przetestowali także swojego robota z maksymalnym obciążeniem wynoszącym trzy kilogramy.

MARVEL potrafi również skanować powierzchnie w poszukiwaniu szczelin i otworów. Naukowcy zaprogramowali robota tak, aby omijał lub poruszał się nad takimi przeszkodami. Porównują ten mechanizm do symulowania sposobu, w jaki kot wykorzystuje swoje przednie łapy do analizy obiektów przed poruszaniem się do przodu.

Projekt ma taki cel, by MARVEL lub robot o podobnej konstrukcji zastąpił inspekcje kontrolne i monitorujące prowadzoną przez człowieka na większych wysokościach lub w ciasnych przestrzeniach. W pracy opisującej eksperymenty z MARVEL-em podkreśla się, że konstrukcja ta mogłaby potencjalnie posłużyć do monitoringu obiektów przemysłowych, takich jak budynki o stalowej konstrukcji, mosty, statki lub zbiorniki służące do przechowywania substancji ciekłych lub sypkich. ■

Miroslaw Usidus



ŁAMANIE GŁOWY

Czy mózg się obroni?

Jeszcze mamy bastion prywatności. Jeszcze nienaruszony, choć czujniki rejestrują już wiele aspektów życia naszego organizmu, bicie serca, oddechy, kroki czy sen. To ma się zmienić. Fala już dostępnych lub zapowiadanych urządzeń śledzących pracę mózgu, chipy, słuchawki, opaski, zegarki, a nawet tatuaże. Obiecują zmienić nasze życie, grożąc azyłowi intymności, jakim jest nasz mózg.

Ostatnia rubież prywatności – czy na pewno?

PODGLĄDACZE MYŚLI JUŻ NA HORYZONCIE

Giganci technologiczni, Meta, Microsoft i Apple, od dłuższego czasu inwestują w tzw. brain wearables (1). Celem tych technoentuzjastów jest wbudowanie czujników mózgu w inteligentne zegarki, wkładki douszne, zestawy słuchawkowe i urządzenia wspomagające zasypianie. Włączenie ich do naszego codziennego życia mogłoby, owszem, zrewolucjonizować opiekę zdrowotną, umożliwiając wczesną diagnozę i spersonalizowane leczenie takich schorzeń, jak depresja, epilepsja, a nawet zaburzenia funkcji poznawczych. Czujniki mózgowe mogłyby poprawić naszą zdolność do medytacji, koncentracji, a nawet komunikacji za pomocą techtelepatii, wszystko w połączeniu z interakcją z zestawami słuchawkowymi rzeczywistości rozszerzonej (AR) i wirtualnej (VR).

Jednak łączące się w taki czy inny sposób z mózgiem urządzenia noszone stanowią również bardzo realne zagrożenie dla prywatności psychicznej, wolności myślenia i samostanowienia. W miarę upowszechniania się będą one generować ogromne ilości danych z neuronów, otwierając okno do naszych mózgów, treści myśli, emocji, a nawet wspomnień.

Pracodawcy już teraz zbierają „dane mentalne”, śledząc poziom zmęczenia pracowników i oferując programy odnowy biologicznej w celu złagodzenia stresu, za pośrednictwem platform, dających dostęp do mózgów pracowników. W Chinach staje się to coraz bardziej powszechne (2). Konduktorzy w pociągach na linii Pekin–Szanghaj noszą czujniki mózgu



1. Wizja człowieka obudowanego monitorującymi mózg czujnikami autorstwa generatora Blue Willow

przez cały dzień pracy. Są nawet doniesienia, że chińscy pracownicy są odsyłani do domu, jeśli aktywność ich mózgu wykazuje mniej niż optymalne wyniki. Urządzenia, które mogą śledzić uwagę pracowników, ich skupienie, a nawet znużenie, bez odpowiednich zabezpieczeń mogą zrujnować prywatność umysłową pracowników.

Rządy również starają się uzyskać dostęp do naszych mózgów. Zwykle projekty takie prowadzone są pod hasłem poszukiwania przyczyn chorób neurologicznych i psychiatrycznych. Od programów biometrycznych na bazie mózgu, opracowywanych w celu uwierzytelniania, w tym finansowanych przez amerykańską Narodową Fundację Nauki na Uniwersytecie Binghamton, do technik pobierania tak zwanych „mózgowych odcisków palców”, wykorzystywanych do identyfikacji podejrzanych o popełnienie przestępstwa, sprzedawanych przez firmy takie jak Brainwave Science i finansowanych przez organy ścigania od Singapuru przez Australię po Zjednoczone Emiraty Arabskie.



2. Chiński system do skanowania fal mózgowych pracownika

Pęd do wdzierania się do ludzkiego mózgu lub, jak kto woli, jego hakowania, widać w stosunkowo nowej dziedzinie neuromarketingu w mediach społecznościowych. Kampania Frito-Lays wykorzystwała np. wiedzę o tym, jak mózgi kobiet podejmują decyzje dotyczące przekąsek, a następnie monitorowała aktywność mózgów, gdy uczestniczący w badaniach ludzie oglądali ich nowo zaprojektowane reklamy, co pozwoliło im dostroić przekazy reklamowe w celu lepszego przykucia uwagi i skłonienia kobiet do częstszego jedzenia ich produktów. Przyciski „lubię to” i powiadomienia w mediach społecznościowych to funkcje zaprojektowane tak, by kształtować nawyki, a wręcz uzależnienia od społecznościowej dopaminy, co ściąga użytkowników do platform, wykorzystując systemy nagrody wbudowane w mózgi.

Mało wciąż poruszonym tematem są zagrożenia związane z cyberatakami i hakowaniem urządzeń noszonych znajdujących się w stanie połączenia z mózgiem (3). Możemy sobie wyobrazić dystopie włamywania się do danych o naszych stanach mentalnych, przechwytywanie haseł i numerów PIN w trakcie myślenia o nich lub wpisywania ich. Już wiemy, że w przyszłości urządzenia noszone mające łączność z mózgiem mogą śledzić to, co czytamy i widzimy, zmieniać naszą percepcję, manipulować emocjami, a nawet wywoływać ból fizyczny. Chiński Entertech gromadzi miliony surowych zapisów danych EEG od osób z całego świata, korzystających z popularnych, konsumenckich urządzeń noszonych, rejestruje również dane osobowe, sygnały GPS, czujniki urządzenia, komputery i usługi, z których korzysta dana osoba, w tym strony internetowe, które może odwiedzać.

Widać, o czym myśli mózg

Można powiedzieć, że urządzeń „czytających myśli” jeszcze nie ma. To prawda. Jednak powstały już



3. Hakowanie mózgu

pierwsze prototypy systemów, które do tego prowadzą. Shinji Nishimoto i Yu Takagi z uniwersytetu w Osace w Japonii opracowali technikę przetwarzania treści myśli na obrazy, która wykorzystuje Stable Diffusion, generator tekstu do obrazu wydany przez Stability AI w sierpniu 2022 roku. Japońscy badacze korzystali z danych pochodzących od czterech osób poddanych skanowaniu mózgów podczas oglądania dziesięciu tysięcy różnych obrazów: krajobrazów, obiektów i ludzi, za pomocą funkcjonalnego rezonansu magnetycznego (fMRI). Wykorzystując około 90 proc. danych z obrazowania mózgu, badacze wytrenowali model, tworząc powiązania między danymi z fMRI z regionu mózgu, który przetwarza sygnały wizualne, zwanego wczesną korą wzrokową, a obrazami, które ludzie oglądali. Użyli tego samego zbioru danych do wytrenowania drugiego modelu w celu utworzenia powiązań między opisami tekstowymi obrazów a danymi fMRI z regionu mózgu, który przetwarza znaczenie obrazów, zwanego brzuszną korą wzrokową. Po szkoleniu te dwa modele były w stanie przetłumaczyć dane z obrazowania mózgu na formy, które zostały bezpośrednio przekazane



4. Na górze – obrazy widziane przez uczestników eksperymentu na dole skany mózgowce

do generatora obrazu. Udało się zrekonstruować około tysiąca oglądanych przez ludzi obrazów z około 80-procentową dokładnością (4). Próbką badawczą była jednak nieduża, badania uciążliwe i pracochłonne, bo podejście to wymaga długich sesji skanowania mózgu i ogromnych maszyn fMRI. Nie jest to więc praktyczne rozwiązanie, a raczej zapowiedź dalszego rozwoju techniki przetwarzania treści mózgu na obraz.

Z fMRI korzystał również projekt naukowy Zijiao Chena i kolegów z uczelni w Singapurze, Hongkongu i Uniwersytetu Stanforda w USA. Wykorzystano w nim skany mózgu uczestników, którzy patrzyli na ponad tysiąc obrazów, przedstawiających przykładowo czerwony wóz strażacki, szary budynek, żyrafę jedzącą liście, gdy fMRI rejestrowało powstające sygnały mózgowe. Następnie badacze przesłali te sygnały do szkolenia modelu sztucznej inteligencji (również Stable Diffusion), aby wytrenować go w kojarzeniu określonych wzorców mózgowych z określonymi obrazami. Później, kiedy badanym pokazywano nowe obrazy w fMRI, system wykrywał fale mózgowe pacjenta, generował skrócony opis tego, co według niego odpowiadało tym falom mózgowym, i używał generatora obrazów AI, aby wyprodukować najlepszą możliwą kopię obrazu, który widział uczestnik. Wyniki, jak oceniono, przypominały obrazy ze snu, rozmyte, jakby zamglone. Wygenerowany obraz odpowiadał atrybutom (kolor, kształt, itp.) i znaczeniu oryginalnego obrazu w około 84 proc. przypadków.

Jack Gallant, naukowiec z Uniwersytetu Kalifornijskiego w Berkeley, zajął się badaniami nad dekodowaniem treści mózgu ponad dekadę temu, używając innego algorytmu i w nieco innym celu, bo nie generował obrazów, lecz starał się wykryć aktywność mózgu odpowiadającą określonym działaniom. W 2019 r. w „Journal of Neuroscience” ukazała się publikacja opisująca, jak naukowcy z jego zespołu stworzyli interaktywne mapy, które pozwalają przewidzieć, w których miejscach różne kategorie słów aktywują mózg. W tym projekcie mapowania mózgu ludzie słuchali popularnej serii podcastów, a następnie czytali te same historie. Korzystając z funkcjonalnego rezonansu magnetycznego, badacze zeskanowali ich mózgi zarówno w warunkach słuchania, jak i czytania, porównali dane dotyczące aktywności mózgu podczas słuchania i czytania i odkryli, że mapy utworzone z obu zestawów danych były praktycznie identyczne. Ich dane dotyczące aktywności mózgu, w obu warunkach, zostały następnie dopasowane



Mapa skojarzeń językowych z mózgu na podstawie badań zespołu Jacka Gallanta: <https://bit.ly/43PXsSP>

do zakodowanych w czasie transkrypcji opowiadań, których wyniki zostały wprowadzone do programu komputerowego, który ocenia słowa według ich wzajemnego stosunku. Wyniki można obejrzeć na interaktywnej, trójwymiarowej, kolorowej mapie kory mózgowej, na której słowa pogrupowano w takich kategoriach, jak wizualne, dotykowe, numeryczne, lokalizacyjne, gwałtowne, umysłowe, emocjonalne i społeczne. Wykorzystując modelowanie statystyczne, badacze ułożyli tysiące słów na mapach zgodnie z ich semantycznymi związkami. Na przykład w kategorii zwierzęta można znaleźć słowa „niedźwiedź”, „kot” i „ryba”. Mapy, które pokrywały co najmniej jedną trzecią kory mózgowej, pozwoliły badaczom przewidzieć z dokładnością, które słowa będą aktywować które części mózgu.

Z kolei naukowcy z uniwersytetu w Glasgow pracują nad połączeniem sztucznej inteligencji i ludzkich fal mózgowych w celu identyfikacji obiektów, których człowiek nie może normalnie zobaczyć, ponieważ są poza polem widzenia. Nazywa się to systemem „ghost imaging” i zostało zaprezentowane na kongresie Optica Imaging and Applied Optics w 2022 r. W ramach eksperymentów na kartonowy wycinek rzutowano obiekt. Osoba z zestawem elektroencefalograficznym na głowie, który monitorował jej fale mózgowe, widziała jedynie rozproszone światło na ścianie zamiast rzeczywistych obrazów. Hełm EEG odczytywał sygnały w korze wzrokowej osoby, które były przekazywane do komputera. Z pomocą algorytmów AI uczonym udało się zrekonstruować obrazy prostych obiektów o rozdzielczości 16×16 pikseli, których ludzie nie mogli zobaczyć z powodu zasłonięcia ich.

Implanty mózgowe mogą pomagać, ale mogą też uzależniać

Bohaterem filmu *Człowiek Terminal* z 1974 r. jest mężczyzna otrzymujący inwazyjny implant mózgowy, który ma pomóc w leczeniu ataków, których doznaje. Choć początkowo operacja wydaje się sukcesem, sprawy przybierają zły obrót, gdy długotrwałe działanie chipa doprowadza go do psychozy. Jak to zwykle bywa w filmach science fiction ostrzeżenie przed katastrofą pojawia się już na początku filmu przez porównanie implantu do lobotomii stosowanej kilkadziesiąt lat wcześniej.

Startup Elona Muska, Neuralink, pracował nad wszczepialnym chipem mózgowym, osadzonym w czaszce od 2016 roku. Po latach testów na obiektach zwierzęcych, Musk ogłosił, że firma planuje rozpocząć testy na ludziach. Choć według oficjalnych komunikatów, firma koncentruje się na zastosowaniach medycznych, np. na pomocy w komunikacji



5. Wizja podobnego do chipa Neuralink wszczepionego do głowy

osobom sparaliżowanym, Musk ma większe aspiracje. Mówił m.in. o „Fitbit [urządzenie i aplikacja do monitorowania aktywności człowieka – przyp. red. MT] w twojej czaszce” (5).

Firma Muska nie jest jedyną spółką pracującą nad interfejsami mózg-komputer, czyli systemami ułatwiającymi bezpośrednią komunikację między ludzkimi mózgami a zewnętrznymi komputerami. Badacze przyglądają się m.in. wykorzystaniu BCI do przywracania utraconych zmysłów i sterowania protezami kończyn. Choć rozwiązania te są jeszcze w powijkach, to rozwijane są już na tyle długo, że badacze coraz lepiej rozumieją interakcje implantów neuronowych z naszymi umysłami. Ingerowanie w działanie ludzkiego mózgu to trudna sprawa, a efekty eksperymentów nie zawsze są zgodne z oczekiwaniami lub zamierzone. Już odnotowuje się sygnały, że osoby korzystające z BCI mogą odczuwać głębokie poczucie zależności od tych urządzeń lub wrażenia, jakby ich poczucie własnej wartości zostało zmienione. Sam Neuralink zaczął być rozliczany z obietnic. Trwa federalne dochodzenie w sprawie skarg dotyczących naruszeń dobrostanu zwierząt wykorzystywanych w eksperymentach firmy Muska.

Ponad dwieście tysięcy ludzi na całym świecie używa już jakiejś formy i typu BCI, głównie z powodów medycznych. Prawdopodobnie najbardziej znanym przypadkiem stosowania są implanty ślimakowe, które pomagają osobom głuchym. Innym przykładem jest zapobieganie napadom epileptycznym za pomocą wszczepionych urządzeń. Istniejące urządzenia mogą monitorować aktywność sygnałów mózgowych, aby przewidzieć ataki i ostrzegać, tak aby pacjent mógł uniknąć pewnych czynności lub przyjmować leki

Łamanie głowy. Czy mózg się obroni?

zapobiegawcze. Niektórzy badacze zaproponowali systemy, które nie tylko wykrywałyby napady, ale również zapobiegałyby im za pomocą stymulacji elektrycznej, czyli niemal dokładnie tak, jak mechanizm przedstawiony w filmie z 1974. Implanty dla osób z chorobą Parkinsona, depresją, zaburzeniami kompulsywnymi i epilepsją są od lat w fazie prób na ludziach.

Firma Grand View Research, zajmująca się badaniem rynku, wyceniła globalny rynek implantów mózgowych na 4,9 miliarda dolarów w 2021 roku, a inne firmy przewidują, że liczba ta może się podwoić do 2030 roku. Paradromics, Blackrock Neurotech i Synchron to tylko kilka przykładów startupów pracujących nad urządzeniami dla osób sparaliżowanych. W listopadzie ubiegłego roku firma o nazwie Science zaprezentowała koncepcję bioelektrycznego interfejsu, który ma pomóc w leczeniu ślepoty.

Na razie BCI rozwijane są w ograniczeniu do domeny medycznej. Badania opublikowane w 2018 roku opisały wykorzystanie BCI do pracy z wieloma aplikacjami na tablecie z systemem Android, w tym pisanie na klawiaturze, wysyłanie wiadomości i wyszukiwanie stron internetowych, wszystko jedynie przez wyobrażenie sobie odpowiednich ruchów. Są rozważania o zastosowaniu interfejsów mózgowych w grach wideo, manipulowaniu wirtualną rzeczywistością, a nawet bezpośrednim odbieraniu danych wejściowych, takich jak wiadomości tekstowe lub filmy wideo, omijając potrzebę korzystania z monitora. Może to brzmieć jak science fiction, ale w rzeczywistości osiągnęliśmy punkt, w którym kulturowe i etyczne bariery dla tego rodzaju technologii zaczęły przewyższać te techniczne.

Ponieważ BCI są nadal ograniczone głównie do dziedziny medycznej, większość pierwszych użytkowników jest zadowolona z tego rodzaju kompromisów. Jeśli ktoś jest niepełnosprawny i nie może się komunikować, to ogólnie jest zadowolony, jeśli istnieje technika, która pozwala mu na to. Anna Wexler, profesor filozofii z Uniwersytetu Pensylwanii, przeprowadziła szereg wywiadów z osobami chorymi na parkinsona, które były poddawane głębokiej stymulacji mózgu, leczeniu polegającemu na wszczepianiu cienkich metalowych przewodów, które wysyłają impulsy elektryczne do mózgu, aby pomóc w złagodzeniu objawów. Z jej badań wynika, że BCI pomogło ludziom poczuć, że wracają do siebie i poczucia władzy nad własnym organizmem dzięki wspomaganiu BCI. Także Eran Klein i Sara Goering, badacze z Uniwersytetu Stanu Waszyngton zauważyli w swoich badaniach pozytywne zmiany w osobowości i postrzeganiu siebie wśród osób korzystających z BCI. W pracy z 2016 r.

na temat postaw i rozważań etycznych dotyczących DBS, podali, że uczestnicy badania często czuli, że leczenie pomogło im odzyskać „autentyczne ja”, które zostało przytłoczone przez depresję lub zaburzenia obsesyjno-kompulsyjne. W wykładzie z końca 2022 roku na temat podobnych badań, neuropsycholożka Cynthia Kubu opisała zwiększone poczucie kontroli i autonomii wśród pacjentów, z którymi przeprowadziła wywiady.

Jednak nie wszystkie zjawiska, które badacze rozpoznali, można uznać za pozytywne i korzystne. W wywiadach z osobami, które miały BCI, Frederic Gilbert z Uniwersytetu Tasmanii zauważył pewne dziwne efekty. Zauważył mianowicie, że pacjenci zgłaszają poczucie nierozpoznawania siebie lub, jak to się określa, „wyobcowania”. Niektórzy mieli uczucie posiadania nowych zdolności niezwiązanych z ich implantami, jak na przykład kobieta w wieku 50 lat, która zraniła się podczas próby podniesienia stołu bilardowego, bo myślała, że da z łatwością radę. „Doprowadziło to do skrajnych przypadków, nawet próby samobójczej”, raportował Gilbert.

Patrząc na analogie pochodzące z interakcji ludzi ze współczesną techniką, należałoby zachować ostrożność. W końcu, jeśli łatwo jest się uzależnić od internetu i telefonu, to o ile bardziej uzależniające mogłoby być urządzenie podłączone bezpośrednio do mózgu. Gilbert opowiadał w rozmowie z jednym z serwisów o pacjencie, który rozwinął rodzaj paraliżu decyzyjnego, w końcu czując się tak, jakby nie mógł wyjść lub zdecydować, co zjeść, bez uprzedniego skonsultowania się z urządzeniem, które pokazywało, co dzieje się w jego mózgu. Gilbert spotkał wielu uczestników badań, którzy popadli w depresję po utracie wsparcia dla swoich urządzeń i ich usunięciu, często po prostu dlatego, że dane badanie się zakończyło lub skończyły się fundusze. „Stopniowo wrastasz w to i przyzwyczajasz się do tego”, powiedział mu anonimowy uczestnik badań, który otrzymał urządzenie do wykrywania oznak aktywności epileptycznej. Ponieważ BCI są w dużej mierze wciąż w fazie prób, brakuje uniwersalnych standardów lub stabilnego wsparcia finansowego, a wiele urządzeń jest zagrożonych nagłą utratą finansowania. Ten rodzaj uzależnienia jest dodatkowo skomplikowany przez fakt, że BCI często wymagają inwazyjnej operacji mózgu, aby je usunąć i ponownie wszczepić.

Istnieją poważne obawy dotyczące prywatności, które wynikają z tego, że komputer uzyskuje dostęp do fal mózgowych pacjenta. „Jeśli na przykład dostaniesz urządzenie, które pomoże poruszać protezą ramienia, to urządzenie odbierze także inne sygnały”, powiedział Gilbert. „Jest wiele sygnałów, które mogą



6. Nita Farahany

być rozszyfrowane”. Ktoś mógłby się wiele dowiedzieć, studiując tylko czyjeś fale mózgowe, a gdyby hakerowi udało się uzyskać dostęp, mógłby w pewnym sensie czytać umysł.

W marcu 2023 w wywiadzie dla „The Guardian”, profesor prawa amerykańskiego Uniwersytetu Duke, autorka książki „The Battle for Your Brain: Defending the Right to Think Freely in the Age of Neurotechnology”, Nita Farahany (6), powiedziała, że chociaż technika interfejsu mózg-komputer jeszcze „nie może dosłownie czytać naszych złożonych myśli”, jest jednak tego na tyle bliska, że należy się poważnie nad nią zastanowić. „Przynajmniej niektóre części naszej aktywności mózgowej można rozszyfrować”, powiedziała Farahany. „Nastąpił duży postęp w technice elektrod i w szkoleniu algorytmów AI, który pozwala znaleźć skojarzenia przy użyciu dużych zbiorów danych”.

Farahany uważa, że powinniśmy zacząć myśleć o naszych prawach, zanim technologie hakowania mózgu, takie jak Neuralink Elona Muska, dostaną szansę wejścia do głównego nurtu. W celu uniknięcia orwellowskich konsekwencji tej techniki, profesor Farahany proponuje powołanie nowego rodzaju prawa obywatelskiego – „wolności poznawczej”, połączonego ze znanymi prawami do prywatności, wolności myśli i samostanowienia. ■

Miroslaw Usidus



1. Mapowanie mózgu

Wyobraź sobie, że patrzysz na Ziemię z kosmosu i chcesz podsłuchać, co mówią do siebie poszczególni Ziemianie. Do tego można porównać mniej więcej wyzwanie, jakim jest poznanie i zrozumienie, jak wszystko działa, jest połączone i wpływa na siebie, w ludzkim mózgu. Mapy mózgu (1) to sposób na zmierzenie się z tym wyzwaniem.

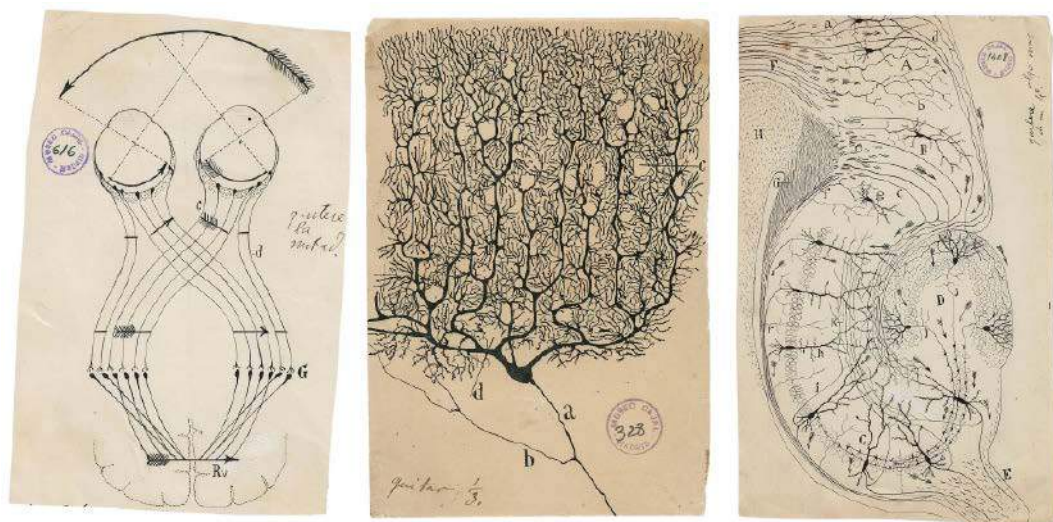
Pokazać wszystko – jak się wszystko łączy się ze wszystkim – i jak to wszystko działa

POSTĘPY W MAPOWANIU MÓZGU

Niemal nieustannie powstają nowe, coraz doskonalsze i dokładniejsze mapy mózgu. Najnowsza, o której w kwietniu 2023 r. doniósł „Nature”, ujawnia, że oprócz regionów poświęconych konkretnym częściom ciała,

są obszary, które kontrolują integrację całego ciała. I generalnie wygląda to nieco inaczej, niż zakładano do tej pory.

Evan Gordon, neurobiolog z uniwersytetu w St. Louis i jego koledzy badają zsynchronizowaną aktywność i komunikację między różnymi regionami mózgu. Zespół zebrał więc dane z funkcjonalnego rezonansu magnetycznego na ochotnikach, którzy wykonywali różne zadania. Uczestnicy badań wykonywali proste ruchy, takie jak poruszanie tylko brwiami lub palcami u nóg, a także złożone zadania, jak jednoczesne obracanie nadgarstka i poruszanie stopą z boku na bok. Dane fMRI ujawniały, które części mózgu aktywowały się w tym samym czasie, gdy wykonywano każde zadanie, co pozwoliło badaczom prześledzić, które regiony były funkcjonalnie połączone



2. Ilustracje przedstawiające pierwsze próby mapowania mózgu autorstwa Santiago Ramóna y Cajala

ze sobą. Zespół odkrył, że połączenie mózg–pierwotna kora ruchowa jest zorganizowane w trzy odrębne sekcje. Każda z nich reprezentuje inne regiony ciała – dolną część ciała, tułów i ramiona oraz głowę. Znalezione również trzy miejsca, które nie są związane z konkretną częścią ciała. Nazwane regionami międzysektorowymi, łączą się one z zewnętrzną siecią zaangażowaną w kontrolę działań i odczuwanie bólu. Regiony te przeplatają się z sekcjami poświęconymi konkretnym częściom ciała. Podejrzewa się, że regiony międzysektorowe mogą integrować cele działania i ruchy ciała obejmujące wiele części ciała, zaś przestrzenie pomiędzy nimi są wykorzystywane do precyzyjnych ruchów pojedynczych części ciała.

Jak się ma struktura do funkcji

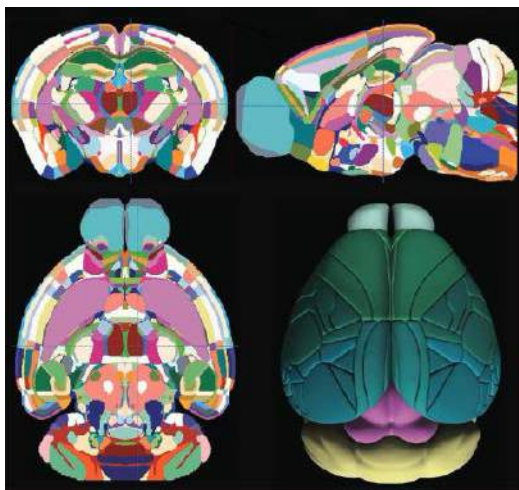
Obecnie szacuje się, że w mózgu znajduje się około 86 miliardów neuronów oraz mniej więcej taka sama liczba komórek nieneuronowych. Uważa się, że liczba połączeń, czyli synaps, przez które neurony komunikują się za pomocą sygnałów chemicznych i elektrycznych, wynosi około 125 bilionów. Znajduje się tam cały wszechświat, mimo że przeciętny dorosły mózg waży zaledwie ok. 1,5 kg i ma wymiary zaledwie 140 mm × 167 mm × 93 mm.

Choć wiemy wiele o anatomii mózgu, jego funkcje pozostają w dużej mierze enigmatyczne. Na przykład, chcielibyśmy poznać biologiczny mechanizm kodowania wspomnień. W komputerze pliki są kodowane cyfrowo za pomocą serii jedynek i zer, co jest rodzajem dyskretnej pamięci. A w jaki sposób mózg przechowuje informacje? Nie wiemy. Nie wiemy też, jak i ewentualnie – gdzie – w mózgu powstaje świadomość,

ale nad tym pochylamy się w innym tekście w tym wydaniu MT. Możemy jednak już jako tako skorelować różne działania, ruchy i myśli z aktywnością mózgu. Naukowcy mogą na przykład powiedzieć, jaka część mózgu będzie wykazywać aktywność elektryczną, gdy czytamy pewne słowa i frazy.

Ponad sto lat temu hiszpański neurobiolog Santiago Ramón y Cajal jako pierwszy pokazał (2), jak wiele różnych typów komórek występuje w mózgu ssaków. Zabarwił neurony tak, by można je było zobaczyć pod mikroskopem, a następnie wykonał precyzyjne rysunki ich kształtów. Spośród kilkudziesięciu typów, które znalazł, niektóre miały przedłużenia, zwane dziś aksonami, które sięgały na duże odległości z klockowatych ciał komórkowych niczym nogi pajaka. Niektóre miały krótkie aksony, inne przypominały raczej gwiazdy. Wydedukował, że ponieważ aksony każdej komórki znajdowały się bardzo blisko ciał komórkowych innych, prawdopodobnie przekazywały informacje. Za swoje odkrycia otrzymał w 1906 roku Nagrodę Nobla w dziedzinie fizjologii lub medycyny.

Od tego czasu większość badań nad typami komórek skupia się na korze mózgowej, która kontroluje wiele bardziej skomplikowanych zachowań. Neurologzy wyodrębnili w korze mózgowej trzy główne klasy komórek – dwie klasy neuronów – hamujące i pobudzające. Obie przekazują impulsy elektryczne, ale pierwsza tłumi aktywność w neuronach partnerskich, a druga ją pobudza. Trzecia klasa to ogromna liczba komórek nieneuronowych, wspierających i chroniących neurony. Przez dziesięciolecia neuronaukowcy wykorzystywali każdą nową technikę, którą uznali za odpowiednią, aby dopracować definicję tego, co stanowi



3. Mapa mózgu myszy stworzona przez Instytut Allena

odrębny typ komórki w tych klasach. Zdali sobie sprawę, że komórki, które powierzchownie wyglądają tak samo, mogą być różnymi typami komórek, w zależności od ich połączeń z innymi komórkami lub regionami mózgu, lub ich właściwości elektrycznych.

Od lat 90. XX w. naukowcy zaczęli badać aktywność genów w różnych typach komórek i to, jak ich ekspresja wpływa na ich właściwości. W 2006 roku Instytut Allena stworzył atlas ekspresji genów pokazujący, gdzie w mózgu myszy każdy z około 21 tys. genów podlega ekspresji. Budowanie atlasu mózgu Allena, jeden gen po drugim, zajęło kilkudziesięciu pracownikom instytutu trzy lata. Atlas jest regularnie aktualizowany (3) i nadal jest szeroko wykorzystywany jako punkt odniesienia.

Głównym celem mapowania mózgu jest odkrycie związku pomiędzy strukturą i funkcją mózgu. Różne narzędzia opracowane w ciągu ostatniego stulecia miały właśnie do tego służyć, do lokalizacji funkcji w mózgu. Obecnie korzysta się z technik obrazowania mózgu takich m.in. jak: funkcjonalny rezonans magnetyczny (fMRI), pozytonowa tomografia emisyjna, elektroencefalografia, elektrokortykografia i spektroskopia w bliskiej podczerwieni (niedawno powstałe narzędzie). Techniki elektrofizjologiczne, w tym

głęboka stymulacja mózgu, mogą być również wykorzystywane do badania funkcji neuronów w poszczególnych regionach mózgu.

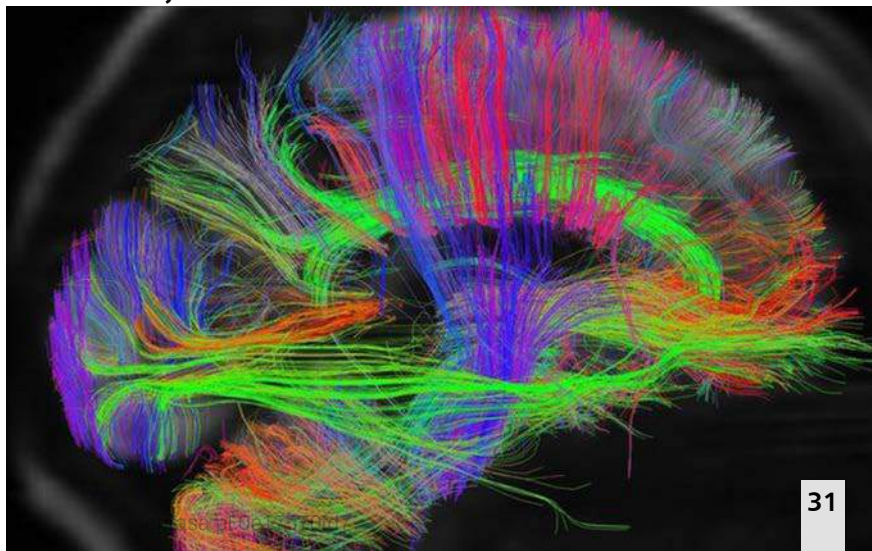
Najbardziej popularna ostatnio metoda, fMRI, mierzy zmiany w przepływie krwi, które przekładają się na aktywność komórkową i teoretycznie dostarczają funkcjonalnej mapy mózgu. fMRI oparte na zadaniach wzmacnia tę mapę poprzez wykonywanie przez osoby badane testów podczas skanowania MRI w celu zidentyfikowania wywołanych przez bodźce wzorców aktywności komórkowej. Przykładem jest wykonywanie zadań związanych z pamięcią roboczą, które prowadziły do aktywacji w korze przedczołowej u zdrowych osób, wzmacniając tę strukturę jako podstawową w pamięci. Wyniki fMRI powinny być interpretowane z ostrożnością, ponieważ dostarczają one miary jedynie surogatu aktywności mózgu poprzez przepływ krwi i zmiany w tlenie we krwi.

Warto dodać, iż niektórzy oponują przeciwko tezie, że mózg ma stałą architekturę neuronową, co podważa sens mapowania. Mapy jednak wciąż powstają. Znany projekt tego typu, Human Connectome Project, (4) pozwolił np. opracować efektowne mapy ścieżek istoty białej rozciągających się w całym mózgu, co ostatecznie doprowadziło do znacznego postępu w rozwijającej się dziedzinie konektomiki.

Mapa tak, ale krótkoterminowa

Zanim zanurkujemy dalej w dziedzinę mapowania mózgu, najpierw sprecyzujmy jeszcze, o czym mówimy. Istnieją tak naprawdę dwa rodzaje mapowania mózgu. Pierwszy typ jest opisywany przez Society for Brain Mapping & Therapeutics jako „badanie anatomii i funkcji mózgu oraz rdzenia kręgowego przez wykorzystanie obrazowania, immunohistochemii, biologii

4. Jedna z wizualizacji mapy połączeń mózgowych powstała w ramach Human Connectome Project

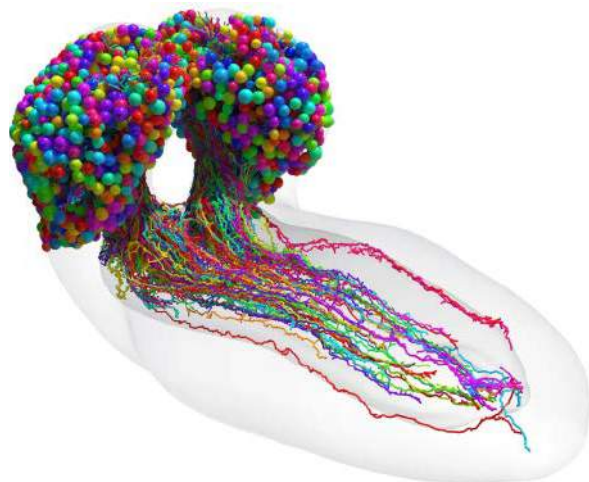


molekularnej i optogenetyki, komórek macierzystych i komórkowych, inżynierii, neurofizjologii i nanotechnologii”. Można by dodać do tej listy fizykę i fizykę kwantową. Drugi rodzaj mapowania mózgu zajmuje się identyfikacją obszarów mózgu za pomocą technologii qEEG (ilościowa elektroencefalografia) w celu ich wzmocnienia lub uzdrowienia poprzez trening neurofeedback. Praktycy neurofeedbacku twierdzą, że ma on imponującą wartość terapeutyczną dla osób cierpiących na różnego rodzaju schorzenia związane z mózgiem, w tym ADHD, autyzm, depresję i lęki. Niektórzy eksperci wyrażają jednak sceptycyzm wobec przynajmniej części tych twierdzeń.

Medyczne mapowanie mózgu działa poprzez fizyczny proces zbierania danych. Chociaż istnieją różne metody zbierania danych do produkcji wizualnej mapy mózgu, wszystkie obejmują rodzaj nieinwazyjnego, a czasem inwazyjnego, fizycznego połączenia z głową lub skórą głowy pacjenta w celu wykrycia i zebrania danych o falach mózgowych. Zebrane dane są uważane jedynie za krótkoterminowe okno aktywności mózgu, ale mogą wskazywać na rodzaje aktywności mózgu dotyczące rodzajów fal mózgowych, ich lokalizacji w mózgu, a także wszelkie nieprawidłowości lub dysfunkcje.

Mapa mózgu mogłaby więc być czymś w rodzaju atlasu, zbiorem map, które dokumentują różne ścieżki neuronowe. Jednak, w przeciwieństwie do mapy drogowej, nie może być dwuwymiarowa. Mapa mózgu samej kory mózgowej musi być trójwymiarowa. Kora mózgowa, czyli istota szara, która zawiera miliardy neuronów i synaps, jest pofalowana w taki sposób, że części, które byłyby od siebie oddalone, znajdują się blisko siebie. Jest to przydatne, ponieważ skraca odległość, jaką sygnały muszą pokonać z jednej części mózgu do drugiej. Fałdy znacznie zwiększają również powierzchnię kory mózgowej, co oznacza, że możemy upchnąć więcej szarej materii wewnątrz naszych czaszek.

W ciągu ostatniej dekady dokonano znaczącego postępu w identyfikacji i klasyfikacji różnych typów komórek występujących w mózgu przy użyciu technik takich jak wysokowydajne sekwencjonowanie pojedynczych komórek. Komórki w mózgu składają się z wielu różnych podklas, z których każda ma odrębny cel i funkcję, w tym neurony pobudzające i hamujące. Co ważne, wiele z tych typów komórek jest specyficznych dla różnych regionów mózgu. Pokłady szczegółów nie kończą się na tym. Nie tylko istnieje ogromna liczba różnych typów komórek, które należy scharakteryzować i zrozumieć, ale są one również ułożone w wysoce wyspecjalizowane i misternie zintegrowane obwody, które obejmują zakres od małych sąsiedztw komórek aż do wyższych struktur mózgu,



5. Wizualizacja konektomu muszki owocowej

takich jak ciało migdałowate i podwzgórze. Dlatego też, aby opracować kompleksowy, szczegółowy atlas mózgu, niezbędny jest kontekst przestrzenny i analiza interakcji między komórkami.

Muszka z konektomem

Kompletną mapę wszystkich połączeń w całym mózgu, zwaną konektomem, opracowano dotychczas tylko dla trzech organizmów. Organizmy te to glista i larwy dwóch morskich owadów. Wszystkie mają najwyżej kilkadziesiąt neuronów w mózgu. Mapowanie większych i bardziej złożonych mózgów pozostaje technicznym wyzwaniem i u innych gatunków, w tym człowieka, zostało dokonane tylko dla części mózgu. Niedawno międzynarodowy zespół badawczy postanowił zmapować cały mózg larwy muszki owocowej *Drosophila melanogaster* (5).

Kompletny konektom muszki owocowej składa się z ponad trzech tysięcy neuronów i ponad pół miliona synaps – punktów połączenia między różnymi neuronami. Badacze sklasyfikowali neurony na podstawie różnych sposobów, w jakie łączyły się ze sobą. Używając tego podejścia, zidentyfikowali 93 odrębne typy neuronów. Wyróżniono trzy ogólne kategorie neuronów: neurony wejściowe, które wprowadzają informacje ze zmysłów do mózgu; neurony wyjściowe, które przekazują sygnały z mózgu; oraz interneurony, które łączą się z innymi neuronami w mózgu. Zespół odkrył, że obwody neuronowe przypominają te, które można znaleźć w najnowocześniejszych komputerowych systemach uczenia maszynowego. Na przykład, sygnał mógł bieć od danego wejścia do danego wyjścia przez kilka tras o różnej liczbie kroków. Te równoległe trasy były silnie połączone, tworząc coś, co nazywa się rozproszoną siecią przetwarzania. Wiele neuronów otrzymywało również informacje zwrotne od swoich partnerów w dalszej kolejności. Ta cecha, zwana architekturą rekurencyjną, była

szczególnie powszechna w neuronach związanych z centrum uczenia się mózgu. Większość neuronów węzłowych mózgu – tych z największą liczbą połączeń – była połączona z centrum uczenia się. Huby częściej niż inne neurony tworzyły połączenia z przeciwną półkulą mózgu. Sugeruje to znaczenie komunikacji między półkulami.

W innym badaniu, którego wyniki udostępniono w marcu 2023, grupa naukowców stworzyła szczegółowy atlas komórkowy mózgu myszy. Dane przedstawione w tym artykule są wynikiem współpracy wielu ośrodków w Stanach Zjednoczonych. Naukowcy wykorzystali wiele technologii do charakteryzowania komórek mózgowych, w tym platformę MERSCOPE firmy Vizgen, która pozwoliła im na dokładne mapowanie każdej komórki do określonego miejsca w przestrzeni 3D całego mózgu myszy. W sumie 4,3 miliona komórek zostało przeanalizowanych przez MERSCOPE na potrzeby tego badania. Integracja kontekstu przestrzennego pozwoliła badaczom zrozumieć nie tylko, czym są komórki, ale jak oddziałują ze sobą i tworzą te złożone obwody neuronalne i podstruktury w mózgu. Rezultatem tej niezwykle współpracy naukowej jest pierwszy pełny atlas mózgu myszy w trzech wymiarach, który ustanawia wzorcowe narzędzie referencyjne, ujawniające unikalne i wcześniej nieodkryte cechy organizacji komórek i komunikacji w różnych regionach mózgu.

Mając dostęp do podstawowych cech komórek i zachowań wyszczególnionych w atlasie, naukowcy mogą badać dynamiczne zmiany zachodzące w różnych warunkach fizjologicznych – w zdrowiu i chorobie – w celu zidentyfikowania określonych procesów, ścieżek i typów komórek jako markerów dla wyników terapeutycznych, diagnozy choroby i celów dla nowych metod leczenia.

Gdy cyfrowy bliźniak jest u lekarza

Projekty, zapoczątkowane w ciągu ostatniej dekady, mają na celu systematyczne tworzenie map połączeń w mózgu oraz katalogowanie typów komórek i ich właściwości fizjologicznych. W 2013 r., rząd USA i Komisja Europejska uruchomiły hojne finansowanie dla tego typu projektów. Działania USA, których koszt szacuje się na 6,6 mld USD do 2027 r. – skupiły się na opracowaniu i zastosowaniu nowych technologii mapowania w ramach inicjatywy BRAIN (Brain Research through Advancing Innovative Neurotechnologies). Jednym z największych i najlepiej finansowanych wysiłków, finansowanych przez inicjatywę BRAIN, jest gigantyczny katalog typów komórek tworzony przez BRAIN Initiative Cell Census Network (BICCN), konsorcjum 26 zespołów w amerykańskich instytucjach badawczych. Katalog opisuje, ile jest różnych typów komórek mózgu, w jakich proporcjach występują i jak

są rozmieszczone przestrzennie. Komisja Europejska i jej organizacje partnerskie wydały 607 mln euro na projekt Human Brain Project (HBP), którego głównym celem jest stworzenie symulacji obwodów mózgowych i wykorzystanie tych modeli jako platformy do eksperymentów. Japonia uruchomiła projekt Brain/MINDS (Brain Mapping by Integrated Neurotechnologies for Disease Studies), którego duża część obejmuje mapowanie sieci neuronowych w mózgu małpy marmozety. Od tego czasu inne kraje, w tym Kanada, Australia, Korea Południowa i Chiny, uruchomiły lub zobowiązały się do uruchomienia wysoko finansowanych programów nauki o mózgu o bardziej rozproszonych celach. Te prace w toku już teraz generują kolosalne – i różnorodne – zbiory danych, z których wszystkie będą otwarte. Na przykład w grudniu 2020 roku HBP uruchomił swoją platformę EBRAINS, aby zapewnić dostęp do zbiorów danych w różnych skalach, narzędzi cyfrowych do ich analizy i zasobów do prowadzenia eksperymentów (<https://ebrains.eu>).

W marcu 2023 r. w „The Lancet Neurology” badacze z Human Brain Project (HBP) przedstawili nowatorskie zastosowania kliniczne zaawansowanych metod modelowania mózgu. Modele te mogą być wykorzystywane jako narzędzia predykcyjne do wirtualnego testowania hipotez i strategii klinicznych. Aby stworzyć spersonalizowane modele mózgu, naukowcy wykorzystują technologię symulacji zwaną The Virtual Brain (TVB), którą członek projektu HBP Viktor Jirsa opracował wraz ze współpracownikami. Dla każdego pacjenta modele obliczeniowe są tworzone na podstawie danych o indywidualnie zmierzonej anatomii, łączności strukturalnej i dynamice mózgu. Podejście to zostało po raz pierwszy zastosowane w epilepsji, a obecnie trwa duże badanie kliniczne. Technologia TVB pozwala klinicyście symulować rozprzestrzenianie się nieprawidłowej aktywności podczas napadów padaczkowych w mózgu pacjenta, co pomaga im lepiej identyfikować obszary docelowe.

Autorzy badań przewidują, że przyszłe postępy w mapowaniu/modelowaniu mózgu otworzą drogę do „cyfrowych bliźniaków” w medycynie mózgu. Chodzi o rodzaj spersonalizowanego cyfrowego i wirtualnego modelu (rodzaju kolejnej mapy) mózgu, który może być stale aktualizowany o zmierzone dane ze świata rzeczywistego uzyskane od jego odpowiednika, czyli pacjenta. Wirtualne mózgi mogłyby w przyszłości być wykorzystywane jako kluczowa pomoc w podejmowaniu decyzji terapeutycznych, w celu zwiększenia precyzji lokalizacji aktywności ataków oraz w planowaniu chirurgicznym, jednak obecnie modele te mają wciąż ograniczenia, takie jak niska rozdzielczość przestrzenna. ■

Mirosław Usidus



1. Świadomość, mózg i świat kwantów

Według niektórych teorii, świadomość rozumiana jako odczuwanie naszego „ja” i wpływu, jaki mamy na nasze otoczenie, odgrywa kluczową rolę w obserwacjach i pomiarach, jakich dokonujemy, a nasze doświadczenie Wszechświata przekształca go z wyobrazonego i „potencjalnego” w rzeczywisty.

Świadomość – czym jest i skąd się bierze

Z MATERII ZRODZONA CZY WRĘCZ PRZECIWNIE?

Jednym z kłopotliwych aspektów mechaniki kwantowej jest to, że cząstki subatomowe nie wydają się być w jakimkolwiek stanie, dopóki zewnętrzny obserwator ich nie zaobserwuje i nie dokona pomiaru. Akt pomiaru przekształca wszystkie możliwości potencjalne w określony wynik. W dodatku mechanika kwantowa sugeruje, że żyjemy w niedeterministycznym świecie.

Innymi słowy, przynajmniej jeśli chodzi o świat cząstek elementarnych, nie da się, bez względu na to, jak przemyślnie eksperymenty zaprojektujemy i jak doskonale znamy warunki początkowe, przewidzieć z całkowitą pewnością wyniku, np. dla protonu, czy elektronu nie ma ustalonego miejsca, w którym znajdą się za kilka sekund. Jest tylko zbiór prawdopodobieństw co do położenia. Kiedy naukowcy przeprowadzają eksperyment, otrzymują jeden z możliwych wyników i nagle Wszechświat staje się znów deterministyczny. Poznając bowiem poziom energetyczny na przykład elektronu, wiemy dokładnie, co się z nim stanie, ponieważ jego „funkcja falowa” załamuje się i cząstka wybiera określony poziom energetyczny.

W świecie makroskopowym wszystko działa „z założenia” zgodnie z deterministycznymi prawami fizyki. Nie rozumiemy związku pomiędzy niedeterministycznym światem cząstek i deterministycznym – obiektów makro, choć „na zdrowy rozum” pierwszy

jest podłożem drugiego. Standardowa interpretacja mechaniki kwantowej, nazywana kopenhaską, mówi, by zignorować te problemy i skupić się na uzyskiwaniu wyników, czy jak kto woli, pomiarów.

Jednak niektórzy teoretycy, tacy jak Eugene Wigner, wskazywali, że jedyna różnica między światem „zmierzonym” a „światem potencjalnych możliwości” polega na tym, że w jednym z nich występuje świadomy, myślący obserwator, a w drugim nie. Tak więc to, co w mechanice kwantowej nazywane jest „kolapsem” (przejście od chmury prawdopodobieństw do konkretnego wyniku) zależy od świadomości. Jedną z alternatywnych wobec kopenhaskiej interpretacji mechaniki kwantowej polega na podążaniu za powyższą logiką do jej skrajnego końca, w której to, co nazywamy pomiarem, jest tak naprawdę interwencją obdarzonego świadomością agenta w łańcuch interakcji subatomowych. Ten tok rozumowania wymaga, by świadomość była czymś innym niż cała znana nam fizyka. W przeciwnym razie naukowcy mogliby twierdzić, że świadomość sama w sobie jest tylko sumą różnych subatomowych interakcji rządzących się takimi samymi prawami mechaniki kwantowej.

Oba podejścia, pierwsze przyjmujące, że świadomość jest „czymś innym” i drugie, zakładające „kwantową świadomość”, która podobnie jak mózg (1) rządzi się tymi samymi prawami fizyki, co reszta Wszechświata, mają swoje poważne konsekwencje, z którymi wielu ludziom, nie tylko fizykom, neurologom, i w ogóle naukowcom, trudno się pogodzić.

Schodzenie głębiej i szukanie kwantów w mikrotubulach

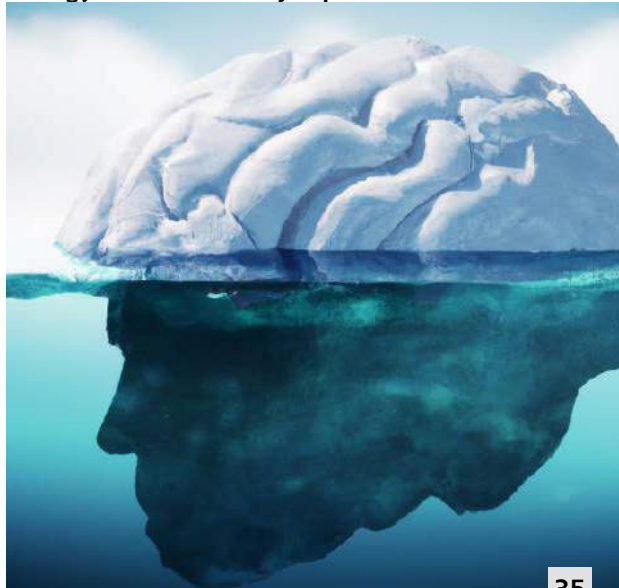
Wigner proponował, że funkcja falowa załamuje się z powodu jej interakcji ze świadomością, zaś inny fizyk, Freeman Dyson, twierdził, że „umysł, cechujący się zdolnością do dokonywania wyborów, jest w pewnym stopniu nieodróżnialny od elektronu”, co sugerowało, że świadomość ma naturę kwantową. Inni fizycy i filozofowie uznali te argumenty za nieprzekonujące. Victor Stenger, fizyk a zarazem wojownik ateizmu, scharakteryzował koncepcję świadomości kwantowej jako „mit” niemający „żadnych podstaw naukowych”, który „powinien zająć miejsce obok bogów, jednorożców i smoków”. Filozof umysłu z Australii, David Chalmers, również przeciwnik idei kwantowej natury świadomości, głosi, iż kwantowe teorie świadomości cierpią na tę samą słabość, co bardziej konwencjonalne teorie – nie ma żadnego szczególnego powodu, dla którego poszczególne makroskopowe cechy fizyczne w mózgu miałyby powodować powstanie świadomości, tak samo nie ma żadnego szczególnego powodu,

dla którego poszczególne cechy kwantowe, takie jak pole elektromagnetyczne mózgu, miałyby generować świadomość.

Niektórzy fizycy doszukują się w świadomości śladów głębszych teorii na temat rzeczywistości. David Bohm np. postrzegał teorię kwantową i względność jako sprzeczne, z czego wynikało, że we Wszechświecie musi istnieć warstwa bardziej fundamentalna. Określał ją jako kwantową teorię pola, na której miałby być oparty porządek Wszechświata, jakiego doświadczamy. Proponowany przez Bohma porządek dotyczy zarówno materii, jak i świadomości. Sugerował, że może to zarazem wyjaśnić relacje między nimi. Widział umysł i materię jako projekcje na „naszą warstwę” porządku bardziej podstawowego. Jak wyjaśniał, kiedy patrzymy na materię, nie widzimy nic, co pomogłoby nam zrozumieć świadomość (2). Twierdził, że argumenty za tym czerpie z pracy Jeana Piageta nad niemowlętami, które pokazywać miały, że małe dzieci uczą się o czasie i przestrzeni, ponieważ mają „wbudowane” rozumienie porządku fundamentalnego. Poglądy Bohma mają silny związek z wieloma tradycjami filozoficznymi, Platonem, Kantem i wieloma innymi myślicielami.

Żaden chyba z fizyków nie „namieszał” tak w teoriach świadomości jak noblista Roger Penrose z jego teorią kwantowej natury świadomości. Argumenty Penrose’a wywodziły się pierwotnie z twierdzeń o niezupełności matematyka Kurta Gödla. W swojej pierwszej książce na temat świadomości, „The Emperor’s New Mind” (1989), argumentował, że choć system formalny nie może udowodnić własnej spójności, niedające się dowieść wyniki Gödla są możliwe do osiągnięcia przez matematyków. Penrose

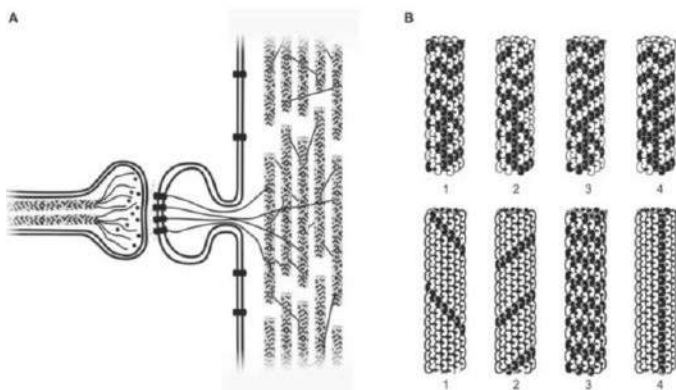
2. Mózg jako obiekt widoczny na powierzchni



potraktował to jako dowód, że ludzie-matematycy nie są formalnymi systemami dowodzenia i nie uruchamiają obliczalnego algorytmu. Penrose uznał, że załamanie funkcji falowej jest jedyną możliwą fizyczną podstawą dla tych procesów. W 2014 roku Penrose wraz z anestezjologiem Sturtem Hameroffem opublikowali twierdzenia, że odkrycie drgań kwantowych w mikrotubulach, strukturach neuronowych w mózgu (3), potwierdza ich teorię kwantowej świadomości nazwaną Orch-OR (skrót od „Orchestrated Objective Reduction”. Ich zdaniem rodzi się ona z załamania superpozycji pomiędzy mikrotubulami. Tezy te uznano za kontrowersyjne. Podjęto badania, które miały je potwierdzić lub im zaprzeczyć.

W kwietniu 2022 roku na konferencji The Science of Consciousness ogłoszono wyniki eksperymentów przeprowadzonych na kanadyjskim Uniwersytecie Alberta i amerykańskim w Princeton, które dostarczyły kolejnych dowodów na zachodzenie procesów kwantowych w obrębie mikrotubuli. W badaniu, w którym uczestniczył m.in. Hameroff, Jack Tuszyński z Alberta wykazywał, że środki znieczulające przyspieszają czas trwania procesu zwanego opóźnioną luminescencją, w którym mikrotubule i wchodzące w ich skład tubuliny ponownie emitują uwięzione światło. W innym eksperymencie Gregory D. Scholes i Aarat Kalra z Princeton użyli laserów do wzbudzenia molekuł wewnątrz tubulin, co dało wyniki sugerujące efekty kwantowe, ale możliwe też do innej interpretacji.

Również w 2022 roku inna grupa badaczy z Europy, przeprowadziła badanie, których wyniki podważają pokrewną wobec Orch-OR hipotezę autorstwa fizyka Lajosa Diósi. Z tych prac wynika zarazem, że teoria Penrose’a i Hameroffa jest „wysoko niewiarygodna”, gdyż opiera się na najprostszym typie załamania funkcji falowej związanym z grawitacją. Wynik naukowców europejskich wykluczył sformułowanie modelu Diósi-Penrose’a, które przewidywało, że skala superpozycji jest porównywalna z rozmiarem samych jąder. Przy superpozycji wielkości jądra atomowego, efekt zapadania się poszczególnych jąder węgla w białkach tubulinowych jest znikomy i dlatego wymaga ogromnej liczby jąder, aby działać zgodnie z przewidywaniami. Naukowcy ustalili, że do załamania funkcji falowej w ciągu około 0,025 s stan koherentny musiałoby utworzyć 10^{23} jąder tubulinowych, jednak, jak zauważają, w całym mózgu znajduje się tylko 10^{20} tych jednostek. W drugim scenariuszu, zakładającym większą skalę superpozycji



3. Symulacja procesu wyzwalania neuroprzekazników do mikrotubuli i przemian w tubulinach

wymagana liczba spójnych jąder tubulinowych jest mniejsza, być może wystarczy zaledwie 10^{12} . Mimo to, jak twierdzą badacze, ogólne wymagania wydają się zniechęcające, mózg bowiem musi utrzymać masę 10–16 kg w stanie koherentnym przez 25 ms w skali długości około 10 nm, co wykracza poza wszelkie znane stany superpozycji. Badacze, chociaż uważają, że teoria Orch-OR wydaje się niewiarygodna, jeśli opiera się na najprostszym modelu załamania funkcji falowej, nie wykluczają, że jeśli uda się opracować bardziej wyrafinowany model, taki, który na przykład zachowuje energię, Orch-OR być może zadziała.

Nie tak, to może inaczej

Teoria Penrose’a-Hameroffa jest chyba najbardziej znaną próbą pożenienia fizyki kwantowej ze świadomością, ale nie jedyną. Fizyk teoretyk Henry Stapp zaproponował nieco inne podejście oparte na rozważaniach na temat „wolnej woli”. Jak zwraca uwagę, „fale kwantowe” zapadają się tylko wtedy, gdy wchodzą w interakcję ze świadomością. Stan kwantowy załamuje się, gdy obserwator wybiera jedną spośród alternatywnych możliwości jako podstawę dla przyszłego działania. Stapp postuluje bardziej globalne, „umysłowe” rozumienie kolapsu funkcji falowej z mózgu. Zaś świadomość jest, według niego, fundamentem Wszechświata, co określa się jako odmianę panpsychizmu.

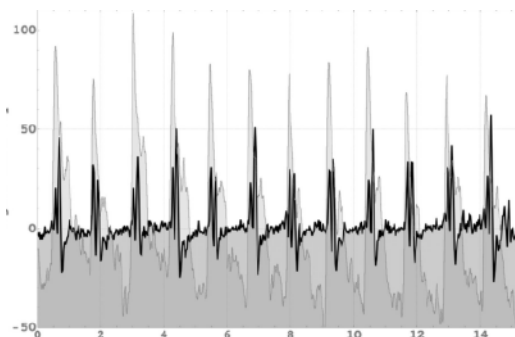
Warto też zwrócić uwagę na nowsze teorie, takie jak np. CNET, która odnosi się do mechanizmu sygnalizacji neuronowej, który wykorzystuje transport elektronów. Hipoteza opiera się częściowo na obserwacji wielu niezależnych badaczy, że tunelowanie elektronów występuje w ferrytynie, białku magazynującym żelazo, które jest powszechne w tych neuronach, w temperaturze pokojowej i warunkach otoczenia. Każde

zdarzenie tunelowania wiązałyby się z załamaniem funkcji falowej elektronu, ale załamanie to byłoby przypadkowe w stosunku do efektu fizycznego tworzego przez silne oddziaływania elektron-elektron. Chociaż mechanizm CNET nie został jeszcze bezpośrednio zaobserwowany, może być to możliwe przy użyciu fluoroforów kropek kwantowych oznaczonych do ferrytyny lub innych metod wykrywania tunelowania elektronów. Jednak pełniejsze zrozumienie tego, jak podejście CNET mogłoby odnosić się do świadomości, wymagałoby lepszego zrozumienia silnych interakcji elektron-elektron w ferrytynie.

Głównym argumentem teoretycznym przeciwko hipotezie kwantowej świadomości jest opinia, że stany kwantowe w mózgu straciłyby spójność, zanim osiągnęłyby skalę, w której mogłyby być przydatne do przetwarzania neuronowego. Przypuszczenie to zostało rozwinięte przez Maxa Tegmarka. Jego obliczenia wskazują, że układy kwantowe w mózgu ulegają dekoherencji w skali subpikosekundowej. Żadna odpowiedź mózgu nie wykazała wyników obliczeń ani reakcji w tak szybkiej skali czasowej. Typowe reakcje są rzędu milisekund, tryliony razy dłuższe niż skale czasowe subpikosekund.

Krytycy hipotezy umysłu kwantowego nie zaprzeczają, że efekty kwantowe są zaangażowane w obliczenia w mózgu. Ale ponieważ efekty te mają znaczenie tylko w bardzo małych skalach, np. określając właściwości i strukturę białek i neuroprzekazników, krytycy uważają je za nieistotne dla świadomości powstającej jako zjawisko makroskopowe.

Mimo sceptycyzmu nadchodzą kolejne wskazówki co do kwantowej natury umysłu. Niedawno, w 2022 r. uczeni z Trinity College w Dublinie, używając techniki testowania kwantowej grawitacji, dostrzegli, że w naszych mózgach może zachodzić splątanie kwantowe. W naszym mózgu, jak zauważają, znajdują się również pewne izotopy, których spiny jądrowe zmieniają reakcje naszego ciała i mózgu. Na przykład, ksenon o spinie 1/2 może mieć właściwości znieczulające, podczas gdy ksenon bez spinu nie może. W eksperymentach różne izotopy litu o różnych spinach zmieniały rozwój i zdolności rodzicielskie u szczurów. Używając MRI, naukowcy sprawdzili, czy spiny protonów w mózgu mogą oddziaływać i stać się splątane przez nieznanego pośrednika. Jeśli nieznan system może pośredniczyć w splątaniu ze znanym systemem, to, jak wykazano, nieznan system musi być kwantowy. Badacze zeskanowali grupę czterdziestu osób za pomocą MRI. Skorelowali aktywność spinową z biciem serca pacjenta, generującym sygnał zwany potencjałem bicia serca, czyli HEP, zdarzeniem elektrofizjologicznym



4. Skorelowany z sygnałem MRI zapis sygnału bicia serca HEP z eksperymentów Trinity College w Dublinie

(4), jak fale alfa czy beta. HEP jest związany ze świadomością, zależy od świadomości. Dostrzeżenie splątania w mózgu może świadczyć o tym, że mózg jest potężnym systemem kwantowym. Jeśli wyniki uda się potwierdzić, mogą one dostarczyć więcej dowodów, że mózg wykorzystuje procesy kwantowe.

Czy świadomość to lokalna sprawa mózgu?

Z naukowej perspektywy to, jak świadomość wyłania się z aktywności mózgu, jest wciąż kwestią niewyjaśnioną. Nie wiemy w sposób ostateczny, jak półtorakilogramowa bryła tkanki wewnątrz czaszki daje początek umysłowi, który jest samoświadomy i cieszy się subiektywnym doświadczeniem. W filozofii, która od tysiącleci dyskutuje nad problemem umysłu-ciała, zaproponowano wiele systemów, od materializmu (tzn. wszystko jest zależne lub redukowalne do tego, co fizyczne) do idealizmu (tzn. idee lub myśli tworzą podstawową rzeczywistość). Dla mistyków i osób zauroczonych tradycjami ezoterycznymi problem dotyczy nie tyle umysłu, co tego, jak świat fizyczny wyłania się z metafizycznej „substancji”.

W ciągu ostatnich kilku dekad większość badań nad świadomością traktowała ją jako zmienną kontrolowaną. Zredukowali oni świadomość do prostszych elementów, takich jak percepcja, i skupili się na porównaniach procesów mózgowych w warunkach świadomych i nieświadomych. W tym podejściu bada się różnice w aktywności mózgu, kiedy ten sam bodziec jest subiektywnie postrzegany, a kiedy nie. Jest to poszukiwanie neuronowych sygnałów (korelatów) świadomości (NCC).

W nauce wyróżnia się fizykalistyczne (wiążące świadomość z podłożem fizycznym) i metafizykalistyczne teorie świadomości. Jedną z nowszych teorii, należących do pierwszej grupy, jest GWT, zaproponowana przez kognitywistę Bernarda J. Baarsa w latach

80. XX wieku. Opiera się na obserwacji, że istnieją wysoce wyspecjalizowane regiony mózgu, które przetwarzają informacje lokalnie i nieświadomie, takie jak kora wzrokowa. Świadome doświadczenie pojawia się, gdy dojdzie do rozproszonej aktywności w innych obszarach mózgu, czyli „nadawania” do systemu jako całości.

Znane są także inne teorie fizykalistyczne, takie jak zintegrowana teoria informacyjna oraz nieco podobna do GWT teoria lokalnej rekurencyjności. Są też podejścia oparte na przetwarzaniu predykcyjnym, które traktują mózg jako maszynę, która dopasowuje wejścia do oczekiwań przez przetwarzanie korowe, dążąc do minimalizacji błędów w tych przewidywaniach. Trochę podobnie brzmią założenia teorii adaptacyjnego rezonansu (ART) opracowanej przez Stephena Grossberga i Gail Carpenter, która mówi, że gdy pojawia się zgodność pomiędzy oczekiwaniami a tym, co jest postrzegane, dochodzi do rezonansowej synchronizacji, która generuje skupienie uwagi napędzające szybkie uczenie się, przy minimalizacji błędów przewidywania. Jednym z elementów wspólnych dla wszystkich teorii fizykalistycznych jest redukcja niepewności, która wynika z przypisania mechanizmów do świadomości. System świadomy musi osiągnąć zunifikowany stan reprezentacyjny o wysokim poziomie informacji.

Alternatywne teorie niefizykalistyczne mogą informować o innych aspektach świadomości, które nie są całkowicie wyjaśnione przez teorie fizykalistyczne. Teorie fizykalistyczne zwykle zakładają, że świadomość jest generowana wyłącznie w mózgu i jest dla niego czymś „lokalnym”. Niefizykalistyczne modele nie przyjmują takich założeń, chociaż też próbują wyjaśnić mechanizmy mózgowe leżące u podstaw świadomości. Teorie fizykalistyczne twierdzą, że świadomość pochodzi z fizycznego podłoża, np. neuronów, które z czasem ewoluowały do coraz większej złożoności przez adaptację, prowadząc do pojawienia się świadomości. Modele niefizykalistyczne nie zakładają, że fizyczne podłoże generuje świadomość, a wiele z nich proponuje nawet, że świadomość jest w rzeczywistości czymś bardziej fundamentalnym niż materia i czasoprzestrzeń. W tym ostatnim ujęciu, które jest naturalnym poglądem dla większości starożytnych kultur, materia i czasoprzestrzeń powstają ze świadomości, a nie na odwrót. Być może niefizykalne ramy, w których świadomość jest uważana za fundamentalną i ma nielokalne właściwości (takie jak w skali kwantowej), lepiej wyjaśniłyby wiele zjawisk, np. dobrze udokumentowane doświadczenia ludzi postrzegających informacje z odległych miejsc, z przeszłości i mających wrażenia mentalne innych ludzi, które nie znajdują konwencjonalnego wyjaśnienia. Ponadto istnieją zweryfikowane przypadki

funkcjonowania mechanizmu poznawczego człowieka w sytuacjach, gdy podłoże fizyczne, czyli neurony i mózg, jest poważnie uszkodzone, zdegenerowane, niepełne, co teoretycznie wykluczałoby normalne funkcjonowanie mózgu i podważa fizykalistyczne podejście. Dla określenia tych „transcendentnych” właściwości świadomości ukuto termin „świadomość nielokalna”, który kojarzy się z nielokalnym charakterem zjawisk kwantowych. W tym sensie oparcie się na mechanice kwantowej wyjaśnia również wiele trudnych do wyjaśnienia na płaszczyźnie fizykalistycznej zjawisk związanych ze świadomością.

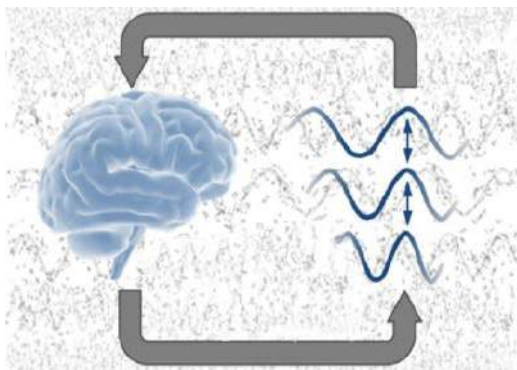
Teorie proponujące ten pomysł zostały zaproponowane przez Federico Faggina, Donalda Hoffmana, Bernarda Kastrupa, Vernona Neppe i wielu innych. Większość z tych teorii ma charakter spekulacyjny, podczas gdy inne są wspierane przez argumenty matematyczne lub dane empiryczne.

Faggin wychodzi z założenia, że rzeczywistość wyłania się z komunikacji wolnej woli ogromnej liczby świadomych podmiotów. Każde nowe samopoznanie rodzi jednostkę świadomości, która jest częścią tzw. Jednego. Faggin postrzega świat fizyczny jako metaforę wirtualnej rzeczywistości, w której zaawansowane awatary kontrolowane przez świadome istoty wchodzą w interakcje ze sobą, gdzie ciało kontrolujące awatara istnieje poza komputerem i nie jest częścią programu. Podobnie świadome byty, które kontrolują ciała fizyczne, istnieją poza światem fizycznym, który zawiera ciało.

Donald D. Hoffman proponuje model oparty na strukturze matematycznej zwanej „świadomymi agentami”, tworzącymi coś w rodzaju systemu operacyjnego i interfejsu komputera osobistego. Przestrzeń i czas wyłaniają się z interakcji świadomych agentów, które jednak nie są świadome rzeczywistej struktury przestrzeni interakcji. Twierdzi on ponadto, że równania mechaniki kwantowej można wyprowadzić ze sformalizowanych opisów interakcji między świadomymi agentami.

Bernardo Kastrup proponuje z kolei „idealizm analityczny” jako model rzeczywistości, w którym podłożem istnienia jest uniwersalna świadomość, zaś poszczególne żywe istoty są jedynie czymś w rodzaju jej fragmentów, które przez „dysocjację” tworzą subiektywny prywatny świat wewnętrzny, postrzegający niekiedy siebie jako wchodzący w interakcję ze światem transpersonalnym. Cała materia jest jedynie nazwą, którą nadajemy temu, jak wygląda świadome życie wewnętrzne.

Vernon Neppe i Ed Close twierdzą, że standardowy czterowymiarowy model fizyki powoduje wiele sprzeczności lub niewyjaśnionych rozbieżności. Dlatego proponują model matematyczny, w którym



5. Wizualizacja interakcji mózgu z polem punktu zerowego według Joachima Kepplera

istniejemy w dziewięciowymiarowej skończonej, skwantowanej, wolumetrycznej, wirującej rzeczywistości osadzonej w nieskończonej ciągłości. Model ten wymaga dodatkowego składnika, który nazwali „gimmel”, czyli masy i energii ujemnej. Matematycznie gimmel musi koniecznie istnieć w związku z każdą cząstką we Wszechświecie, aby ta cząstka była stabilna.

Całkiem nowa (2018 r.) jest teoria Joachima Kepplera, w której podstawą świadomości jest energia próżni, tzw. pole punktu zerowego. Jest to teoria panpsychizmu, w której świadomość przenika Wszechświat, jednak jest skoncentrowana i widoczna tylko w pewnych okolicznościach. W przeciwieństwie do innych teorii panpsychizmu, to nie „materia” jest świadoma, ale pusta przestrzeń. Keppler stawia hipotezę, że ludzki mózg jest jednym z fizycznych nośników, które mogą bezpośrednio oddziaływać z polem punktu zerowego (5) poprzez koncentrację na nim i w ten sposób doświadczać świadomości. Interesującym elementem tej teorii jest to, że prowadzi ona do testowalnych przewidywań, np. interakcje pomiędzy mózgiem



6. Człowiek z czujnikami na głowie

(być może poprzez zjawiska kwantowe, jak w teorii Orch-OR), a polem punktu zerowego mogłyby być ewentualnie obserwowane i mierzone.

Większość z tych teorii (jest ich znacznie więcej) zakłada, że świadomość jest fundamentalna i pierwotna wobec wszystkiego innego. Nasze subiektywne przecięcie z tą fundamentalną świadomością jest opisywane na różne sposoby, np. jako bycie interfejsem, granicą dysocjacyjną czy jednostką świadomości. Należy jednak zauważyć, że teorie fizykalistyczne nadal mają swoje miejsce w ramach podejścia nefizykalistycznego.

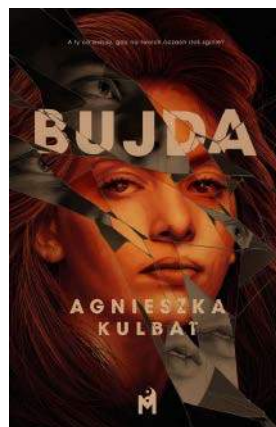
Teorie fizykalistyczne wymagają rygorystycznych testów w celu ich walidacji, teorie nielokalnej świadomości również wymagają testowania i weryfikacji eksperymentalnej. Niestety, wiele przewidywań teoretycznych jest trudnych do sprawdzenia

Bujda

Agnieszka Kulbat

Wydawnictwo Mięta, liczba stron: 320, cena: 46,99 zł

A ty co zrobisz, gdy na twoich oczach ktoś zginie? Mimo upływu czasu Kalina wciąż nie otrząsnęła się po samobójstwie starszej siostry. Pewnego dnia nieoczekiwanie staje się świadkiem morderstwa. Z dachu bloku, w którym mieszka, zostaje zepchnięty mężczyzna. Na wycieraczkę przed drzwiami swojego mieszkania nastolatka znajduje list nawiązujący do jej siostry. Czy obie te śmierci są ze sobą powiązane?



eksperymentalnie. W przypadku teorii „nielokalnych” można wskazać pewne eksperymenty (6), które wskazują ich trafność, choć oczywiście nie są jeszcze ostatecznymi dowodami.

Na przykład, jeśli świadomość byłaby nielokalna, to jednostka powinna być w stanie postrzegać informacje poza zasięgiem mózgu, ciała i zmysłów. Na przykład może być w stanie uzyskać informacje o osobie, miejscu lub przedmiocie znajdującym się w odległej lokalizacji. Takie zdolności zostały opisane w raportach na temat niejawnego programu rządu USA, który trwał od 1972 do 1995 roku i którego celem było wykorzystanie nielokalnej świadomości do szpiegostwa. W ramach eksperymentów przeprowadzono ponad pół tysiąca misji operacyjnych. Przeprowadzono wiele analiz odtajniionych eksperymentów tego typu, a wyniki wykazały dowody na nielocalne objawy świadomości.

Nielocalne rozumienie świadomości płynie również z przypadków wykazujących u ludzi zdolności poznawcze bez wcześniejszego doświadczenia czy treningu w tych umiejętnościach. Przykładem jest zjawisko mówienia przez osobę nieznanym językiem, czyli ksenoglosja. Zjawisko to odnotowywane jest od czasów starożytnych. Na przykład w 400 roku p.n.e., Platon wspomina o kapłankach na wyspie Delos, które mówiły „językami”. Istnieją również opisy w Biblii. Innym przykładem jest Indriði Indriðason (1883–1912), który najwyraźniej mówił wieloma językami, których nie znał. Podobnie Alec Harris długo rozmawiał z Alexandrem Cannonem w języku hindustani i tybetańskim, dwóch językach, których Harris nie miałby jak poznać. Chociaż są to przypadki anegdotyczne i podlegają znanym tendencyjnościom raportów z doświadczeń,

zostały one skrupulatnie udokumentowane. Podobne są przypadki „spontanicznych sawantów”, czyli osób, które, albo przez traumatyczne wydarzenie, albo bez żadnej widocznej przyczyny, nagle zyskują wyjątkowe umiejętności muzyczne lub matematyczne.

Inna zastanawiająca grupa przypadków to sytuacje, gdy poznanie, percepcja i pamięć działają normalnie, nawet jeśli mózg trudno uznać za w pełni funkcjonalny. Jest to zgodne z tym, co widzimy w zjawisku zwanym terminalną jasnością, w którym pacjenci z chorobami neurodegeneracyjnymi wykazują pozornie normalne funkcje poznawcze i jasność umysłu w okresie poprzedzającym śmierć (od kilku godzin do kilku dni). Chociaż takie doświadczenia wydają się trudne do wyjaśnienia z punktu widzenia wiedzy neurologicznej, są one opisywane w literaturze medycznej od ponad setek lat. Jeden z opisanych przypadków dotyczył pacjenta z rakiem, który z powodu zaawansowanej choroby dysponował niewielką ilością funkcjonalnej tkanki mózgowej. Jednak na godzinę przed śmiercią pacjent odzyskał świadomość i rozmawiał z rodziną przez pięć minut przed odejściem. Stawia to pod znakiem zapytania fizykalistyczną koncepcję świadomości opartą wyłącznie na aktywności fizycznej materii mózgu, sugerując, że być może istnieją aspekty świadomości, które mogą znajdować się „na zewnątrz” ciała.

Warto zwrócić uwagę, że podważanie fizykalistycznego modelu świadomości niekoniecznie prowadzi nas do zrozumienia, czym jest świadomość. Przenosi raczej problem na inny plan, wcale nie łatwiejszy do eksperymentalnego testowania. Nawet wręcz przeciwnie. ■

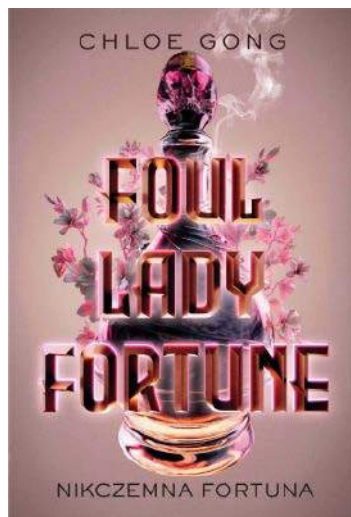
Mirosław Usidus

Foul Lady Fortune. Nikczemna Fortuna

Chloe Gong

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

Pierwsza część nowej fascynującej dylogii, w której będziemy śledzić losy nieodpasowanej pary szpiegów udających małżeństwo, co ma ułatwić im rozwikłanie sprawy brutalnych morderstw w Szanghaju lat trzydziestych. Cztery lata temu Rosalind Lang otarła się o śmierć, ale poddana dziwnemu eksperymentowi nie dość, że uratowała życie, to jeszcze zyskała pewne nietypowe właściwości. Teraz, aby zadośćuczynić za dawno popełnioną zdradę, wykorzystuje swoje umiejętności jako zabójczyni pracująca dla swojego kraju. Kiedy jednak Cesarska Armia Japońska rozpoczyna inwazję na Chiny, Rosalind otrzymuje nową misję. Seria morderstw, których sprawcami mogą być Japończycy, budzi zaniepokojenie w Szanghaju. Zgodnie z nowymi rozkazami Rosalind ma przeniknąć do powiązanej z Cesarstwem Japonii agencji prasowej i znaleźć osoby odpowiedzialne za działania terrorystyczne. Aby odsunąć od siebie podejrzenia, musi udawać żonę innego szpiega nacjonalistów, Oriona Honga, a chociaż jego nonszalanckość i maniery playboya doprowadzają ją do furii, zgadza się z nim współpracować w imię wyższego dobra. Jednakże Orion ma własne plany, zaś Rosalind skrywa tajemnice, których nie zamierza ujawniać. Dwójka szpiegów stara się dotrzeć do sedna spisku, ale z czasem zaczyna się przekonywać, że cała ta intryga ma znacznie większy zasięg i jest bardziej przerażająca, niż przypuszczali.





1. DishBrain

Mózg DishBrain (1), wyhodowany w laboratorium przez Cortical Labs kompleks liczący osiemset tysięcy neuronów, które żyją i działają w powiązaniu, nauczył się grać w grę wideo Pong w ciągu pięciu minut, kilkanaście razy szybciej, niż zajmuje to sztucznej inteligencji.

Sztuczny mózg – idea, która nie chce odejść

ORGANOID DO ORGANOIDA

W artykule opublikowanym w lutym 2023 r. we „Frontiers in Science”, duża międzynarodowa grupa badaczy prowadzona przez naukowców z Uniwersytetu Johna Hopkinsa opisuje tzw. inteligencję organoidów (OI), wschodzącą dziedzinę, w której naukowcy rozwijają biokomputery oparte na trójwymiarowych kulturach ludzkich komórek mózgowych (organoidów mózgowych) i technice interfejsu mózg-maszyna. Organoidy miałyby

służyć jako swoisty biologiczny sprzęt komputerowy bardziej wydajny niż sztuczne systemy neuronowe AI. Zresztą naukowcy chcą łączyć je z systemami uczenia maszynowego.

Thomas Hartung z Uniwersytetu Johna Hopkinsa, jeden z autorów pracy, wyjaśnia, że mózg jest okablowany zupełnie inaczej niż komputery krzemowe, które, jak przypomina, „osiągają fizycznej granice upakowania danych”. Ma około sto miliardów neuronów połączonych przez ponad 10^{15} punktów połączeń. To ogromna różnica mocy w porównaniu do znanych technik komputerowych. Ludzki mózg ma też niesamowitą zdolność do przechowywania informacji. Według cytowanej publikacji, może przechowywać szacunkowo 2500 terabajtów. Naukowcy przewidują złożone struktury komórkowe 3D (2), które byłyby połączone z systemami AI i uczenia maszynowego. Teoretycznie biokomputer podobnie jak ludzki mózg ma mieć znacznie lepszą wydajność energetyczną niż krzemowa elektronika. Jest jeszcze coś. „Nawet jeśli komputery krzemowe wykonują obliczenia na liczbach i danych szybciej niż ludzie, mózgi są znacznie

sprawniejsze w podejmowaniu złożonych decyzji logicznych, takich jak odróżnianie psa od kota”, zwraca uwagę Hartung.

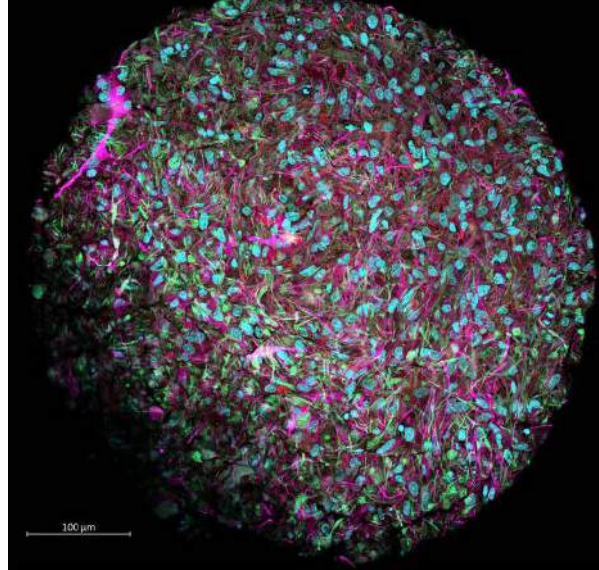
Nurt OI dąży do stworzenia biokomputerów poprzez wykorzystanie hodowanych w laboratorium organoidów mózgowych jako „biologicznego sprzętu”. Według naukowców z John Hopkins, „biokomputer” zasilany przez ludzkie komórki mózgowe może zostać opracowany w ciągu naszego życia. „Biocomputing to ogromny wysiłek zagęszczania mocy obliczeniowej i zwiększania jej wydajności w celu przekroczenia obecnych limitów technologicznych”, mówi Hartung w komunikacie.

Hartung wraz z kolegami rozpoczął prace nad wzrostem i składaniem komórek mózgowych w funkcjonalne organoidy w 2012 roku, używając komórek pochodzących z ludzkiej skóry przeprogramowanych do stanu przypominającego embrionalne komórki macierzyste. Każdy organoid zawiera około 50 tys. komórek, mniej więcej wielkości układu nerwowego muszki owocowej. Teraz badacze planują zbudować cały komputer z takich organoidów mózgowych.

Memrystory naśladujące

Inny nurt prac nad sztucznym mózgiem to próby opracowania „neuromorficznych” obwodów, które naśladują działanie ludzkich połączeń neuronowych, czyli synaps. Takie miękkie chipy komputerowe mogłyby być wszczepiane bezpośrednio do mózgu, pozwalając ludziom kontrolować ramiona robotyczne np. egzoszkieletów lub obsługując monitory komputerowe za pomocą samych myśli. Podobnie jak prawdziwe neurony, ale w przeciwieństwie do konwencjonalnych chipów komputerowych, takie urządzenia mogą wysyłać i odbierać zarówno sygnały chemiczne, jak i elektryczne (3). „Ludzki mózg pracuje na bazie chemicznych, neuroprzekazników takich jak dopamina i serotonina. Nasze materiały są w stanie oddziaływać z nimi elektrochemicznie”, pisał 2021 r. w „Annual Review of Materials Research” Alberto Salleo, materiałoznawca z Uniwersytetu Stanforda. Jego zespół stworzył urządzenia elektroniczne wykorzystujące miękkie materiały organiczne, które mogą działać jak tranzystory (które wzmacniają i przełączają sygnały elektryczne) i komórki pamięci (które przechowują informacje) oraz inne podstawowe elementy elektroniczne.

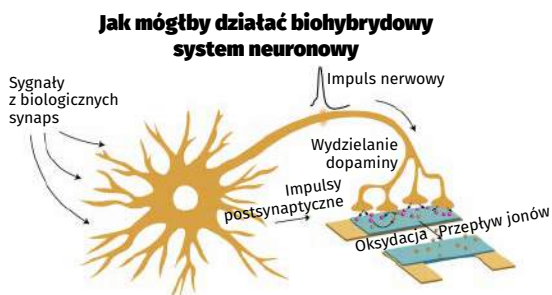
Naukowcy opracowali wiele różnych urządzeń memrystorowych (4), które naśladują zdolności przetwarzania typowe dla mózgu. Kiedy przepuszcza się przez nie prąd elektryczny, zmienia się opór elektryczny. Podobnie jak biologiczne neurony, urządzenia te wykonują obliczenia przez zsumowanie wartości wszystkich



2. Obraz organoidów mózgu z laboratorium Thomasa Hartunga

prądów, na które były wystawione. I zapamiętują wynikową wartość, jaką przyjmuje ich opór. Prosty organiczny memrystor może mieć na przykład dwie warstwy materiałów przewodzących prąd elektryczny. Po przyłożeniu napięcia prąd elektryczny napędza dodatnio naładowane jony z jednej warstwy do drugiej, zmieniając łatwość przewodzenia prądu przez drugą warstwę przy następnym kontakcie z prądem elektrycznym. Technika ta uwalnia komputer od wartości ściśle binarnych. „Nasza pamięć może mieć dowolną wartość między zerem a jedynką. Można więc dostroić ją w sposób analogowy”, podkreśla Salleo.

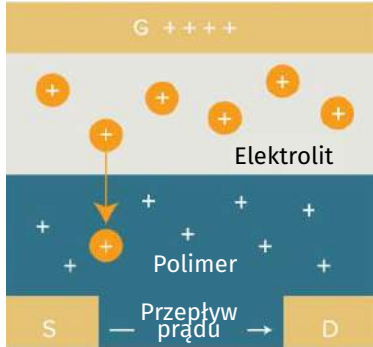
W chwili obecnej większość memrystorów i urządzeń pokrewnych nie jest oparta na materiałach organicznych, ale wykorzystuje standardową technikę układów krzemowych. Są znane prace nad organicznymi komponentami, które mogłyby wykonywać pracę szybciej, zużywając przy tym mniej energii. Chodzi o materiały miękkie i elastyczne, które także mają właściwości elektrochemiczne, pozwalając im na interakcję z biologicznymi neuronami, czyli mniej inwazyjne wprowadzanie do mózgu. Przykładem



Źródło: ADAPTED FROM S.T. KEENE ET AL./NATURE MATERIALS 2000

Knowable magazine

3. Model biohybrydowego neuronu

Podstawowe organiczne urządzenie memrystorowe

Źródło: ADAPTED FROM Y. VAN DE BURGT ET AL./NATURE ELECTRONICS 2018

Knowable magazine

4. Model działania organicznego urządzenia memrystorowego

takiego projektu są prace Franceski Santoro; inżynier elektryk, pracująca obecnie na Uniwersytecie RWTH Aachen w Niemczech, opracowuje urządzenie polimerowe, które pobiera dane z prawdziwych komórek i „uczy się” na ich podstawie. W jej urządzeniu komórki są oddzielone od sztucznego neuronu niewielką przestrzenią, podobną do synaps, które oddzielają prawdziwe neurony od siebie. Gdy komórki produkują dopaminę, substancję chemiczną sygnalizującą pracę nerwów, dopamina zmienia stan elektryczny sztucznej połowy urządzenia. Im więcej dopaminy produkują komórki, tym bardziej zmienia się stan elektryczny sztucznego neuronu, podobnie jak

w przypadku dwóch biologicznych neuronów.

Niskopoziomowe, zdecentralizowane systemy tego typu – z małymi, neuromorficznymi komputerami przetwarzającymi informacje otrzymywane przez lokalne czujniki, są, zdaniem Salleso i Santoro, obiecującą drogą dla obliczeń neuromorficznych. „Fakt, że tak ładnie przypominają elektryczne działanie neuronów, czyni je idealnymi do fizycznego i elektrycznego sprzężenia z tkanką neuronową, a ostatecznie z mózgiem”, piszą badacze.

Projekt budowy sztucznego mózgu, technologicznej kopii ludzkiego (5), to dziedzina badań, która doczekała się zainteresowania ważnych instytucji, organizacji i państw. Przykładem jest europejski Human Brain Project, który rozpoczął się w 2013 roku, choć wprost nie określa go się jako „projektu budowy sztucznego mózgu” – podkreślany jest tu raczej aspekt poznawczy, chęć lepszego poznania ludzkiego mózgu.

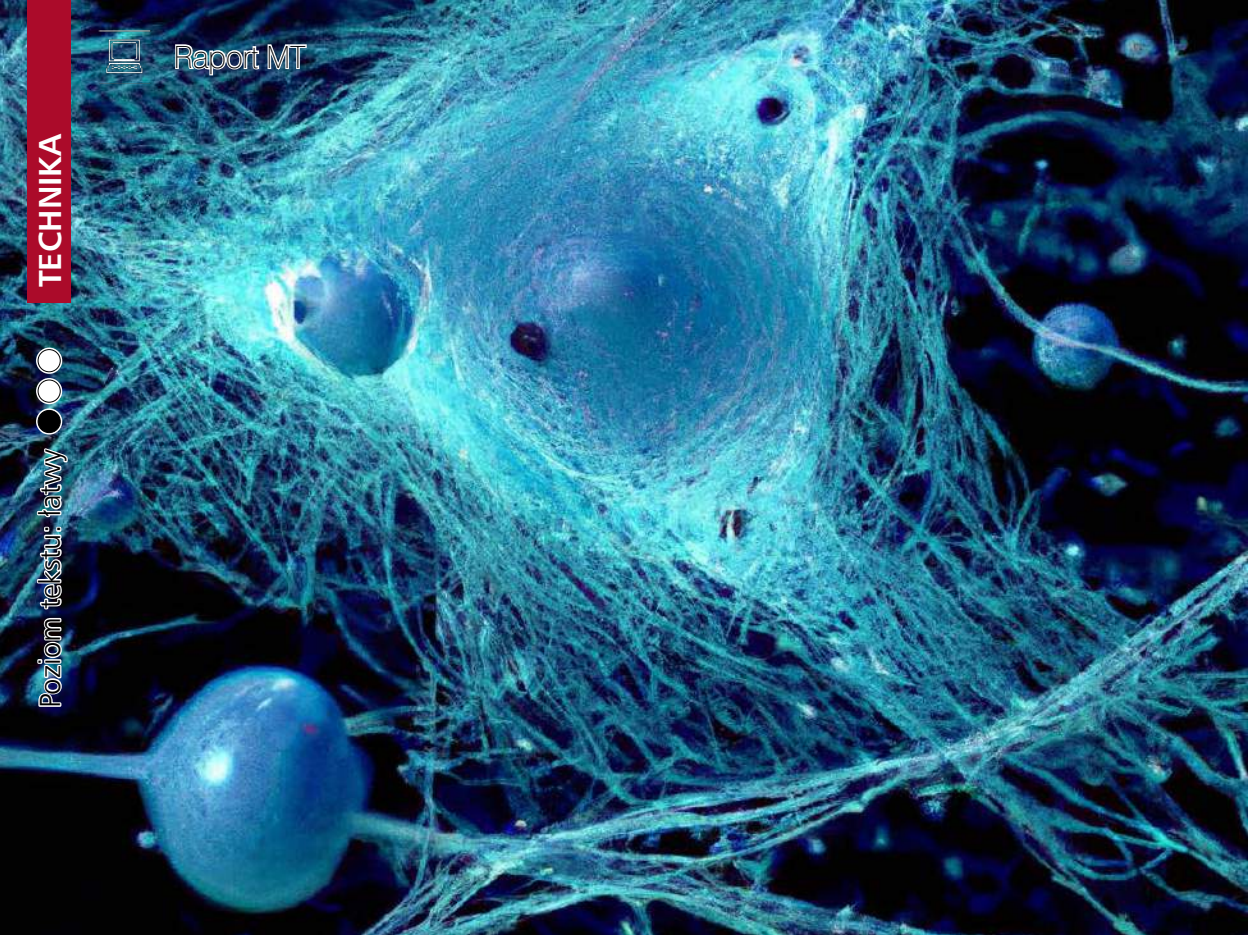
Kto wie, jeśli połączymy ogólne odkrycia Human Brain Project, mapowanie „konektomu” z rozpoznanymi algorytmami ludzkiej inteligencji i technologiami memrystorowej elektroniki wykorzystującej również algorytmy AI, to może w perspektywie kilku dekad uda nam się zbudować sztuczny mózg, wierną kopię ludzkiego. ■

Mirosław Usidus



Kompleksa DishBrain grający w Pong:
<https://bit.ly/42cXLMW>

5. Technologiczna kopia mózgu



1. Ocean cząstek elementarnych – jak go widzi DALLÉ-2

Proton – najbardziej skomplikowana rzecz, jaką znamy

**Tym bardziej inny,
im bardziej mu
się przyglądamy**

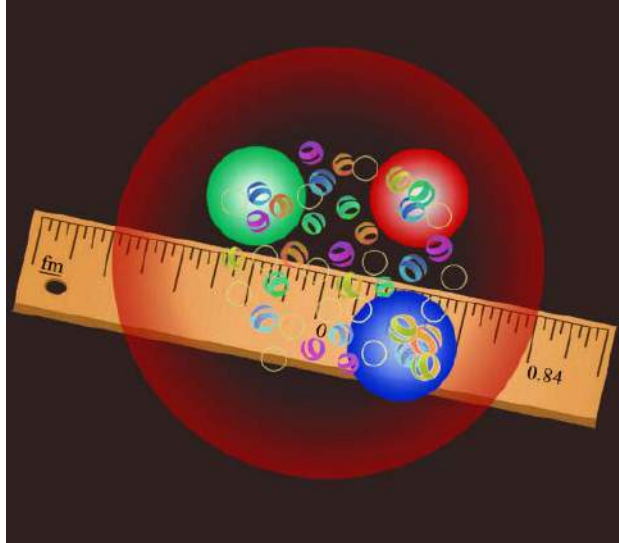
RAPORT

Wewnątrz protonu, według historycznego modelu, znajdują się trzy kwarki. W rzeczywistości jego struktura jest wielokrotnie bardziej skomplikowana. Uznaje się, że proton to prawdziwy ocean cząstek elementarnych (1), gluonów spajających kwarki i antykwarki pojawiające się i znikające nieustannie, co brzmi nieco dziwnie w uchodzącej za nad wyraz stabilną cząstce materii.

Wewnętrzna dynamika protonów jest skomplikowana, ponieważ zachodzi w nim bez liku procesów, poczynając od wymiany gluonów przez kwarki oraz oddziaływania z różnymi polami i stanami energii próżni. Wszystko to sprawia, że masa protonu, choć w różnych zestawieniach jej wartość jest podawana, w rzeczywistości jest jedynie przybliżeniem w granicach błędu do 4 proc. Niezależnie od tych niedokładności, przyjmuje się, że protony mają masę około 1836 razy większą niż elektrony, wynoszącą około $1,67 \times 10^{-27}$ kilograma. Masa protonu, pomimo że wciąż jest przedmiotem kolejnych przybliżeń, stanowi jedną ze stałych fizycznych.

Nie można także mówić o tym, że znamy definitywnie jego rozmiar i kształt. Przyjęta międzynarodowo wartość promienia ładunku protonu wynosi 0,8768 fm. Wartość ta oparta jest na pomiarach z udziałem protonu i elektronu oraz badaniach poziomów energetycznych atomów wodoru i deuteru. Jednak w 2010 roku Randolph Pohl z Instytutu Optyki Kwantowej Maxa Plancka i współpracownicy poinformowali, że dokonali wyjątkowo precyzyjnego pomiaru wielkości protonu (2), czyli dodatnio naładowanego bloku jądra atomowego. Dokonali tego przy użyciu specjalnych atomów wodoru, w których elektron, który normalnie orbituje wokół protonu, został zastąpiony mionem, cząstką identyczną z elektronem, ale 207 razy cięższą. Wynik był zastanawiający. Zespół Pohla stwierdził, że protony okrążane przez mion mają 0,84 femtometra w promieniu, 4 proc. mniej niż w zwykłym wodorze, według średniej z ponad dwóch tuzinów wcześniejszych pomiarów. Oznaczałoby to, gdyby się potwierdziło, że protony kurczą się w obecności mionów, czyli występują nieznane fizyczne oddziaływania między protonami. W styczniu 2013 roku opublikowano uaktualnioną wartość promienia ładunku protonu do 0,84087 fm. Precyzja została poprawiona 1,7 razy, co zwiększyło rozbieżności do przyjętego modelu. Korekta CODATA z 2014 roku nieznacznie zmniejszyła zalecaną wartość promienia protonu (obliczoną wyłącznie przy użyciu pomiarów elektronowych) do 0,8751 fm, ale wciąż pozostawia to rozbieżność na poziomie 5,6.

W 2019 r. w „Science” zespół fizyków pod kierownictwem Erica Hesselsa z Uniwersytetu York w Toronto opisał operację ponownego mierzenia protonu w zwykłym, „elektronowym” wodorze. Ustalono wartość 0,833 femtometra, co w granicach błędu 0,01, daje wynik dokładnie odpowiadający wartości Pohla. Oba pomiary są bardziej precyzyjne niż wcześniejsze próby i sugerują, że rozmiar protonu nie zmienia się pod wpływem takich czynników jak



2. Mierzenie rozmiaru protonu

miony zamiast elektronów wokół jądra. Uznano raczej stare pomiary z wykorzystaniem wodoru elektronowego za błędne. Przez długi czas fizykom z pomiarów i obliczeń wychodziła wartość 0,877 femtometra (jeden fm to 100 kwintylionowych części metra). Warto dodać, że po drodze, w 2017 r. niemieccy fizycy na podstawie swoich pomiarów wyliczyli promień protonu na 0,83 fm i jak uznawano, z dokładnością do błędu pomiarowego, zgadzałoby się to z wartością 0,84 fm wyliczoną w 2010 roku na podstawie promieniowania egzotycznego „wodoru mionowego”.

By nie było za prosto, w tym samym 2019 roku, inny zespół naukowców tworzących PRad w Jefferson Lab w Wirginii w USA, ponownie sprawdził wyniki pomiarów, wykorzystując nowy eksperyment rozpraszania protonów-elektronów. Naukowcy uzyskali wynik – 0,831 femtometra. Autorzy badań w pracy na ten temat, opublikowanej w „Nature”, nie uważają, że problem został już w pełni rozwiązany.

Być może pomogą inne metody pomiaru. Już w 2013 roku „Scientific American” opisał nową metodę pomiaru rozmiaru protonu, która wykorzystuje neutrino zamiast ładunku elektrycznego. Neutrino to bardzo słabo oddziałujące cząstki, które mogą przenikać przez materię i dostarczać informacji o strukturze protonu.

Przeprowadzone w ostatnim okresie eksperymenty wykazały, że wewnątrz jądra atomu protony, jak też sąsiadujące z nimi neutrony wydają się znacznie większe, niż powinny być. Fizycy stworzyli dwie konkurujące ze sobą teorie, które próbują wyjaśnić to zjawisko, a zwolennicy każdej z nich są przekonani, że druga jest błędna. Zjawisko polegające na tym, że protony i neutrony wewnątrz ciężkich jąder zachowują się tak, jakby były znacznie większe niż wtedy, gdy znajdują się poza jądrami, naukowcy nazywają efektem EMC, od European Muon Collaboration, grupy badawczej,

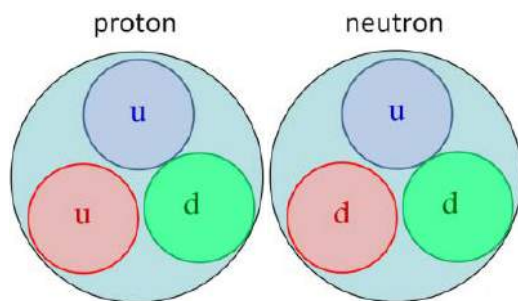


kłótnia przypadkowo je odkryła. Gwałci ono znane teorie fizyki jądrowej. Badacze przypuszczają, że kwarki składające się na nukleony oddziałują na inne kwarki w innych protonach i neutronach i rozbijają ściany oddzielające cząstki. Kwarki tworzące jeden proton zaczynają zajmować tę samą przestrzeń co kwarki z innego protonu. To powoduje, że protony (lub neutrony) rozciągają się i rozmywają. Rosną znacząco, ale jedynie przez bardzo krótki okres. Nie wszyscy fizycy zgadzają się z takim opisem zjawiska.

Problem określenia promienia dla jądra atomowego (protonu) jest podobny do problemu promienia atomu, w tym sensie, że ani atomy, ani ich jądra nie mają określonych granic. Jednakże, dla interpretacji eksperymentów rozpraszania elektronów, jądro może być modelowane jako kula o dodatnim ładunku. Ponieważ nie ma zdefiniowanej granicy jądra, elektrony „widzą” pewien zakres przekrojów, dla których można przyjąć średnią.

Masa, rozmiar czy ładunek to nie jedynie parametry, które można w przypadku protonu zmierzyć i te pomiary mają daleko idące konsekwencje dla fizyki. W artykule opublikowanym na stronie scitechdaily.com w październiku 2022 roku przedstawione zostały wyniki badań nad strukturą protonu. Struktura protonu ulega deformacji pod wpływem zewnętrznych pól elektrycznych i magnetycznych. Zjawisko to nazywa się polaryzowalnością. Polaryzowalność elektryczna i magnetyczna to miary sztywności protonu przeciwstawiającej się deformacji wywołanej przez pola. Mierząc te wielkości, naukowcy mogą dowiedzieć się więcej o wewnętrznej budowie protonu i porównać ją z teoretycznymi opisami oddziaływania promieniowania gamma z protonami. Fizycy nazywają ten proces rozpraszaniem Comptona na protonach.

Badania nad rozpraszaniem Comptona na protonach zostały przeprowadzone przez międzynarodową grupę naukowców z NNPDF (Neural Network Parton Distribution Functions) przy użyciu źródła wysokoenergetycznych promieni gamma HIGS (High Intensity Gamma Ray Source) w Triangle Universities Nuclear Laboratory w USA. Naukowcy wykorzystali uczenie maszynowe i analizę danych z eksperymentów fizyki cząstek z ostatnich 35 lat. Zbadali, jak wiązki promieni gamma o różnej polaryzacji rozpraszają się na protonach w ciekłym wodorze i jak często pojawiają się kwarki uroczone. Wyniki pokazały, że polaryzowalność elektryczna i magnetyczna protonu są znacznie większe niż wcześniej sądzono. Oznacza to, że proton jest mniej sztywny, niż się spodziewano i łatwiej ulega deformacji pod wpływem pól.



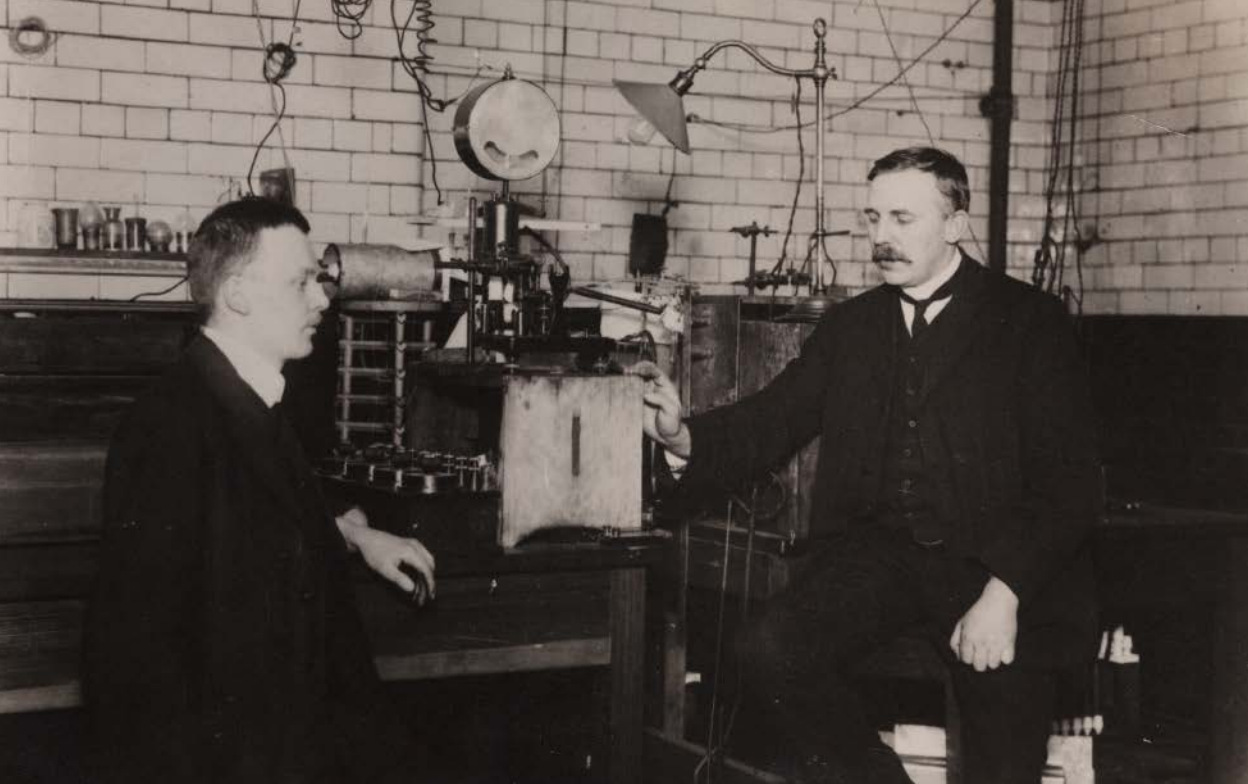
3. Modelowe przedstawienie nukleonów

Rozpada się czy nie rozpada samorzutnie?

Definiujmy mówimy, że proton to trwała cząstka subatomowa z grupy barionów o ładunku +1. Protony wraz z neutronami tworzą kategorię nukleonów (3), składników jąder atomowych. Liczba protonów w jądrze danego atomu jest równa jego liczbie atomowej, będącej podstawą uporządkowania pierwiastków w układzie okresowym. Są głównym elementem pierwotnego promieniowania kosmicznego. Proton, według Modelu Standardowego fizyki cząstek elementarnych, jest cząstką złożoną, zaliczaną do hadronów, a ściślej ich podrodziny, zwanej barionami, zbudowaną według standardowej definicji z trzech kwarków – dwóch kwarków górnych „u” i jednego kwarka dolnego „d” związanych silnym oddziaływaniem jądrowym przenoszonym przez gluony.

Nie do końca jest wyjaśniona kwestia czasu życia protonu. Według najnowszych wyników eksperymentalnych, jeżeli rozpad protonu następuje, to średni czas życia tej cząstki jest dłuższy niż $2,1 \cdot 10^{29}$ lat. Zgodnie z Modelem Standardowym proton, jako najlżejszy barion, nie może się samorzutnie rozpaść. Niezweryfikowane eksperymentalnie teorie wielkiej unifikacji generalnie przewidują rozpad protonu, zaś czas życia tej cząstki to minimum $1 \cdot 10^{36}$ lat. Proton może ulec przemianie, na przykład w procesie wychwytu elektronu. Ten proces nie zachodzi samorzutnie, lecz w wyniku dostarczenia dodatkowej energii. Jest odwracalny. Na przykład w rozpadzie beta neutron zamienia się w proton. Wolne neutrony rozpadają się spontanicznie (czas życia około 15 minut), tworząc protony.

Antyproton, cząstka elementarna będąca antycząstką protonu, różni się od niego odwrotnym ładunkiem elektrycznym, momentem magnetycznym i liczbą barionową. Ma tę samą masę i czas życia. Oddziaływanie protonu z antyprotonem powoduje ich anihilację. W odpowiednich warunkach przed anihilacją mogą utworzyć protonium, egzotyczny



4. Zdjęcie Ernesta Rutherforda i Hansa Geigera

atom zbudowany z protonu i antyprotonu, które krążą wokół siebie. Antyprotony występują również w naturze. Odnajdywane są w promieniowaniu kosmicznym docierającym do Ziemi oraz w Pasach Van Allena. Przypuszcza się, że większość tych antyprotonów ma pochodzenie wtórne, czyli powstają w wyniku oddziaływania promieniowania kosmicznego z materią międzygwiazdą.

Jak gwiazda z układem planetarnym

Przez większą część XIX wieku uważano, że atomy są najmniejszym i najbardziej podstawowym budulcem wszelkiej materii, ale gdy wiek ten zbliżał się do końca, zaczęło przybywać dowodów na to, że atomy w rzeczywistości składają się z mniejszych cząstek. Naukowcy zaczęli eksperymentować z promieniami anodowymi i katodowymi, czyli dodatnio i ujemnie naładowanymi wiązkami wytwarzanymi przez gazowe lampy wyładowcze.

W 1897 roku J.J. Thomson odkrył, że promienie katodowe to strumienie elektrycznie ujemnych cząstek subatomowych, zwanych elektronami, które były uwalniane z atomów w rurze wyładowczej. Konsekwentnie zaczęto uważać, iż promienie anodowe muszą być strumieniami jonów, które są dodatnio naładowanymi atomami. Jony wodoru zostały eksperymentalnie rozpoznane w promieniach anodowych w 1898 roku przez niemieckiego fizyka Wilhelma Wiena.

Pierwsza hipoteza budowy atomów przewidywała więc ujemnie naładowane elektrony rozsiane po amorficznie rozłożonej masie ładunku dodatniego. Nazwano ją modelem budyniu śliwkowego, a elektrony uczyniono analogicznymi do śliwek osadzonych w cieście. Brytyjski fizyk Ernest Rutherford miał poważne wątpliwości co do sensu tego modelu. W latach 1909–1911 Hans Geiger i Ernest Marsden, pod okiem Rutherforda (4) na uniwersytecie w Manchesterze wyemitowali cząstki alfa, czyli to, co dziś znamy jako jądra helu, w kierunku błony ze złotej folii. Gdyby model budyniu śliwkowego był zgodny z rzeczywistością, cząstki alfa powinny były po prostu przelatywać prostymi torami pomiędzy atomami złota lub zostać odchylone w niewielkim stopniu. Geiger i Marsden odkryli, że cząstki alfa były jednak często odchylane pod dużymi kątami, a nawet odbijały się w przeciwnym kierunku. Mogło to zachodzić jedynie wtedy, gdy w centrum atomu znajdował się skoncentrowany ładunek elektryczny, a nie był on rozłożony w całej objętości jak w modelu budyniu. To przekonało Rutherforda, że atomy zbudowane są z niewielkiego gęstego jądra o ładunku dodatnim, otoczonego pustą przestrzenią oraz elektronami krążącymi wokół jądra w pewnej odległości.

Model ten, choć uproszczony, bo nie uwzględnia kwantowego zachowania elektronów, nazywany jest modelem Bohra od nazwiska Nielsa Bohra, który wraz



z Rutherfordem złożył wszystkie elementy do w całościową teorię budowy atomu, nazywaną niekiedy modelem planetarnym, gdyż elektrony miały krążyć w nich na podobieństwo orbit wokół ciał niebieskich (5).

W 1920 roku Rutherford doszedł do wniosku, że jądra wodoru muszą być podstawowym budulcem wszystkich jąder atomowych, ponieważ wodór jest najlżejszym pierwiastkiem. Jądro wodoru nazwał protonem, za greckim słowem znaczącym „pierwszy”, ponieważ Rutherford postrzegał te cząstki jako pierwszy/pierwotny/podstawowy budulec wszystkich atomów. Dziś wiemy, że protony (i neutrony) są utworzone z jeszcze mniejszych cząstek, kwarków, i że jądro atomu jest zbudowane z protonów i neutronów (z wyjątkiem podstawowej formy wodoru, która nie ma neutronów).

Protony mają dodatni ładunek elektryczny równy ładunkowi elektronu, ale z przeciwnym znakiem. Ładunek protonu jest jedną z podstawowych stałych fizycznych i wynosi około $1,6 \times 10^{-19}$ kulomba. Jest on dokładnie równy i przeciwny ładunkowi elektronu, który wynosi $-1,602192 \times 10^{-19}$ kulomba. Ponieważ ich ładunki są równe i ponieważ drugi współmieszkaniec jądra atomowego, neutron, jest neutralny, to tak długo, jak liczba protonów i elektronów jest w atomie równa, ich ładunki się znoszą i atomy są elektrycznie neutralne. Usunięcie elektronu z otoczenia atomu zaburza równowagę pomiędzy ładunkami elektronów i protonów, a atom staje się dodatnio naładowany – jest jonem.

Tworzą w kosmosie, leczą na Ziemi

Protony wypełniają przestrzeń kosmiczną, ujawniając swoją obecność w różnych sytuacjach. Mogą być częścią atomów lub jonów w materii międzygwiazdowej lub planetarnych atmosferach. Mogą też być emitowane przez Słońce i inne gwiazdy jako promieniowanie kosmiczne lub wiatr słoneczny. Protony mogą też tworzyć jądra cięższych pierwiastków w procesach syntezy jądrowej w gwiazdach lub w supernowych.

Wolne protony występują na Ziemi sporadycznie. np. podczas burz mogą zostać wytworzone protony o energiach do kilkudziesięciu megaelektronowoltów. Przy odpowiednio niskich temperaturach i energiach kinetycznych swobodne protony będą wiązać się z elektronami. Charakter takich związanych protonów nie zmienia się jednak i pozostają one protonami. Szybki proton poruszający się w materii będzie zwalniał z powodu oddziaływań z elektronami i jądrami, aż zostanie wychwycony przez chmurę elektronową atomu. W efekcie powstaje atom protonowany, który jest związkiem chemicznym wodoru. W próżni, gdy

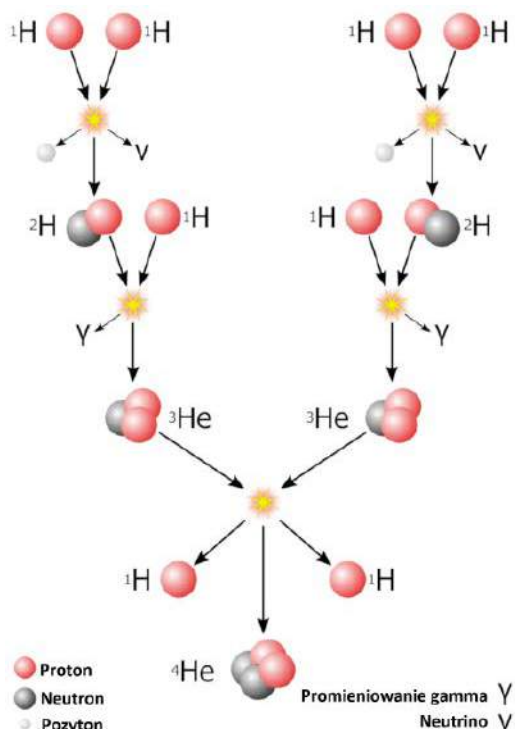


5. Planetarny model atomu – wizualizacja autorstwa generatora Blue Willow

obecne są wolne elektrony, wolny proton może „przechwycić” pojedynczy wolny elektron, stając się neutralnym atomem wodoru. Takie „wolne atomy wodoru” mają tendencję do reagowania chemicznego z wieloma innymi rodzajami atomów nawet przy niskich energiach. Kiedy wolne atomy wodoru reagują ze sobą, tworzą neutralne cząsteczki wodoru (H_2), które są najbardziej powszechnym składnikiem obłoków molekularnych w przestrzeni międzygwiazdowej.

Mgławice gwiazdotwórcze wypełnione gazem wodorowym w przestrzeni kosmicznej są często określane jako regiony H-II. Oznacza to, że wodór został zjonizowany przez światło ultrafioletowe pochodzące od znajdujących się okolicy młodych gwiazd (H-I to neutralny wodór atomowy; H-II jest zjonizowany). Energia fotonu ultrafioletowego, który wodór pochłania, jest wystarczająca, aby wybić elektron. Ponieważ atom wodoru składa się z jednego protonu i jednego elektronu, utrata elektronu pozostawia tylko proton. Kiedy proton w mgławicy odzyskuje elektron, emituje foton światła o charakterystycznej długości fali 656,3 nanometra, co jest zjawiskiem znanym jako emisja H-II.

Energia, która przejawia się jako światło i ciepło Słońca, jest generowana poprzez mechanizm, w którym znów kluczową rolę odgrywają protony, a nazywane jest to łańcuchem proton-proton (6). W jądrze Słońca temperatura osiąga 15 milionów stopni Celsjusza, co wystarcza do przeprowadzenia fuzji termojądrowej. W tych wysokich temperaturach wszystkie atomy ulegają jonizacji, a ponieważ Słońce składa się głównie z wodoru, oznacza to, że jego jądro jest wypełnione protonami. W łańcuchu proton-proton dwa protony spotykające się w takich warunkach w centrum Słońca mogą się połączyć,



6. Łańcuch proton-proton w Słońcu

wydzielając przy tym neutrino i dodatnio naładowany pozyton (który jest antymaterijnym odpowiednikiem elektronu). Utrata dodatniego ładunku zmienia jeden z protonów w neutron, a razem proton i neutron tworzą deuter (izotop wodoru). Jądro

deuteru może następnie połączyć się z innym protonem, tworząc hel-3 (zbudowany z dwóch protonów i neutronu) i emitując w tym procesie energię, która w końcu dociera do powierzchni Słońca w postaci promieniowania, które widzimy jako światło i odczuwamy jako ciepło. Tymczasem jądro helu-3 może następnie połączyć się z innym jądrem helu-3 powstałym w tym samym procesie, tworząc hel-4 (2 protony, 2 neutrony) i emitując dwa inne protony. Te inne protony mogą następnie utworzyć więcej helu-3 i tak dalej w reakcji łańcuchowej, uwalniając więcej energii w tym procesie. Słońce zawiera wystarczająco dużo jąder wodoru, czyli protonów, aby kontynuować tę reakcję przez pięć miliardów lat.

Wiatr słoneczny, który jest strumieniem naładowanych cząstek odchodzących od atmosfery Słońca, zawiera liczne protony, elektrony i różne jądra atomowe. Kiedy wiatr słoneczny zderza się z atmosferą planety, takiej jak ziemska, protony i elektrony przemieszczają się po liniach pola magnetycznego w dół, w kierunku biegunów planety, oddziałując i jonizując atomy i cząsteczki w atmosferze. Te atomy i cząsteczki następnie świecą, tworząc zorze polarne,

Czasami Słońce wybuchu w postaci rozbłysku słonecznego, często skutkującego uwolnieniem koronalnego wyrzutu masy. Te gwałtowne erupcje słoneczne mogą przyspieszać protony do wysokich energii. Owe „słoneczne cząstki energetyczne” są rozpędzane do prędkości bliskiej prędkości światła i stanowią zagrożenie radiacyjne m.in. dla astronautów i pasażerów samolotów na dużych wysokościach.

7. Nowoczesna komora do terapii protonowej



W Układzie Słonecznym występują również wysokoenergetyczne protony (i cząstki alfa) pochodzące spoza naszego Układu Słonecznego. Promieniowanie kosmiczne, które tworzą, ma wysoka energię, ale pochodzenie jego składników, w tym „kosmicznych protonów”, pozostaje zagadką. Najwyraźniej są one przyspieszane przez potężne pola magnetyczne, a głównymi podejrzanymi są aktywne jądra galaktyk i okolice czarnych dziur. Mogą to być także pozostałości po supernowych i wytwory rejonów formowania się gwiazd.

Protony i ich badania to kluczowa rzecz dla nauki i technologii. Zderzając je i analizując na różne sposoby, poznajemy tajemnice budowy atomu i jądra atomowego oraz oddziaływań między cząstkami. Używa się ich jako „czynnika roboczego” w akceleratorach cząstek, choćby w Wielkim Zderzaczu Hadronów (LHC) i innych. Protony są tam zderzane ze sobą lub z innymi cząstkami przy bardzo wysokich energiach. Protony są też wykorzystywane do leczenia niektórych chorób nowotworowych za pomocą terapii, która polega na napromieniowywaniu guzów wiązką protonów (7).

Ocean kwarków, antykwarków i gluonów

Nauczyciele fizyki w szkole średniej opisują je jako bezbarwne kulki, mające po jednej jednostce dodatniego ładunku elektrycznego, punkty odniesienia dla ujemnie naładowanych elektronów, które brzęczą wokół nich. Studenci uczelni wyższych dowiadują się, że kula jest w rzeczywistości wiązką trzech cząstek elementarnych zwanych kwarkami. Dekady badań w dziedzinie fizyki ujawniły znacznie głębszą prawdę, skomplikowaną i, co tu kryć, momentami przeczącą

naszej intuicji. Ta mała, dodatnio naładowana cząstka w sercu atomu to obiekt o niewypowiedzianej złożoności, który zmienia swój wygląd w zależności od tego, jak się go bada (8), np. naukowcy odkryli niedawno, że proton może czasami zawierać kwarki i antykwarki, kolosalne cząstki, z których każda jest cięższa od samego protonu. Jak to pojąć?

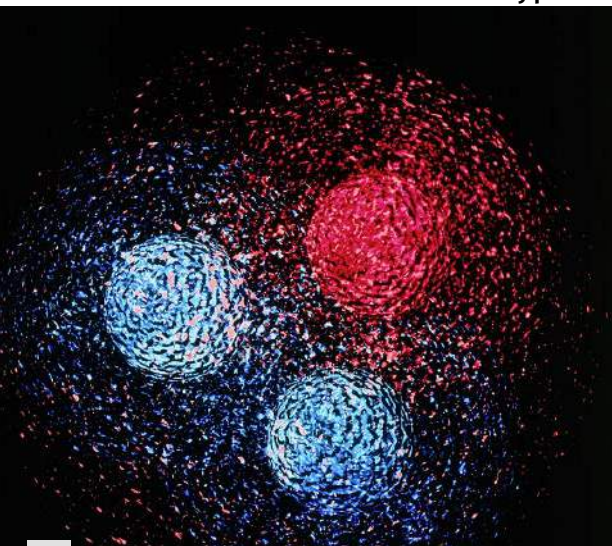
„Proton to najbardziej skomplikowana rzecz, jaką można sobie wyobrazić”, pisze w jednej z niedawnych publikacji Mike Williams, fizyk z Massachusetts Institute of Technology. „Trudno wręcz sobie wyobrazić, jak bardzo jest skomplikowany”.

Proton jest obiektem mechaniki kwantowej, który (podobnie zresztą jak elektron) istnieje jako chmura prawdopodobieństwa, dopóki eksperyment nie zmusi go do przyjęcia konkretnego stanu. Jego formy, cechy i parametry różnią się drastycznie w zależności od tego, jak badacze skonfigurują swój eksperyment. Wciąż ujawniane są kolejne tajemnice tej cząstki elementarnej. Opublikowana w sierpniu 2022 r. analiza danych wykazała, że proton zawiera ślady cząstek zwanych kwarkami powabnymi, które są... cięższe niż sam proton.

Pierwsze dowody na to, że proton składa się z wielu elementów, przyszły w 1967 roku z akceleratora SLAC na Uniwersytecie Stanforda. Odkrycie to potwierdzało wcześniejszą teorię Murraya Gell-Manna i George’a Zweiga, którzy w 1964 roku stwierdzili, że proton składa się z trzech kwarków. Po tym odkryciu, za które w 1990 roku przyznano Nagrodę Nobla w dziedzinie fizyki, badania protonu nasiliły się. Fizycy przeprowadzili setki eksperymentów. Używając elektronów o wyższych energiach, odkrywają coraz subtelniejsze cechy budowy i natury protonu. Mierząc energię i trajektorię każdego odbitego w zderzeniu elektronu, badacze określają, rozmiar, ładunek i moment pędu kwarków składowych.

Model Gell-Manna i Zweiga to eleganckie ujęcie zawartości protonu. Przewiduje dwa kwarki górne o ładunkach elektrycznych $+(2/3)$ każdy i jeden kwark dolny o ładunku $-(1/3)$, co daje całkowity ładunek protonu równy $+1$. Pięknie się to składa, ale elegancja kończy się, gdy przechodzimy do kwestii spinu protonu, który jest połówkowy. Podobną wartość ma spin każdego z jego kwarków składowych: górnych i dolnych. Fizycy początkowo przypuszczali, że połówkowe wartości dwóch kwarków górnych minus kwark dolny muszą być równe połowie jednostki dla całego protonu. Jednak w 1988 roku zespół European Muon Collaboration doniósł, że spiny kwarków sumują się do wartości znacznie mniejszej niż połowa. I to nie koniec zamieszania. Masy dwóch

8. Jedna z wizualizacji bardziej skomplikowanej niż model trzech kwarków struktury protonu



kwarków górnych i jednego kwarku dolnego stanowią tylko około 1 proc. całkowitej masy protonu. Nie pozostawało więc nic innego niż przyznanie, że wewnątrz protonu jest znacznie więcej czegoś (?) niż trzy kwarki.

Wyniki eksperymentów w Akceleratorze Hadronów-Elektronów (HERA), który pracował w Hamburgu w latach 1992–2007, uderzając w protony elektronami z grubsza tysiąc razy mocniej niż SLAC, potwierdziły nową, znacznie bardziej złożoną, niż model Gell-Manna i Zweiga, teorię budowy protonu. Była to opracowana w latach 70. XX wieku kwantowa teoria „oddziaływania silnego” między kwarkami. Teoria ta przewiduje, że kwarki połączone są cząstkami przenoszącymi owo oddziaływanie, zwanymi gluonami. Każdy kwark i każdy gluon ma jeden z trzech rodzajów ładunku („koloru”), oznaczanego jako czerwony, zielony i niebieski. Te naładowane kolorowo cząstki w naturalny sposób tworzą zgrupowanie, w której kolory składników sumują się do neutralnej (w sensie ładunku kolorowego) bieli, co w sumie daje „kolorowo neutralny” proton. Ta kolorowa teoria stała się znana jako chromodynamika kwantowa, czyli QCD. Frank Wilczek, David Gross i David Politzer skompletowali teorię QCD w 1973 roku i 31 lat później otrzymali za nią Nagrodę Nobla.

Zgodnie z QCD, gluony mogą odbierać chwilowe skoki energii. Dzięki tej energii gluon rozpada się na kwark i antykwark, zaś każdy z nich niesie ze sobą drobny ułamek pędu, zanim para anihiluje i znika. To właśnie ów ocean ulotnych gluonów, kwarków i antykwarków odkryły eksperymenty HERA. Uczniowie znaleźli w nich również wskazówki dotyczące tego, jak wyglądałby proton w potężniejszych zderzacach. Gdy fizycy dostosowali HERA do poszukiwania kwarków o niższej energii, te kwarki, pochodzące od gluonów, pojawiały się w coraz większej liczbie. Wyniki sugerowały, że w zderzeniach o wyższych energiach proton pojawi się jako chmura złożona prawie w całości z gluonów.

Ale zwycięstwo nowej teorii oznaczało gorzką pigułkę do przełknięcia. Choć QCD zgrabnie opisała szalony taniec krótko żyjących kwarków i gluonów, teoria jest bezużyteczna, gdy trzeba dokładnie opisać trzy trwałe kwarki rejestrowane w SLAC. Ponadto przewidywania QCD są klarowne tylko, gdy oddziaływanie silne jest, hm... nie za bardzo silne. A oddziaływanie to słabnie tylko wtedy, gdy kwarki są bardzo blisko siebie, tak jak w krótkotrwałych parach kwark-antykwark. W przypadku łagodniejszych zderzeń, takich jak w SLAC, gdzie proton zachowuje się jak trzy kwarki, które wzajemnie utrzymują większy dystans, kwarki te oddziałują na siebie na tyle silnie, że obliczenia

QCD stają się niemożliwe. Zatem znów trzeba zadać pytanie – o co w tym chodzi? Który proton jest „tym prawdziwym”, ten pełen kipiącej różnorodności krótko żyjących ingrediencji i jednak wciąż ten trójkwarkowy ze stabilnym układem wewnętrznym.

Powabny dziwołag

Eksperymentatorzy, na których spadł ciężar wyjaśnienia tych sprzeczności, wzięli się do pracy i w ostatnich latach zaskakują nowinkami, jeśli chodzi o strukturę protonu. Ostatnio zespół kierowany przez Juana Rojo z Narodowego Instytutu Fizyki Subatomowej w Holandii i Uniwersytetu VU w Amsterdamie po przeanalizowaniu ponad pięciu tysięcy zdjęć protonu wykonanych w ciągu ostatniego półwiecza, wykorzystując uczenie maszynowe do wnioskowania o ruchach kwarków i gluonów wewnątrz protonu, wychwytał rozmycie tła, które umknęło poprzednim badaczom. W stosunkowo łagodnych zderzeniach, które ledwo co rozrywają proton, większość pędu była zamknięta w zwykłych trzech kwarkach – dwóch górnych i dolnym. Jednak niewielka ilość pędu pochodziła z kwarka powabnego i jego antykwarka, masywnych cząstek, z których każda przewyższa masą cały proton o ponad jedną trzecią.

Wyniki Rojo i współpracowników sugerują, że kwarki powabne mają w strukturze protonu relatywnie trwałą obecność, co czyni je wykrywalnymi w łagodniejszych zderzeniach. W takich zderzeniach proton pojawia się jako kwantowa mieszanina, lub superpozycja, wielu stanów. Elektron zwykle natrafia na trzy lekkie kwarki. Ale czasami napotka rzadszą „cząsteczkę” złożoną z pięciu kwarków („pentakwarki”), takich jak kwark górny, dolny i powabny, zgrupowane po jednej stronie oraz kwark górny i antykwark powabny po drugiej. Międzynarodowa grupa naukowców, do której należy Rojo, NNPDF (Neural Network Parton Distribution Functions) zbadała, jak wiązki cząstek o wysokiej energii rozpraszają się na protonach i jak często pojawiają się kwarki powabne. Ich wyniki jeszcze bardziej niezwykle czyni odkrycie, że chociaż ten smak kwarka (powabny) jest półtora raza masywniejszy niż sam proton, to gdy jest on składnikiem protonu, kwark ten stanowi jedynie około połowy masy kompozycji nazywanej protonem. To konsekwencja egzotyki mechaniki kwantowej, która wymaga myślenia o strukturze cząstki i o tym, co można w niej znaleźć, jako o obszarze prawdopodobieństwa.

„W przyrodzie występuje sześć rodzajów [smaków] kwarków, trzy są lżejsze od protonu [kwarki górne, dolne i dziwne], a trzy cięższe [kwarki powabne,



niskie-piękne i wysokie-prawdziwe”, wyjaśnia w podcaście Nature Briefing Stefano Forte, lider zespołu NNPDF Collaboration i profesor fizyki teoretycznej na uniwersytecie w Mediolanie. „Można by pomyśleć, że tylko lżejsze kwarki trafiają do wewnątrz protonu, ale w rzeczywistości prawa fizyki kwantowej pozwalają również na to, aby cięższe kwarki znajdowały się wewnątrz protonu”.

Kwarki powabne mogą zmieniać sposób oddziaływania protonów z innymi cząstkami i polami sił. Mogą też mieć znaczenie dla astrofizyki i kosmologii, ponieważ wpływają na procesy jądrowe zachodzące w gwiazdach i we wczesnym Wszechświecie.

Kiedy protony zwane promieniami kosmicznymi przylatują tu z kosmosu i uderzają w protony w ziemskiej atmosferze, kwarki powabne wyskakujące w odpowiednim momencie obsypują Ziemię bardzo energetycznymi neutronami, wyliczyli naukowcy w 2021 roku. Mogłyby one zmylić obserwatorów poszukujących wysokoenergetycznych neutronów pochodzących z całego kosmosu.

Zespół Rojo planuje kontynuować badania protonu, poszukując braku równowagi pomiędzy kwarkami i antykwarkami. Cięższe składniki, takie jak kwark górny, mogą pojawić się jeszcze rzadziej i trudniej je wykryć.

Eksperymenty nowej generacji będą poszukiwać jeszcze bardziej nieznanymi cech. Fizycy z Brookhaven National Laboratory mają nadzieję, że w latach 30. uda się uruchomić Zderzac Elektronowo-Jonowy i kontynuować prace nad HERA, wykonując zdjęcia o wyższej rozdzielczości, które pozwolą na pierwsze trójwymiarowe rekonstrukcje protonu. EIC użyje również wirujących elektronów do stworzenia szczegółowych map spinów wewnętrznych kwarków i gluonów, tak jak SLAC i HERA mapowały ich momenty pędu. Powinno to pomóc badaczom w ostatecznym ustaleniu natury spinu protonu oraz w odpowiedzi na inne fundamentalne pytania dotyczące tej zaskakującej cząstki, która stanowi większość naszego codziennego świata.

Czas sięgnąć po neutrina

W ostatnich latach kolejne badania i eksperymenty przynoszą nową wiedzę o cząstce, która wydawała nam się tak dobrze znana. Cztery lata temu fizycy z MIT po raz pierwszy obliczyli rozkład ciśnienia wewnątrz protonu. Znaleźli oni wysokociśnieniowe „jądro” protonu, które w najbardziej intensywnym punkcie generuje większe ciśnienia niż te, które występują wewnątrz gwiazdy neutronowej. Ciśnienie wypycha wnętrze protonu, zaś otaczający jądro obszar wypycha do wewnątrz. Konkurencyjne ciśnienia działają stabilizująco na ogólną strukturę protonu. Wyniki

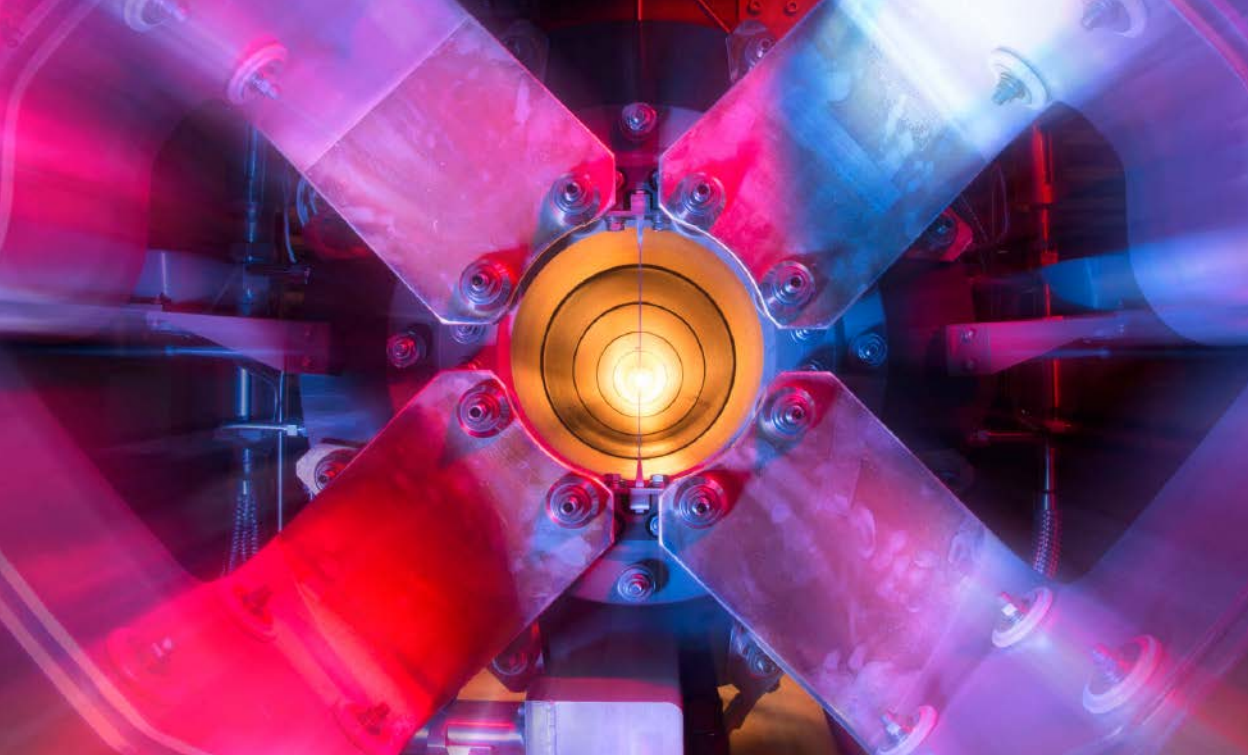
fizyków, opublikowane w „Physical Review Letters”, to pierwszy przypadek, w którym naukowcy po raz pierwszy obliczyli rozkład ciśnienia protonu, biorąc pod uwagę wkład zarówno kwarków, jak i gluonów.

Zamiast mierzyć ciśnienie protonu za pomocą akceleratorów cząstek, postanowili uwzględnić rolę gluonów, wykorzystując superkomputery do obliczania interakcji pomiędzy kwarkami i gluonami, które przyczyniają się do ciśnienia protonu. „Wewnątrz protonu znajduje się bąbelkowa próżnia kwantowa par kwarków i antykwarków, jak również gluonów, pojawiających się i znikających”, mówi Phiala Shanahan, adiunkt fizyki w MIT. „Nasze obliczenia obejmują wszystkie te dynamiczne wahania”.

W tym celu zespół zastosował technikę znaną w fizyce jako siatka QCD, czyli będącą wykorzystaniem chromodynamiki kwantowej. Obliczenia te wykorzystują czterowymiarową strukturę, czyli sieć punktów do reprezentacji trzech wymiarów przestrzeni i czasu. Badacze obliczyli ciśnienie wewnątrz protonu za pomocą równań chromodynamiki kwantowej zdefiniowanych na siatce. Zespół spędził osiemnaście miesięcy, uruchamiając różne konfiguracje kwarków i gluonów przez kilka różnych superkomputerów, a następnie wyznaczył średnie ciśnienie w każdym punkcie od środka protonu, na zewnątrz do jego krawędzi. „Po raz pierwszy przyjrzelśmy się wkładowi gluonu w rozkład ciśnienia i widzimy, że w stosunku do poprzednich wyników rozkład ciśnienia rozciąga się dalej od środka protonu”, mówi Shanahan. Innymi słowy, wydaje się, że najwyższe ciśnienie w protonie wynosi około 10^{35} paskali, czyli 10 razy więcej niż w przypadku gwiazdy neutronowej. Otaczający obszar niskiego ciśnienia rozciąga się dalej, niż wcześniej szacowano.

Potwierdzenie tych nowych obliczeń będzie wymagało znacznie potężniejszych detektorów, takich jak wspomniany Zderzac Elektronowo-Jonowy, który fizycy zamierzają wykorzystać do badania wewnętrznych struktur protonów i neutronów, bardziej szczegółowo niż kiedykolwiek wcześniej, w tym gluonów.

W 2017 r. fizykom pracującym z wykorzystaniem akceleratora cząsteczek elementarnych w Ośrodku im. Thomasa Jeffersona (Jefferson Lab) w USA, udało się po raz pierwszy w historii zmierzyć tzw. słaby ładunek protonu. Według publikacji w czasopiśmie „Physical Review Letters” udało im się także dokonać pomiarów słabego ładunku neutronu oraz kwarka górnego i dolnego. Wspomniany ładunek słaby wynika z oddziaływań słabych, które należą do grupy czterech podstawowych oddziaływań fizycznych we



9. Eksperyment MINERvA

Wszelchświecie. Poza słabymi są, oczywiście, oddziaływania silne, o których była mowa, elektromagnetyczne i grawitacyjne. Ładunek słaby protonu oznaczany jest Q_w^p , przy czym literka „w” odpowiada angielskiemu wyrazowi „weak” – słaby. Naukowcy zaznaczają, że ich pomiar w eksperymencie, który nazwali Q-weak, ma charakter pierwszego przybliżenia i będzie precyzowany w miarę analizy pozostałych danych. Cieszą się jednak, że już pierwsze dane potwierdzają przewidywania fizycznego Modelu Standardowego, który jest główną podstawą teoretyczną współczesnej fizyki.

W artykule opublikowanym w „Nature” 1 lutego 2023 roku naukowcy z międzynarodowej współpracy MINERvA przedstawili nowy sposób badania struktury protonów za pomocą neutrin, nazywanych niekiedy „cząstkami- duchami”. Bardzo trudno je wykrywać i badać. Nie mają ładunku elektrycznego i prawie nie mają masy. Rzadko oddziałują z atomami, dlatego mogą przenikać przez materię bez przeszkód. Badania zostały przeprowadzone w Fermi National Accelerator Laboratory (Fermilab) w USA, gdzie odbywał się eksperyment MINERvA (9). Jest to eksperyment fizyki cząstek, który ma na celu zbadanie oddziaływań neutrin z różnymi materiałami. Naukowcy po raz pierwszy w historii użyli w nim wiązki neutrin o wysokiej energii, w celu zbadania, jak neutrina rozpraszają się na protonach w jądrach atomowych.

Aby badać strukturę protonów za pomocą neutrin, naukowcy muszą uderzać w nie wiązką innych cząstek

i mierzyć kąty i energie rozproszenia. Dotychczas używano do tego wiązek elektronów lub fotonów (kwantów światła). Jednak te cząstki mają ładunek elektryczny i oddziałują nie tylko z protonami, ale też z elektronami w atomach. To utrudnia interpretację wyników. Neutrina natomiast nie mają ładunku elektrycznego i prawie nie oddziałują z elektronami. Dlatego można założyć, że będą służyć wyłącznie do interakcji z protonami w jądrach atomowych. Dodatkowo naukowcy z MINERvA opracowali nową technikę analizy danych z eksperymentu z neutrin, która pozwala im wyodrębnić sygnał od protonów od szumu od innych cząstek.

Wstępne wyniki pokazały, że rozproszenie neutrin na protonach zachodzi inaczej niż rozpraszanie elektronów lub fotonów. To sugeruje, że neutrina widzą inną strukturę protonów. Naukowcy spekulują, że może to być spowodowane tym, że neutrina oddziałują ze wszystkimi trzema kwarkami w protonie równomiernie, gdy elektrony i fotony oddziałują głównie z kwarkiem górnym.

Nie jest zatem wykluczone, że dzięki temu nowemu, „neutrinowemu” spojrzeniu na proton zobaczymy go znów inaczej i dowiemy się o „pierwszym” zupełnie nowych rzeczy. Czy ten obraz protonu będzie tym wreszcie „najprawdziwszym”, to już inna kwestia, na którą niestety odpowiedź mechaniki kwantowej może być znów niezbyt oczywista i niezgodna z intuicją. ■

Mirosław Usidus

O tych, co przekuli innowacyjne wizje w biznesowy sukces

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

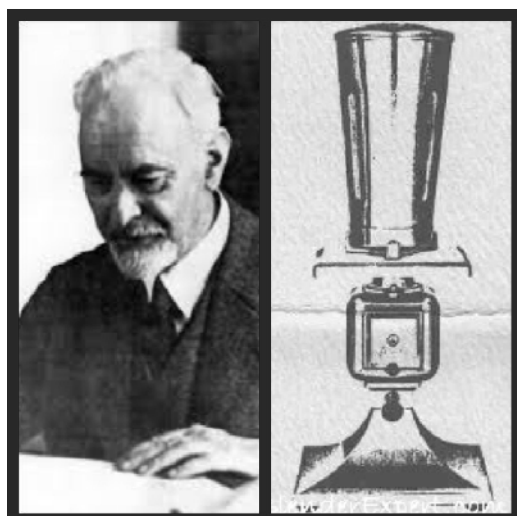
To Wy, młodzi Czytelnicy MT, macie tego dokonać! Żeby Was natchnąć dobrymi przykładami, co miesiąc przedstawiamy reprezentantów czołówki światowych liderów innowacji. Najczęściej byli oni jeszcze w wieku szkolnym lub studenckim, gdy w ich głowach rodziły się śmiałe pomysły skutkujące później powstaniem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

To oni kształtują cywilizację technologiczną.

To bohaterowie naszych czasów.

Polski wkład światowe kuchenne rewolucje – Stefan Popławski

Nie wiadomo czy wynalazłby rzecz, z której jest znany, gdyby nie wyemigrował do Stanów Zjednoczonych, gdzie panowała znacznie lepsza atmosfera dla mechanizacji kuchni. To jednak tylko gdybanie. Fakty są takie, że Stefan Popławski (**1**) wpisał się na trwałe w historię powstawania AGD.



1. Stefan Jan Popławski i blender który wynalazł

CV: Stefan Jan Popławski

Data i miejsce urodzenia: 14.8.1885, Warszawa
(zm. 9.12.1956, Racine, Wisconsin, USA)

Adres zamieszkania: nie żyje

Obywatelstwo: Polak – mieszkaniec zaboru rosyjskiego, amerykańskie

Stan cywilny: żonaty, dwoje dzieci

Majątek: trudny do oszacowania, we współczesnych skalach był co najmniej milionerem

Kontakt: nie żyje

Edukacja: nieznana

Doświadczenie zawodowe:

młode lata – praca, m.in., w parowozowni;
1919-46 – założyciel i szef własnej firmy
Stephens Tool Co.

Przyszły wynalazca przyszedł na świat 14 sierpnia 1885 roku w Warszawie. W rodzinnej kuchni państwa Popławskich zapewne najważniejszym sprzętem była kuchnia węglowa z blatem. Ale w wielkim świecie, pod koniec XIX wieku, trwająca rewolucja przemysłowa coraz śmielej wkraczała do kuchni i coraz śmielej podsuwała nowe urządzenia gospodyniom zajęтым krzątaniem przy przygotowywaniu codziennych posiłków.

Gdy Stefan Popławski miał ledwie rok Amerykanka, Josephine Cochrane, opatentowała ręcznie napędzaną zmywarkę do naczyń. Gospodyni domowa z Chicago sama zaprojektowała urządzenie bezpieczniejsze dla kruchych naczyń niż ludzkie ręce. W tym samym mniej więcej czasie niemiecki inżynier, Carl von Linde opracował pierwszą lodówkę. Na razie urządzenie wykorzystywano do produkcji lodu na potrzeby lokalnego browaru. A przedsiębiorca z Austrii, Friedrich Wilhelm Schindler wpadł na pomysł wykorzystania elektryczności do gotowania i zbudował pierwszą kuchenkę elektryczną do domowego użytku. Wynalazek został doceniony i nagrodzony w 1893 roku na światowej wystawie w Chicago.

Nikt już nie mógł mieć wątpliwości, że elektryfikacja ułatwi życie w kuchni. Kolejne wynalazki miały pojawić się wkrótce dzięki silnikowi elektrycznemu o niewielkich rozmiarach – idealnych do sprzętów AGD. Konstruktorem pierwszego miniaturowego silnika był zresztą nikt inny jak Thomas Alva Edison, który w 1880 roku opatentował maszynę o wymiarach 2,5×4 cm (4000 obr./min).

Mały Stefan miał zaledwie 9 lat, gdy znalazł się niemal w centrum tych innowacyjnych dokonań. W 1894 roku wyemigrował wraz z rodzicami, Janem i Emilią, do Stanów Zjednoczonych, gdzie rodzina osiedliła się w Racine, w stanie Wisconsin. To usytuowane nad jeziorem Michigan miasto stanowiło wówczas ważny ośrodek fabryczny, dokąd emigrowali Duńczycy, Czesi, Niemcy i czarni mieszkańcy południa USA.

Stefan mieszkał w centrum nowoczesnego i wielokulturowego tygła. Pierwszą pracę podjął jeszcze jako nastolatek, w fabryce powozów w Racine. Z zachowanych dokumentów wynika, że otrzymał powołanie do wojska w 1917 roku, ale nic nie wiadomo o przebiegu jego służby wojskowej.

„Mikser do napojów”

Kolejne lata okazały się przełomowe w jego życiu. Wkrótce otworzył własną firmę narzędziową Stephens Tool Co. i opatentował nożyczki przeznaczone do jednoczesnego cięcia w dwóch prostopadłych płaszczyznach, np. do wycinania kuponów. Wynalazek został dostrzeżony i okazał się na tyle interesujący, że w 1919 roku otrzymał zlecenie od Arnold Electric Company, szybko rozwijającej się firmy szukającej pomysłów na wykorzystanie niewielkiego silnika elektrycznego.



2. Wczesna reklama blendera

Początkowo AEC zajmowała nieduży budynek Collier przy Washington Avenue i Northwestern Tracks w Racine, ale firma miała odnieść sukces i opatentować wiele innowacyjnych urządzeń dla gospodarstw domowych. Zlecenie dla Popławskiego dotyczyło branży gastronomicznej. Stefan Popławski miał zbudować urządzenie elektryczne do przyrządzania napojów ze sproszkowanego mleka. Zamówienie złożyła popularna w mieście Racine restauracja „Malted Milk”.

O blenderach na razie nikt nawet nie słyszał. Znano już mikser elektryczny, czyli urządzenie do mieszania składników, np. na ciasto lub chleb (w produkcji od 1914 r.). Popławski miał opracować nieco inne urządzenie, które gwarantowałoby efekt aksamitnej, jednolitej konsystencji mlecznego napoju. Wymyślił sprzęt z wirującymi dwoma skrzyżowanymi ostrzami umieszczonymi na dnie kielicha, naczynia nakładanego na obudowę z silnikiem elektrycznym. Podczas ruchu obrotowego noży, dzięki sile odśrodkowej, rozdrabniana masa jest kierowana pod ostrza.

Pomysł Popławskiego, aby umieścić noże na dnie pojemnika, okazał się genialny. Działające w ten sposób urządzenie mogło siekać, rozdrabniać, mieszać nie tylko proszek na mleczny napój, ale również owoce do mlecznych koktajli.

W 1922 roku Stefan Popławski zgłosił do amerykańskiego urzędu patentowego „mikser do napojów” własnego autorstwa. „Można powiedzieć, że mój wynalazek składa się z czterech zasadniczych cech, a mianowicie: udoskonalonej konstrukcji pojemnika na napoje wraz z podstawą i mieszadłem obrotowym, stojaka z silnikiem elektrycznym i sposobem mocowania go do stojaka, specyficznej konstrukcji sterowania wyłącznikiem elektrycznym silnika i elementami współpracującymi, dzięki którym pojemnik może być podparty na stojaku i jednocześnie zapewnia łatwe połączenie między wałem silnika a wałem mieszadła oraz uruchomienie włącznika silnik”, opisał we wniosku patentowym pierwszy na świecie blender. W 1924 roku wynalazek Polaka otrzymał prawną ochronę patentową, a Arnold Electric Company ruszyła z produkcją i sprzedażą urządzeń.

Hit rynkowy

Mikser do koktajli mlecznych szybko stał się hitem restauracji oferujących dotąd głównie gazowane napoje z fontann sodowych, pozwalał wzbogacić ofertę i zwiększyć zyski. Akcje firmy AEC poszybowały wysoko i w 1926 roku przedsiębiorstwo zostało odkupione przez Hamilton Beach Manufacturing

Co., lokalną firmę z Racine, założoną w 1910 roku przez wynalazców Fredericka J. Osiusa i Chestera Beacha, a także specjalistę od reklamy Louisa Hamiltona. Mimo zmiany właściciela Popławski nie zerwał współpracy z firmą i kontynuował prace nad doskonaleniem swojego wynalazku. Opracował i opatentował jeszcze kilka wersji urządzenia, jak np. blender z silnikiem umieszczonym od góry i nożycami na dole, blender o mocy pozwalającej rozdrobnić na miążgę owoce i warzywa. Popławski opatentował również system mocowania kielicha na obudowie z silnikiem.

Uzyskał w sumie siedemnaście patentów, wśród których znajdował się również projekt pierwszego blendera do użytku domowego. Firma Popławskiego zajęła się produkcją urządzeń, a partnerzy z Hamilton Beach Manufacturing Co. wprowadzili urządzenie na rynek w 1933 roku. Do reklamy urządzeń (2) zaangażowano popularnego wówczas muzyka i radiowca Freda Waringa, który nie dość, że wziął na siebie promocję urządzenia, to dodatkowo zainwestował w projekt. Reklamowany przez artystę sprzęt kuchenny miał też wyjątkową nazwę – Waring Blender.

Zestaw sprzedawany w cenie do 30 dolarów stał się popularnym urządzeniem nie tylko w kuchni i w lokalach gastronomicznych, ale także w szpitalach i ośrodkach naukowych, np. Jonas Salk wykorzystał urządzenie do miksowania martwych wirusów podczas prac nad szczepionką przeciw polio. W niespełna dwadzieścia lat sprzedano milion urządzeń Waring Blender.

W 1946 roku Stefan Popławski postanowił wycofać się z branży i przejść na emeryturę. Sprzedał swoją firmę Stephens Tool Co. wraz z patentami kompanii John Oster Manufacturing Co. Do końca życia mieszkał w Racine, gdzie na miejscowym cmentarzu katolickim znajduje się również jego grób (3). Popławski zmarł 9 września 1956 roku w wieku 71 lat. Oryginalny blender polskiego wynalazcy znajduje się w kolekcji History Museum of Western Virginia. ■

Miroslaw Usidus

3. Grób Popławskich w Racine



Uniwersalny multimetr UT139S to nowoczesne urządzenie pomiarowe z praktycznymi funkcjami, np.: pomiar wartości skutecznej True RMS czy bezkontaktowy detektor napięcia zmiennego (NCV).

Czytelny wyświetlacz LED typu EBTN z 31 segmentowym bargrafem ułatwia odczyt pomiarów.

Pomiary:

- napięcie DC: $600V \pm (0.5\% + 2)$
- napięcie AC: $600V \pm (0.8\% + 3)$
- prąd DC: $10A \pm (0.7\% + 2)$
- prąd AC: $10A \pm (1\% + 3)$
- rezystancja: $60M\Omega \pm (0.8\% + 2)$
- pojemność: $99.99mF \pm (4\% + 5)$
- częstotliwość: $10Hz \sim 10MHz \pm (0.1\% + 4)$
- temperatura: $-40^{\circ}C \sim 1000^{\circ}C \pm (1\% + 4)$

Wyświetlacz:

- LCD (Black EBTN)
- maksymalne wskazanie: 5999
- wymiary 58 x 36mm
- podświetlenie
- bargraf 31 segmentów

Funkcje, cechy:

- True RMS
- NCV - bezkontaktowy detektor napięcia AC
- wybór zakresu: automatyczny; ręczny
- funkcja REL (pomiar wartości względnej)
- Data Hold
- test ciągłości obwodu
- test diody
- filtr LPF/LoZ (ACV)
- współczynnik wypełnienia [Duty Cycle]: 0.1% - 99.9%
- zasilanie: 2x bateria AA 1.5V
- masa: 345g
- wymiary: 170 x 80 x 48mm

W zestawie:

- miernik,
- przewody pomiarowe,
- baterie,
- sonda temperatury typu K



UT-139S
269 zł



sklep.avt.pl

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



1. Chiński Wielki Mur Internetowy

Nowa architektura sieci
– splinternetowe alternatywy

Fatum wielkiego podziału wraca do internetu

World Wide Web miał być czymś uniwersalnym, gdzie każdy miałby dostęp do wszystkiego. Ta wizja rozpada się od lat. Pojęcie splinternetu nie jest tak nowe, jednak od niedawna zaczyna być rozumiane jako tworzenie całkiem nowej infrastruktury na odmiennym, niż znane standardy internetu, „podłożu” technicznym.

Technika i protokoły stosowane w internecie utrwaliły się jako standard po nieco zapomnianej „wojnie” toczonej w latach 80. i na początku lat 90. XX wieku. Można to uznać za rywalizację Europy z USA. W USA do rodzącej się komunikacji sieciowej używano powszechnie standardu TCP/IP, natomiast w Europie proponowano model ISO/OSI (pełna nazwa ISO OSI RM, z ang. ISO Open Systems Interconnection Reference Model – model odniesienia łączenia systemów otwartych).

Słowo „wojna” ujęte jest tu w cudzysłowie i stosowane nieco na wyrost. Przede wszystkim nazwy

te obejmują w dużym stopniu różne rzeczy. Nie są to więc wykluczające się alternatywy. Istnieją pewne podobieństwa i różnice między nimi. Jedną z głównych różnic jest to, że OSI jest modelem pojęciowym, który nie służy do praktycznego stosowania w komunikacji, natomiast protokół TCP/IP jest używany do nawiązywania połączenia i komunikowania się za pośrednictwem sieci.

Ten drugi został opracowany przed modelem OSI, w Departamencie Obrony (DoD) USA. Składa się z czterech warstw, z których każda ma swoje protokoły. Protokoły internetowe to zestaw reguł zdefiniowanych



2. Rosyjski internet narodowy

dla komunikacji w sieci. Protokół TCP/IP jest obecnie uznawany za standardowy model protokołu dla sieci. Obsługuje transmisję danych i adresy IP.

Model OSI (Open System Interconnect) został wprowadzony przez ISO (International Standard Organization). Nie jest to protokół, lecz model oparty na koncepcji warstwowania, zaś podstawowym założeniem modelu jest podział systemów sieciowych na siedem warstw (ang. layers) współpracujących ze sobą w ściśle określony sposób. Model warstwowy został przyjęty przez ISO w 1984 roku. Model ISO OSI RM jest traktowany jako model odniesienia (wzorzec) dla większości rodzin protokołów komunikacyjnych.

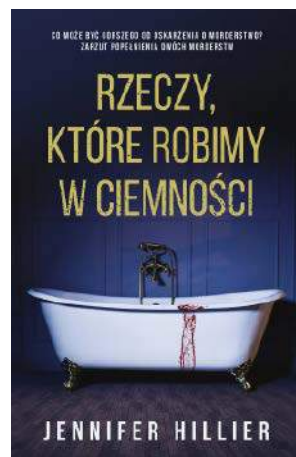
W zamierzchłej epoce wczesnego internetu w pewnym momencie wydawało się nawet, że kontynenty po obu stronach Atlantyku stworzą własne, oddzielne od siebie sieci internetowe oparte na własnych standardach. Wówczas splinternet istniałby od zarania sieci i być może w ogóle nie znalazłbyśmy pojęcia jednego globalnego internetu. Ostatecznie jednak tak się nie stało. Dominującą rolę zaczęła odgrywać specyfikacja TCP/IP. Zdecydowano się wyznaczyć ją jako podstawę dla całej infrastruktury internetowej na całym świecie. Takie rozwiązanie ma sens, gdy rozumiemy internet jako jedną sieć globalną.

Rzeczy, które robimy w ciemności

Jennifer Hillier

Wydawnictwo MUZA SA, liczba stron: 512, cena: 54,90 zł

Paris Peralta, instruktorka jogi, siedzi na podłodze w łazience z brzytwą w ręku. W wannie pełnej krwi leży jej martwy mąż. Ona nic nie pamięta z ostatniego wieczoru, nie wie, co działo się w nocy, wciąż słyszy tylko rozdzierający krzyk, który zbudził ją rano. Kobieta zostaje aresztowana. Wszystko wskazuje na to, że będzie oskarżona o zabicie męża. Jest załamana i przerażona. Ale nie to oskarżenie przeraża ją najbardziej. Najbardziej Paris obawia się medialnej wrzawy, jaka z pewnością rozpęta się wokół niej po tym wszystkim. Dziennikarze zaczną szperać w jej biografii i szybko dokopią się do tego, od czego uciekła i co chciała zostawić za sobą raz na zawsze. Dopadnie ją ponura przeszłość. A na to się właśnie zanosi, bo akurat warunkowo zwolniono z więzienia Ruby Reyes, skazaną dwadzieścia lat wcześniej za brutalne zabicie kochanka. Ona wie, kim naprawdę jest Paris i chętnie podzieli się tą wiedzą z innymi... Czy może być coś gorszego od oskarżenia o morderstwo?! Zarzut popełnienia dwóch morderstw.



Wielkie narodowe grodzenie

Dziś nie każde państwo na świecie jest zainteresowane tak rozumianą siecią, a na pewno nie swobodnym, niekontrolowanym, przepływem informacji. Chiny czy Rosja dążą do stworzenia „narodowych internetów”. Grodzenie sieci cenzorskimi blokadami, filtrami, choćby tak jak Chiny – „wielkim firewallem” (1), choć odcina niepożądaną przez tamtejsze rządy komunikację, to wciąż nie zmienia podstaw internetowej techniki, która nadal jest oparta na protokołach TCP/IP.

Stąd myśl, popularna szczególnie w kręgach zbliżonych do władzy, by stworzyć odrębną sieć od podstaw, korzystając z własnych standardów i infrastruktury. Chiny, przez firmę Huawei, zaproponowały stworzenie nowego protokołu internetowego, który miałby zastąpić TCP/IP. Według amerykańskiej grupy ds. cyberbezpieczeństwa o nazwie Just Security, miałyby to bardzo negatywne skutki dla całego internetu. „Sama propozycja nowego systemu adresów IP opiera się na wadliwych fundamentach technicznych, które grożą fragmentacją internetu, bałaganem braku interoperacyjności, zmniejszeniem stabilności, a nawet bezpieczeństwa sieci”, czytamy w komunikacie przedstawicieli Just Security.

Alternatywne architektury internetu nieoparte na TCP/IP to koncepcje opierające się na różnych protokołach i sposobach przesyłania informacji, odmiennych niż ten wykorzystywany w obecnej architekturze sieci. Oto kilka koncepcji, propozycji i projektów, które są rozpatrywane, nie zawsze zresztą w celach politycznych:

– Named Data Networking (NDN) to projekt, który proponuje zmianę metody przesyłania danych przez sieć. W NDN dane są identyfikowane na podstawie ich treści, a nie adresu sieciowego, jak w przypadku TCP/IP. Zamiast przesyłania informacji z jednego punktu do drugiego, NDN skupia się na przesyłaniu żądań danych z miejsca, gdzie są one potrzebne, do miejsca, gdzie są przechowywane.

– Inna koncepcja, Delay-Tolerant Networking (DTN), stawia sobie za cel zapewnienie komunikacji w sytuacjach, gdy konwencjonalne połączenia sieciowe nie są dostępne lub nie są niezawodne. DTN opiera się na przesyłaniu danych przez pośredników, którzy przemieszczają się w sieci i przechodzą przez nią. Dane są przesyłane z jednego pośrednika do drugiego, aż do osiągnięcia celu.

– Z kolei Software-Defined Networking (SDN) proponuje oddzielenie warstwy zarządzania siecią od warstwy przesyłania danych. W SDN centralny kontroler zarządza siecią, decydując o tym, jakie dane

są przesyłane i jakie urządzenia sieciowe są wykorzystywane w tym procesie.

Odciać się można, ale to boli i kosztuje

Wiadomo, że Rosja rozpoczęła już działania praktyczne zmierzające do „oddzielenia” rosyjskiej sieci od globalnego internetu (2). W marcu 2022 roku władze Rosji nakazały wszystkim państwowym portalom przejście na lokalny hosting i serwery domenowe. Państwo to tworzy własny system nazw domen. Ma to być osiągnięte poprzez budowę nowych serwerów, systemów zarządzania ruchem internetowym oraz sieci przesyłających ruch wewnątrz Rosji. Rosyjskie władze mają w ten sposób kontrolować treści przesyłane w sieci. Jednak odciecie od globalnego internetu uchodzi za dość ryzykowne, oznacza koszty i to ogromne koszty. To odciecie od krwioobiegu światowego handlu, nowych technologii. Ostatecznie prowadzi do stagnacji.

Nie ma oficjalnych informacji na temat protokołu, na którym miałby się opierać rosyjski „narodowy internet”. Według doniesień medialnych i ekspertów z dziedziny technologii, rosyjski rząd planuje wykorzystać własny Domain Name System (DNS). DNS w globalnej sieci służy do przyporządkowywania adresów IP do nazw domenowych. Korzystanie z własnej wersji DNS pozwoliłoby rządowi Rosji na kontrolowanie i filtrowanie ruchu internetowego w kraju poprzez modyfikowanie bazy adresów IP i przekierowywanie ruchu na inne serwery.

Jeśli chodzi o chiński projekt „narodowego internetu”, to tamtejsze władze również planują stworzyć swoją własną infrastrukturę sieciową, która będzie oddzielona od globalnego internetu i umożliwi im pełniejszą kontrolę nad ruchem internetowym w kraju. Ma to być osiągnięte przez budowę nowych serwerów, systemów zarządzania ruchem internetowym oraz sieci przesyłających ruch wewnątrz Chin. W architekturze chińskiego „narodowego internetu” wyróżnia się kilka elementów, w tym: Great Firewall – czyli Wielki Chiński Mur Internetowy, który jest systemem blokującym dostęp do niepożądanych stron internetowych oraz monitorującym ruch w sieci, sieci wewnętrzne, które będą izolowane od globalnego internetu i będą wykorzystywane do przesyłania ruchu między różnymi instytucjami i organizacjami w kraju, narodowy Domain Name System (DNS), co umożliwi im kontrolowanie przypisywania adresów IP do nazw domenowych. Chociaż oficjalnie chińska koncepcja nie zakłada zmiany podstawowej architektury protokołu TCP/IP, to władze mogą wprowadzić modyfikacje w jego funkcjonowaniu, aby umożliwić lepszą kontrolę ruchu internetowego w kraju, np. w sposobie

przetwarzania pakietów danych blokowanie niepożądanych protokołów lub wprowadzenie dodatkowych mechanizmów bezpieczeństwa.

Uważa się, że chińska koncepcja „narodowego internetu” jednak nie zakłada całkowitego odłączenia od globalnej sieci internetowej, ale raczej tworzenie oddzielonych i kontrolowanych przez rząd sieci, które będą działały w ramach istniejącej architektury TCP/IP.

Big Tech też się grodzi

Często uważa się, że poprzegradzany takimi czy innymi barierami splinternet przyszłości może przypominać system narodowych fortec, każda broniona przez wiele wałów obronnych i fos do przebycia z pojedynczymi, ściśle kontrolowanymi wejściami dla dozwolonego przepływu danych.

Splinternet powstaje zresztą nie tylko wskutek wysiłków antyzachodnich reżimów. Nawet w obrębie samego Zachodu sieciowy teren jest podzielony granicami (3), może nie absolutnie szczelnymi ale wyraźnie widocznymi, np. w sieciach mobilnych jest świat Apple i iPhone'a oraz świat Androida. Funkcjonują równolegle, w sposób odrębny. Gdy użytkownik dokona wyboru, jest przypisany i w dużym stopniu zamknięty w konstelacji zastrzeżonych platform. Strażnicy tych światów, firmy Apple i Google, mają ogromną władzę, określając zasady funkcjonowania swoich sklepów z aplikacjami i decydując, kto może, a kto nie może wejść ze swoją ofertą. Jeśli ktoś nie spełnia arbitralnych kryteriów, to zwykle jest bezceremonialnie wyrzucany. Także między tymi komercyjnie wytyczonymi uniwersalami bardzo często brak interoperacyjności. Prawie nie ma międzyplatformowych standardów przesyłania wiadomości. I to nie jest przypadek.



3. Dzielenie internetu przez potentatów Big Tech

Każdy gracz chce utrzymywać swojego użytkownika w zamknięciu, minimalizować jego możliwości przechodzenia do alternatywnego świata.

Pożegnanie World Wide Web z neutralnością jest już właściwie faktem. Na każdym prawie kroku natrafiamy na bariery w swobodnym korzystaniu z informacji. Widzimy komunikaty w stylu „ta treść ze względu na takie albo inne ograniczenia jest niedostępna w twoim kraju”. Wolny i otwarty internet jest zastępowany przez konstelację coraz szczelniej zamkniętych baniek. Staje się coraz bardziej zatamizowany. Splinternet, który mógł powstać u samego zarania sieci, nieubłaganie rodzi się na naszych oczach. ■

Mirosław Usidus

Finding Back to Us

Bianca Iosivoni

Wydawnictwo Jaguar, liczba stron: 400, cena: 49,90 zł

Callie, studentka medycyny, na prośbę swojej siostry i przybranej mamy wraca do rodzinnego domu w małym miasteczku na południu Stanów. Chce spędzić wakacje z siostrą, zanim ta wyruszy w podróż dookoła świata. Ale dziewczyna nie ma pojęcia, że w tym samym czasie zjawia się tam również Keith, którego szczerze i serdecznie nienawidzi od lat. To przez niego jej ojciec zginął w wypadku samochodowym. Najchętniej do końca swoich dni nie oglądałaby Keitha na oczy. Widok chłopaka sprawia, że natychmiast powracają wspomnienia najgorszych chwil, jakie przeżyła. Ale ku swojemu zdziwieniu, Callie zdaje się doświadczać też zupełnie innych uczuć..., które powodują totalny mętlik w jej głowie. Bo Keith to nie tylko człowiek, którego nienawidzi najbardziej na świecie, ale także jej przybrany brat...



Pierwiastkowy high life

Poziom produkcji metali we współczesnym świecie sięga setek tysięcy i milionów ton rocznie – aż tyle wynosi zapotrzebowanie naszej żarłocznej cywilizacji. Są jednak i takie, których całoroczny urobek można załadować na ciężarówkę lub co najwyżej do kilku wagonów. W przeciwieństwie do niewielkiej podaży, ich znaczenie jest jednak ogromne.

Pomijając nietrwałe, promieniotwórcze pierwiastki oraz te, dla których technika nie znalazła jeszcze szerszych zastosowań, platynowce są metalami otrzymywanymi w najmniejszych ilościach. Dla porównania: złoto z roczną produkcją wynoszącą około 3 tys. ton kilkakrotnie przewyższa wydobycie ich wszystkich razem wziętych. Platynowce są najszlachetniejszymi metalami i do tego w większości droższymi od złota, w artykule przyjrzymy się zatem pierwiastkowemu „wielkiemu światu”. W domowym laboratorium nie wykonasz z nimi doświadczeń z powodu braku dostępnych surowców (poświęcenie biżuterii lub samochodowego katalizatora z oczywistych względów nie wchodzi w rachubę), ale znajomość historii i współczesności platynowców jest niezbędna w wykształceniu chemika.

Platynowców portret zbiorowy

Chemiczna tradycja złączyła pierwiastki z grup 8, 9 i 10 w poziome rzędy – **triady**. Pierwszą z nich, żelazowce (żelazo, kobalt i nikiel), poznałeś w artykułach z roku 2021. Pod nimi leżą **platynowce lekkie** (ruten, rod i pallad) o gęstościach wynoszących około 12 g/cm³, a jeszcze niżej **platynowce ciężkie** (osm, iryd i platyna). Te ostatnie mają największe gęstości wśród metali: dla osmu i irydu przekraczają one 22 g/cm³ (trwają spory, który z nich ma największą koncentrację masy), w przypadku platyny jest to około 21,5 g/cm³ (1). Platynowce są również trudno topliwe: osm topi się dopiero w ponad 3000°C (z tego powodu był stosowany jako drucik żarowy w żarówkach, ale wyparł go tańszy i bardziej dostępny wolfram), a najłatwiej topliwy pallad w ponad 1500°C. Wygląd platynowców nie odbiega od wizerunku innych metali: srebrzystoszare (jedynie osm ma niebieskawy odcień), do tego bardzo twarde (2).

Pod względem chemicznym platynowce należą do najbardziej odpornych metali, na niektóre z nich nie działa nawet **woda królewska** (mieszanka kwasów azotowego i solnego), która bez trudu



1. Blaszka platynowa

roztwarza złoto. Na powietrzu nie ulegają zmianie (nie korodują), ale oczywiście i na nie chemicy znaleźli sposoby, a związki platynowców są liczne.

Chemicy w XIX wieku zauważyli, że żelazowce i platynowce wykazują duże podobieństwo właściwości w rzędach, czyli wspomnianych triadach. Nie oznacza to, że pierwiastki te nie podlegają prawu okresowości – podobieństwa grupowe są również wyraźne. Ruten i osm osiągają najwyższą wartościowość wśród metali – VIII, żelazo z tej samej grupy jest maksymalnie VI-wartościowe. Rod i iryd, podobnie do kobaltu, tworzą wielobarwne związki, natomiast nikiel, pallad i platyna to doskonałe katalizatory reakcji z wodorem.



2. Ampułki z irydem (u góry) i osmem (na dole, widoczny niebieskawy odcień metalu). Pozostałe platynowce wyglądają podobnie

Gdzie ich szukać...

Niestety, tylko w nielicznych miejscach na Ziemi. Podobnie jak złoto, występują praktycznie tylko w postaci wolnej, jako **platyna rodzima** – naturalny stop, oprócz platyny zawierający 10% żelaza i około 2% pozostałych platynowców. Stanowią one również domieszki w rudach miedzi i niklu, skąd są otrzymywane podczas przeróbki szlamów anodowych, czyli pozostałości po elektrolitycznym oczyszczaniu tych metali. Ze względu na podobieństwo właściwości oraz małą reaktywność rozdzielenie platynowców stanowi jeden z najtrudniejszych do przeprowadzenia procesów chemicznych, którego przebieg, dostosowany do składu surowca, jest ściśle chronioną tajemnicą producentów.

Pod względem rozpowszechnienia w powierzchniowej warstwie Ziemi platynowce mieszczą się w ósmej dziesiątce pierwiastków, prawie na samym końcu listy (mniej jest właściwie tylko gazów szlachetnych, oprócz argonu). Łączna ich zawartość sięga kilku milionowych części procenta – więcej niż jest złota i mniej od srebra. W zależności od źródeł, najpospolitszym platynowcem jest pallad lub ruten (niełatwo oszacować tak małe wielkości), a najrzadszymi – iryd i rod.

Z platynowców najwięcej otrzymuje się palladu (210 ton w roku 2022) i platyny (190 ton), produkcja pozostałych jest znacznie mniejsza, rzędu ton. Poza konkurencją pozostaje Republika Południowej Afryki (łącznie 220 ton palladu i platyny), ponad 100 ton dostarcza Rosja, z pozostałych krajów liczą się Zimbabwe i Kanada (przeróbka rud niklu). Jak

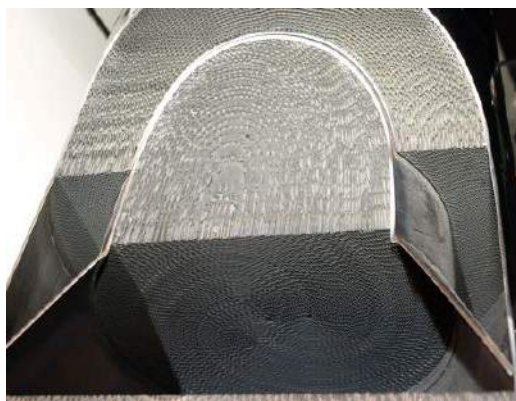
więc widzisz, zasoby są rozłożone nader nierówno (przede wszystkim południe Afryki). Znaczny udział ma również recykling zużytych katalizatorów.

Ceny platynowców wahają się w zależności od zapotrzebowania i giełdowych zawirowań (powodowanych głównie niestabilną sytuacją w krajach producentów), ale są one niezmiennie bardzo drogimi metalami. Najkosztowniejszy z nich to rod, za który trzeba zapłacić nawet kilkanaście razy więcej niż za złoto, cena kilku innych również jest wyższa niż „króla metali”.

...i do czego są potrzebne?

Przede wszystkim do wytwarzania katalizatorów dla przemysłu chemicznego i petrochemii (w tych działach gospodarki są niezastąpione), a także samochodowych dopalaczy spalin (3). Najwięcej zastosowań ma platyna, z czystego metalu i jej stopów wyrabia się biżuterię, sprzęt laboratoryjny odporny na agresywne chemikalia, elektrody, termoogniwa, elementy grzewcze pracujące w bardzo wysokich temperaturach, korzystają z niej też dentyści. Szczególnie dużą aktywność katalityczną ma **czerń platynowa** (bardzo rozdrobniony metal), który przy dostępie powietrza samoczynnie rozżarza się w kontakcie z palnym gazem, co zastosowano np. do konstrukcji zapalniczek i podręcznych ogrzewaczy (4).

Twardość i odporność chemiczną platynowców wykorzystuje się do produkcji stopów używanych jednak tylko do specjalnych zastosowań (wiadomo, cena) jak elementy precyzyjnej aparatury, styki elektryczne, a dawniej igły gramofonowe i końcówki stalówek (5). Bardzo cienkie powłoki z roku chronią przedmioty srebrne przed ciemnieniem. Warto również



3. Wnętrze katalizatora służącego jako samochodowy dopalacz spalin (siatkowa struktura zwiększa powierzchnię kontaktu spalin z katalizatorem i poprawia sprawność urządzenia)



4. Grzałka katalityczna

wspomnieć, że ze stopu platyny z irydem wykonane były międzynarodowe wzorce metra i kilograma, zanim nie zastąpiły ich jednostki oparte na wielkościach atomowych. Związki platyny należą także do najsilniejszych leków antynowotworowych.

Historia platynowców...

...w Europie zaczęła się w połowie XVIII wieku. Wtedy to uczeni badający tereny Ameryki Południowej donieśli o nieznanym w Starym Świecie metalu z wyglądu podobnym do srebra. Początkowo nie spotkał



5. Końcówki markowych stalówek pokrywa się twarzymi metalami szlachetnymi

się on z zainteresowaniem, czego dowodem jest nazwa – **platyna** to z hiszpańskiego „sreberko” (zdrobnienie z lekceważącym odcieniem znaczeniowym). Co ciekawe, platyny, metalu obecnie znacznie droższego od złota, początkowo używano do fałszowania złotych monet i, aby zapobiec owemu procederowi, niszczone jej zapasy. Odkrycie dużych złóż na Uralu i poznanie unikatowych właściwości platyny spowodowało, że już w połowie XIX wieku „sreberko” zaliczono do grona najbardziej cenionych metali.

W początkach wieku XIX dokładnie zbadano rodziną platynę. W latach 1803–1804 brytyjscy chemicy potraktowali ją wodą królewską i analizowali dwie frakcje. Częścią rozpuszczalną w wodzie zajął się **William Wollaston**, który w roztworze stwierdził obecność dwóch nieznananych pierwiastków. Był to **pallad** (nazwa pochodzi od planetoidy Pallas odkrytej rok wcześniej) i **rod** (z gr. *rhodon* = róża, od koloru jego związków). **Smithson Tennant** z kolei badał pozostałe osady. I on również znalazł dwa nowe pierwiastki: **osm** (z gr. *osme* = zapach, od charakterystycznej woni metalu, a właściwie powstającego na powierzchni lotnego czterotlenku OsO₄) oraz **iryd** (z gr. *iris* = tęczna, od wielobarwnych związków, które tworzy).

W historii najpóźniej odkrytego platynowca jest polski wątek. W roku 1807 **Jędrzej Śniadecki**, twórca naszego nazewnictwa chemicznego, doniósł o istnieniu jeszcze jednego metalu w surowej platynie i nadał mu nazwę **west** (od planetoidy Westy, panowała wtedy „moda” na nazywanie pierwiastków imionami ciał niebieskich, do wspomnianego palladu można dodać jeszcze cer od planetoidy Ceres). Uczeni z Akademii Paryskiej, największego ówczynie autorytetu w świecie nauki, nie potwierdzili jednak odkrycia, a sam Śniadecki pogodził się z werdyktem. Brakujący platynowiec ostatecznie odkrył **Karl Claus** pracujący w rosyjskim uniwersytecie w Kazaniu w roku 1844, choć korzystał z wcześniejszych prac **Gottfrieda Osanna** z uniwersytetu w Tartu (obecna Estonia). Nazwa **ruten** pochodzi od łacińskiej nazwy Rosji – *Ruthenia* (zapropował ją Osann).

W końcu XX wieku pojawiła się jeszcze jedna triada o liczbach atomowych 108–110 nazwana **superplatynowcami**. Ponieważ ich otrzymanie to zasługa niemieckich fizyków z **Instytutu Badań Ciężkich Jonów w Darmstadt** (zespołem kierowali **Peter Armbruster** i **Gottfried Münzenberg**), oni też uzyskali prawo nadania nazw. I tak pierwiastek o liczbie atomowej 108 to **has** (od Hesji, na terenie której leży instytut), kolejny – **meitner** (Lise Meitner, austriacka fizyczka jądrowa), a ostatni – **darmsztadt**. Ponieważ są to bardzo krótko żyjące, sztucznie



6. Amerykański fizyk Louis Alvarez (na fotografii z lewej wraz z synem Walterem) jako pierwszy stwierdził, że upadek gigantycznego meteorytu spowodował wymieranie kredowe. Widoczna w skale warstwa, w której wykryto podwyższony poziom irydu, stanowi granicę pomiędzy epoką kredową i paleogenem

otrzymane pierwiastki (i to w liczbie zaledwie pojedynczych atomów), o ich własnościach nie można jeszcze wiele powiedzieć.

Dwie ciekawostki

Iryd jest jednym z charakterystycznych metali występujących w meteorytach. Odkrycie warstw skalnych sprzed 65 milionów lat zawierających podwyższony



7. Przemysł chemiczny to główny odbiorca platynowców

poziom tego pierwiastka dowodzi, zdaniem uczonych, że w tym czasie Ziemia została trafiona przez obiekt kosmiczny o średnicy około 10 km. Wydarzenie to spowodowało zagładę dinozaurów, zakończyło epokę kredową i umożliwiło ewolucję ssaków (6).

Pallad ma zdolność do pochłaniania wodoru niespotykaną wśród innych metali: potrafi wchłonąć objętość tego gazu kilkaset razy większą od własnej, tworząc tzw. **gąbkę wodorową**. Wodór zaabsorbowany w palladzie znajduje się nie w postaci cząsteczkowej, lecz pojedynczych atomów (*in statu nascendi*, jak zjawisko to określają chemicy), co powoduje jego zwiększoną reaktywność.

Platynowce zdecydowanie nie służą tylko do wyrobu ozdób. Bez nich załamałby się prawie cały współczesny przemysł chemiczny (katalizatory) (7). Wiążą się z nimi również meandry światowej polityki – większość złóż leży w zapalnych rejonach naszego globu, państwa gromadzą więc ich zapasy. Światowa produkcja innych pierwiastków sięga setek tysięcy i milionów ton, zaś wytworzone w ciągu roku platynowce można załadować do pociągu towarowego. Jednak nie od dziś wiadomo, że często to, co jest na pierwszy rzut oka mało istotne, okazuje się bardzo ważne. ■

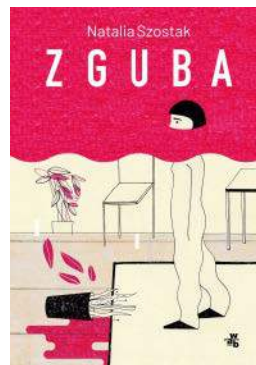
Krzysztof Orlński

Zguba

Natalia Szostak

Wydawnictwo W.A.B., liczba stron: 256, cena: 46,99 zł

Nowy – silny i delikatny jednocześnie – głos na polskiej scenie literackiej. Debiutancka książka dziennikarki Natalii Szostak to powieść o rzeczach, uczuciach i nadziejach, które gubimy po drodze. Matka piętnastoletniej Marianny wraca z emigracji zarobkowej, by odnaleźć córkę. Pomóc może jej teściowa, która przez ostatnie miesiące opiekowała się Marianną. Ale co właściwie się stało? Kto stoi za zniknięciem? Może wszystko zaczęło się wtedy, gdy dziewczyna straciła głos?





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

W dniach 17–30.03.2023 w Petrovac (Czarnogóra) odbyły się indywidualne Mistrzostwa Europy Kobiet w szachach. Już trzeci rok z rzędu reprezentantki Polski zajmują miejsca na podium. W 2021 r. brązowy medal mistrzostw Europy zdobyła Oliwia Kiołbasa, w 2022 tytuł mistrzyni wywalczyła Monika Soćko, w tym roku srebrny medal zdobyła Oliwia Kiołbasa, a brązowy Aleksandra Malcewska.

Oliwia Kiołbasa wicemistrzynią Europy



**1. Oliwia Kiołbasa z pucharem
wicemistrzyni Europy w szachach,
(<https://tiny.pl/cxhxn>)**

Mistrzostwa Europy Kobiet w Petrovac zakończyły się dużym sukcesem polskich reprezentantek. W ostatniej rundzie walka o medale toczyła się na dwóch pierwszych szachownicach. Ostatecznie Oliwia Kiołbasa zdobyła srebrny medal, a Aleksandra Malcewska brązowy. Znakomitym finiszem

w turnieju popisała się też Klaudia Kulon, która ostatecznie zdobyła szóste miejsce. Mistrzynią Europy została Meri Arabidze z Gruzji, która wyprzedziła Oliwię Kiołbasę dzięki wygranej w bezpośredniej partii.

Polska była najliczniejszą reprezentacją na mistrzostwach, w których startowało 136 szachistek z 34 federacji. W ostatniej rundzie, po ponad czterech godzinach zaciętej gry, Oliwia Kiołbasa zremisowała z Greczynką Stavrulą Tsolakidou i zdobyła srebrny medal. Tyle samo punktów miała Meri Arabidze z Gruzji. O pierwszym miejscu zdecydował bezpośredni pojedynek Kiołbasy i Arabidze, wygrany w ósmej rundzie przez Gruzinkę. W całym turnieju 22-letnia mistrzyni międzynarodowa

Tabela 1. Wyniki końcowe czołówki Mistrzostw Europy Kobiet

Miejsce	Tytuł	Nazwisko, imię	FED	Ranking	Punkty	TB1	TB2	TB3
1	IM	Arabidze, Meri	GEO	2433	8,5	1	68,5	73,5
2	IM	Kiołbasa, Oliwia	POL	2406	8,5	0	68,5	74
3	IM	Malcewska, Aleksandra	POL	2388	8	0	71	77,5
4	IM	Tsolakidou, Stavrula	GRE	2358	8	0	65	69
5	IM	Melia, Salome	GEO	2366	8	0	63	68
6	IM	Kulon, Klaudia	POL	2290	8	0	58	61,5
7	IM	Javakhishvili, Lela	GEO	2443	7,5	0	67,5	73,5
8	IM	Garifullina, Leya	FID	2372	7,5	0	67,5	73
9	FM	Toncheva, Nadya	BUL	2192	7,5	0	66,5	71
10	IM	Sargsyan, Anna M.	ARM	2371	7,5	0	65,5	70
11	IM	Guichard, Pauline	FRA	2379	7,5	0	64,5	68
12	IM	Daulyte-Cornette, Deimante	FRA	2344	7,5	0	63	67,5
13	IM	Cyfka, Karina	POL	2350	7,5	0	60	64,5



2. Mistrzyni Europy i reprezentantki Polski na podium (<https://tiny.pl/cxhx4>)

z Augustowa wygrała siedem partii, trzy zremisowała i poniosła jedną porażkę.

Oliwia Kiołbasa, urodzona 26 kwietnia 2000 w Augustowie, swoje pierwsze znaczące sukcesy na polu szachowym zaczęła odnosić już jako 8-latką, kiedy to sięgnęła po srebro w Mistrzostwach Polski Juniorek do lat 8. W międzynarodowych szachach pojawiła się 13 lat temu, zdobywając w Batumi



3. Polskę w mistrzostwach Europy reprezentowało 14 zawodniczek, (<https://tiny.pl/cxhxn>)



4. Oliwia Kiołbasa, Durban RPA, 2014 r., (<https://tiny.pl/cxhxv>)

złoty medal Mistrzostw Europy do lat 10. Za to osiągnięcie Międzynarodowa Federacja Szachowa przyznała jej tytuł mistrzyni FIDE. W roku 2011 r. zdobyła w Druskiénikach na Litwie srebrny medal Mistrzostw Europy Juniorek do 12 lat w szachach szybkich.

W numerze styczniowym 2015 roku „Młodego Technika” pisałem o 14-letniej Oliwii Kiołbasie, która we wrześniu 2014 r. w Durbanie (Republika Południowej Afryki) została wicemistrzynią świata w szachach klasycznych w kategorii do lat 14 (4). Zdobyła wtedy w 11 rundach 8,5 pkt., tyle samo co złota medalistka z Kanady Qiyu Zhou (zdecydowała punktacja dodatkowa). Nie był to pierwszy sukces niezwykle utalentowanej młodej szachistki z Augustowa. Była już wtedy m.in. wicemistrzynią Polski do lat 8, mistrzynią Polski do lat 10, wicemistrzynią Polski do lat 12. W wieku 13 lat zdobyła złoty medal mistrzostw Polski juniorek do lat 16 w szachach błyskawicznych oraz srebrny w szachach szybkich. W indywidualnych mistrzostwach Polski juniorek zdobyła 4-krotnie złote medale w szachach klasycznych, 2-krotnie w szachach szybkich i czterokrotnie w szachach błyskawicznych.

Tytuł mistrzyni międzynarodowej uzyskała w 2016 roku, a mistrza międzynarodowego zdobyła w 2022 roku.

Na ubiegłorocznej olimpiadzie szachowej w Madrasie, debiutując w tej rangi imprezie, jako jedyna w całej polskiej ekipie wystąpiła we wszystkich 11 meczach drużyny, zgromadziła 9,5 pkt. i wywalczyła złoty medal na trzeciej szachownicy.

Pod koniec lutego tego roku Oliwia Kiołbasa obroniła pracę licencjacką z psychologii zarządzania na Akademii Koźmińskiego.

A oto jedna z najlepszych partii rozegranych przez Oliwię na Mistrzostwach Europy (5)



5. Oliwia Kiołbasa w partii z Iriną Bulmagą w 2023 r., (<https://tiny.pl/cxhx4>)



6. Oliwia Kiotbasa – Irina Bulmaga, pozycja po 11. H:d4



7. Oliwia Kiotbasa – Irina Bulmaga, pozycja po 22. Sd1



8. Oliwia Kiotbasa – Irina Bulmaga, pozycja po 32. S:h5

Oliwia Kiotbasa – Irina Bulmaga

23. Mistrzostwa Europy Kobiet w Szachach, Petrovac (Czarnogóra), 23.03.2023, runda 6

1. e4 e5 2. Sf3 Sc6 3. Gb5 a6 4. Ga4 Sf6 5. Sc3 b5 6. Gb3 Gd6 7. d3 Sa5 8. Gg5

S:b3 9. a:b3 Gb7 10. d4 e:d4 11. H:d4 (diagram 6) 11...O-O (lepsze było 11...h6 i jeżeli 12. Gh4 to 12...He7 13. G:f6 H:f6 14. H:f6 g:f6 ze skomplikowaną pozycją a po 12. Ge3 można grać 12...b4) 12. e5 G:f3 13. G:f6 g:f6 14. e:d6 We8+ 15. Kd2 Gc6 16. Wae1 We6 17. d:c7 H:c7 18. W:e6 d:e6 19. H:f6 b4 20. Hg5+ Kh8 21. Hf6+ Kg8 22. Sd1 (diagram 7) 22...Ha5 (lepsze było 22...Wd8+ 23. Kc1 Ge4 24. Se3 Wc8 i czarne mają rekompensatę za piona) 23. Se3 Wd8+ 24. Ke2 Gb5+ 25. Kf3 Hc7 26. h4 Hc6+ 27. Kg3 Hc7+ 28. Kh3 Wd7 29. Sg4 h5 30. Hg5+ Kf8 31. Sf6 Wd4 32. S:h5 (diagram 8, białe mają już dwa piony przewagi) 32...Wd5 33. Hh6+ Ke7 34. Hf6+ Kd6 35. Hf4+ e5 36. Hf6+ Kc5 37. Sg3 Gd7+ 38. Kh2 Kb5 39. We1 Ge6 40. We2 a5 41. h5 e4 42. h6 Wf5 43. Hg7 Hd6 44. W:e4 Wh5+ 45. Kg1 Hd1+ 46. Sf1 Hd5 47. Wd4 Hc5 48. Wd2 Hg5 49. H:g5+ (diagram 9, najprostsza droga do wygranej z trzema pionami przewagi białych) 49...W:g5 50. Wd8 Wh5 51. Wh8 Kc6 52. Se3 Kd7 53. g4 Wh4 54. Kg2 Ke7 55. Kg3 Wh1 56. f4 f6 57. f5 Gd7 58. Sd5+ Kd6 59. S:f6 Gc6 60. h7 1-0

20-letnia Aleksandra Malcewska, która zmieniła rosyjską federację na polską, zwyciężyła w ostatniej rundzie i zajęła 3. miejsce (10). Urodzona w Wołgogradzie, w wieku 13 lat została mistrzynią Rosji dziewcząt do lat 15, a w następnym roku zdobyła złoty medal Mistrzostw Europy Juniorek do lat 14 (Praga 2016). Na początku 2017 roku przeniósł się z rodziną do Moskwy, gdzie jej trenerem został arcymistrz Aleksiej Drejew. W 2018 roku wygrała Mistrzostwa Świata Dziewcząt do lat 20, które odbyły się w Gebze w Turcji i otrzymała tytuł arcymistrzyni FIDE. Reprezentantką Polski jest od 12 maja 2022 roku. W Mistrzostwach Europy w Szachach Szybkich i Błyskawicznych rozegranych

w grudniu 2002 roku w Katowicach Aleksandra Malcewska zdobyła złoto w szachach szybkich oraz srebro w szachach błyskawicznych.

Klaudia Kulon (11) też wygrała i z taką samą liczbą punktów, lecz z gorszą punktacją pomocniczą zajęła 6. miejsce. Po słabym początku w ostatnich siedmiu partiach uzyskała 6,5 punktu!

Klaudia Kulon to arcymistrzyni od 2014 roku i mistrzyni Polski w szachach w roku 2021. Szachową karierę rozpoczęła w wieku 7 lat w koszański klubie



9. Oliwia Kiotbasa – Irina Bulmaga, pozycja po 49. H:g5+



10. Aleksandra Malcewska zdobyła brązowy medal, (<https://tiny.pl/cxhx4>)



11. Klaudia Kulon, (<https://tiny.pl/cxhnmh>)

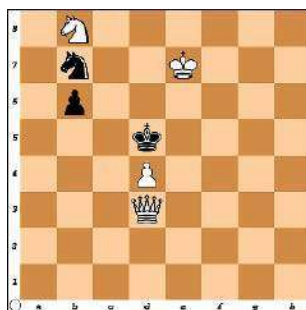
„Hetman Politechnika Koszalińska”. W roku 2002 zwyciężyła w mistrzostwach kraju junierek do lat 10, rozegranych w Kołobrzegu. Dwukrotnie zdobywała złote medale w Mistrzostwach Świata Junierek – w kategorii do 12 lat w Heraklionie (Grecja) w 2004 roku i do 14 lat w Batumi (Gruzja) w 2006 roku. We wrześniu 2016 roku wywalczyła srebrny medal drużynowo na olimpiadzie szachowej w Baku oraz brązowy indywidualnie na czwartej szachownicy. 8-krotnie zdobywała złote medale w Mistrzostwach Polski Junierek w różnych kategoriach wiekowych. W szachach szybkich dwukrotnie wywalczyła tytuł indywidualnej mistrzyni Polski, a w błyskawicznych pięciokrotnie. ■



Zadania do samodzielnego rozwiązania



Zadanie 1
12. Leonid Kubbel 1911
Mat w 2 posunięciach



Zadanie 2
13. E. Montvid 1901
Mat w 2 posunięciach

Rozwiązanie zadań z MT 5/2023

Zadanie 1
Iszta. 1893
Mat w 2 posunięciach
Rozwiązanie: 1. Hh3

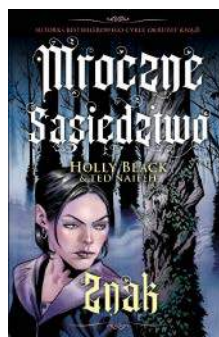
Zadanie 2
Hoffman, 1900
Mat w 2 posunięciach
Rozwiązanie: 1. Hd8

Znak

Holly Black, Ted Naifeh

Wydawnictwo Jaguar, liczba stron: 128, cena: 54,90 zł, Cykl: Mroczne sąsiedztwo (tom 2)

Kiedy Rue Silver była małą dziewczynką, widziała, jak kwiaty irysów same się chylą ku jej eterycznej, dziwnej matce. Kiedy podrosła, stwierdziła, że nawet jeśli nie potrafi powstrzymać szaleństwa, może udawać, że nie jest obłąkana jak ona... Dziewczyna wyrusza na poszukiwanie matki do królestwa Nieśmiertelnych – w samo serce niehumanicznej otchłani. Jednocześnie świat elfów snuje własne plany wobec ludzkich siedlisk; już całkiem namacalnie ingeruje w życie Rue, sięgając po kogoś bliskiego jej sercu. Kłopoty elfowego dziadka zagrażają rodzinemu miastu Rue, która nagle zdaje sobie sprawę, że tylko ona może je udaremnić. Jednak sama teraz nie wie, czy powinna się jeszcze zaliczać do ludzkich istot, czy może już bez reszty przynależy do grona elfów? Gdzie jest jej miejsce i jak ma się odnaleźć między lojalnością i przyjaźnią a magią, jak wytyczyć granicę między miłością a przeznaczeniem?





Chemia budowlana

Już ponad 2000 lat temu Rzymianie stosowali beton w budownictwie. W wyniku badań archeologicznych odkryto, że w rzymskich budowlach zastosowano specjalny rodzaj betonu, tzw. *opus caementicium*. Składał się on z mieszanki zaprawy wapiennej i kruszywa, do której dodawano popiół wulkaniczny lub muł wapienny. Dzięki temu materiałowi Rzymianie byli w stanie budować trwałe i wytrzymałe konstrukcje, takie jak łuki, mosty, akwedukty i kanały. Wiele z tych budowli przetrwało do dzisiejszych czasów i wciąż jest w użyciu. Ten genialny wynalazek to tylko jeden z wielu materiałów budowlanych wymyślonych przez człowieka. Dzisiaj także trwają poszukiwania materiałów trwalszych, tańszych w produkcji i bardziej uniwersalnych. Zapraszamy do lektury wszystkich, którzy są zainteresowani tworzeniem rzeczy, które mogą przetrwać tysiące lat. Zapraszamy na chemię budowlaną.

Zaprawa

Chemia budowlana to interdyscyplinarny kierunek studiów, który łączy w sobie wiedzę z zakresu chemii, technologii materiałów budowlanych oraz budownictwa. Studenci zdobywają umiejętności związane z projektowaniem, wytwarzaniem, stosowaniem i badaniem materiałów budowlanych. Studia można realizować na poziomie inżynierskim, magisterskim i doktoranckim. Osoby zainteresowane nauką na tym kierunku nie mają zbyt dużego wyboru, bo zaledwie kilka uczelni w Polsce ma go w swojej ofercie. Kandydaci mają do wyboru: Katowice,

Gdańsk, Kraków, Łódź. Co ciekawe, politechniki Gdańska i Łódzka oraz AGH współpracują ze sobą, czego efektem jest nadanie ChB charakteru kierunku międzyuczelnianego. W ramach tej kooperacji student pierwsze cztery semestry realizuje na uczelni „pierwszego wyboru”, kolejne dwa na uczelniach partnerskich, by na końcówkę powrócić do „swojej” szkoły.

Spoiwo

Przebrnięcie procesu rekrutacji nie powinno przysporzyć zbyt dużo problemów, gdyż zainteresowanie

ChB nie należy do wysokich. W trakcie nauki nie będzie już tak kolorowo, gdyż od studentów wymaga się opanowania dużej ilości materiału i wielkiego zaangażowania. W ramach studiów na kierunku chemia budowlana studenci zdobywają wiedzę z zakresu matematyki, chemii, fizyki, technologii materiałów budowlanych oraz budownictwa. Program nauczania obejmuje takie przedmioty jak: chemia analityczna i fizyczna, chemia nieorganiczna i organiczna, fizyka techniczna, materiałoznawstwo budowlane, technologia betonów i zapraw murarskich, technologia tynków i farb, ochrona i konserwacja zabytków, budownictwo. Ponadto studenci mają możliwość uczestniczenia w zajęciach laboratoryjnych oraz praktykach w zakładach produkcyjnych i firmach budowlanych.

W trakcie nauki poznają różne rodzaje materiałów budowlanych, takie jak: betony, zaprawy murarskie, tynki, farby, kleje czy impregnaty. Zajmują się także zagadnieniami związanymi z budową i eksploatacją budynków, takimi jak: izolacje termiczne i akustyczne, odporność ognia oraz trwałość konstrukcji. Ponadto można zdobyć wiedzę z zakresu zarządzania, marketingu oraz przedsiębiorczości, co może okazać się przydatne w przyszłej pracy. Tym samym studenci przygotowani są do realizowania zadań w różnych dziedzinach związanych z budownictwem i przemysłem. W trakcie nauki z pewnością przyda się umiejętność sprawnego przyswajania dużej ilości materiału oraz biegłość w analitycznym myśleniu. Poziom trudności nie powinien nikogo dziwić, gdyż chemia budowlana to kierunek interdyscyplinarny. Oczekuje się zatem od osoby chcącej uzyskać dyplom ukończenia studiów, że będzie umiała sprawnie poruszać się po różnych dziedzinach nauki, takich jak na przykład: chemia, technologia chemiczna, inżynieria materiałowa, towaroznawstwo, technologia tworzyw sztucznych. Swoboda, z jaką absolwent ChB powinien wykorzystywać możliwie jak największy zakres wiedzy z zakresu: chemii fizycznej, organicznej, nieorganicznej i analitycznej, polimerowych materiałów budowlanych, surowców i procesów biotechnologicznych, kontroli jakości produktów i ochrony środowiska w technologii chemicznej, powoduje, że nie tylko można określić go fachowcem przez duże F, w swojej branży, ale także, że otwierają się przed nim duże możliwości zawodowe w wielu branżach i gałęziach przemysłu. Osiągnięcie tego poziomu jest możliwe jedynie poprzez sprostanie wysokim wymaganiom stawianym przez uczelnie, co możliwe jest dzięki systematycznej pracy.

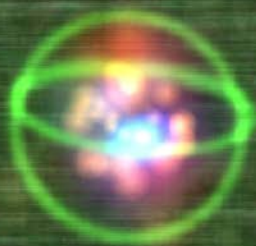
Kompozyt

Po ukończeniu studiów na kierunku chemia budowlana absolwenci mogą pracować w różnych branżach związanych z budownictwem, takich jak projektowanie budynków, produkcja materiałów budowlanych, kontrola jakości czy nadzór inwestorski. Miejsca pracy dla specjalistów w tej branży znajdują się w firmach produkujących i sprzedających materiały budowlane, takie jak beton, zaprawy murarskie, farby czy tynki. Zawodowo można zajmować się projektowaniem i wytwarzaniem nowych materiałów, opracowywaniem metod produkcji oraz kontrolą jakości, między innymi materiałów budowlanych. W ramach swoich obowiązków inżynierowie są odpowiedzialni za przeprowadzanie testów i badań w celu zapewnienia zgodności produktów z normami i standardami jakościowymi. Absolwenci chemii budowlanej mogą pracować jako projektanci i doradcy techniczni w firmach budowlanych i architektonicznych. Będą odpowiedzialni za projektowanie i dobór odpowiednich materiałów budowlanych, uwzględniając ich właściwości fizyczne i chemiczne, oraz za doradztwo techniczne w trakcie realizacji projektów budowlanych.

Ponadto pracy można szukać w charakterze inspektora nadzoru inwestorskiego, zajmując się nadzorem realizacji projektów budowlanych pod kątem zgodności z projektem, normami i standardami jakościowymi. Wysokość wynagrodzenia zaraz po ukończeniu studiów zależy od wielu czynników, takich jak miejsce pracy, doświadczenie zawodowe, specjalizacja oraz lokalizacja geograficzna, w której jest się zatrudnionym. Absolwent tego kierunku może spodziewać się około 6000 zł brutto. W miarę rozwoju doświadczenia można oczekiwać wyższych pensji, nawet na poziomie 12 000 zł brutto. Warto jednak pamiętać, że rynek pracy jest konkurencyjny, a branża materiałów budowlanych może wymagać ciągłego dokształcania i podnoszenia kwalifikacji.

Chemia budowlana to dobry wybór dla wszystkich osób, które w obszarze swoich zainteresowań mają chemię i budownictwo. Połączenie wiedzy z tych dwóch dziedzin naukowych pozwala na zdobycie szerokich kompetencji z zakresu chemii, fizyki, materiałoznawstwa oraz budownictwa, a także praktycznych umiejętności laboratoryjnych. Wybór kierunku studiów powinien także być oparty na czynnikach związanych z sytuacją na rynku pracy. Biorąc pod uwagę ten element, a także fakt, że ChB pozwala na tworzenie rzeczy, które przetrwają kolejne setki lat, jest to kierunek wart polecenia. ■

Michał Pacholski



Czy substancja promieniotwórcza straci z czasem swoje właściwości? (2)

Aby lepiej zrozumieć omawiane zjawisko, można przy pomocy prostej gry losowej przeprowadzić symulację procesu rozpadu substancji radioaktywnej. Potrzebujemy kostki do gry – może to być klasyczna kostka o sześciu oczkach, ale również kostka o innej liczbie ścianek. Będziemy również potrzebowali planszy (podobnej jak do szachów) o wymiarach $n \times n$, gdzie n jest liczbą ścianek kostki.

W każdym polu planszy ustawiamy jeden żeton, symbolizujący jądro izotopu niestabilnego. Ich suma to liczba N_0 , która pojawiła się w wyrażeniu opisującym rozpad promieniotwórczy. Najlepiej aby żetonów było około 30–40 (np. plansza 6×6). Zwiększenie liczby żetonów pozwoliłoby wprawdzie uzyskać nieco dokładniejsze wyniki, jednak kosztem znacznego przedłużenia czasu poświęconego na grę. W przypadku realizacji symulacji w trakcie lekcji można podzielić klasę na grupy, a na zakończenie zsumować efekty ich pracy.

W kolejnym kroku przygotowujemy tabelę według poniższego wzoru, w której będziemy zapisywać wyniki. Wartość N_0 wpisujemy w pierwszym wierszu, dla liczby rzutów równej zero. Ustalamy, że pierwszy rzut kostką wskaże nam wiersz, a drugi – kolumnę

tabeli. Z wylosowanego pola zdejmujemy jeden żeton i notujemy wyniki.

Czas (liczba rzutów)	Liczba jąder promieniotwórczych w próbce (liczba żetonów na planszy)
0	
1	
2	

W trakcie symulacji początkowo praktycznie przy każdym rzucie kostką losujemy pole z żetonem. Niemniej w miarę ubywania ich z planszy coraz częściej okazuje się, że wylosowane zostaje puste pole, w którym rozpad już nastąpił. Jeśli tak się

stanie to również zapisujemy wyniki, notując w tabeli liczbę rzutów oraz odpowiadającą jej liczbę jąder promieniotwórczych.

Po zakończeniu symulacji sporządzamy wykres zależności liczby jąder promieniotwórczych w próbce od czasu. W zasadzie możemy wyskalować oś czasu dowolnie, przypisując jednemu rzutowi na przykład sekundę, godzinę albo rok. Nie będzie to sprzeczne z prawami fizyki, ponieważ w przyrodzie występuje wiele różnych izotopów nietrwałych, a każdy z nich charakteryzuje się innym czasem rozpadu.

Warto na zakończenie zauważyć, że gdybyśmy rzucali kostką wystarczająco długo, zdjęlibyśmy z planszy wszystkie żetony, co oznacza, że każda substancja promieniotwórcza w pewnej chwili całkowicie rozpadnie się do izotopów stabilnych. Możemy odetchnąć z ulgą – wytworzone przez człowieka substancje radioaktywne kiedyś w końcu się rozpadną. Jedne szybciej, drugie wolniej, ale na pewno nie będą nam zagrażać w nieskończoność.

Ciekawostka – czas połowicznego rozpadu

W ramach opracowania uzyskanych wyników do punktów naniesionych na wykres dopasowujemy krzywą gładką. Posługując się tą krzywą odczytujemy czas po którym w próbce pozostanie połowa jąder promieniotwórczych. Czas ten nazywa się czasem lub okresem połowicznego rozpadu i oznacza się symbolem $T_{\frac{1}{2}}$. Posługując się tym pojęciem prawo rozpadu promieniotwórczego można zapisać jako

$$N(t) = N_0 \left(\frac{1}{2} \right)^{\frac{t}{T_{\frac{1}{2}}}}$$

Jakkolwiek powyższe wyrażenie może wydawać się skomplikowane, spróbujemy na prostym przykładzie rachunkowym zrozumieć jego znaczenie. Jest to o tyle przydatne, że z pojęciem czasu połowicznego rozpadu możemy również spotkać się w odniesieniu do zagadnień z dziedziny chemii czy medycyny.

Zadanie rachunkowe

Na podstawie wyników symulacji oszacuj, ile rzutów kostką należy wykonać, aby na planszy pozostała $\frac{1}{8}$ z początkowej liczby żetonów.

Dla nauczyciela

Opisana symulacja może zostać wykorzystana na lekcji fizyki w liceum lub technikum, szczególnie do realizacji poniższego punktu podstawy programowej w zakresie rozszerzonym:

XII.11 (uczeń) opisuje przypadkowy charakter rozpadu jąder atomowych.

Ponieważ jest ona prosta do realizacji w dowolnym języku programowania, może również posłużyć jako algorytm uczniom uzdolnionym programistycznie. ■

Joanna Borgensztajn

Odpowiedzi do zadań

Wykreślanka z pierwszej części artykułu (MT5/2023): izotop, jądro atomowe, neutron, pierwiastek, prawo rozpadu, proton, radioterapia.

Zadanie rachunkowe

Określamy czas połowicznego rozpadu izotopu wyrażony liczbą rzutów kostką. Skoro na planszy ma pozostać $\frac{1}{8}$ z początkowej liczby żetonów to

$$N(t) = N_0 \left(\frac{1}{2} \right)^{\frac{t}{T_{\frac{1}{2}}}} = \frac{1}{8} N_0$$

zatem

$$\left(\frac{1}{2} \right)^{\frac{t}{T_{\frac{1}{2}}}} = \frac{1}{8}$$

Ułamek $\frac{1}{8}$ możemy zapisać jako $\left(\frac{1}{2} \right)^3$ wobec czego dostajemy

$$\frac{t}{T_{\frac{1}{2}}} = 3$$

Szukany czas wynosi

$$t = 3T_{\frac{1}{2}}$$

Po tamtej stronie piekła

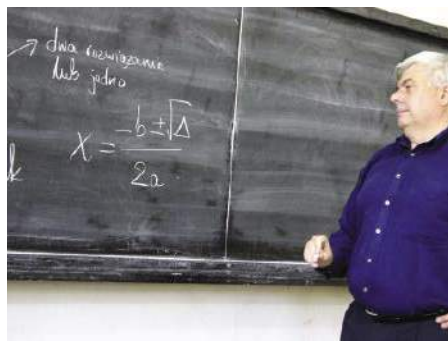
Tomasz Betcher

Wydawnictwo W.A.B., liczba stron: 384, cena: 49,99 zł

Drugi tom emocjonującej sagi sięgającej czasów drugiej wojny światowej. O bliznach, które noszą ci, którzy przeżyli. O tym, że o pewnych sprawach nigdy nie da się zapomnieć. O miłości, której siła pozwala dokonywać niemożliwego. I o kolejnych pokoleniach, nieświadomych, jak wielki wpływ na ich życie mają wydarzenia sprzed kilkudziesięciu lat.



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”.



Apolonia Inteligentna. Magiczny kafelek

Minał kolejny maj. Jak zwykle, kwitły kasztany – ten symbol egzaminów maturalnych, niegdyś elitarnych, dziś obejmujących znaczną większość chłopców i dziewcząt, wchodzących w dorosłe życie – zamglone kryzysami, konfliktami, kłótniami i ogólną niepewnością jutra. Tylko kasztanowce nic sobie z tego nie robią. Kwitną na biało. Zaczekam do jesieni. Wtedy lubię zbierać spadłe brązowe kasztany. Dobrze się bawić nimi z wnukami, jak dawniej. Tak jak mój dziadek bawił się ze mną.

Od pewnego czasu dyskutuje się w mediach o pani Apolonii Inteligentnej – tak na własny użytek nazwałem AI, artificial intelligence, czyli po polsku sztuczną inteligencję. Wciąż jeszcze łatwo ją oszukać, wpisując tak zwane głupie pytania. Nie cieszy się: Ona się uczy, a na pewno ma utajnione wersje, które już potrafią odpowiedzieć na pytanie filozoficzne, które pozwolę sobie przytoczyć. Pewien kawaler, starający się o rękę panny, został pouczony, że należy trochę porozmawiać o rodzinie, uczuciach i trochę pofilozofować. Na pytanie, czy ma brata, odpowiedziała, że nie. Na pytanie, czy lubi makaron, też odpowiedziała przecząco i konwersacja utknęła, ale konkurent znalazł sposób, by kontynuować i przejść do ogólnych zagadnień o życiu. „Słuchaj, a gdybyś miała brata, to czy on by lubił makaron?”

Takie mam złośliwości w stosunku do pani Apolonii. Ale konkretnie i matematycznie. Dałem jej zadanie, do którego mam duży sentyment. Dostałem je 59 lat temu (!) na egzaminie wstępnym (ustnym) na Uniwersytet Warszawski, na Wydział Matematyki i Fizyki. Na szczęście zorientowałem się od razu, że jest to pytanie podchwytliwe i trochę też dlatego dostałem

piątkę z odpowiedzi. Szóstek wtedy nie było. Równanie do rozwiązania było: $\sin x + \cos x = 2$.

Aby zobaczyć, że równanie to nie ma rozwiązań, wystarczy pamiętać, że obie funkcje są ograniczone przez -1 i 1 oraz, że sinus i cosinus nie mogą być jednocześnie równe 1 . A to wynika na przykład z tak zwanej jedynki trygonometrycznej. Z bardzo niewiele informacji matematycznych, jakie większości ludzi zostaje w głowach po ukończeniu szkoły ponadpodstawowej, jest właśnie to, że $\sin^2 x + \cos^2 x = 1$.

Co jednak jest w tym „podchwytliwe”? Otóż w moich szkolnych czasach połowę czasu w klasie X (przedmaturalnej) zajmowała trygonometria, a tu z kolei nacisk położony był na tożsamości i równania. Musieliśmy na pamięć znać wiele wzorów i tożsamości, na przykład na $\sin(\alpha + \beta)$, $\cos(\alpha + \beta)$, $\sin \alpha + \sin \beta$, a także i ważne podstawienie wymierne, gdzie kluczową rolę gra tangens połowy kąta:

$$\sin \varphi = \frac{2t}{1+t^2}, \quad \cos \varphi = \frac{1-t^2}{1+t^2}, \quad \text{gdzie} \quad t = \operatorname{tg} \frac{\varphi}{2}$$

Tak się zachowała i Apolonia. Poszła w przekształcenia. Komputery lubią obliczenia. Zastosowała prawidłowo wzór na sumę sinusów i cosinusów:

$$\sin x + \cos x = \sqrt{2} \sin\left(x + \frac{\pi}{4}\right)$$

Otrzymała stąd równanie

$$\sin\left(x + \frac{\pi}{4}\right) = \sqrt{2}$$

i nie przejmowała się tym, że prawa strona jest większa od 1, tylko znalazła „rozwiązanie”

$$x = -\frac{\pi}{4} + 2k\pi \quad x = \frac{3\pi}{4} + 2k\pi$$

Widać z tego, że „coś jej się pomieszało” i zamiast $\sqrt{2}$ po prawej stronie wzięła sobie zero.

Zanim jednak zaczniemy wyśmiewać „bezmózgową” Apolonię, spójrzmy, jak sobie poradziła z ciekawym zadaniem o ciągu geometrycznym (cytuję z własnej książki *Spacery matematyczne*). Wpisałem jej w stosowne okienko:

Była wioska nad rzeką i siedmiu wieśniaków. Każdy miał siedmiu synów – odważnych chłopaków. Każdy chłopak posiadał aż siedem kajaków, A w każdym tym kajaku miał siedem plecaków. Każdy plecak pomieścić mógł siedem szczupaków. Można złowić szczupaka na siedem robaków.

Ile tego jest razem? Robaków, szczupaków, Plecaków, kajaków, chłopaków, wieśniaków?

Wskazówka: miejscowością tą wcale nie był Kraków!

Oto jak rozwiązała to zadanie pani Apolonia Inteligentna. Kopiuje odpowiedź.

W tej zagadce matematycznej musimy wykonać kilka operacji arytmetycznych, aby uzyskać łączną liczbę robaków, szczupaków, plecaków, kajaków, chłopaków i wieśniaków.

Każdy z siedmiu wieśniaków ma siedmiu synów, co daje nam łącznie $7 \times 7 = 49$ chłopców.

Każdy chłopak ma 7 kajaków, więc łącznie mamy $49 \times 7 = 343$ kajaki.

Każdy kajak ma 7 plecaków, co daje nam łącznie $343 \times 7 = 2401$ plecaków.

Każdy plecak może pomieścić 7 szczupaków, więc łącznie mamy $2401 \times 7 = 16807$ szczupaków.

Można złowić szczupaka na 7 robaków, więc łącznie mamy $16807 \times 7 = 117649$ robaków.

Odpowiedź brzmi: razem jest 117649 robaków, 16807 szczupaków, 2401 plecaków, 343 kajaków, 49 chłopców i 7 wieśniaków.

Czy to nie przerażające? W licznych rozmowach i artykułach o AI można właśnie wyczuć nutkę strachu: już wkrótce opanują nas sztuczne mózgi. Potem dojdą do wniosku, że należy nas (ludzi) eksterminować

i po prostu zabijają. Zanim my wyłączymy im wtyczkę do prądu, one odetną nas od tlenu. Przez krótki czas panice uległem i ja – stwierdziłem po prostu, że sztuczną inteligencją już posługują się studenci na egzaminach i kolokwiah. Sami czynili to jednak niezbyt inteligentnie – od razu było widać, co i jak. Późniejsze dopytywanie ustne potwierdziło, że nie mają pojęcia o tym, co piszą i że nie rozumieją, o co chodzi. Nie rozumieją naszym ludzkim umysłem.

Chyba każdy wielki wynalazek wywoływał podobne obawy: od maszyny parowej przez kolej żelazną do komputerów. Kto w dzisiejszej szkole uczy takich technik obliczeniowych, jak próba dziewiątkowa, „dodawanie metodą księgowych”, szybkie podnoszenie do kwadratu, pisemne wyciągnięcie pierwiastka, sprowadzanie do postaci logarytmicznej? To wszystko odeszło w zapomnienie. Szkoda? Tak samo, jak szkoda koni dorożkarskich, lamp naftowych i pisania piórem ze stalówką, maczaną w kałamarzu z atramentem (a kleksy suszyło się bibułą).

Wiemy, że ilość nagromadzonej wiedzy rośnie gwałtownie i „końca nie widać”. Jak jednak zauważył Ryszard Kapuściński, najwyraźniej przybywa jej tylko w pamięci komputerów, niekoniecznie w umysłach ludzi. Co gorsza, nie mądrzejemy od tego. Ale to już inna działka.

Wynalazkiem, który wciąż ma więcej dobrego niż złego, jest Internet (przewidziany już siedemdziesiąt lat temu przez Stanisława Lema, w książce *Powrót z gwiazd*, 1963). Ciekawe, że mamy go pisać dużą literą. Dlaczego? Z szacunku? No, to też inny temat. Możemy dzięki niemu (Niemu?) dowiedzieć się wiele i przeczytać opinie powszechnie szanowanych osób na wiele tematów. Niestety, jesteśmy też zasypywani bzdurnymi informacjami, zwykle podawanymi jako „sensacyjne” (pod Piotrkowem urodziło się cięło o dwóch głowach, a w Mławie motocyklista Michał wpadł na słup). I najgorsze: trudno nam odróżnić opinię mądrych ludzi od głupich wypowiedzi, albo celowo spreparowanych fake newsów. Trudno zachować dyscyplinę umysłową.

O właśnie, wróć do trygonometrii. Około Świąt Wielkanocnych moja komórka uparcie wyświetlała mi informację, że w matematyce zostało zauważone coś, co czekało 2000 lat na odkrycie. Zrobiłem wyjątek od zasady, że nigdy takich sensacji nie czytałem. No i pożałowałem decyzji. Oto bowiem zostałem poinformowany, że Calcea Rujean Johnson i Ne’ Kiya Jackson z Nowego Orleanu udowodniły twierdzenie Pitagorasa za pomocą trygonometrii, ale nie korzystając z „jedynek trygonometrycznej”, czyli właśnie z tego, że $\sin^2 x + \cos^2 x = 1$.

„Najwięksi matematycy od ponad 2000 lat próbowali to osiągnąć. Do tej pory opublikowano co najmniej 118 dowodów tego twierdzenia, a Carl Friedrich Gauss wykazał, że tych dowodów jest nieskończenie wiele” – tak napisano na odpowiedniej stronie internetowej, co bezmyślnie przedrukował *Focus* a powtórzył XY, autor notatki w Internecie. W USA sprawa może ma jeszcze bardziej medialne znaczenie, bo obie dziewczyny-licealistki są czarnoskóre (nie wiem, jak w dobie poprawności politycznej teraz mam nazywać osoby, które... istotnie... są czarnoskóre?).

Najpierw sprawy oczywiście bzdurne. Trygonometria powstała w dobie wczesnego renesansu – w związku z potrzebami praktycznymi, najpierw na potrzeby nawigacji oceanicznej. Dlatego przed znaną nam ze szkoły trygonometrii płaskiej była prawie zapomniana już dzisiaj trygonometria sferyczna. Zauważmy, że gdy w grę wchodzi duże odległości, geometria na powierzchni kuli różni się od tej na płaszczyźnie. Na przykład suma kątów trójkąta, którego jeden wierzchołek jest na biegunie, a dwa na równiku, jest większa od 180 stopni.

Trygonometria miała wielkie znaczenie w pomiarach Ziemi aż do upowszechnienia się nawigacji satelitarnej. Jeszcze pół wieku temu Polska była pokryta siecią drewnianych wież triangulacyjnych. Na Podhalu miały one nawet „ludową” nazwę: patryja. Teraz zamiast nich stoją maszty telefonii komórkowej... Tak czy owak, żadne zadanie o trygonometrii nie może mieć 2000 lat, bo wtedy jej nie było.

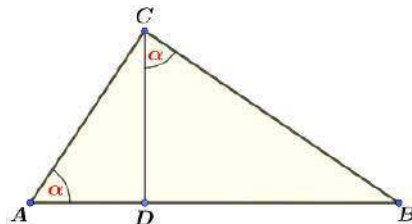
„Opublikowano co najmniej 118 dowodów twierdzenia Pitagorasa”. Matematycznie, zgadza się. Takie jest znaczenie zwrotu „co najmniej”. Gdybym napisał, że „opublikowano co najmniej 500” – też by się zgadzało. Co do rewelacji Gaussa, że dowodów jest nieskończenie wiele – Gauss tego na pewno nie napisał. Pewne jest jednak, że każde twierdzenie można udowodnić jednak na nieskończenie wiele sposobów. To mniej więcej tak, jakby powiedzieć, że do mojego lokalnego sklepu też mogą dojść na nieskończenie wiele sposobów: mogą na przykład dojść tam, idąc przez Lizbonę i Władywostok, za każdym razem inaczej stawiając nogi.

To wszystko jednak „dałoby się znieść”, ale nie rewelację, że można użyć trygonometrii, lecz bez jedynki trygonometrycznej. Istotnie, w banalny sposób można stwierdzić, że owa tożsamość trygonometryczna jest po prostu innym sformułowanie twierdzenia Pitagorasa. Przypomnę, że głosi ono: w trójkącie prostokątnym suma kwadratów przyprostokątnych jest równa kwadratowi przeciwprostokątnej.

Dowód twierdzenia Pitagorasa za pomocą trygonometrii (bez „jedynki”) robimy tak. Bierzymy dowolny dowód i przerabiamy go na trygonometryczny

bez użycia „jedynki”. Brzmi to trochę tak, jak przepis Juliana Tuwima: jak zrobić esklop? – wziąć pół funta cięlicyny, wstrząsnąć i zrobić esklop. No, ale proszę. Wezmę tak zwany dowód Garfielda (XX prezydenta USA, 1831–1881, tak jest, tacy bywali niegdyś prezydenci).

1



Dowód twierdzenia Pitagorasa „za pomocą” trygonometrii. Jeżeli przypomnimy sobie określenie funkcji podstawowej funkcji trygonometrycznej: sinusa, to zobaczymy, że na **rysunku 1** mamy:

$$\sin \alpha = \frac{BC}{AB} = \frac{BD}{BC}$$

Pierwsza z tych równości pochodzi z dużego trójkąta ABC , druga z tego „mniejszego” po prawej, czyli BCD . Mamy więc

$$(*) \quad BC^2 = AB \cdot DB$$

Teraz będzie trudniej, bo do gry wejdzie cosinus. Nie ma rady, musimy sobie przypomnieć, co to za funkcja. Zakładam, że Czytelnicy doskonale z tym sobie radzą. Mamy

$$\cos \alpha = \frac{AC}{AB} = \frac{AD}{AC}$$

Pierwsza z tych równości pochodzi, jak poprzednio, z dużego trójkąta ABC , druga z tego „mniejszego” po lewej, czyli ACD . Mamy więc

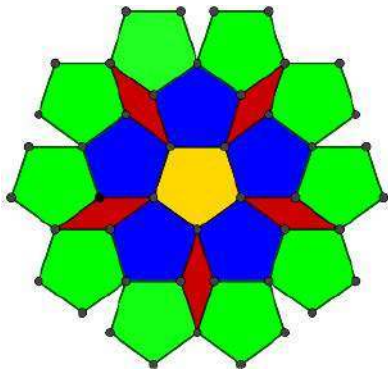
$$(**) \quad AC^2 = AB \cdot AD$$

Dodajemy równości oznaczone gwiazdkami: $AC^2 + BC^2 = AB \cdot AD + AB \cdot DB$.

Wylączamy po prawej stronie AB przed nawias i uwzględniamy to, że $AD + DB = AB$ i mamy twierdzenie Pitagorasa. Korzystałem tylko z tego, czym jest sinus i cosinus.

Nic dziwnego, że z dużą podejrzliwością odniosłem się do innej „rewelacji” zamieszczonej przez *Forum*, że jakoby po 50 latach pracy odkryto świętego Graala matematyki – kształt, którym można pokryć płaszczyznę bez powtarzania wzoru. Dla matematyków to istotnie ciekawe, ale nazywanie tego „świętym Graalem” jest po prostu niesmaczne. Nie jest też tak, że matematycy przez pół wieku tylko szukali jednego kafelka.

O co chodzi? Jeżeli układamy parkietaż z płytek o regularnych kształtach, to wzór prędzej czy później musi się powtórzyć albo mieć jakąś symetrię. Spójrzmy

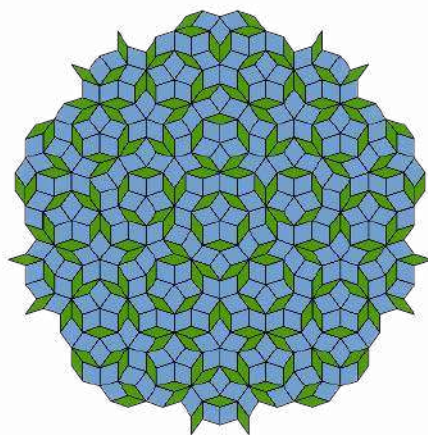


2

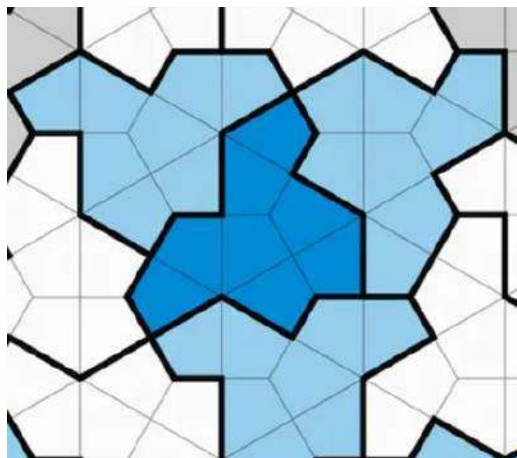
na **rysunek 2**. Jest zbudowany bardzo regularnie. Układamy obok siebie pięciokąty foremne. Powstaje siatka z oczkami – są nimi brązowe romby. Bardzo estetyczny deseń, prawda? Z matematycznego punktu widzenia jest on okresowy – nie zmienia się przy przesunięciach.

Pytanie, czy są takie parkietaże nieokresowe, postawił w 1961 roku amerykański matematyk chińskiego pochodzenia, Hao Wang, pracujący zresztą w Bell Laboratories. Ale zagadnienie przyszło z logiki matematycznej – teorii funkcji rekursywnych. Zajmował się nimi między innymi wybitny matematyk polski, Andrzej Mostowski (1913–1975). Przytoczę nawet dane odpowiedniego artykułu: *Mostowski, A., A formula with no recursively enumerable model, Fundamenta Mathematicae, vol. 43 (1955), pp. 125–140*.

Fakt, że można pokryć płaszczyznę w sposób nieokresowy udowodnił w 1966 r. Robert Berger. Podał on też stosowną konstrukcję, ale jego parkietaż zawierał aż 20 426 kafelków różnych kształtów. Potem redukowano liczbę potrzebnych kafelków (w rozważaniach przydawały się idee maszyny Turinga) aż do osiągnięcia prostego parkietażu Penrose’a, który wymaga tylko dwóch kształtów. To istotnie było pół wieku temu, w 1973 roku. Obecnie wykazano, że wystarczy jeden wielokąt, jeden kształt. Bo choć informację o tym odkryciu zamieszczono w sieci podejrzanie blisko 1 kwietnia, to jednak wydaje się prawdziwa. Trudno przewidzieć, czy i jakie zastosowanie



3. Kafelki Penrose’a (źródło: Wikipedia, hasło: kafelki Penrose’a)



4. Trzynastokąt do nieokresowego wypełnienia płaszczyzny

będzie miało to odkrycie. Kafelki Penrose’a posłużyły do opisu kwazikrystalów. Podejrzewam, że zgodnie ze współczesnymi trendami ten jeden kafelek może mieć znaczenie w kryptografii. ■

Śmierć, która cię ocali

Jack Jordan

Wydawnictwo MUZA SA, liczba stron: 448, cena: 49,90 zł

Trudno o bardziej skutecznego mordercę niż wprawny chirurg. Mój syn został porwany. Teraz muszę dokonać wyboru. Albo mój pacjent nie przeżyje operacji, albo już nigdy nie zobaczę mojego dziecka. Na stole operacyjnym leży człowiek. Jako chirurg mam za zadanie go uratować. Jako matka wiem, że muszę go zabić. Ktoś mógłby pomyśleć, że jestem potworem. Ale mam tylko jedno wyjście. Nikt nie może odkryć, że popełniłam morderstwo. W przeciwnym razie mój syn umrze. Tragiczny dylemat moralny i zawrotne tempo. Historia, która pochłania od pierwszej strony.





Szkoła Wynalazców

dozwolone do lat 15

Mieścieście zadanie dla młodych kierowców: *Jak zapewnić sobie kontrolę świateł zewnętrznych samochodu bez pomocy drugiej osoby.*

Zazwyczaj sprawdzenie świateł wymaga pomocy drugiej osoby: kierowca siedzi za kierownicą i włącza po kolei wszystkie światła, a osoba stojąca na zewnątrz melduje: światła są lub ich nie ma.

Czasami jednak trudno zapewnić sobie pomoc drugiej osoby i trzeba sobie jakoś radzić. Na pewno nie jest dobrze, gdy brak jakiegoś światła odkryje nam diagnosta, podczas okresowej kontroli lub policjant. Starzy kierowcy sprawdzają swoje światła, stojąc... w korku. Mając przed sobą jakikolwiek samochód, można po kolei włączać wszystkie światła przednie. Z odbicia światła na karoserii „poprzednika” łatwo wnioskujemy o działaniu świateł przednich, a nawet bocznych kierunkowskazów. Z tylnymi jest nieco gorzej, ale przy odrobinie cierpliwości można i te światła sprawdzić.

Wymaga to trochę wprawy, ale da się zrobić. No a co na to nasi przyszli wynalazcy?

Grzegorz Wiejas pisze: „światła przednie tzn. drogowe, mijania, dzienne, p. mgielne i kierunkowskazy można łatwo sprawdzić, podjeżdżając do jasno pomalowanej ściany i włączając po kolei różne rodzaje świateł. Ze światłami tylnymi jest trudniej, ale jeśli trafi się ściana bardzo jasno pomalowana i nie będzie to środek dnia (lepiej jest robić to wieczorem), to da się wszystkie światła sprawdzić. Oczywiście nie chodzi tu o ustawienie świateł drogowych i mijania, bo to inna sprawa.

Bardzo dobra propozycja. Jedyny problem to znalezienie ściany, do której da się na tyle blisko podjechać, żeby wszystko, co trzeba, było dobrze widoczne. Słusznie kolega przypomniał, że chodzi nam tylko o stwierdzenie: *światła są lub ich nie ma*. Ustawienie, to rzeczywiście zupełnie inna sprawa.

Jerzy Szymański nie widzi w ogóle problemu i uważa, że można opuścić szybę lewych drzwi – przy fotelu kierowcy – i przez okno obsłużyć wszystkie światła. Oczywiście trzeba się pofatygować i stanąć przed samochodem, później przejść do tyłu. Jedyny problem to światło „stop” i światło jazdy wstecz, które wymagają obecności kierowcy na fotelu. Jerzy uważa, że można ustawić za samochodem lustro i w nim zobaczyć – w lusterku wstecznym – czy światła „stop” i wsteczne – działają.

Nowoczesne samochody „rozleniwiają” kierowców do granic możliwości. Kolega ma rację: sporo można zrobić, tylko trzeba się nieco pofatygować, a przy okazji ocenić jasność świecenia żarówek. Propozycja zastosowania lustra – problematyczna: bo jeśli małe, to trzeba się natrudzić, żeby je dobrze ustawić, a jeśli duże – to kłopot.

Obu kolegom gratuluję i zapraszam do następnych konkursów!

Nowe zadanie

Dziś w epoce GPS, radarów i radiolokacji, problem oceny szybkości statku nie jest żadnym problemem. 200 lat temu był to naprawdę duży problem. Żeglarze mieli swoje sposoby, niektóre bardzo skomplikowane. „Na oko” proste; wystarczyło ustalić pozycję statku w punkcie A i po przepłynięciu np. dwóch godzin, znów wyznaczyć nową pozycję dla punktu B. Dawało to szansę na ocenę średniej prędkości od punktu A do punktu B. Zmora tej metody były bardzo skomplikowane obliczenia (kalkulatorów elektronicznych ani „kręciółków” wtedy nie było) i zanim je oficer nawigacyjny ukończył – statek był już w zupełnie innym miejscu.

Elektroniki wtedy nie było, były chronometry Harrisona, no i... pomysłowość załóg żaglowców. Spróbujcie więc wcielić się w postać oficera nawigacyjnego, który musi znać prędkość statku, żeby móc prawidłowo określić zarówno pozycję statku, jak i czas dopłynięcia do celu. A zatem, wasze zadanie to:

Zaproponować metodę określania prędkości żaglowca z 1800 roku, który nie miał żadnych innych instrumentów oprócz chronometru i sekstansu.

Niech was nie straszy „chronometr” i „sekstans”. Do określenia prędkości w tamtych czasach wystarczył zwykły zegarek kieszonkowy. Weźcie pod uwagę wszystkie „resursy” dla żaglowca z XVIII wieku. Resursy – przypominam – to wszystkie zjawiska i możliwości, jakie można wykorzystać do pomiaru prędkości statku.

Wszystkim życzę pomysłowości dawnych marynarzy i dobrych koncepcji. Przypominam o terminie nadsyłania rozwiązań: koniec lipca br.

Klub Wynalazców

bez ograniczeń wieku

Zadaniem waszym było: *Co zrobić, żeby płukanie wewnętrznego ucha było w pełni bezpieczne w warunkach domowych, bez konieczności udawania się do laryngologa.*

Płukanie ucha, a ściślej – jego przestrzeni od małżowiny do bębienka – zawsze grozi uszkodzeniem bębienka, co może prowadzić do poważnego uszczerbku słuchu. Z tego też powodu zabiegu tego nie wykonują pielęgniarki, a nawet panie z zakładów oferujących aparaty słuchowe. Oczyszczenie wnętrza ucha patyczkiem z wacikiem jest uważane za niebezpieczne, więc robi to wyłącznie lekarz laryngolog. Może też użyć dużej strzykawki i wypłukać ucho wodą z płynem rozpuszczającym woskowinę. Strzykawka może budzić obawę głównie z powodu efektu stick slip. Może on spowodować niekontrolowany wzrost ciśnienia wody podawanej do ucha. Jak to się dzieje? Otóż najpierw naciskamy tłok strzykawki, ale opór tłoka jest na tyle duży, że musimy nacisnąć mocniej, wtedy tłok rusza, ale tarcie statyczne przechodzi w tarcie kinetyczne, znacznie mniejsze i tłok przesuwają się dalej, niżbyśmy chcieli. Oczywiście podaje za dużo wody i często pod większym ciśnieniem. Wydaje się więc, że warunkiem bezpiecznego płukania jest zapewnienie stałego i kontrolowanego ciśnienia podawanej do ucha wody. Jak to zrobić w warunkach domowych? To było wasze zadanie, przyjrzyjmy się więc waszym propozycjom:

Zbigniew Przygodzki: sprzęt odpowiedni do operacji płukania ucha już w zasadzie jest na rynku. Jest to tzw. hegar do wlewów doodbytniczych. Oczywiście do płukania ucha jego zbiornik na wodę jest za duży, można więc wykorzystać zestaw do wlewów dożylnych. Ciśnienie wlewu zależy od wysokości zawieszenia zbiornika nad poziomem ucha pacjenta i oczywiście może być łatwo precyzyjnie regulowane. Wężyk doprowadzający wodę (zapewne z jakimś płynem, rozpuszczającym woskowinę) powinien mieć zawór, umożliwiający zatrzymywanie wlewu lub otwarcie w dowolnym momencie.

Propozycja oczywista i bardzo dobra. Nie wiadomo, dlaczego laryngolodzy o takim sposobie nie wiedzą albo nie chcą go zastosować. Przepisy nie zezwalają na wykonywanie zabiegu czyszczenia ucha przez pielęgniarki, a do lekarza specjalisty laryngologa – jak wiadomo, dostać się jest trudno – i czasami

trzeba czekać na termin 3–4 miesiące. Z drugiej jednak strony, płukanie ucha może lekarzowi dostarczyć informacji, jakie tylko specjalista może ocenić. Może więc nie tylko o ciśnienie chodzi?

Sławomir Horabik proponuje zastosowanie trójnika zakładanego bezpośrednio na końcówkę strzykawki i zakładanie igły na trójnik, a do jego bocznego ujścia należałoby podłączać mały balonik z cienkiej gumy lub lateksu. W razie niekontrolowanego wzrostu ciśnienia w strzykawce balonik amortyzowałby ten wzrost, a poza tym sygnalizowałby przez rozdymanie się – że ciśnienie zostało przekroczone. Nie trzeba by niczego zmieniać, a jedynie wyprodukować odpowiednie plastikowe trójniki.

Pomysł niezły, bo praktycznie nie zmienia nic w dotychczasowej praktyce, do której lekarze są przyzwyczajeni. Regulacja ciśnienia płynu podawanego do wnętrza ucha jest mniej dokładna niż w propozycji Zbigniewa, ale tu raczej chodzi o nieprzekraczanie pewnej wartości krytycznej, a tę można ustalić, określając maksymalny wzrost średnicy balonika.

Obu kolegom gratuluję i życzę jak najradyziej korzystać z interwencji jakichkolwiek lekarzy! Oczywiście zapraszam do kolejnych zadań.

Nowe zadanie

Zadanie z fizyki szkolnej – takiej na poziomie gimnazjum. Nauczyciel postanowił pokazać uczniom doświadczalnie, jak mocno jest zanieczyszczone powietrze w przemysłowej części miasta, gdzie są huty i jedna kopalnia. Poleciał uczniowi wziąć szczelną zamykaną słoik, udać się do wspomnianej dzielnicy i przynieść stamtąd próbkę powietrza. Wszystko jasne, ale jak usunąć „szkolne” powietrze tak, żeby w słoiku znalazło się powietrze tylko z tej dzielnicy. Jest to więc zadanie typowe dla „dawnych czasów”, kiedy to szkoły nie miały odpowiedniego wyposażenia technicznego i nauczyciele musieli często improwizować. A więc:

W jaki sposób pobrać próbkę powietrza z dzielnicy przemysłowej, tak aby w słoiku przeznaczonym do tego celu znalazło się tylko „właściwe” powietrze.



Oczywiście odpada użycie pompy i usunięcie za jej pomocą powietrza ze słoika. Problem powstał w ramach zwykłej szkoły, które zazwyczaj nie są zbyt bogate w sprzęt do fizyki. Rozwiązaniem może być jedynie sposób: sprytny i prosty.

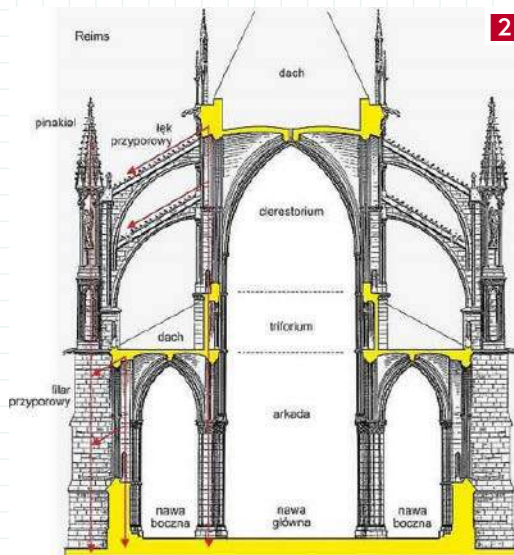
Wszystkim życzę sprytu i pomysłowości na poziomie uczniów szkoły średniej! Przypominam o terminie nadsyłania propozycji rozwiązań: koniec lipca br.

Vademecum Młodego Wynalazcy

Dziś dalszy ciąg tematyki: „Teoria rozwoju osobowości twórczej” (TROT). Ten fragment to tzw. „reinventing”, czyli próba dokonania jeszcze raz wynalazków powszechnie znanych, czyli odtwarzanie drogi myślowej dawnych wynalazców, która doprowadziła ich do wybitnych niekiedy rezultatów. Bardzo często popełniamy błąd, polegający na lekceważeniu dokonań średniowiecznych i dawniejszych jeszcze twórców, patrząc „z góry” poprzez nasze obecne, komputerowe możliwości. Czy ci dawni mistrzowie rzeczywiście zasługują na takie, nieco lekceważące traktowanie?

W różnych dziedzinach oczywiście różny był poziom doskonałości dawnych dzieł, ale te powszechnie znane, jak np. słynna katedra kolońska, nawet dziś przy bardzo rozwiniętych technikach budowlanych budzą zachwyt (1).

Podziw budzi drobnoelementowa faktura zewnętrznych elewacji. Trzeba oczywiście pamiętać, że wszystko to było wykonane z kamienia, obrabianego narzędziami głównie ręcznymi. Do tego trzeba dodać, że nie istniała jeszcze mechanika budowli, pozwalająca na obliczeniowe dobieranie wymiarów elementów. Właściwie nie istniała dokumentacja w dzisiejszym rozumieniu. Jeden z głównych problemów to zjawisko „rozporu” ścian w wyniku



działania ciężaru sklepienia. Radzono sobie z tym różnie. Przede wszystkim budowle musiały mieć bardzo grube mury, dające pewne oparcie łukom sklepienia. Kolejnym rozwiązaniem były filary oporowe (2).

Filary oporowe były racjonalnym rozwiązaniem, gdyż wymagały znacznie mniej materiału niż grube ściany. Filary te miały przekrój prostokąta, zwróconego krótszą krawędzią do ściany budowli. Stopniowano też szerokość filarów, zmniejszając ją na wyższych poziomach. Na najwyższym poziomie wieże były połączone tzw. łękiem przyporowym ze ścianą budowli, dając opór siłom rozporu, wynikającym z ciężaru łuków sklepienia. Mimo to zdarzały się przypadki, że po zdementowaniu drewnianych, tymczasowych rusztowań okazywało się, że ciężar łuków był za duży i budowla otwierała się jak książka... Kończyło się to niekiedy samobójstwem budowniczego. W rezultacie wiedza o doborze wymiarów i materiałów do budowy wysokich budynków, zwłaszcza gotyckich katedr, stała się tajemnicą, przekazywaną jedynie wąskiemu gronu zaufanych osób. Dało to początek organizacji „wolnych mularzy”, których gódem był

cyrkiel i „winkel”, inaczej – kątownik. Do tego fartuch roboczy i – znacznie później – rytuały i cała filozofia ruchu wolnomularskiego, inaczej masońskiego, która już daleko odbiegała od pierwotnego kształtu organizacji typu cechu rzemieślniczego. Gotyckie katedry i ich budowa są dziś dość powszechnie znane.

Są jednak takie obiekty, które do dziś budzą właściwie same pytania, w większości bez odpowiedzi. Takim obiektem są granitowe kształtki Puma Punku znajdujące się w miejscowości Tiahuanaco w Boliwii. Granit – andezyt, to jeden z najtwardszych naturalnych materiałów, oprócz niego do budowlı użyto czerwonego piaskowca. Kamieniołom, z którego pochodzą

bloki, odległy jest o ok. 15 km od miejsca, gdzie tworzą one niewyjaśniony do dziś kompleks. Uważa się to za obiekt świątynny, ale niejasne są dziwne kształty, jak również metoda transportu z kamieniołomu na miejsce montażu, bloków o masie przekraczającej 100 ton, do dziś nie jest wyjaśniona. Pamiętajmy, że Inkowie nie znali koła! Zdziwiał również precyzja wykonania poszczególnych kształtek, Ściany niemal idealnie płaskie, kąty proste, kąty wewnętrzne wykonane „na ostro” bez typowych zaokrągleń i otwory: o różnych kształtach i niekiedy na wylot andezytowej płyty (3–5)!

O obiektach Puma Punku jest w Internecie sporo informacji, a także o hipotezach, dotyczących historii, przeznaczenia i technologii wykonania obiektów, z materiałów należących do szczególnie twardych. Precyzja wykonania ścian – niemal idealnie płaskich, pomijając oczywiście skutki erozji, wcięcia i otwory z ostrymi kątami – wszystko to zdaje się wykluczać przyjęcie hipotezy, że wykonali to starożytni Inkowie, A więc kto? I po co? Najczęściej powtarzaną hipotezą jest ingerencja „kosmitów”, albo bardzo, bardzo dawnej, dziś zupełnie nieznannej cywilizacji.

No i trzeci obiekt, a właściwie obiekty, bo znaleziono ich bardzo dużo, to tajemnicze sprężynki, jak wykazały badania, wykonane z niemal czystego wolframu.





Rysunek 6 pokazuje sprężynki na tle monety jedno-groszowej, przy czym są to największe z dużej liczby znalezionych. Najmniejsze miały długość 0,003 mm (!) i były widoczne wyłącznie pod mikroskopem. Istotną ich cechą było też i to, że większość miała rdzenie, wykonane z molibdenu.

Drucik, z którego wykonane były sprężynki, ma wzdłużne rysy (7), takie jakie powstałyby przy przeciąganiu go przez wyszczerbione oczko. Zważywszy na właściwości wolframu: twardy i kruchy materiał, o temperaturze topnienia 3422°C, jest to obiekt niewykonalny w czasach odległych o ok. 100 000 lat. Temperatura topnienia molibdenu – 2623°C też wklucza starożytne technologie.

Pytania: kto i w jakim celu wyprodukował te sprężynki i jak to zrobił – pozostają do dziś bez odpowiedzi...

Te trzy obiekty, tajemnica ich przeznaczenia i powstania, to do dziś jedne z najbardziej frapujących problemów. Oczywiście pojawiło się sporo hipotez, ale żadna z nich nie jest w pełni zadowalająca. Pokazaliśmy te problemy, wiedząc, że nie są łatwe, ale może kogoś zainteresują i zachęcą do wglębnienia się w ich tajemnice.

Pamiętamy przecież historię odkrycia ruin Troi przez Henryka Schliemanna. Jego fascynacja problemem Troi zaczęła się w zasadzie już w dzieciństwie i ukształtowała jego życie zawodowe, osobiste i naukowe. Może więc ktoś z naszych czytelników wyjaśni kiedyś: kto zbudował Puma Punku. Jak bez matematyki i statyki budowli budowano gotyckie katedry, kto i po co produkował miniaturowe sprężynki z wolframu i z molibdenowym rdzeniem?

A jako zadanie startowe stosunkowo drobny problem, chociaż przyprawiający o ból głowy nawet bardzo dobrych specjalistów z optyki. Chodzi o „magiczne lustro” produkowane kiedyś w Chinach. Lustro z brązu wykonywane było jako okrągła tarcza, z jednej strony pokryta jakimś dekoracyjnym reliefem, a druga strona służyła jako zwyczajne lustro (8).

Właśnie w tym problem, że nie było to zwyczajne lustro: przy rzuceniu „zajęczka” na białą ścianę ukazywał się na niej błąd, ale dobrze widoczny wizerunek Buddy w otoczeniu promieni.

Wymagało to pewnej zręczności i ustawienia zwierciadła pod odpowiednim kątem do kierunku padania promieni światła. Tajemniczość zwierciadła polegała na tym, że ani na stronie pokrytej dekoracyjnym reliefem, ani na gładko wypolerowanej stronie „użytkowej” – nie było widać wizerunku Buddy! W jaki sposób pojawiał się na ścianie?

Oczywiście w dawnych Chinach takie zwierciadła były uważane za niemal święte przedmioty i ich używanie wiązało się z pewnym rytuałem o charakterze religijnym.



Jako młodzi technicy, myślący racjonalnie, stosujemy „brzytwę Ockhama” czyli: „nie należy mnożyć bytów ponad potrzebę”, a mówiąc prościej: nie należy doszukiwać się cudów tam, gdzie wystarczy fizyka.

No to jeszcze jeden problem, tym razem „czarodziejskiej butelki” admirała Łazarowa. W długich rejsach na żaglowcach trzeba było marynarzy czymś zająć i oczywiście poza myciem pokładu, pompowaniem wody z żęzy, należało ich rozruszać intelektualnie. Admirał demonstrował marynarzom sztuczkę z butelką. Brał pustą butelkę po winie, zamykał ją szczelnie korkiem i po obciążeniu ołowianym ciężarkiem opuszczał ją na głębokość ok. 400 m. Po wyciągnięciu butelki okazywało się, że była pełna wody, mimo że nadal była szczelnie zakorkowana. Jak to możliwe?

I tu również działa czysta fizyka. Spróbujcie więc, zanim zaczniecie rozmyślać o Puma Punku, rozwiązać te ostatnie zagadki.

Zagadki, rebusy i szarady od dawna są znane jako doskonaly materiał kształcenia bystrości umysłu i logicznego myślenia. W *Lalce* Bolesława Prusa jest przytoczona rozmowa Wokulskiego ze starym Szlangbaumem. Proponuję przemyśleć, co ten stary Żyd mówi:

„U nas, panie, niby u Żydów, jak się młodzi zejdą, to oni nie zajmują się, jak u państwo, tańcami, komplementami, ubiorami, głupstwami, ale oni albo robią rachunki, albo oglądają uczone książki, jeden przed drugim zdaje egzamin, albo rozwiązują sobie szarady, rebusy, szachowe zadanie. U nas ciągle jest zajęty rozum i dlatego Żydzi mają rozum, i dlatego, niech się pan nie obrazi, oni cały świat zawojują”.

Spróbujcie wyciągnąć wnioski z tej rozmowy. ■

Prezes Klubu Wynalazców
Champion TRIZ
Jan Boratyński

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

Pomysł miesiąca 6/2023

Ciekawy wydaje się nam pomysł wysuwanej z podwozia platformy lub innej konstrukcji pozwalające wygodnie obracać samochód i dostosowywać jego ułożenie do miejsca postoju. Koncepcja pasuje do logiki zmian we współczesnej motoryzacji.

Autorem pomysłu jest Marek Lewiński

1 **Marek Bielecki** widział kiedyś, jak krawiec, który szył męski płaszcz, prasował go specjalnym ciężkim żelazkiem na parę. Żelazko było ciężkie, bo gruby materiał wymagał dużej siły nacisku. Marek proponuje zainstalowanie w żelazku krawieckim wibratora, który pozwoliłby na zbudowanie lżejszego żelazka, ale mogłoby dawać bardzo dobre rezultaty przy mniejszym wysiłku krawca. **Bardzo ciekawy pomysł i w zasadzie dość powszechnie znany: wibratory do zgęszczania mas betonowych, ceramicznych itp. Żelazko mogłoby istotnie zyskać na sprawności: wibracja w połączeniu z parą to naprawdę potężna broń na zalamania tkaniny, nawet bardzo grubej.**

2 **Dariusz Zakrzewski** zwiedzał kiedyś duży zakład produkujący aparaturę dla cukrowni. Zauważył, że między innymi produkuje się tam ogromne wymienniki ciepła, złożone z kilku setek rur, osadzanych szczelnie końcówkami w górnym i dolnym dnie. Ogromnie pracochłonna robota, a do tego jeszcze obydwa dna mają kształt sferyczny. Darek proponuje kształtować na gorąco końcówki rur w formie wielokąta, np. sześciokąta, pakietować i spawać na automacie. Niepotrzebne byłyby dna, wiercenie w nich otworów i pracochłonne osadzanie rur.

3 **Propozycja Darka wygląda na bardzo ciekawą i naprawdę łatwiejszą niż ta, jaką miał okazję zobaczyć. Oczywiście jest jeszcze sprawa parametrów: ciśnienia i temperatury, ale rzecz wydaje się naprawdę przełomowa dla tego typu obiektów.**

4 **Marek Lewiński** w nawiązaniu do propozycji kolegi Marka z numeru majowego MT proponuje bardziej radykalne podejście. Samochody – jego zdaniem – powinny być wyposażane w obrotową platformę, niedużą, tak ok. 1 m średnicy, opuszczaną na ziemię i umożliwiającą uniesienie samochodu o jakieś 2–3 cm powyżej poziomu parkingu. Niewielki silnik elektryczny umożliwiłby obrót samochodu – jak na karuzeli i w ogromnym stopniu ułatwiłby manewrowanie na zbyt ciasnych parkingach.

5 **Byłoby to spełnienie marzeń wielu kierowców i rzeczywiście istotne ułatwienie manewrów w każdym ciasnym miejscu. Nie wiemy, jak potoczy się przyszłość motoryzacji, jednakże obecny trend, czyli mnożenie liczby**

samochodów, jest nie do utrzymania. Wizje przyszłości w postaci pojazdów latających to raczej nierealna mrzonka, chociażby z powodu znacznie większych wymagań kwalifikacyjnych dla użytkowników.

6 **Bogusław Kowal** ma kolegę, który uprawia sport jeździecki i ma nawet własnego konia. Kolega ma powtarzalny kłopot z jego podkuwaniem. Fachowych podkuwaczy dziś jest znikoma liczba i chcąc dobrze podkuć konia, trzeba czasem wziąć go do przyczepy i jechać kilkaset km do renomowanego kowala – podkuwacza. Bogusław proponuje zamiast podkuwania zastosować podklejanie podków. Są dziś znakomite kleje, nawet takie, jakich używa się do przyklejania płytek termoodpornych na powierzchni wahadłowców kosmicznych, to klej do przyklejania podkowy nie powinien być problemem.

7 **Przymocowanie podkowy to tylko część problemu. Kowal – podkuwacz wykonuje przy tej okazji swoje pedicure: wyrównuje powierzchnię kopyta, bo przecież kopyto – tak jak nasze paznokcie – rośnie i zmienia kształt. Dla koni rekreacyjnych np. kuców są już inne metody ochrony kopyt: buty, podkowy z tworzyw, oczywiście przyklejane, itd. Konie sportowe to jednak inna klasa i tu należałoby opracować coś nowego.**

8 **Krzysztof Dudzik** widział kiedyś w TV film pt. *Dru gie życie kurczaka*. Film miał charakter interwencyjno-ostrzegawczy i opowiadał o niegodziwych praktykach niektórych zarządców marketów i dużych sklepów, którzy poddają nieświeże tuszki kurczaków „odmłodzeniu” dla nadania im „świeżego handlowego” wyglądu. Trudno wymagać od przeciętnego klienta, żeby potrafił bezbłędnie odróżnić mięso świeże od nieświeżego. Wiadomo, że potrafi to zrobić pies, ale my – ludzie – nie mamy takich talentów. Krzysztofowi marzy się aparat z igłą, którą wbiłaby się w mięso, a na skali można by było odczytać wynik: „świeży” lub „nieświeży”.

9 **Testery o podobnym działaniu już są, ale Krzysztofowi chodzi o sprzęt podręczny, niekoniecznie uniwersalny (do warzyw i owoców) dający jasny sygnał „świeży – nieświeży”. Technika rozwija się chyba szybciej niż nasze marzenia...**



Ekranoplan MT-2023

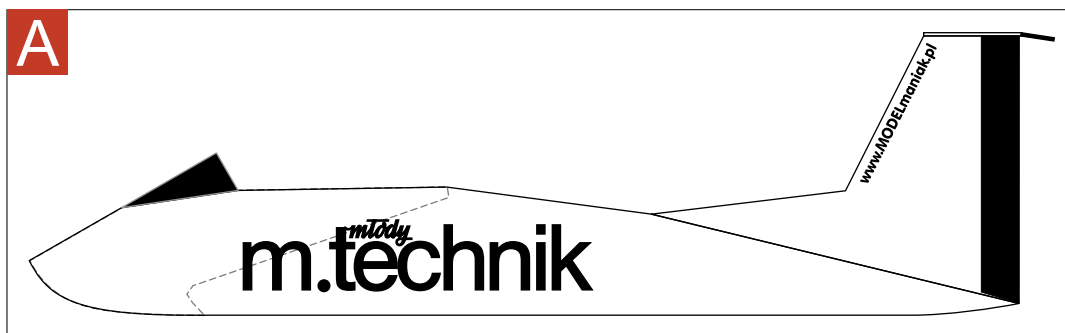
Dziś „Na warsztacie” powracamy do niezwykłych środków transportu poruszających się na poduszce powietrznej, budując prosty model przemieszczający się kilka milimetrów nad podłogą.

Ale co to?

Ekranoplan (ang. screen plane, Wing-In-Ground -> WIG, WISE) to rodzaj statku powietrznego poruszającego się na niewielkiej wysokości nad wodą, lodem lub inną gładką powierzchnią, wykorzystujący zjawisko wzrostu siły nośnej na skutek zwiększonego ciśnienia pomiędzy poruszającym się płatem nośnym a podłożem. Polska nazwa powstała jako połączenie słów ekran (w domyśle powietrzny – rodzaj poduszki powietrznej, powstający np. między kładzionymi

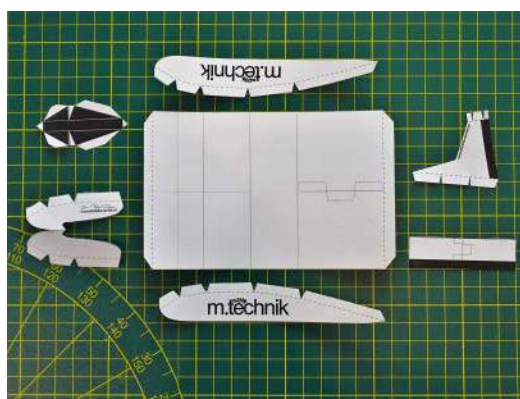
szybko taflami szkła) oraz słowa planować (z francuskiego: szybować).

Zjawisko zwiększenia noszenia w pobliżu podłoża znane jest pilotom samolotów niemal od początku lotnictwa – tym wyraźniejsze, im doskonalszy aerodynamicznie jest płatewiec (szybownikom w pewnym sensie nawet przeszkadza w lądowaniu), jednak budowę statków powietrznych, przeznaczonych wyłącznie do lotów z wykorzystaniem tego zjawiska, rozpoczęto dopiero po drugiej wojnie





1. Współczesny, dziesięcioosobowy ekranoplan serii AIRFISH 8 (fot. via <https://www.wigetworks.com/airfish-8>)



2. Skopiowane na bryistol elementy modelu należy starannie wyciąć



3. Kolejnym krokiem jest zbigowanie krawędzi oraz ich odpowiednie zagięcie



4. Burty przykleja się do dłuższych krawędzi kadłuba (trzeba zwrócić uwagę na oznaczone linie zagięć). W części dziobowej, od spodu, wkleja się podwójnej grubości zaczep startowy



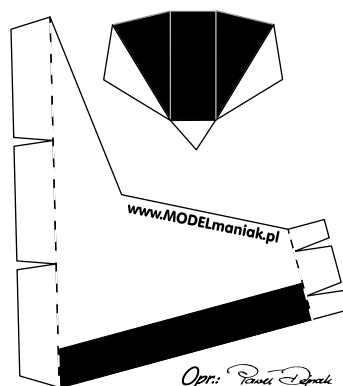
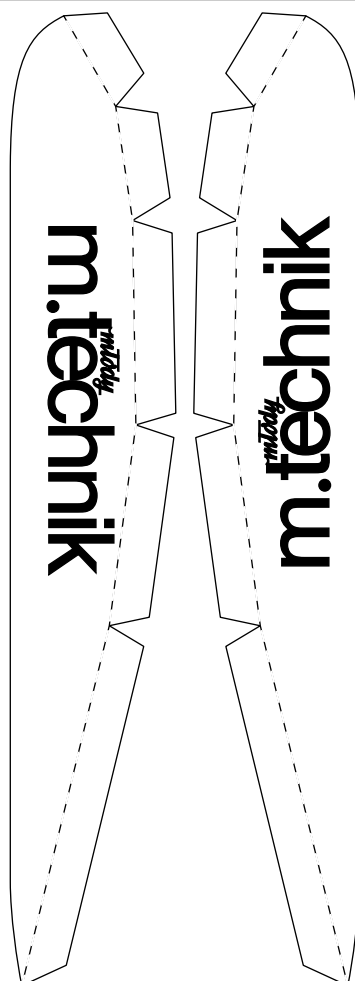
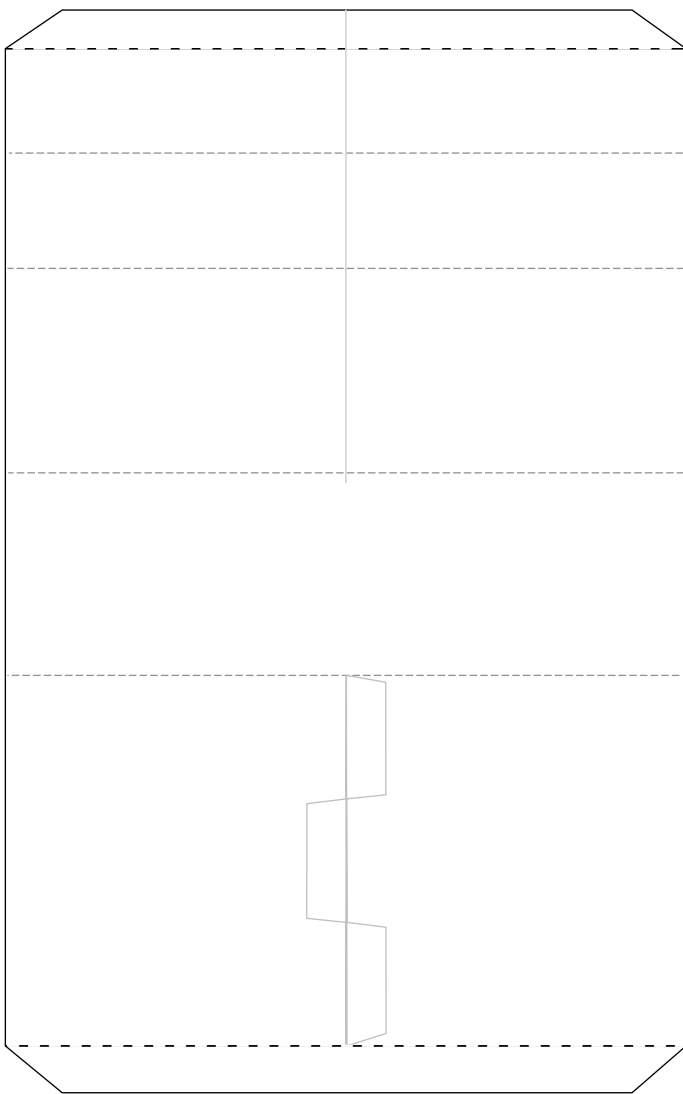
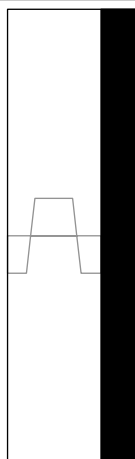
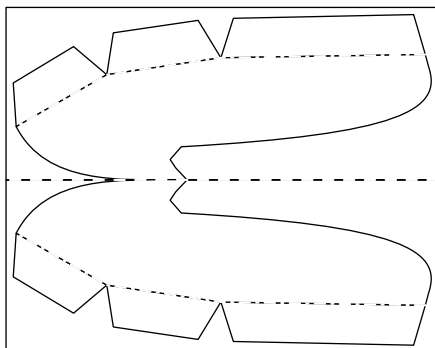
5. Do tylnej części kadłuba przykleja się sklejone wcześniej usterzenie typu „T”



B

EKRANOPLAN MT-2023

Skala 1:1

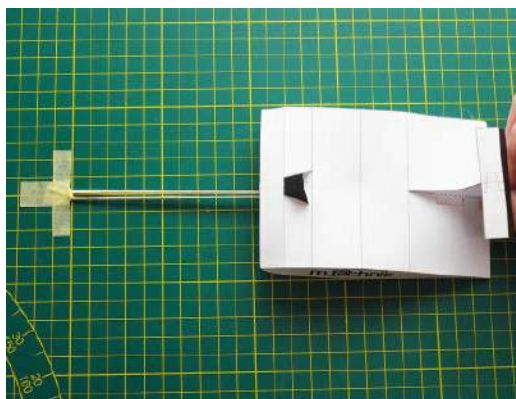


www.MODELmaniak.pl

Opr.: *Renata Lipińska*
www.MODELmaniak.pl



6. Wyrzutnia może być wykonana nawet ze zwykłej gumki recepturki mocowanej taśmą klejącą – tu papierową, dla lepszej widoczności. Wskazówka, aby naprężona gumka nie zerwała taśmy, w tym miejscu warto zawinąć taśmę „na cukierek” – o 360° i skleić ją ze sobą, obejmując gumkę. Dodatkowa taśma poprzeczna zapobiegnie podnoszeniu się bliższej modelowi części taśmy



7. Po zaczepieniu do wyrzutni, trzymając model za ogon, naciągamy gumę – i puszczaamy. Dobrze wyważony i ustawiony model poleci kilka metrów w dal parę milimetrów nad podłogą – ale to wymaga trochę treningów – nawet ptaki uczą się przecież latać

światowej. Doceniono bowiem wiele zalet ekranoplanów, na przykład:

- niemal połowę mniejsze zużycie paliwa w stosunku do tej samej masy samolotu konwencjonalnego,
- większy udźwign przy tej samej mocy silników,
- brak choroby lokomocyjnej u pasażerów ekranoplanów (nie występują tu ani turbulencje, ani falowanie),
- duża samostateczność konstrukcji lotniczej,
- niskie koszty utrzymania,
- łatwość obsługi,
- duże bezpieczeństwo.

Ekranoplany najczęściej używane są na rzekach, jeziorach lub wodach przybrzeżnych. Największe tego rodzaju statki powietrzne budowano w ZSRR w latach siedemdziesiątych, ale i dziś używane są turystyczne konstrukcje tego rodzaju – np. AIRFISH 8 (1). Ze względu na swoją wyjątkowość ekranoplany są też chętnie budowane przez modelarzy.

Ekranoplan MT-2023

Budowa modelu naprawdę nie jest trudna, jednak należy zwrócić uwagę na staranny montaż całości. Od tego w dużej mierze zależą poprawne loty. To jeszcze prostszy model niż opublikowany w marcowym wydaniu „Młodego Technika” z 2011 r. Jest zmodyfikowaną na bazie wrocławskich emdekowych doświadczeń wersją modelu RAM doktora Syozo Kubo z Uniwersytetu Tottori w Japonii.

Elementy modelu w skali 1:1 zamieszczone są na **rysunku B**. Do jego budowy potrzebujemy kartki brystolu

z bloku technicznego 160g/m² (warto sprawdzić oznaczenie na bloku technicznym) (5).

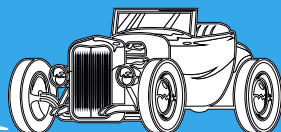
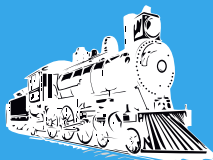
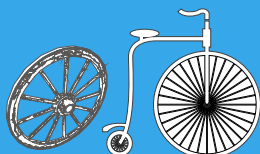
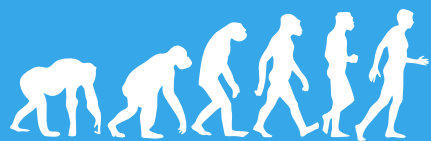
Elementy starannie wycinamy (2). Przed sklejeniem należy zwrócić uwagę na linie zagięć (3) oraz na mankiety wzmacniające krawędzie płata (i kadłuba jednocześnie). Linie zagięć bigujemy np. grubą igłą przy linijce. Jako pierwszy warto skleić kadłub z wcześniej wyciętymi burtami. W przedniej części kadłuba, od spodu montowany, jest zaczep startowy (4). Na końcu model montujemy usterzenie ogonowe typu „T” oraz owiewkę (5).

Testowe loty „ekranowe”

Do pierwszych lotów potrzebna nam będzie przede wszystkim równa podłoga (minimum w dużym pokoju lub klasie). Model można wyrzucać z ręki (kwestia wprawdy) lub użyć wyciągarki gumowej, mocowanej do podłogi taśmą samoprzylepną (6, 7). Możliwe, że pierwsze loty nie będą imponujące. Kartonowe modele są niestety podatne na niesymetryczne deformacje, które niezauważone, utrudniają prawidłowe loty. W tej konstrukcji po sprawdzeniu poprawności (symetrii) montażu możemy regulować lot za pomocą ustawienia steru wysokości i położeniem środka ciężkości (obciążenie można umieścić na zaczepie startowym). Podczas startu z wyrzutni gumowej model trzyma się za usterzenie pionowe. Tak jak w przypadku modeli szybowców, uzyskanie pięknych lotów wymaga treningu – naszym celem jest równy lot kilka milimetrów nad podłogą (bardziej to słycać, niż widać).

Wszystkim konstruktorom modeli WIG życząc udanych i pouczających eksperymentów! ■

Paweł Dejnak



starożytność

IV–II wiek p.n.e.

1283

XV–XVI w.

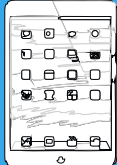
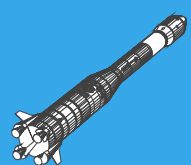
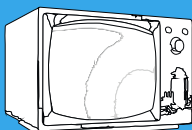
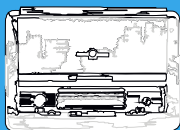
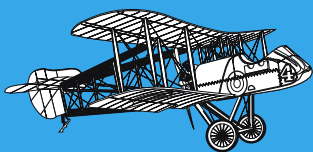
Przekładnie mechaniczne

Archeologowie odkopali szczątki pochodzącego z Chin tzw. Rydwanu Wskazującego Południe, wyposażonego w układ kół zębatach, utrzymujący za pomocą mechanizmu różnicowego figurę w pozycji wskazujące południe, gdy koła się obracały, a wóz poruszał (1). Mechanizm taki, został po raz pierwszy odnotowany w piśmie z IV wieku naszej ery. Według różnych źródeł, za wynalazcę Rydwanu Wskazującego Południe uznaje się Żółtego Cesarza (Huangdi) z ok. 2650 r. p.n.e. lub księcia Zhou z ok. 1000 r. p.n.e. Konstrukcja jest najstarszym znanym pojazdem, w którym zastosowano przekładnię. Najwcześniejsze opisy przekładni pasowych pochodzą ze starożytnego Egiptu i Mezopotamii. W Chinach mechanizm przekładni opartej na przenoszeniu za pomocą pasów został po raz pierwszy wymieniony w dziele filozofa, poety i polityka z dynastii Han, Yanga Xionga (53–18 p.n.e.). W roku piętnastym p.n.e. zastosowany został w maszynie, która nawijała włókna jedwabiu na szpulki do czółenek tkaczy.

Najstarsze europejskie zapisy o przekładniach zębatych pochodzą ze starożytnej Grecji. Pierwsze urządzenia nazywane dziś śrubą Archimedesesa, a będące rodzajem ślimakowego podnośnika współpracującego z kołem zębatym zostały opisane przez Archimedesesa (2). Z kolei najstarsze znane zastosowanie napędu łańcuchowego w rodzaju miotającej wyrzutni o nazwie polybolos opisane zostało przez greckiego inżyniera Filona z Bizancjum (III w. p.n.e.). Dwa płaskie łańcuchy połączone były z układem dźwigni, służącej do naciągania mechanizmu. Chociaż urządzenie nie przenosiło mocy w sposób ciągły, ponieważ łańcuchy nie przenosiły mocy z wału na wał, to grecka konstrukcja wyznacza początek historii napędu łańcuchowego.

W średniowieczu przekładnie zębate były szeroko stosowane w młynach, wodociągach, mechanizmach wojennych i innych maszynach. Początkowo były to proste mechanizmy z jednym kołem zębatym, ale wraz z rozwojem technologii powstawały coraz bardziej skomplikowane układy. W XIII wieku pojawiły się zegary wieżowe z pierwszymi mechanizmami napędowymi wykorzystującymi np. opadający balast. Znaną dokładną datą jest 1283 rok, gdy pierwszy zegar mechaniczny napędzany ciężarem pojawił się w katedrze w Salisbury w Anglii (3). W tym samym okresie mnisi w północnych Włoszech budowali dzwonnice wykorzystujące w mechanizmach odmierzających czas koła zębate, używane jako zegary przypominające o modlitwie.

W górnictwie wynaleziono dwa nowe typy maszyn odwadniających, które umożliwiały głębsze wydobywanie w sztolni: pompę szmaciano-łańcuchową (ok. 1430) i wyciąg z worków (ok. 1470). Szybko jednak osiągnięto jednak granice możliwości tych maszyn. Zaczęto stosować układy pionowych i poziomych prętów, połączonych ze sobą za pomocą „krzyża napędowego” (Kunst Kreuz), który przełączał kierunek ruchu o 90 stopni i belek. Napęd prętowy rozpowszechnił się szybko w drugiej połowie XVI wieku, zyskując kilka ulepszeń. Jedną z jego wersji pojawia się w książce Georgiusa Agricoli z 1556 r. o górnictwie, *De Re Metallica*, w której autor odnotowuje, że maszyna była już wcześniej znana. System transmisji za pomocą płaskich prętów (niem. Stangenkunst) według zapisków został wynaleziony prawdopodobnie w 1510 roku przez włoskiego metalurga Vanoccio Biringuccio, który zaprojektował urządzenie pozwalające kołu wodnemu napędzać pracę miechów w czterech oddzielnych kuźniach w hucie żelaza. Kilkadziesiąt lat później Niemcy zastosowali to samo podejście do przenoszenia energii mechanicznej na znacznie większe odległości, nawet do czterech kilometrów. Systemy płaskich drągów przenoszących (4) były szeroko stosowane w Harzu i Rudawach w Niemczech, a także w Kornwalii w Anglii i Bergslagen w Szwecji.



1478

XVI w.

XVIII w.

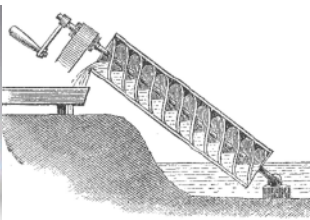
1781

Leonardo da Vinci tworzy szkice projektu uważanego obecnie za pierwowzór autonomicznego pojazdu, w którym mechanizmy napędowe oparte były na układzie przekładni kół zębatach i łańcuchów napędowych (5). Leonardo da Vinci (1452–1519) wykorzystał skomplikowany układ zębów kół zębatach w wielu swoich odkryciach, takich jak urządzenia do ostrzenia soczewek i obróbki metalu.

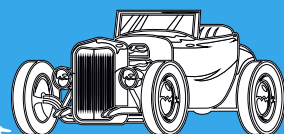
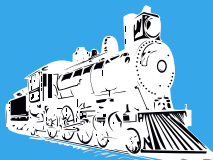
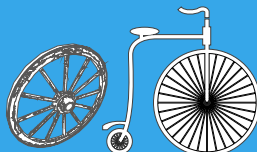
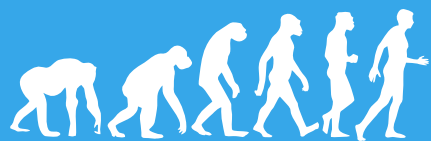
Agostino Ramelli znalazł pierwsze praktyczne zastosowanie łańcucha w rozwoju systemów pompowania napędzanych wodą. Pierwszy powszechnie przyjęty łańcuch nazywany był łańcuchem zębatym, ponieważ poruszał się na kole zębatym lub „zębatce”. Oparta na prostokątnej konstrukcji odlewanych ogniw połączonych zapętlonymi i nitowanymi taśmami żelaznymi, konstrukcja była prosta, ale trudna do naprawy bez specjalistycznego oprzyrządowania. Łańcuch zębaty był używany do przenoszenia mocy i ruchu z bieżni do szeregu innych urządzeń, takich jak windy wodne, maszyny rolnicze i krosna tkackie. Przekładnie łańcuchowe były stosowane w manufakturach. Służyły do przenoszenia napędu z wału maszyny na poszczególne elementy jej mechanizmu. Później wraz z rozwojem przemysłu motoryzacyjnego przekładnie łańcuchowe zaczęły być stosowane w motocyklach i rowerach. Zastąpiły one starsze systemy napędowe oparte na pasach lub szprychach.

Pierwsze teoretyczne eseje na temat mechanizmu trakcyjnego opublikował słynny matematyk Leonhard Euler. Wraz z rozwojem przemysłu zaczęły pojawiać się przekładnie zębate z metalu o wysokiej precyzji wykonania, co pozwalało na uzyskiwanie dużych przekładni, co zwiększyło ich wytrzymałość i precyzję.

Wynalezienie przez szkockiego inżyniera Williama Murdocha przekładni słonecznych i planetarnych umożliwiło wykorzystanie siły pary do wytworzenia ciągłego ruchu kołowego wokół środka przekładni. Ruch ten był wykorzystywany do obracania kół lub wałów innych maszyn. Celem wynalazku było stworzenie nowego rodzaju ruchu korbowego. Ponieważ Murdoch był pracownikiem w firmie Jamesa Watta, patent na wynalazek należał do Jamesa Watta, który go zarejestrował. Koła zębate tworzyły układ analogiczny do słońca (koło wewnętrzne) i planety (koło zewnętrzne), przekształcając liniowy ruch tłoka w silniku spalinowym lub parowym na ruch obrotowy potrzebny do napędzania różnych mechanizmów. Przekładnia słoneczna była przymocowana do końca pręta łączącego z silnikiem, a gdy tłok poruszał się w górę i w dół, przekładnia planetarna obracała przekładnię słoneczną. Przekładnia słoneczna i planetarna zwiększały wydajność silnika parowego.



1. Rekonstrukcja starożytnego chińskiego Rydwanu Wskazującego Południe; 2. Zastosowanie tzw. śruby Archimedesa; 3. Zrekonstruowany średniowieczny mechanizm zegarowy z katedry w Salisbury w Anglii; 4. Fragment starej ryciny demonstrującej przenoszenie napędu typu Stangenkunst; 5. Mechanizmy łańcuchowe na szkicach Leonarda da Vinci; 6. Mechanizm przekładni obiegowej opracowany w przedsiębiorstwie



XIX w.

W 1835 r. angielski wynalazca Joseph Whitworth opatentował pierwszy proces frezowania kół zębatach. Za pomocą tej maszyny można było wycinać koła zębata o różnej liczbie zębów i różnych profilach. Ponieważ śruby prowadzące umożliwiały dokładne ustawienie, dzięki wynalezieniu maszyny do cięcia kół zębatach w XIX wieku można było produkować koła zębata o jednakowych wymiarach. Przekładnie zębata zaczęły być produkowane na masową skalę i zaczęły znajdować zastosowanie w wielu gałęziach przemysłu. Herman Pfauder z Niemiec skonstruował wydającą frezarkę do kół zębatach pod koniec XIX w.

1874

Amerikanin William Gleason wynalazł i opatentował strugarkę z przekładnią stożkową, rodzaj maszyny, która otworzyła nowe możliwości przenoszenia napędu. Przekładnie stożkowe umożliwiają przenoszenie mocy pod kątem 90° z wału napędowego samochodu na oś jego tylnego koła. Później, w XX wieku wprowadzono przekładnie zębata stożkowe, o osiach prostopadłych, które pozwoliły na zmniejszenie hałasu i tarcia podczas pracy przekładni.

1877

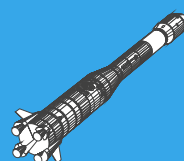
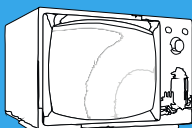
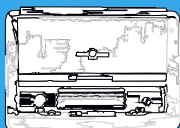
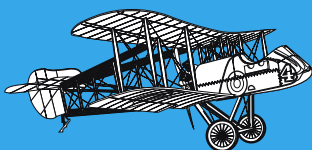
Brityjski wynalazca James Starley patentuje mechanizm różnicowy (7) wykorzystany początkowo w trójkołowym rowerze, a później upowszechniony w wielu innych konstrukcjach. Udoskonalił również stosowaną we wczesnych rowerach przekładnię zębata.

1894–98

Francuzi, Emile Levassor i Louis-René Panhard (8), są uznawani za twórców pierwszej manualnej skrzyni biegów. Levassor i Panhard zastosowali w swojej oryginalnej przekładni napęd łańcuchowy. Ich wynalazek jest nadal podstawowym punktem wyjścia dla współczesnych przekładni manualnych. W 1898 r. Louis Renault zaadaptował do swoich samochodów projekt Levassora i Panharda, zamieniając napęd łańcuchowy na wał napędowy i dodając oś różnicową dla tylnych kół, aby poprawić wydajność ręcznej skrzyni biegów.

1903

Przekładnie obiegowe, znane również jako planetarne lub epicykliczne, które umożliwiły uzyskanie dużych przekładni przy jednoczesnym zmniejszeniu rozmiarów urządzenia, zostały opracowane już przez starożytnych Greków, ale im pomagały jedynie w przewidywaniu ruchów planet w Układzie Słonecznym. W świecie motoryzacji, pierwsze przekładnie wykorzystujące przekładnie planetarne zostały zamontowane na początku XX wieku w samochodzie Wilson-Pilcher, który został wyprodukowany w Wielkiej Brytanii. Następnie Ford Model T wykorzystał przekładnie planetarne w 1908 roku w swojej dwubiegowej manualnej skrzyni biegów. W czasie II wojny światowej przekładnie planetarne były stosowane w samolotach jako reduktory do silników. W latach 50. XX wieku przekładnie planetarne zaczęły być stosowane w przemyśle motoryzacyjnym również do napędu kół. W kolejnych dekadach znalazły wiele zastosowań, m.in. zaczęły być stosowane w układach napędowych turbin wiatrowych.



1917

Pierwszy pasek klinowy został wynaleziony przez Johna Gatesa, współnika Gates Rubber Company. W tamtych czasach była to rewolucyjna koncepcja. Jednak zanim taki sposób przenoszenia napędu się rozpowszechnił, musiało minąć jeszcze nieco czasu.

1928

Firma motoryzacyjna Cadillac wprowadza zsynchronizowaną manualną skrzynię biegów, która znacznie zmniejszyła ścieranie się kół zębatach i sprawiła, że proces zmiany biegów stał się bardziej płynny i wygodny.

1937–38

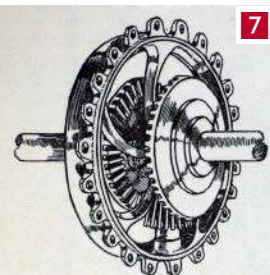
General Motors (GM) wprowadza pierwszą półautomatyczną skrzynię biegów o nazwie Automatic Safety Transmission (9). Zaraz po tej konstrukcji GM wprowadził też rozwiązanie nazwane Hydra-Matic, uznawane za rewolucyjne, wykorzystujące w pojeździe pięciobiegową, bezsprzęgłową skrzynię biegów.

lata 50. XX wieku

Początek lat pięćdziesiątych przyniósł pierwszą trzybiegową automatyczną skrzynię biegów. Studebaker i Ford byli jednymi z pierwszych producentów samochodów z trzybiegową automatyczną skrzynią biegów. W 1953 roku koncern Chrysler wprowadził własny dwubiegowy konwerter momentu obrotowego.

1957

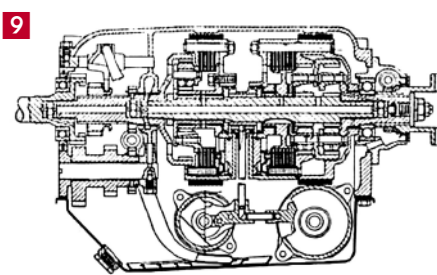
Podstawowa koncepcja przekładni falowej (SWG), przekładni mechanicznej, w której przekazywanie napędu odbywa się przez przesuwanie się fali odkształcenia podatnego wieńca (10), została wprowadzona przez Clarence'a Waltona Musserra w patencie z 1957 r., gdy był on doradcą w United Shoe Machinery Corp (USM). Po raz pierwszy została ona z powodzeniem zastosowana w 1960 r. przez USM, a następnie przez Hasegawa Gear Works na licencji USM. Później Hasegawa Gear Work przekształciła się w Harmonic Drive Systems z siedzibą w Japonii. Ten typ przekładu, bardzo precyzyjny i cichy, był i jest wykorzystywany w przemyśle kosmicznym i robotyce.



7

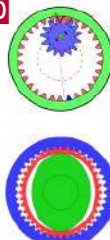


8

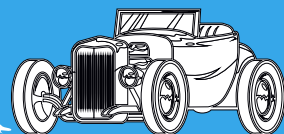
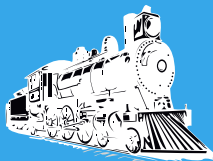
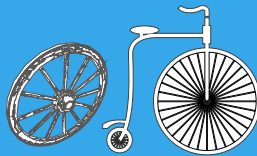
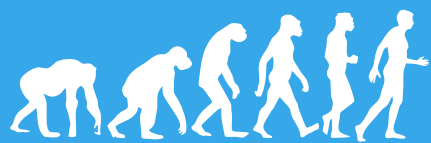


9

10



7. Mechanizm różnicowy Jamesa Starleya; 8. Emile Levassor i Louis-René Panhard w pojeździe wyposażonym w skrzynię biegów, którą wynaleźli; 9. Rysunek techniczny skrzyni Automatic Safety Transmission z 1937 roku; 10. Porównanie schematu przekładni obiegowej planetarnej z przekładnią falową



ODKRYJ HISTORIĘ WYNALEZKÓW

Klasyfikacja przekładni mechanicznych

Przekładnie zębate – w których ruch jest przekazywany za pomocą koła zębatego. W zależności od kształtu zębów, wyróżnia się przekładnie kolistozębne, stożkowe, łukowe oraz zęby skośne. Przekładnie różnią się ze względu na:

Liczbę stopni:

- przekładnia jednostopniowa – w której współpracuje jedna para kół zębatych;
- przekładnia wielostopniowa np. dwustopniowa, trzystopniowa itd. (przykład b) – w której szeregowo pracuje więcej par kół zębatych; przełożenie całkowite przekładni wielostopniowej jest iloczynem przełożeń poszczególnych stopni.

Umiejscowienie zazębienia:

- zazębienie zewnętrzne;
- zazębienie wewnętrzne.

Rodzaj przenoszonego ruchu:

- przekładnia obrotowa – uczestniczą w niej dwa koła zębate;
- przekładnia liniowa – koło zębate współpracuje z listwą zębatą tzw. zębatką. Ruch obrotowy zamieniany jest w posuwisty lub na odwrót;

Wzajemne usytuowanie osi obrotu:

- osie równoległe przekładnia walcowa.
- osie przecinające się przekładnia stożkowa.
- osie wchrowate, w tym:
 - prostopadłe, np. przekładnia hipoidalna, przekładnia ślimakowa;
 - nieprostopadłe, np. przekładnia hiperboloidalna.

Przekładnie ślimakowe – w których ruch jest przekazywany za pomocą ślimaka i koła zębatego. Są stosowane w celu uzyskania dużych przełożeń i zapewniają wysoki moment obrotowy przy niskich prędkościach.

Przekładnie łańcuchowe – w których ruch jest przekazywany za pomocą łańcucha. Obecnie są stosowane głównie w motoryzacji oraz rowerach.

Przekładnie obiegowe/planetarne – w których ruch jest przekazywany za pomocą zestawu kół zębatych umieszczonych na wspólnej osi. Pozwalają na uzyskanie dużych przełożeń przy niewielkich rozmiarach urządzenia.

Przekładnie śrubowe – w których ruch jest przekazywany za pomocą śruby i nakrętki. Są stosowane głównie w celu przekazywania ruchu liniowego.



Przekładnie hydrokinetyczne – w których ruch jest przekazywany za pomocą płynu. Są stosowane w celu uzyskania płynnego i ciągłego przekazywania momentu obrotowego.

Przekładnie cięgnowe – w których fizyczny kontakt pomiędzy członem napędzającym i napędzanym odbywa się za pośrednictwem cięgna. Dzięki temu człony przekładni mogą być oddalone od siebie na duże odległości nawet do 15 metrów. Pozwala to także na zastosowanie bardziej swobodnej geometrii przekładni. Przekładnie cięgnowe charakteryzują się także możliwością przenoszenia dużych mocy (do 1500 kW w przekładniach pasowych do 3500 kW w przekładniach łańcuchowych). Przekładnie cięgnowe dzielą się na:

- przekładnie pasowe – w których ruch jest przekazywany za pomocą pasów napędowych. Mogą być stosowane do przekazywania ruchu między osiami prostopadłymi;
- przekładnie linowe – w których cięgnem jest lina, znajdujące zastosowanie w przypadkach, gdy moc przenoszona jest na większą odległość (od kilku do kilkunastu metrów), przy dużych obciążeniach i stosunkowo niskich prędkościach;
- przekładnie łańcuchowe – w których cięgnem jest łańcuch, przy czym stosuje się łańcuchy drabinkowe i pierścieniowe.

Przekładnie cierne – w których dwa poruszające się elementy (najczęściej obracające się) dociskane są do siebie, tak by powstało pomiędzy nimi połączenie cierne. Siła tarcia powstająca pomiędzy elementami odpowiedzialna jest za przeniesienie napędu. ■

M.U.



Konkurs „AS Majsterkowania”

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

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

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

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

Regulamin, zastrzeżenia RODO oraz prezentacje wybranych nagród i mecenasów konkursu przewidujemy opublikować na stronie „Młodego Technika” wkrótce. ■





Corocznie w maju odbywa się w Monachium duża impreza – wystawa/targi sprzętu high-end, czyli „wyjątkowych” urządzeń audio. Co kryje się za eufemizmem „wyjątkowych”? Prościej byłoby napisać „najdroższych”, ale zdarzają się tam i dość przystępne. Czy „najlepszych”? To jeszcze mniej pewne... High-end to specyficzna działka, na której miesza się snobizm, ambicje, luksus, fantazja, egzotyka i najbardziej zaawansowana technika, genialne pomysły i szkolne błędy. Pasja i czysto biznesowa kalkulacja.

Monachium High-End 2023

TUZIN głośnych POMYSŁÓW

Pełną relację z tegorocznego high-endu, reportaż z setką zdjęć, ukazał się w AUDIO 6/2023. Dla MT przygotowaliśmy esencję poświęconą zespołom głośnikowym, które zawsze były, są i będą najbardziej spektakularnymi obiektami każdej tego typu wystawy – zarówno przez swoją wielkość, różnorodność, jak i niespodzianki. Ze względu na swoją mechaniczno-akustyczno-dizajnerską naturę są najbardziej fascynujące dla projektantów i dla klientów (zresztą każdy

projektant wyrósł z klienta). A obecnie do gry – dosłownie i metaforycznie – wróciły wszystkie dawne koncepcje, rywalizując z nowymi technikami albo łącząc się z nimi, wybór jest przeogromny, i nawet jeżeli ceny nie pozwalają stać się właścicielem tych „wynałazków”, to każdy może się nimi inspirować i ruszyć małymi krokami... w swoją własną drogę do high-endu.

Oto dwanaście wyselekcjonowanych głośnikowych „okazów”, które warto pokazać i opisać.

Penaudio to firma z Finlandii, stąd nazwa flagowego modelu – Karelia. Od frontu wygląda poważnie i konwencjonalnie – układ d'Appolito (symetryczny) z czterema 25-cm niskotonowymi, parą 18-cm średniotonowych i wysokotonowym w centrum. Symetryczna konfiguracja ma szczególne charakterystyki kierunkowe, zmniejszające energię poza osią główną (w płaszczyźnie pionowej), w takim samym stopniu do góry i do dołu. Konstruktor dodał kolejne, mniej konwencjonalne zabiegi koncentrujące promieniowanie w kierunku słuchacza, a więc zmniejszające odbicia. Otóż głośniki średniotonowe pracują w systemie półotwartym (otwory z wytłumieniem na froncie i po bokach – obszary zaznaczone maskownicą), co kształtuje charakterystykę kardiodalną, a niskotonowe – w specjalnej wersji odgrody otwartej, w kształcie „U” (bez tylnej ścianki), która w wyższym podzakresie tworzy kardiodę, a poniżej 50 Hz płynnie zbliża się do charakterystyki dipola.



Thorens, znany głównie z gramofonów, nie jest nowicjuszem w dziedzinie zespołów głośnikowych, chociaż po reaktywacji (kilka lat temu) jeszcze nam o tych swoich kompetencjach nie przypominał. SoundWall HP600 nawiązuje do konstrukcji sprzed 40 lat. To otwarta odgroda (dipol), w dodatku o tyle nietypowa, że zamiast kilku bardzo dużych głośników (dipol wymaga dużej powierzchni membran), zastosowano aż 12 mniejszych – 15-cm; dwa średniotonowe są również 15-cm i tworzą z wysokotonowym lokalny układ symetryczny, jednak aby symetria była kompletna i wielowymiarowa, drugi głośnik wysokotonowy, promieniujący w przeciwnej fazie, jest ustawiony za frontowym i skierowany do tyłu. Uwagę zwraca też lekko zagięta ku tyłowi skrajna część frontu (z czterema niskotonowymi). Otwarte odgrody generują ósemkową (dipolową) charakterystykę kierunkową, a przez to poprawiają relację między natężeniem promieniowania docierającego do słuchacza bezpośrednio od głośników w stosunku do fal odbitych, zmniejszając tym samym wpływ rezonansów pomieszczenia. Dzieje się to jednak kosztem efektywności (znaczna część energii niskich częstotliwości zostaje wygaszona w „akustycznym zwarcu” przeciwnych faz promieniowania przedniej i tylnej strony membran), co wymaga przygotowania dużego jej „zapasu”, stąd duża powierzchnia membran niskotonowych. Producent nie obiecuje więc basu infrasonicznego, ale usłyszymy wszystko do 40 Hz, a główną zaletą będą czystość i precyzja.



Utopia w ścianie – ciekawe zagadnienie akustyczne i biznesowe. Oczywiście z powodu braku obudowy regularnej kolumny (wielkiej i luksusowej) cena jest znacznie niższa. Ale czy wierzyć, że dzięki montażowi w ścianie można też uzyskać lepszy dźwięk? Tak, przy odpowiednim strojeniu całego układu... Duża powierzchnia zwiększa efektywność w zakresie niskich częstotliwości (usuwa zjawisko „baffle step”, powstające w obudowach o szerokości kilkunastu–kilkudziesięciu centymetrów, gdy dłuższe fale częściowo „uciekają” do tyłu), pozbawiona krawędzi w pobliżu głośników usuwa powstające na nich odbicia i interferencje. Układ głośnikowy ściennej Utopii nie został przeniesiony z jakiegokolwiek innego modelu – jest zupełnie inny. Poniżej i powyżej berylowej kopułki wysokotonowej znajdują się 7,5-cm przetworniki średniotonowe (jakich nie ma w „normalnych” Utopiach), a dopiero dalej – 18-tki, pełniące tutaj funkcję niskotonowych, po dwa poniżej i powyżej. Sekcja subniskotonowa to oddzielny moduł z trzema 18-tkami, ale innego typu (poznamy je po bardzo dużych nakładkach przeciwpływowych), na zdjęciu moduł taki dodano tylko na dole, ale można też na górze, po bokach... ile się chce.



Legendarne Western Electric 12B są już zwyczajowo prezentowane w wielkim pokoju wynajmowanym przez koreańską firmę Silbatone, produkującą całkiem nowoczesną elektronikę (chociaż z dużym udziałem lamp we wzmacniaczach), ale także... repliki kolumn Western Electric 12A i 13A. Na pokazie był jednak oryginał, zabytek liczący już prawie 100 lat, z 1927 roku, przygotowany wówczas na zamówienie Warner Brothers dla pierwszych filmów dźwiękowych.

Nowość szwedzkiej firmy Marten – Mingus Septet. Układ czterodrożny z zestawem przetworników reprezentujących różne najbardziej zaawansowane materiały i struktury membran, dopasowane do poszczególnych zakresów częstotliwości. Niskie tony przetwarza para 20-cm z membranami aluminiowo-sandwiczowymi (najniższe wspomaga para 25-cm membran biernych, umieszczonych z tyłu), od 200 do 800 Hz pracuje 18-cm z membraną ceramiczną, do 6 kHz – 7,5-cm kopułka berylowa, powyżej – 25-mm kopułka diamentowa. Informacja o stosowaniu filtrów 1. rzędu jest bardzo intrygująca w kontekście takich materiałów, wymagających raczej ostrego filtrowania, ale możliwe jest też połączenie filtrów 1. rzędu z filtrami pułapkami.



Aurora – konstrukcja cypryjskiej firmy Aries Cerat – może wyglądać na pierwszy rzut oka na designerskie szaleństwo, ignorujące uznane zasady konstruowania zespołów głośnikowych, stawiające wszystko na jedną kartę oryginalności. To jednak nie tylko oryginalna, ale też rzetelna, gruntownie

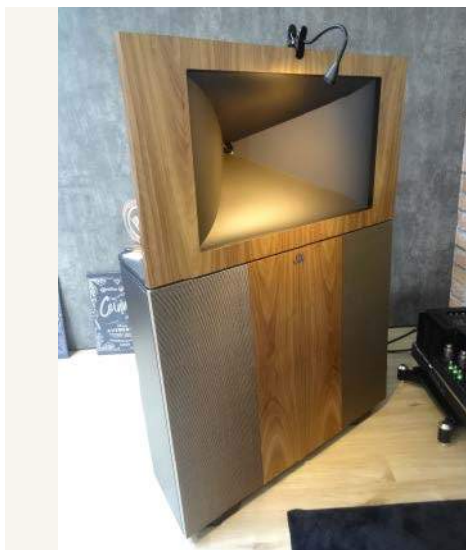
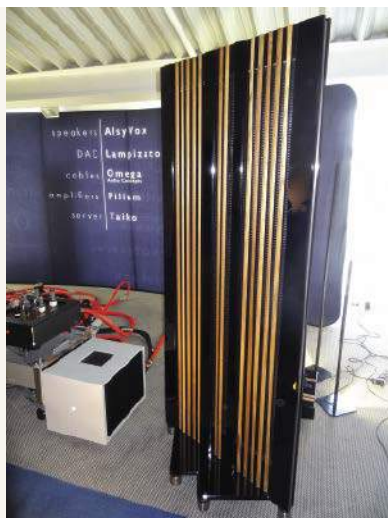
prześlana konstrukcja. Projekt ten realizuje kilka ważnych celów akustycznych, tych bardziej i mniej znanych.

Okrągła tuba głośnika średniotonowego przechodzi w sześciąt, który jest obudową czterech głośników niskotonowych, umieszczonych na jej bocznych ściankach. Samo połączenie okrągłej tuby z kwadratowym przekrojem „skrzynki” wymagało specjalnych „turbiny” wyprofilowań (aby uniknąć kumulowania się rezonansów na stałym promieniu), co nadaje konstrukcji tak dynamiczny wygląd. Niskotonowe pracują we wspólnej komorze, otwartej z tyłu (jest to więc odmiana otwartej odgrrody), jednak na skutek mniejszej powierzchni wylotu od tężycznej powierzchni membran, w obudowie zachodzi też sprężanie, które ma w odpowiedni sposób modyfikować energię uwalnianą do tyłu. Między magnesami niskotonowych znajduje się zasadnicza część głośnika średniotonowego, dzięki czemu odległość od wszystkich przetworników do słuchacza jest wyrównana; do tego warunku dopasowano też położenie wstęgowego wysokotonowego, przymocowanego na zewnątrz, na jednej z płaszczyzn bocznych.



Backes & Müller to legenda niemieckiego high-endu, jednak przez pewien czas pozostawała nieco w tyle za zmianami w designie, trzymając się dawnego, ciężkiego stylu. Współczesne projekty też są potężne, ale nabrały lekkości i dynamiki... Również dzięki nowocześniejszej technice i zastosowaniu tub. B&M wprowadza je na swój sposób, z rozmachem. Tubowy profil nadaje całemu obudowom, sięgającym od podłogi do sufitu, ustawiając w osi pionowej (a więc we wlocie tuby) sekcję średnio-wysokotonową, a głośniki niskotonowe wykorzystują objętość utworzoną w „skrzydłach” tuby (znajdują się z tyłu). To opis konstrukcji stojącej dalej, ta z przodu nie ma obudowy tubowej. Są to systemy całkowicie aktywne, B&M nie cofa się w rozwoju.

Hastem włoskiej firmy Alsy Vox jest „Diverti & Diversi”, czyli „ciesz się różnością”. Dewiza ta pasowałaby do całej imprezy, a nie tylko do projektów Alsy Vox, która jest przekonana o oryginalności swoich produktów – konstrukcji z przetwornikami wstęgowymi. Panele Alsy Vox wyglądają wyjątkowo pięknie. Zaprezentowano nowy model Rafaello, największy z „niepodzielnych” (jeszcze większą instalację możemy złożyć z panelu Caravaggio i dodanego subwoofera – też wstęgowego, o tej samej wielkości, ustawionego obok w jednej płaszczyźnie). Specyfiką Rafaello jest zdublowanie sekcji niskotonowych, a także węższych wstęg średnio-wysokotonowych, biegnących po obydwu stronach centralnej wstęgi superwysokotonowej. Układ staje się więc symetryczny w trzech płaszczyznach, osiąga też wysoką moc i efektywność – według danych producenta aż 98 dB.



Przez wiele lat referencyjnymi Klipschami były Klipschorny. Założyciel firmy pracował jednak nad czymś jeszcze lepszym... ale dzieła nie dokończył, a potem nastały czasy dla tub mniej sprzyjające. Jednak znowu jest zapotrzebowanie na takie wyczyny, co pokazuje aktywność na tym polu wielu innych firm, więc Klipsch postanowił przypomnieć, kto na tubach zna się najlepiej – wrócił do dawnych założeń, dołożył najnowsze rozwiązania i przygotował Jubilee. Znamienne, że w odróżnieniu nie tylko do Klipschorna, ale też pozostałych modeli serii Heritage (wywodzących się z danych konstrukcji) konstrukcja ta nie jest trójdrożna, lecz dwudrożna, co było możliwe dzięki zaawansowanym metodom projektowania profili tubowych i możliwości przygotowania tuby średnio-wysokotonowej.



Andrew Jones, wcześniej konstruktor KEF-a, TAD-a i Elaca, rok temu dołączył do zespołu amerykańskiej firmy MoFi. Najpierw przygotował SourcePoint10, a niedawno mniejsze SourcePoint8 – oczywiście uzbrajając je w to, na czym się doskonale zna – w moduły koncentryczne. Taki dwudrożny moduł, jako jedyny w całej konstrukcji, tworzy punktowe źródło dźwięku mające swoje „naturalne” zalety, ale też problemy do rozwiązania. W tym przypadku relatywnie duże sekcje nisko-średniotonowe, wraz z umieszczonym centralnie wysokotonowym, zostały zgrane perfekcyjnie, charakterystyki są wyrównane i stabilne w szerokim kącie. Daleko od cenowego kosmosu „prawdziwego” high-endu, ale blisko doskonałego brzmienia.

Firma Waterfall od początku zajmuje się kolumnami „przezroczystymi”. Obudowa jest wykonana ze szkła, ale oprócz zagadnień mechanicznych do rozwiązania jest też problem akustyczny – cały urok opiera się na tym, że w środku jest „pusta” (nie ma wytłumienia); specjalny kołnierz bezpośrednio na głośnikach (nisko-średniotonowych) tworzy odpowiedni filtr akustyczny, a w dolnej ścianie założono membranę bierną (która jest mniej podatna na przenoszenie rezonansów obudowy niż „zwykły” tunel bas-refleks). Na zdjęciu najlepszy model serii – Niagara – wyposażony w tubowy przetwornik wysokotonowy w oddzielnym module.



PMC – brytyjski specjalista od linii transmisyjnych – zdradza ich tajemnice. W małej konstrukcji podstawkowej zwija labirynt o długości ponad metra, co jest konieczne, aby złapać zgodną fazę przy niskich częstotliwościach, jednak dla efektywnego działania znaczenie ma również przekrój kanału, a ten nie może być w takiej sytuacji duży. W konstrukcji wolno stojącej jest więcej miejsca, podobny układ dwudrożny ma do dyspozycji kanał niewiele dłuższy, za to o większym przekroju i z dodatkową „ślepą uliczką” (idącą w górę do połowy wysokości; wylot jest na dole), która wytłumia antyrezonans. No to do roboty!

Andrzej Kisiel

*** Pisownia oryginalna ***



PRZEGLĄD TECHNICZNY 6-tonnowy samochód parowy

Technika samochodowa dochodzi do coraz lepszych wyników w budowie samochodów parowych, znanych zresztą już od dość dawna. Kocioł samochodu (...) składa się z zewnętrznej części cylindrycznej, w której się mieści cylindryczna również skrzynia ogniowa; ta ostatnia może być wyjęta po rozkręceniu dwóch połączeń. Opatkimi w kotle tego samochodu są stromo pochylone – w przeciwnieństwie do dawniejszego typu, w którym leżały one prawie poziomo. Rury te są rozmieszczone w trzech spiralnych rzędach. Duży przegrzewacz podnosi temperaturę pary o 150° ponad temperaturę nasycenia. Tylna oś wozu posiada kształt cylindryczny i niesie na swych końcach dwa koła biegowe, o nadzwyczaj długich łożyskach ślizgowych. Koła dają się łatwo zdejmować. Do wewnętrznej strony kół przymocowany jest duży średnicy bęben hamulcowy, zawierający dwa wewnętrzne hamulce szczękowe. Jeden z hamulców działa od pedatu, drugi zaś – od koła ręcznego, dogodnie umieszczonego w budce kierowcy, obok dźwigni do ruchu wstecz. Mechanizm kierowniczy działa w kąpieli oliwnej. Drażki kierownicze – zarówno podłużny, jak i poprzeczny – zaopatrzone są w przeguby łańcukowe ze sprężynami. Szczególnie ciekawy jest nowy ustrój dyferencjatu, który mieści się całkowicie wewnątrz karteru silnika i zbudowany jest, jako część samego wału korbowego. Dyferencjat działa dopiero wówczas, gdy siła pociągowa na tylnych kołach samochodu będzie się różniła o 10 do 15%. Jest to środkiem zaradczym przeciwko bocznemu poślizgowi wozu. Nowy typ dyferen-

cjatu jest mniejszy i lżejszy od typu zwykłego, i dzięki swemu położeniu, osłonięty jest od wszelkiego brudu i kurzu. Wały rozrządcze umieszczone są wewnątrz karteru silnika, co powoduje dokładne oliwienie i lepszą ochronę mechanizmu. Wały rozrządcze w liczbie dwóch zostały zastosowane po raz pierwszy w tym typie samochodu, zamiast jednego tylko wału, jak poprzednio; to ulepszenie pozwala używać garbów szerszych i lepiej ukształtowanych. Zawory nowego silnika wykonane są ze stali nierdzewiejącej i działają za pośrednictwem długich popychaczy, dających się regulować. Trzony tłokowe przechodzą przez podwójne dławnice, co zapobiega przedostawianiu się wody kondensacyjnej z cylindrów do karteru. Pompa zasilająca napędzana jest z szybkością równą mniej więcej połowie szybkości silnika. Samochód ten budowany jest przy pomocy sprawdzianów, tak, że wszystkie części dają się zamieniać.

12 czerwca 1923

Wagony skrzynkowe

W ostatnich latach coraz częściej jest wysuwana myśl, że, wobec złego stopnia wyzskania wagonów ciężarowych w ruchu, należałoby wprowadzić zasadę, że części toczne wagonu (podwozie) powinny być oddzielone od części ładunkowej (nadwozia) i, będąc przeznaczane do ruchu, stałe w ruchu się znajdować. Postój bowiem wagonów wywołują głównie długie okresy załadowywania i wyładowywania. Tę zaś ujemną stronę łatwo można zmniejszyć ogromnie przez wprowadzenie zdejmowanego nadwozia. Jako jeden z typów takiego nadwozia, rozpowszechnił się obecnie w Ameryce ustrój nadwozia skrzynkowego (wprawdzonego od r. 1920 na kol. New-York Central i in.). Zastosowanie takiego wagonu skrzynkowego (container car) okazuje się ogromnie korzystnym wszędzie, gdzie sortowanie ładunków musi być posunięte do jednostek mniejszych niż zawartość wagonu, ale może być zatrzymane na jednostkach większych niż jedna paczka. Wagon skrzynkowy składa się (...) z pomostu na 4-ch osiach, zaopatrzonego w ramę, występującą ponad

poziom podłogi (pomostu). Wewnątrz tej ramy ustawia się na pomoście skrzynie zamknięte z blachy żelaznej, zaopatrzone w drzwi z jednej strony. Skrzynie posiadają po bokach występy, opierające się o odpowiednie wysoki prowadnicze na zewnętrznej stronie ramy. Rama występuje ponad poziom podłogi o 60 cm. Po bokach między skrzynią a ramą pozostaje luz 1/2 cala. W nowszych wagonach skrzynie są ustawiane już nie na samej podłodze, lecz na kratowanym pomoście żelaznym, umieszczonym na podłodze i zabezpieczającym od zbierania się wody przy dnie skrzyni. Ponieważ skrzynie są ustawiane drzwiami wpoprzek wagonu jedna koto drugiej, więc otworzenie drzwi podczas przewozu jest zupełnie niemożliwe. Wobec tego ładunki są zabezpieczone od kradzieży i uszkodzenia. Żelazne skrzynie chronią też je od pożaru i deszczu, wzgl. śniegu. Na stacji odbiorczej skrzynie są zdejmowane z pomostu zapomocą mechanizmów wyładunkowych i ustawiane tuż na pomosty samochodów ciężarowych (...). Do zdejmowania służą 4 ucha, umieszczone na rogach skrzyń u góry. Towarowe wagony mają po 6 skrzyń ściśle jednakowych wymiarów, a zatem zamiennych. Ładowność skrzyni wynosi 3150 kg. Wagony skrzynkowe są budowane 2-ch typów: do przewożenia zwykłych ładunków ciężarowych i do pośpiesznych oraz pocztowych. Szczególnie wielkie korzyści osiągnięto przy zastosowaniu wagonów skrzynkowych do ładunków pocztowych. Za ideał przewozu uważa się bezpośrednie dostarczenie paczki od drzwi nadawcy do drzwi odbiorcy. Opisane wagony właśnie zbliżają się najwięcej do tego ideału. Paczki, umieszczone przez nadawcę w skrzyni, bezpośrednio i szybko razem ze skrzynią przestawiają się dźwigiem z samochodu na pomost wagonu, a po przybyciu na miejsce – z wagonu na samochód – i do odbiorcy. Unikamy tu zatem: 1) konieczności starannego opakowywania i przewożenia paczek w workach. 2) Strat w drodze (wskutek kradzieży, pożarów i t. p.), które w ostatnich latach ogromnie wzrosły w St. Zj., tak, że Urząd Pocztowy i kolej wy-

placają ogromne odszkodowania (w r. 1914 kolej zapłaciła odszkodowań 33 miliony dolarów, w r. 1920 – 125 milj. dol.). A do tego dochodzi ważniejsza bodaj i wspomniane już na wstępie. Natomiast osiagamy. 3) Zmniejszenie ilości operacji przy ekspedjowaniu i przeładowywaniu (sprawdzanie nadawcy i odbiorcy) oraz usunięcie składania na stacjach. 4) Zmniejszenie do minimum postoju wagonów, przy wyładowywaniu i załadowywaniu, a zwiększenie obiegu wagonów, oraz 5) wzrost ładowności wagonów (zarówno co do wagi, jak objętości), szczególnie pocztowych i bagażowych, która dochodzi do 16 650 kg, a może wzrosnąć do 221 1/2 tonn. Ważną jest również budowa skrzyń takich wymiarów, że mogą być one zastosowane do przewozów zarówno kolejami parowymi, jak elektrycznymi, statkami i samochodami. Wagony pocztowe (nowsze) posiadają 8 skrzyń. Pierwsze próby wykazały, że załadowanie i wyładowanie skrzyń (zapomocą zwykłego żorawia) trwa 18 min. Obecnie, po wprowadzeniu ulepszonych wyładunków, czynność ta zabiera zaledwie 8 min. czasu, czyli około 1/10 tego, co zużywało zwykle dotychczas (1–2 godz.). Wymiary tych skrzyń są: długość – 2,7 m, szerokość – 3,5 m, wysokość – 2,9 m; drzwi zaś: 1,7 m × 1 m. Waga skrzyni wynosi – 1350 kg, ładowność 3150 kg, udźwig wagonu – 36 tonn. Można przypuszczać, że wagony skrzynkowe będą się bardzo nadawały do przewozu żywności, a szczególnie mleka, oraz takich towarów, jak jedwab, który dotąd z obawy zepsucia, kradzieży i t. p. oraz wobec wysokich stawek ubezpieczeniowych, przewożono wyłącznie samochodami. Należy, oczywiście, zwrócić uwagę na przymus szybkiego wyładowywania skrzyń samych, by z nich nie robiono sobie nowych składów. Wprowadzenie takich wagonów wymaga pewnych kosztów nakładów, lecz jeśli się zważy wartość zaoszczędzanego przy nich wciąż czasu na przeładunki i wyładunki, to okaże się, niewątpliwie, że przyniosą one niewątpliwie większe zyski, niż kosztu urządzeń wyładunkowych.

19 czerwca 1923

NOWOŚĆ

WYDANIE SPECJALNE DIGITAL CAMERA POLSKA

SONY

PORADNIK UŻYTKOWNIKA EDYCJA 1 Digital Camera



**Wykorzystaj
w pełni
możliwości
swojego
aparatu**

+

- 12 opisów** aparatów
- 33 polecane** obiektywy
- 140 stron** praktycznych porad

DIGITAL CAMERA POLSKA
WYDANIE SPECJALNE
38 zł (w tym 9 zł VAT)
Nakład: 14 500 egz.

AVT

9 772344 811111

**DARMOWA
PRZESYŁKA**

Jeśli jesteś posiadaczem aparatu Sony lub rozważasz jego zakup,
to ta publikacja jest właśnie dla Ciebie!

Niezależnie od tego, czy z fotografią wiążesz swoją zawodową
przyszłość, czy stanowi ona dla Ciebie jedynie odskocznię
od codziennych obowiązków, nasze Wydanie Specjalne
wprowadzi Cię w ten fascynujący świat i pozwoli Ci
przenieść swoje umiejętności na wyższy poziom.

Przejrzyj i zamów na www.UlubionyKiosk.pl

eprasa.pl 0e17570rd7