



Tu przejrzysz
i kupisz ten numer

NEWS 24/7
przełóżaj codziennie
na swoim smartfonie

mlody **m.technik**

Ciekawi świata są zawsze młodzi



TAJEMNICA ZIEMI

Nasza obca planeta

RAPORT: Wielki Kosmiczny Teleskop Jamesa Webba
Start następcy Hubble'a i wielkie nadzieje

ISSN 0462-9760 Indeks 365408



9 770462 976212 12 >
cena: **11,90 zł** (w tym 8% VAT)



Active Reader

Zapraszamy do udziału w nieustającym konkursie **Active Reader**.

Nagrody rozdajemy **codziennie**.

Zapamiętaj!

Uczestnik **Active Reader** zbiera punkty na swoim koncie i w każdej chwili może „zapłacić” swoimi punktami za nagrody wybrane z listy publikowanej na:

www.mlodytechnik.pl/active-reader-nagrody

Wybrane nagrody wysyłamy wraz z najbliższą przesyłką prenumerat. Zbierasz punkty na koncie osobistym i w każdej chwili możesz sobie „kupić” za te punkty dowolne nagrody (wycenione w punktach). Wysyłka nagród i aktualizacja stanu dorobku punktowego na Twoim koncie odbywa się raz w miesiącu, podczas wysyłki prenumerat.

Stan swojego konta możesz sprawdzać na stronie:

www.mlodytechnik.pl/active-reader-ranking

Tylko Prenumeratorzy „Młodego Technika” mogą brać udział w Konkursie **Active Reader**.

Zbieraj punkty i zgarniaj nagrody

Do konkursu **Active Reader** można przystąpić w każdej chwili, wysyłając e-mail na adres: **activerreader@mt.com.pl** o treści: „Zgłaszam swój udział w konkursie Active Reader. Jestem prenumeratorem „Młodego Technika”. Mój numer prenumerat...”

TYLKO PRENUMERATORZY „Młodego Technika” mogą brać udział w konkursie **ACTIVE READER**.

Punkty otrzymuje się za różne formy aktywności:

Listy 30 pkt. za każdy opublikowany w „Młodym Techniku” list/wpis z facebookowego fanpage’a MT.

Pomysły 30 pkt. za każdy pomysł opublikowany w „Młodym Techniku”, w rubryce „Pomysły genialne, zwiariowane i takie sobie”.

Konkurs futurystyczny 30 pkt. za ciekawą wizję futurystyczną opublikowaną w „Młodym Techniku”, w rubryce „Pomysły genialne, zwiariowane i takie sobie”.

Na warsztacie 100 pkt. za wykonanie modelu wg projektu publikowanego w rubryce „Na warsztacie” i przesłanie jego zdjęć na e-mail: **activerreader@mt.com.pl**. Przypominamy, że projekty można wysłać maksymalnie do **trzeciego numeru wstecz!**

Klub/Szkoła Wynalazców N x 10 pkt. liczba punktów N uzyskanych w Rankingu Klubu Wynalazców lub Rankingu Szkoły Wynalazców pomnożona razy 10.

Facebook 30 pkt. za wpis merytorycznie istotny dla „Młodego Technika”, opublikowany w wydaniu drukowanym (w rubryce Listy).

MiniQuiz 10 pkt. za każdą poprawną odpowiedź przesłaną na e-mail: **activerreader@mt.com.pl**

Chemia 20 pkt. za zdjęcia i krótki opis przeprowadzonych doświadczeń chemicznych i przesłanie na e-mail: **activerreader@mt.com.pl**

Temat numeru, temat artykułu 50-100 pkt.

Zapraszamy do wspólnego kształtowania planu tematycznego kolejnych wydań MT. Zgłaszajcie na adres: **redakcja@mt.com.pl** propozycje tematów artykułów, które chcielibyście przeczytać w MT, w szczególności zagadnienia, które nadają się na temat numeru, opracowany w postaci zbioru artykułów. Jeśli w ciągu jednego roku od Twojego zgłoszenia w „Młodym Techniku” pojawi się artykuł lub temat numeru zgodny z Twoją propozycją, to otrzymasz punkty w AR:

1. **temat numeru** – 100 pkt.

2. **artykuł** – 50 pkt.

Do zgłaszanych tematów należy dołączyć krótkie objaśnienie (do 140 znaków), co powinien zawierać proponowany przez Ciebie artykuł.

Inne X pkt. Udział w konkursach nieregularnych, ogłaszanych *ad hoc* w poszczególnych numerach ma wycenę punktową, określaną indywidualnie dla każdego konkursu.

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

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

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

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

• **Redaktor Wydania:**
Wojciech Marciniak

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

• **Konsultacja graficzna:**
Małgorzata Jabłońska

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

• **Kontakt z redakcją:**
e-mail: mt@mt.com.pl
<http://www.mlodytechnik.pl>
<http://facebook.com/magazynMlodyTechnik>

• **Prenumerata w Wydawnictwie AVT**
www.ulubionykiosk.pl
tel. 22 257 84 22
e-mail: prenumerata@avt.pl

• **Prenumerata w RUCH S.A.**
www.prenumerata.ruch.com.pl
lub tel. 801 800 803, 22 717 58 59
e-mail: prenumerata@ruch.com.pl

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



Temat okładkowy

Mamy tyle map. I tworzymy wciąż nowe. Mapy powierzchni lądu, oceanów, wnętrza naszej planety. Wszystko jest dzięki rozwojowi techniki coraz dokładniejsze i wszechstronne. Jednocześnie nie przestaje kołatać w nas myśl, że nasza wiedza jest wciąż nader skromna.

Nasz dom pełen tajemnic

Powiedzieć o Ziemi, że jest to „obca planeta”, jeśli mamy w głowie porównania z innymi planetami z Układu Słonecznego, lub egzoplanetami krążącymi wokół odległych gwiazd, to zapewne przesada. Niemniej, choć planeta nasza jest znacznie mniej nieznaną niż dalekie światy, to jednak mówienie o Ziemi „obca” nie jest pozbawione sensu.

Przez wypowiedzi i publikacje, do których odwołujemy się w tym numerze „Młodego Technika” przewija się znany motyw, w różnych wersjach ale sprowadzający się do takiej oto tezy, że lepiej zbadaliśmy i więcej wiemy o niejednym zjawisku i obiekcie we Wszechświecie niż o wielu miejscach na planecie Ziemia, np. o głębiach oceanicznych lub o wnętrzu Ziemi. Na Księżycu stopę postawiło kilkunastu ludzi – do najgłębszego miejsca ziemskich oceanów dotarło trzech, przy czym o „stawianiu stopy” nie było tam mowy. We wnętrzu Ziemi udało nam się wgłębić na ułamek ułamka jej promienia i tym co znajduje się w środku naszej planety wnioskujemy raczej na podstawie pośrednich danych niż cokolwiek widzieliśmy, dotknęliśmy i eksplorowaliśmy.

Mówiąc o nieznanach planetach, mówimy też o Ziemi

Dzięki rozwijającym się technicznym narzędziom, wysyłanym w kosmos ale również np. w głębie oceanów, coraz dokładniej i wszechstronniej mapujemy powierzchnię Ziemi. Ta jej część, która jest pokryta wodą

wciąż jest znana niezbyt szczegółowo. Można m.in. spotkać opinię, że wcale nie jest jeszcze przesądzone ostatecznie, że to dno Rowu Mariańskiego jest najgłębszym punktem, że tak naprawdę należy dodawać mówiąc o tym miejscu – „z tego co wiemy”, co podkreśla wciąż mocno ograniczony zasób ludzkiej wiedzy na ten temat.

Także wiedza o początkach naszej planety i jej najstarszej historii, choć wydaje nam się, że poukładaliśmy sobie to w dość spójne modele, nieustannie wstrząsana jest przez odkrycia, które, jak to się zwykle w raportach medialnych pisze, „zmuszają do zweryfikowania istniejących teorii”. Od niedawna wiemy np., że formy życia pojawiły się na Ziemi bardzo wcześnie, zaskakująco szybko po tym jak uformowała się gorąca kula materii tworząca pierwotną Ziemię. Czy to znaczy, że życie jest importem z zewnątrz? Dobre pytanie, na które nie ma równie dobrej odpowiedzi.

W tym kontekście mówienie o Ziemi „obca planeta” nabierałby jeszcze więcej sensu. Tak czy inaczej nauka ma jeszcze wiele do zrobienia i zrozumienia, jeśli chodzi o tę błękitną kulę, która nazywamy naszym domem.

Mirosław Usidus

DO
50%

**TANIEJ
W PRENUMERACIE
DLA SZKÓŁ
I PLACÓWEK
OŚWIATOWYCH!**

ROCZNA PRENUMERATA
DRUKOWANA W PROMOCJI
DLA SZKÓŁ I PLACÓWEK
OŚWIATOWYCH KOSZTUJE
99,90 ZŁ, ROCZNY DOSTĘP
ONLINE – 57,00 ZŁ.

SZCZEGÓŁY NA
WWW.ULUBIONYKIOSK.PL/
PRENUMERATA/SZKOLNA

PRENUMERATA – TO SIĘ OPŁACA!
SZCZEGÓŁY NA STR. 62

STAŁY KONKURS

Active Reader

Supernagrody!

Szczegóły na stronie 2

KSIAŻKI
GRY
PŁYTY
MODELE

NARZĘDZIA
SPRZĘT
AKCESORIA



Tajemnice Ziemi. Jak mało wiemy o jej początkach, historii formowania się, kształtowania, wpływach z zewnątrz i o tym jak powstało życie. Mimo że sięgamy gwiazd, to, co leży bezpośrednio pod naszymi stopami, we wnętrzu Ziemi jest nadal wiedzą tajemną. Podobnie z wiedzą o głębinach oceanicznych. Jak się okazuje błękitny glob to dla nas wciąż nieznana, obca, można powiedzieć, planeta.



Spis treści

Temat numeru: Tajemnica Ziemi:

Nasza obca planeta

- 26 • Kartografia, jakiej nigdy dotąd nie mieliśmy. W centrum każdej mapy jest człowiek
- 32 • Seria bombardowań w sprzyjających warunkach. Nieznana historia planety Ziemia
- 39 • Podróż do wnętrza Ziemi. Nieznany świat pod skorupą
- 45 • Przygoda w otchłani dopiero się zaczyna. W głębinach

Technika

- 8 Info Zoom
- 18 Dodaj do obserwowanych
- Horizonty mgłą spowite
- 17 • Czy kosmiczne górnictwo ruszy wreszcie na poważnie? Czekając na księżycowe zagłębie
- 20 • Dlaczego jest tak wiele zapowiedzi budowy reaktorów termojądrowych? Fuzja obietnic
- 23 • Naukowe zabawy światłem. Tajemnice kwantowej walizki
- 50 Raport MT: Wielki Teleskop Jamesa Webba. Kosmiczne oko na rzeczy dotąd nie widziane
- 63 Nasi idole – liderzy innowacji: Nie oszukuj, bo wynajdę coś, co ci się nie spodoba – Almon Brown Strowger

m.technik

- 66 e-Technologie: Krzywdą Niebiańskiego Emu. Starlink – nowe oblicze internetu czy utrapienie?

Szkola

- 69 Matematyka za ludzką twarzą: Klocki, monety, kasztany i Encyklopedia Ciągów Liczbowych
- 73 Koniec i co dalej: Hasło. Pożegnanie z 123456?
- 76 Chemia inna niż w szkole: Na cześć złośliwych duchów (2)
- 80 Edukacja przez szachy: Włodzimierz Schmidt – legenda polskich szachów
- 84 MT studiuje: Akustyka
- Klub i Szkoła Wynalazców
- 86 • Szkoła Wynalazców, dozwolone do lat 15
- 87 • Klub Wynalazców, bez ograniczeń wieku
- 88 • Vademecum Młodego Wynalazcy
- 91 Pomysły genialne, zwariowane i takie sobie
- Na warsztacie
- 92 • Elektronika dla Ciebie: Klaskacz 12 V na jedno lub dwa kłaśnięcia
- 94 • Efektowna skarbonka
- Odkryj historię wynalazków
- 100 • Motocykle
- 104 • Rodzaje motocykli

Hobby

- 106 Fotografia: Kreatywne fotografowanie
- 108 Akademia audio: Nowoczesne soundbary (2)
- 2 Konkurs: Active Reader
- 3 Od wydawcy
- 6 Listy, Facebook
- 62 Prenumerata
- 105 Sędziwy Technik – 100 lat temu prasa pisała

25

Nasza obca planeta

Start następcy
Hubble'a i wielkie
nadzieje

50

List miesiąca

nagroda: 30 punktów AR

Szczegóły na stronie 2

Przyszłość języków programowania

Zainteresowałam mnie artykuł na temat języków i platform programistycznych przyszłości, który ukazał się w październikowym wydaniu „Młodego Technika”. Jest to temat ciekawy dla mnie do tego stopnia, że postanowiłem dodać swoje trzy grosze.

Spędzam dużo czasu na identyfikowaniu i analizowaniu rosnących trendów, aby nadążyć za popytem w branży. W mojej ocenie AI to nie przejściowa moda. Sztuczna inteligencja będzie nieuchronnie wkraczać we wszystko, co robimy. Jednym z obszarów, którym jestem szczególnie zafascynowany, są rozwiązania, które można by obrazowo przedstawić tak: wprowadzasz swój luźno zdefiniowany pomysł na projekt do narzędzia, które tworzy na jego podstawie interfejs użytkownika, makietę, ekrany.

Co jest potrzebne? Czego należy uczyć się, aby wziąć udział w rozwoju AI? Wymieniłbym takie języki jak Python, R, Lisp, Prolog i Java.

Kolejna wielka rzecz to rzeczywistość rozszerzona. Koszty rozwiązań AR spadają, co jest ogólnie znakiem, że jesteśmy blisko szerszej adopcji. Prawdopodobnie będziemy świadkami coraz częstszego przejmowania AR przez urządzenia mobilne, ponieważ te dwa elementy w naturalny sposób współpracują ze sobą. AR staje się bardziej popularne przed VR. Zarówno Apple, jak i Google wydały swoje własne platformy programistyczne dla AR. Google rozwija ARCore framework z Java, jako głównym językiem programowania. Apple oferuje ARKit Framework z językami Swift lub C obiektywnym.

Kolejny interesujący dla szukających swoich szans w programowaniu obszar to wirtualna rzeczywistość. Chociaż nie widzimy jeszcze wielu projektów VR, jest to bez wątpienia fascynujący obszar. Jeśli technologia ta przyjmie się wśród szerszej publiczności, może wprowadzić szalenie innowacyjne zmiany do naszego codziennego życia. Przewidywanie takich zmian jest oczywiście bardzo trudne. Podobnie jak w przypadku AR, nie możemy sobie wyobrazić, dokąd nas to zaprowadzi, nie wiemy też, czy i jak dojrzeje.

Czego powinna się nauczyć osoba zainteresowana programowaniem dla VR? Na pewno jest to JavaScript i Java a także języki C++ i C#.

Kolejny obiecujący obszar to internet rzeczy, IoT. Od niedawna możemy mówić o eksplozji tego rodzaju rozwiązań. Rozwój sieci inteligentnych urządzeń był wolniejszy niż oczekiwano z powodu problemów z komercjalizacją danych IoT. Rozwój jest też silnie uzależniony od nowych technik telekomunikacyjnych, przede wszystkim upowszechniania 5G.

Z punktu widzenia umiejętności programistycznych dla osoby myślącej o tworzeniu rozwiązań dla IoT, na pewno cenna będzie znajomość języków Python i JavaScript.

Kuszącym dla adeptów programowania jest również świat rozwiązań typu Blockchain, łańcuchów blokowych, kryptowalut i bezpiecznych protokołów transakcyjnych

Czego warto nauczyć, aby rozwijać się na Blockchain? Znowu przede wszystkim Pythona, a także C++, JavaScript a także czegoś takiego jak zorientowany obiektowo język programowania do pisania inteligentnych kontraktów – Solidity.

Jeśli chodzi o tajniki zarządzania wielkimi zasobami danych (big data) to znajomość Pythona i JavaScript na pewno będzie cenna. Jest zapotrzebowanie na programistów biegłych w obsłudze opartej na Java platformy Hadoop.

O programowaniu komputerów kwantowych nie wspominam bo to zupełnie inna „zabawa”. Niewątpliwie jednak w dłuższej perspektywie umiejętności tego rodzaju mogą okazać się bardzo cenne.

Są jeszcze inne zjawiska i tendencje w świecie programowania, które należy brać pod uwagę, myśląc o karierze w tej dziedzinie. Należy do nich np. postępująca redukcja redundancji. Jest programowanie po stronie klienta i jest programowanie po stronie serwera. To rozdzielenie jest kosztowne i prowadzi do wielu redundancji, czyli marnotrawienia pracy programistów. Obecnie dąży się do połączenia tych dwu sfer, co programiści myślący o specjalizacji, powinni brać pod uwagę. Jest całkiem realne, że w niedalekiej przyszłości, gdy zbudujesz aplikację mobilną lub webową, automatycznie wygeneruje ona dla siebie niezbędną stronę serwerową.

Powstają wciąż nowe języki programowania i będą powstawać. Jednocześnie do zaprogramowania aplikacji

będzie potrzebne coraz mniej linii kodu i wiedzy. Penetracja rynku programistycznego przez nie-programistów będzie rostała, a płace niektórych specjalistów będą spadały. Należy więc dobrze zastanowić się nad tym, co chcemy robić.

W dłuższej perspektywie, gdy dostępna będzie większa moc obliczeniowa, będziemy mówić raczej o szkoleniu komputerów niż o ich programowaniu. Programowanie może być ograniczone do pisania gier lub sterowników urządzeń. Z drugiej strony rośnie rola architektów sieci neuronowych AI, być może jednak nie na zawsze.

Niektórzy uważają, że gwoździem do trumny programowania będą obliczenia kwantowe. Ponieważ wtedy rzeczywistością staną się komputery samouczące się i kreatywne.

Ostatecznie wszystko to co nazywamy obecnie programowaniem, żmudne klecenie kodu itd. odejdzie do la-musa. To nie znaczy, że wiedza i umiejętności programistów nie będą już potrzebne. Będą ale z pewnością nie będzie chodziło o znajomość kodu. Nacisk będzie kładziony na umiejętności tworzenia architektury, projektowania rozwiązań biznesowych, kreatywnych. Tu doświadczeni programiści mogą wciąż wygrywać, jeśli właściwie zrozumieją kierunek zmian.

Eryk Szymanderski, Szczecin

Problemy z procesorami do samochodów

Opisane przez „Młodego Technika” problemy z półprzewodnikami i chipami to dość złożone kłębowisko. W przypadku układów do samochodów kwestie te są starsze niż pandemia. Już lata temu producenci krzemowi przestawili się na bardziej rentowne produkty i niechętnie produkują tańsze układy. W sytuacji spiętrzenia zamówień problem się tylko pogłębił.

Tak czy inaczej, do producentów półprzewodników ustawiła się długa kolejka. Przemysł półprzewodnikowy tak naprawdę nigdy nie dogonił popytu.

Producenci samochodów przewidywali spadek sprzedaży i anulowali zamówienia (wyjątkiem była Toyota, która przyjęła nieco inną politykę). Tymczasem pandemia spowodowała wzrost popytu na elektronikę użytkową (powodując braki w prawie wszystkich produktach, ponieważ producenci mieli kłopoty z nadżeniem za zamówieniami). Sprzedaż samochodów wzrosła wcześniej niż oczekiwano. Kiedy producenci samochodów zaczęli składać zamówienia, znaleźli się na końcu i tak już dłuższych niż wcześniej kolejek.

Epidemia Covid nie była przyczyną spowodował niedoborów. Po prostu wyzwoliła łańcuch zjawisk, które były prawdziwą przyczyną. Podstawowa struktura systemu i łańcuch dostaw był i bez pandemii słaby. Przewróciłoby się to przy pierwszej dużej fali, która się pojawiła a akurat pojawił się Covid.

Teraz o różnicy pomiędzy Toyota a innymi firmami. Japoński producent rozpoznał dostawy



półprzewodników jako krytyczny element podatny na zakłócenia już po trzęsieniu ziemi w Japonii w 2011 r. Od tego czasu firma dbała o to, by posiadać wystarczające zapasy, tak, aby zawsze pokryć braki z powodu opóźnień w realizacji zamówień. Znana metoda produkcyjna Toyoty (system produkcji just in time) polega na eliminowaniu nadmiernych zapasów, ale nie wszystkich zapasów. Należy utrzymywać rezerwy niezbędne do tego, aby utrzymać fabrykę w ruchu. Jest to coś, co firmy naśladujące bardzo udany system zarządzania Toyoty całkowicie zignorowały, czytając tylko pierwszą stronę opisu systemu Toyot, ignorując zaś resztę literatury na ten temat.

Zatem Toyota przygotowywała się na kryzys od dekady. I gdy przyszedł, miała znacznie mniejsze problemy, niż inne firmy samochodowe. Kryzys w zaopatrzeniu w krzem nie przeminie łatwo i szybko. A w branży motoryzacyjnej łączy się z innymi przełomowymi zmianami, co może skończyć się dla tego sektora całkowitym przemebłowaniem, ale to już trochę inna historia.

Stanisław Żuk, Wambierzyce

Od Redakcji

Autorów opublikowanych listów, którzy są prenumeratorami MT, nagradzamy płytami z najwyższej półki. Mamy ponad 100 tytułów wspaniałych albumów muzycznych.

Prosimy Autorów listów, aby z zestawu „Płyty z najwyższej półki”, publikowanej w każdym wydaniu miesięcznika „Audio”, wybrali płytę dla siebie i napisali do redakcji (e-mail: redakcja@mt.com.pl) list zawierający: tytuł wybranej płyty (Autor Listu miesiąca ma prawo do nagrody w postaci 3 płyt wybranych z ww. listy); numer prenumeratora MT.

Wybraną płytę wyślemy wraz z przesyłką najbliższego numeru MT.





FIZYKA

Ile jest informacji we Wszechświecie?

Sześć z osiemdziesięcioma zerami – to liczba bitów informacji, jakie zawiera cały dający się zaobserwować Wszechświat – ogłosił fizyk Melvin Vopson z Uniwersytetu w Portsmouth w Wielkiej Brytanii po przeprowadzeniu oszacowań ilości informacji, których nośnikiem jest pojedyncza cząstka elementarna. Po wyliczeniu tego było już z górki, bo wystarczyło pomnożyć szacunkową wartość przez liczbę cząstek w znanym nam Uniwersum.

Wykorzystując teorię informacji Claude'a Shannona, Vopson oszacował, że każda cząstka zawiera 1,509 bitów informacji. Teoria ta wiąże entropię, wielkość niepewności w systemie, z informacją: Wartość informacyjna wiadomości jest miarą tego, jak bardzo niepewność jest zredukowana przez tę wiadomość. Różne rodzaje wiadomości mają różne wartości. Aby otrzymać szacunkowe dane na temat tego, jak wiele informacji przechowują, uczony zastosował obliczenia entropii do masy, ładunku i spinu protonów, neutronów (i tworzących je kwarków) oraz elektronów.

Choć gigantyczny, to jednak wynik mnożenia tych oszacowań dał Vopsonowi wynik niższy niż wcześniejsze szacunki, które często brały pod uwagę zakładane rozmiary całego Wszechświata (w rzeczywistości nieznane). Brytyjski fizyk wykluczył z rachunku też niektóre cząstki, np. antycząstki czy bozony przenoszące oddziaływania, przyjmując zasadę, że mierzy wyłącznie to co daje się obserwować i jest stabilne. Praca na temat jego obliczeń ukazała się w czasopiśmie „AIP Advances”.



FIZYKA

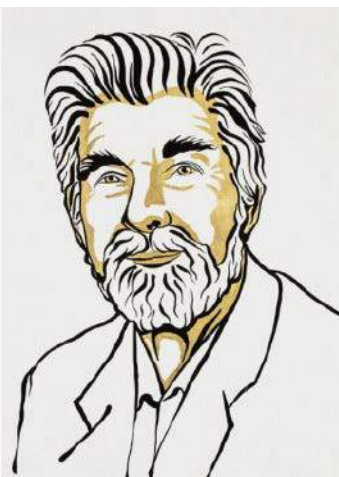
Tornado z atomów

Pierwsze w historii atomowo-cząsteczkowe tornado udało się stworzyć fizykom z Uniwersytetu Kalifornijskiego w Berkeley. Dokonali tego, wysyłając prostą wiązkę atomów helu przez kratownicę z małymi szczelinami o rozmiarach ok. 600 nanometrów, i wykorzystując zasady mechaniki kwantowej do wprowadzanie wiązki w stan wirowania. Wiązka atomów rozprzaskłała się w siatkę, wyginając się tak bardzo, że tworzyła wir, który zakrzywiał się w przestrzeni.

Wirujące atomy docierały następnie do detektora, który wykazywał istnienie wielu wiązek, rozproszonych w różnym stopniu, z różnymi momentami pędu. Naukowcy zauważyli również mniejsze, jaśniejsze pierścienie zaklinowane wewnątrz trzech zauważonych centralnych wirów. Uznali to za ślady tzw. ekscymerów helu, cząsteczek powstających, gdy jeden energetycznie wzbudzony atom helu przykleja się do innego.

Badacze mają teraz w planach zderzanie wiązek atomowych heli z fotonami, elektronami i atomami innych pierwiastków, aby zobaczyć co z tego wyniknie i jak to wpłynie na powstające „tornado”. Jeśli eksperymentach tych wiry zachowywać się w inny sposób, to zdaniem Yaira Segeva, fizyka z Berkeley, który uczestniczył w badaniach, otwierają się możliwości konstruowania nowego typu mikroskopów do obserwacji efektów i cząstek subatomowych.

124 słowa ma „najprostszy język świata”, Toki Pona, zaprojektowany przez kanadyjską tłumaczkę i esperantystkę Sonję Lang na początku XXI wieku.



NAGRODY NOBLA

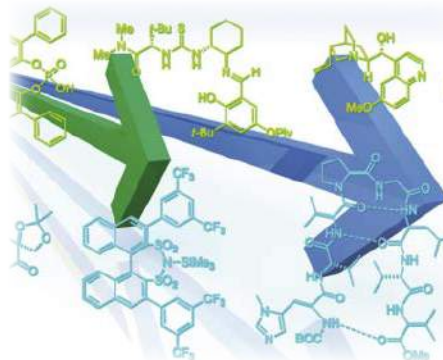
Fizycy nagrodzeni za modele klimatu a chemicy za organokatalizę

Fizyczny Nobel trafił w tym roku do trzech uczonych – Klausa Hasselmann, Syukuro Manabe, i Giorgio Parisiego. Zostali oni przez Komitet Noblowski uhonorowani za badania nad fizycznym modelowaniem klimatu, pomiarami jego zmienności w zależności od poziomu emisji gazów cieplarnianych i oszacowania dotyczące poziomu globalnego ocieplenia.

Zależność pomiędzy zawartością dwutlenku węgla w atmosferze ziemskiej a wzrostem temperatury przy powierzchni Ziemi badał Syukuro Manabe, pracujący na Uniwersytecie Princeton meteorolog i klimatolog japońskiego pochodzenia. Był pionierem stosowania komputerów do przeprowadzania symulacji klimatu Ziemi. Fizycznym Noblem podzielił się z Klausem Hasselmannem, niemieckim oceanografem i klimatologiem, autorem modelu zmienności klimatycznej nazwanej „modelem Hasselmanna”. Giorgio Parisi, jedyny definicyjny fizyk w gronie tegorocznych laureatów, otrzymał nagrodę za powiązanie z modelami klimatycznymi „odkrycie wzajemnego oddziaływania

nieuporządkowania i fluktuacji w systemach fizycznych od skali atomowej do planetarnej”.

Nagroda Nobla w dziedzinie chemii przypadała w tym roku Benjaminowi Listowi i Davidowi MacMillanowi za prace nad rozwojem nowego rodzaju procesów organokatalitycznych, które służą do syntezy cząsteczek organicznych, bez konieczności sięgania po drogie i toksyczne nierzadko dla środowiska katalizatory metaliczne. List jest niemieckim uczonym, który pracuje na Uniwersytecie w Kolonii, współpracując również z Instytutem Maksa Plancka. Opracowane przez niego procesy organokatalizy, wykorzystującej związki węgla, wodoru, siarki i innych substancji niemetalicznych wykorzystywane były od lat w badaniach akademickich i przemysłowych na całym świecie. David William Cross MacMillan to brytyjski chemik współpracujący obecnie m.in. z Uniwersytetem Princeton. Opracowane przez jego zespół techniki organokatalizy asymetrycznej znalazły zastosowanie w wytwarzaniu różnego rodzaju produktów naturalnych.





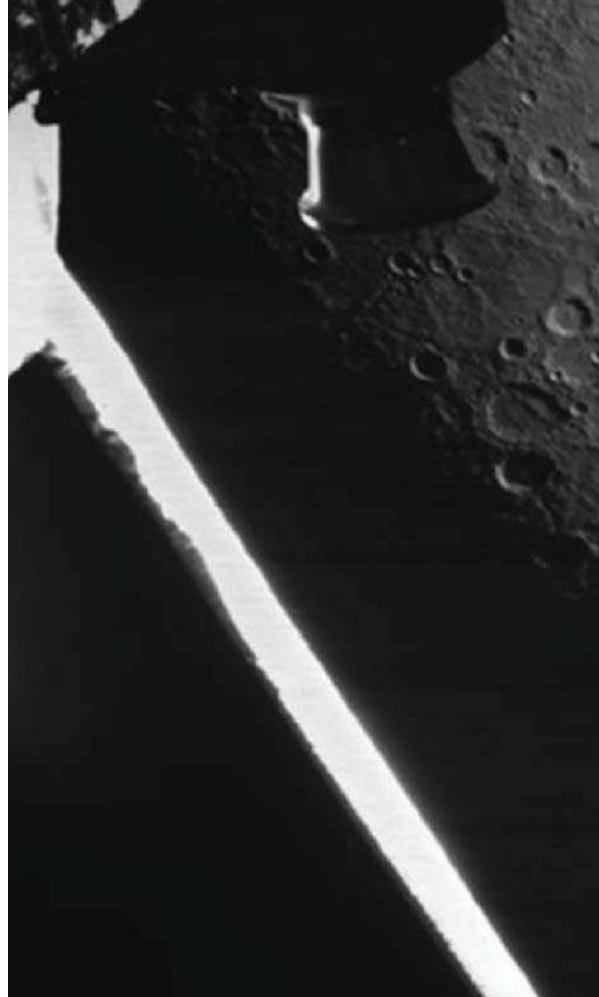
SIEĆ

Połowa świata używa mobilnego internetu

Organizacja GSMA w cyklicznie publikowanym raporcie „State of the Internet Connectivity Report” podała, że ponad połowa światowej populacji korzysta obecnie z mobilnego Internetu. Stanowi to znaczący wzrost w porównaniu z jedną trzecią populacji z dostępem do tej sieci w 2014 roku. Obecnie jedynie sześć proc. ludzkości nie ma dostępu do żadnej sieci komórkowej. Sześć lat temu była to ok. jedna czwarta mieszkańców Ziemi.

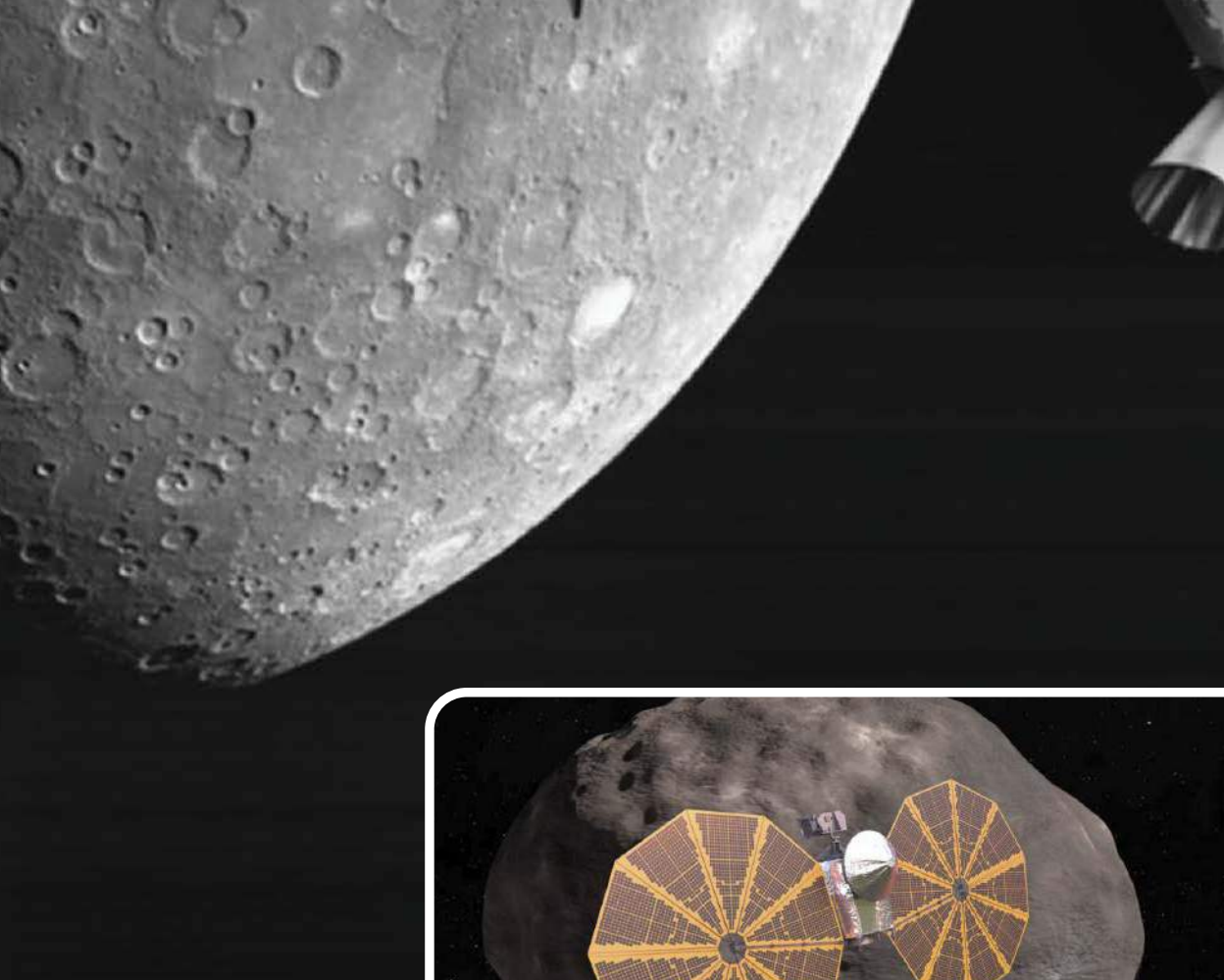
Zachodzi, co widać w podanych odsetkach, rozbieżność pomiędzy poziomem dostępu do sieci a rzeczywistym korzystaniem z niej. Spośród 3,8 mln osób, które nie korzystają z mobilnych łączów szerokopasmowych, tylko 450 mln mieszka na obszarach, na których nie ma żadnej sieci komórkowej. Jest to tłumaczone w raporcie brakiem dostatecznej świadomości istnienia takich usług, umiejętności cyfrowych a także zbyt wysoką dla pewnych grup ich ceną.

93 proc. ludności świata pozbawionej połączenia z siecią i ponad 98 proc. ludności pozostającej bez zasięgu mieszka w krajach o niskim i średnim poziomie dochodów. Dla GSMA jest to szczególnie niepokojące, ponieważ z powodu braku stacjonarnej infrastruktury telefony komórkowe są podstawowym środkiem komunikacji dla wielu ludzi w krajach rozwijających się.



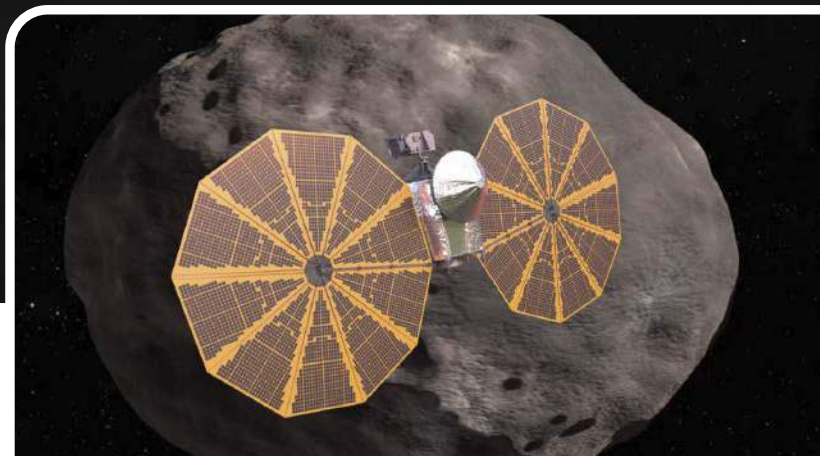
Do Merkurego po prawie trzech latach lotu dotarła sonda Europejskiej Agencji Kosmicznej i Japońskiej Agencji Eksploracji Aerokosmicznej (JAXA) nazwana – BepiColombo. Statek zbliżył się do planety na odległość momentami mniejszą nawet niż 200 kilometrów. Aparatura pokładowa wykonała zdjęcia powierzchni najbliższej Słońcu planety Układu Słonecznego. Bliski przełot nad Merkurym nie oznacza jeszcze, że sonda (a właściwie dwie sondy, bo w ramach misji wysłano połączone urządzenia – europejskie i japońskie) zajęła cykliczną orbitę wokół tej małej planety. To jeden z szeregu manewrów przygotowujących wejście na orbitę Merkurego do 2025 roku.

Z pierwszymi zdjęciami Merkurego wykonanymi przez BepiColombo zbiegł się w czasie start Lucy, sonda NASA, której celem jest orbita Jowisza a głównym zadaniem – badanie pozostałości po wczesnym okresie istnienia Układu Słonecznego. Niestety już na samym początku misja napotkała problemy techniczne, gdyż jeden z dwóch kolistych paneli słonecznych, choć



EKSPLORACJA KOSMOSU

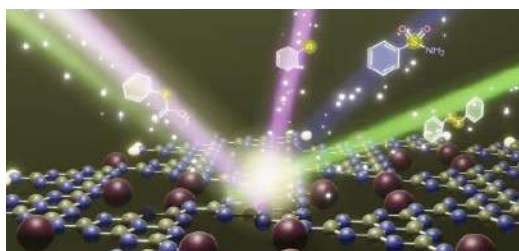
Bliskie spotkanie z Merkurym i misja do Trojan



rozwinął się po starcie zgodnie z oczekiwaniami, to nie zamknął się poprawnie. Warto jednak zaznaczyć, że oba panele działają.

W ciągu dwunastu lat trwania misji Lucy ma eksplorować świat planetoid trojańskich, jak nazywa się dwie grupy obiektów krążących wokół Słońca po orbitach zbliżonych do orbity Jowisza. Ponieważ orbita, po której krążą jest silnie stabilizowana przez interakcję grawitacyjną Słońca z Jowiszem, uczeni sądzą,

że ciała te są tam „zamknięte” od bardzo dawna, będąc nienaruszonymi „skamielinami” wczesnego Układu Słonecznego z czasów formowania się planet. Według planu sonda Lucy ma spotkać się z główną grupą obiektów trojańskich na przełomie 2027/28 roku. Ma po dotarciu na miejsce wykorzystać swoje bogate oprzyrządowanie do badania obiektów gabarytów kilku kilometrów lub większych, szczegółowo analizując ich kształty, strukturę, cechy powierzchni, skład i temperaturę.



CHEMIA

Reakcje chemiczne sterowane barwami światła

Badaczom z niemieckiego Instytutu Maxa Plancka udało się opracować cechujący się trwałością „inteligentny fotokatalizator”, który potrafi rozróżniać kolory światła (niebieski, czerwony i zielony) i w odpowiedzi umożliwia zaprogramowaną w nim specyficzną reakcję chemiczną. „Nasz inteligentny fotokatalizator działa jak drogowskaz, który otwiera jedną konkretną ścieżkę w odpowiedzi na światło o określonym kolorze”, wyjaśnia w publikacji Jewgienia Markuszyna, główna autorka opracowania na ten temat w periodyku pt. „Angewandte Chemie International Edition”.

Fotokatalizatory to materiały, które wykorzystują energię światła słonecznego lub światła LED do wywołania pożądanej reakcji chemicznej. Ten opracowany w niemieckim instytucie pozwolił m.in. na syntezę sulfonamidów, związków stosowanych jako antybiotyki w leczeniu infekcji bakteryjnych, w sposób kontrolowany. „Szczególną cechą sposobu działania nowego fotokatalizatora jest to, że możemy kontrolować selektywność reakcji chemicznej, włączając żarówkę o odpowiednim kolorze”, wyjaśnia poglądowo dr Markuszyna.

Złożone obiekty biologiczne, takie jak ludzkie oko czy najnowocześniejsze kamery w urządzeniach elektronicznych, rozpoznają barwy światła. Fotokatalizator z Instytutu Plancka jest odpowiedzią na znacznie większe wyzwanie, jakim jest opracowanie znacznie mniejszych systemów tego rodzaju, „inteligentnych cząsteczek”, składających się z kilkudziesięciu atomów, które nie tylko rozróżniałyby kolory światła (niebieski, czerwony i zielony), ale także wykonywałyby pewną „zaprogramowaną” czynność, która zależy od określonej barwy światła.



WODA

Hydrożelowa tabletka czyści wodę skuteczniej niż znane metody

Litr wody rzecznej w ciągu godziny potrafi oczyścić i uzdatnić do picia, według publikacji, która ukazała się na łamach czasopisma „Advanced Materials”, tabletka hydrożelowa opracowana przez naukowców z Uniwersytetu Teksańskiego w Austin. Wrzucana po prostu do pojemnika z wodą pastylka potrafi zlikwidować 99,999 procenta drobnoustrojów.

Tabletki działają poprzez generowanie nadtlenu wodoru, który współpracuje z cząsteczkami węgla aktywnego, zabijając bakterie przez zakłócenia ich metabolizmu. Autorzy metody twierdzą, że w procesie tym nie powstają żadne szkodliwe produkty uboczne. Chociaż do tej pory przeprowadzono jedynie testy na małą skalę, nowe hydrożele powinny być również dość łatwe do używania w większych skalach. Materiały i procesy są niedrogie i proste, a pastylki można formować w najróżniejsze kształty i rozmiary, aby dopasować je do konkretnych zastosowań.

Oczyszczacze hydrożelowe mogłyby również znaleźć zastosowanie w doskonaleniu innych technik, np. destylacji słonecznej, która działa przez skupienie ciepła słonecznego do odparowania wody i zebrania jej w innym pojemniku, z zanieczyszczeniami pozostawionymi w odseparowanym zbiorniku. Sprzęt do takiej destylacji może być zatykany przez mikroby. Nowy hydrożel mógłby temu zaradzić. Badacze chcą rozszerzać zakres neutralizowanych przez substancję zanieczyszczeń.



CYBERWOJNA

Złośliwy kod DNA – broń hakerów przyszłości

Wyniki badań prezentowane na konferencji USENIX Security przez grupę naukowców z Uniwersytetu Waszyngtońskiego wykazują (a jest to pierwszy taki przypadek), że możliwe jest zakodowanie złośliwego oprogramowania w fizycznych łańcuchach DNA, co doprowadzić może do sytuacji, w której sekwencjonowanie genów w trakcie ich analizy przekształca powstałe dane w program uszkadzający system sekwencjonowania genów i przejmujący kontrolę nad komputerem.

„Wiemy, że jeśli przeciwnik ma kontrolę nad danymi, które komputer przetwarza, może potencjalnie przejąć kontrolę nad tym komputerem”, mówił Tadayoshi Kohno, profesor informatyki z Uniwersytetu w Waszyngtonie, który kierował projektem, porównując tę technikę do tradycyjnych ataków hakerów, którzy umieszczają złośliwy

kod na stronach internetowych lub w załącznikach poczty elektronicznej. Chociaż atak tego rodzaju jest jeszcze mało użyteczny dla prawdziwego szpiega lub przestępcy, to zdaniem badaczy może stać się z czasem bardziej prawdopodobny, ponieważ sekwencjonowanie DNA staje się coraz powszechniejsze, wydajne a ponadto realizowane przez zewnętrzne podmioty we wrażliwych systemach komputerowych.

Jeśli hakerom rzeczywiście udałoby się przeprowadzić tego typu atak, to zdaniem badaczy mogliby oni potencjalnie uzyskać dostęp do cennej własności intelektualnej, albo też zaszkodzić analizom genetycznym, np. kryminalistycznym testom DNA. Z drugiej strony np. firmy mogłyby umieszczać złośliwy kod w DNA genetycznie modyfikowanych produktów, traktując to jako sposób na ochronę tajemnic handlowych.



BIOTECHNOLOGIA

Wytwarzanie skrobi z dwutlenku węgla

Jak wynika z publikacji, która ukazała się w periodyku „Science”, chińscy naukowcy z powodzeniem przeprowadzili syntezę skrobi z dwutlenku węgla. Jeśli ich metoda potwierdzi swoją praktyczną wartość w większej skali, to może oznaczać rewolucję, zarówno w rolnictwie jak również w wysiłkach na rzecz osiągnięcia neutralności emisyjnej w gospodarce.

Ma Yanhe, dyrektor Instytutu Biologii Przemysłowej w Tianjin powiedział mediom, że w jego placówce udało się zaprojektować proces jedenastostopniowej reakcji chemicznej, będącej ścieżką wiązania dwutlenku węgla i syntezy skrobi. Proces, jak wyjaśnia chiński uczony, przypomina nieco warzenie piwa. CO_2 jest w nim redukowany do metanolu, który może być następnie przetwarzany na skrobię.

Roczna produkcja skrobi w bioreaktorze o objętości jednego metra sześciennego w nowym procesie stanowi ekwiwalent produkcji z jednej trzeciej hektara uprawy kukurydzy. Według chińskich badaczy, dzięki temu procesowi produkcji węglowodanów można zaoszczędzić nawet 90 proc. gruntów ornych i ogromne ilości wody, zasobów normalnie zużywanych w produkcji rolnej. Choć proces uprawy, np. kukurydzy, również wiąże się z pochłanianiem CO_2 w procesach fotosyntezy, to jednak efektywność naturalnego procesu jest szacowana na jedynie 2 proc. Naukowcy mają nadzieję, że ich metoda będzie znacznie bardziej wydajna.



GADŻETY

Widać przez ubranie, że masz wiadomość

PocketView to nowego typu inteligentny wyświetlacz LED skonstruowany przez specjalistów z Uniwersytetu Waterloo. Może funkcjonować jako samodzielny układ albo jako zintegrowany element istniejących już urządzeń i systemów lub też rzeczy, które dopiero powstaną. Z prezentacji wynika, że noszony pod tkaniną ubrania skutecznie przenika ja światłem wyświetlając proste komunikaty i obrazy.

System ma obecnie postać płaskiego urządzenia z układem diod LED o wysokiej intensywności na frontowej powierzchni, które mieści się w kieszeni ubrania użytkownika. Wyświetlacz tego typu może być wbudowany w równe urządzenia, np. trackery do ćwiczeń fizycznych, lub może funkcjonować jako łatwy do przeglądania, podłączony za pośrednictwem Bluetooth dodatkowy wyświetlacz dla gadżetu głównego, takiego jak smartfon.

Pomysł polega na tym, że gdy użytkownik potrzebuje powiadomień o przychodzących e-mailach lub wiadomościach tekstowych, PocketView wyświetla odpowiednie powiadomienia w formie symboli o niskiej rozdzielczości, które przeświecają przez ubranie. Po potwierdzeniu, powiadomienia te mogą być odrzucone poprzez dwukrotne dotknięcie wyświetlacza przez tkaninę. Dodatkowo, poprzez pojedyncze dotknięcie urządzenia, użytkownicy mogą przełączać się pomiędzy wyświetlaczami, które oferują dane takie jak czas, warunki pogodowe, statystyki fitness lub nawigacyjne strzałki w lewo/prawo, które wskazują, w którą stronę należy skręcić.



ROBOTY

Terminator na czterech nogach

Robot Vision 60, który zaprojektowała go i zbudowała amerykańska firma Ghost Robotics przy wodzi na myśl skojarzenia z cyklem filmowym „Terminator”. Poruszająca się na nogach, w podobny sposób do znanych „robotycznych psów” Boston Dynamics, jest wyposażona w układ broni podobny do karabinu – Special Purpose Unmanned Rifle (SPUR) – osadzony na kadłubie. Jest wyposażony w 30-krotny zoom optyczny, kamerę termowizyjną do celowania w ciemności. Ma efektywny zasięg do 1200 metrów.

Firma zaprezentowała również innego czteroźnego robota o nazwie Minitaur, którego przeznaczeniem jest patrolowanie obiektów wojskowych a także, jako opcja, wsparcie dla saperów lub strażaków. Roboty te mogą być sterowane zdalnie za pomocą aplikacji w smartfonie, albo w trybie autonomicznym, bez bieżącej ingerencji operatora.

W zeszłym roku 325 Dywizjon Sił Bezpieczeństwa w Bazie Sił Powietrznych Tyndall na Florydzie stał się pierwszą jednostką w amerykańskich siłach zbrojnych, która zaczęła wykorzystywać czworonożne roboty w regularnych operacjach. Zastosowano je m.in. do patrolowania granic bazy, poruszania się po bagnistych terenach itp. Choć zwiad jest jednym z najbardziej oczywistych zastosowań tego rodzaju robotów, producenci powoli eksperymentują z innymi funkcjami użytkowymi. Oprócz rejestracji zdalnego obrazu wideo i mapowania, maszyny te mogą być wykorzystywane jako mobilne wieże komórkowe, do rozbijania bomb lub wykrywania substancji chemicznych, biologicznych, radiologicznych i jądrowych.

319 000 000 000 000 000 bitów,
czyli 319 terabitów,
na sekundę wynosi najnowszy rekord szybkości transmisji danych
w światłowodzie, pobity przez japońskich inżynierów latem 2021 roku.



ROBOTY

◆ Singapurska Agencja Nauki i Technologii (HTSTA) ogłosiła start projektu rozmieszczania na ulicach miasta policyjnych robotów Xavier, zaprogramowanych do wykrywania drobnych wykroczeń, np. nieprawidłowo zaparkowanych rowerów, palenia papierosów w miejscach publicznych oraz wykrywanie naruszeń przepisów dotyczących COVID-19, w tym braku maseczki, co, jak się zauważa, zakłada wyposażenie „robocopów” w jakąś formę techniki rozpoznawania twarzy. ◆ Według „Forbesa” amerykańskie wojsko rozpoczęło testy pierwszej generacji zrobotyzowanych pojazdów bojowych, Robotic Combat Vehicles (RCV), zbudowanych przez QinetiQ North America i Pratt Miller, zaś główna faza ich rozwoju inżynieryjno-produkcyjnego planowana jest na połowę 2023 roku. ◆

TECHNIKA NUKLEARNA

◆ W rdzeniu reaktora MARIA w Świerku została uruchomiona z powodzeniem sonda ISTHAR (Irradiation System for High-Temperature Reactors) zbudowana i zaprojektowana w Departamencie Eksploatacji Obiektów Jądrowych polskiego NCBJ, pozwalająca na napromienianie próbek materiałowych w temperaturze 1000°C, czyli w warunkach jakie panują w reaktorach HTGR, koncepcyjnych, chłodzonych gazem reaktorach jądrowych IV generacji z moderatorem grafitowym i chłodziwem gazowym. ◆ Duński startup Seaborg Technologies poinformowała o planach i pozyskaniu funduszy na budowę małych pływających po morzu reaktorów jądrowych, opartych na technice mieszania paliwa nuklearnego z solami kwasu fluorowodorowego. ◆

ELEKTRONIKA

◆ Komisja Europejska zapowiedziała wprowadzenie przepisów mających na celu ustanowienie jednej wspólnej ładowarki do smartfonów oraz innego sprzętu elektronicznego w 27 krajach Unii Europejskiej. ◆ Zespół badaczy z Uniwersytetu Northwestern w Illinois skonstruował „latającego chipa” o rozmiarach ziarenka piasku czyli ok. milimetra, który, jeśli wierzyć doniesieniom mass mediów, jest najmniej-

szą latającą strukturą stworzoną przez człowieka. ◆

TECHNIKA ORBITALNA

◆ Jak podała agencja Xinhua, Chiny planują wystrzelić na niską orbitę okołozemską konstelację trzydziestu sześciu satelitów, których zadaniem ma być zbieranie danych o możliwych klęskach żywiołowych i próba ich przewidywania, przy czym odbywać się to ma przez wysyłanie na Ziemię obrazów o wysokiej rozdzielczości, gdyż system ma rejestrować geologiczne deformacje powierzchni Ziemi na poziomie milimetra na wczesnym etapie, co pomoże w przewidywaniu osunięć ziemi, osiadania gruntu i innych klęsk żywiołowych. ◆ Samuel Reid, szef firmy Geometric Energy Corporation, poinformował, że jego firma nawiązała współpracę SpaceX, której celem jest uruchomienie niewielkiego satelity reklamowego, wyposażonego w ekran do wyświetlania reklam wykupionych na specjalnych aukcjach oraz w „kij do selfie”, za pomocą którego można będzie pstrykać zdjęcia tych reklam z naszą planetą w tle, przy czym, wbrew obawom wyrażanym w mediach społecznościowych, rzecz cała będzie zupełnie niewidoczna z powierzchni Ziemi. ◆

BIONIKA

◆ Samsung i Uniwersytet Harvarda opublikowały w „Nature Electronics” pracę naukową z opisem techniki zapisu struktury i aktywności mózgu ludzkiego na chipie elektronicznym, np. mapy połączeń neuronowych w mózgu, która mogłaby być skopiowana do macierzy nanoelektrod. ◆ Zespół naukowców z włoskiej uczelni Sant’Anna opracował niewielkiego, dostarczającego insulinę, składającego się z dwu części – wewnętrznego, sterowanego bezprzewodowo (Bluetooth) dozownika insuliny wszczepianego chirurgicznie do organizmu pacjenta oraz z magnetycznej kapsułki wypełnionej insuliną, która jest połykana, niczym pigułka, przemieszcza się do urządzenia i instalowana jest za pomocą magnetycznego mechanizmu, zaś potem, gdy jest już pusta, kontynuuje podróż w dół układu pokarmowego, aż organizm ją wywali. ■ *M.U.*



1. Wizja osady górniczej NASA na Księżycu

Czy kosmiczne górnictwo ruszy wreszcie na poważnie?

Czekając na księżycowe zagłębienie

Pozaziemskie projekty wydobywcze znów są na tapecie. Tym razem brzmią poważniej niż pierwsza ich fala sprzed kilku lat. Nie tylko dlatego, że za rzecz bierze się NASA. Komplikacje techniczne, biznesowe i prawne z tym związane nie przestały być aktualne.

We wrześniu tego roku brytyjski magazyn „This Is Money” ujawnił szczegóły projektu mini-reaktora do wypraw kosmicznych, nad którym pracuje firma Rolls-Royce. Choć takie reaktory prawdopodobnie wkrótce już trafią na pokłady statków kosmicznych, to zwykle myśli się o nich jako o jednostkach służących do napędu. Jak wiadomo jednak, reaktor może działać długo. Dlaczego więc nie wykorzystać go jako źródła mocy, gdy misja kosmiczna doleci do celu? Jeśli myśli się

o wierceniach, wydobywaniu i transportowaniu urobku z powrotem, to sens wykorzystania tego rodzaju urządzeń jeszcze bardziej rośnie.

Rolls-Royce ma duże doświadczenie w opracowywaniu napędu atomowego do okrętów podwodnych dla Royal Navy. Ekspertki zwracają uwagę, że w istocie różnice pomiędzy reaktorami napędzającymi okręty morskie a hipotetycznymi kosmicznymi nie są wielkie. W obu środowiskach muszą funkcjonować w podobnych środowiskach (bez dostępu



2. Wydobycie na asteroidzie

do powietrza) i stawia się przed nimi te same wymagania (mają być niezawodne, generując dużo energii), a także, oczywiście służą do zasilania całości misji a nie jedynie do napędu.

Chłodzenie i tumany pyłu – w kosmosie te problemy nie są banalne

Być może więc reaktory Rolls-Royce zasilą kosmiczny przemysł wydobywczy. Najpierw jednak trzeba rozwiązać szereg innych problemów i ograniczeń. I tym od pewnego czasu intensywnie zaczęła zajmować się NASA. Pod szyldem księżycowego programu Artemis, NASA finansuje startupy dla misji noszącej nazwę – Sustained Lunar Exploration & Development. Zgodnie z jej założeniami tym razem plany NASA to już nie tylko nie tylko wysłanie człowieka tam i z powrotem na Księżyc, ale także założenie stałej bazy eksploracyjnej i jej rozwoju na Księżycu. Wymaga to czasu, nakładów i myślenia o tym, w jaki sposób sięgnąć po znajdujące się na miejscu, na powierzchni i pod powierzchnią Księżyca, zasoby naturalne, co zmniejszyłoby koszty transportu wszystkiego z Ziemi (1).

W grupie firm, którym NASA przyznała granty na badania jest Pioneer Astronautics, która podaje, że będzie w stanie zbierać księżycowy regolit i przetwarzać go w użyteczny tlen. Regolitem nazywa się osady gleby księżycowej, które powstały na przestrzeni miliardów lat w wyniku ciągłych uderzeń meteoroidów w powierzchnię

Księżyca. Naukowcy szacują, że grubość warstwy księżycowego regolitu sięga w niektórych miejscach 4-5 metrów pod powierzchnię, a w starszych obszarach górskich nawet do 15 metrów. Uwięzione są w nim m.in. gazy, w tym tlen. I o ten gaz na potrzeby misji i stałego pobytu ludzi na Księżycu chodzi.

Wydobywanie regolitu i pozyskiwanie zeń tlenu to dwa odrębne zadania. O ile przetwarzanie np. wody, o której wiemy, że jest jej pod postacią lodu sporo w glebie księżycowej, to problem głównie energetyczny i wspomniane reaktory przewiezione z Ziemi mogą tu bardzo się przydać, o tyle wykopywanie i transportowanie mas urobku to górnictwo, które oczywiście wymaga energii, ale w warunkach, w których nie ma takiego ciśnienia i atmosfery jak na Ziemi, napotyka nieznane na naszej planecie problemy, np. wiemy, że różnego rodzaju świdry czy inne narzędzia robocze nagrzewają się podczas pracy. W warunkach ziemskiej atmosfery chłodzi je do pewnego stopnia powietrze, choć można też użyć różnego rodzaju chłodziw. Bez atmosfery rozgrzane części metalowe nie mają do czego oddać ciepła. Można to rozwiązać częściowo zamkniętymi układami chłodzenia, ale muszą one być bardziej skomplikowane. I w końcu tak czy inaczej ostrza i nabieraki muszą mieć zewnętrzny kontakt z materiałem wydobywanym. Innym problemem, dość oczywistym przy pracach wydobywczych przy zastosowaniu dużej mocy i nakładu energii jest wyrzucanie kopanego materiału w warunkach braku grawitacji. Chmury gruzu

i pyłu nie opadają tak łatwo i szybko na Księżycu, na asteroidach zaś – właściwie wcale (2).

Oznacza to, że wszelka kosmiczna górnictwa działalność będzie musiała stosować dodatkowe systemy zabezpieczania terenów wydobywania, znaleźć sposób na chłodzenie i konserwację sprzętu, w warunkach, które z górnictwem ziemskim nie mają wiele wspólnego.

Wydobywać wolno, ale – czy urobek może być czyjąkolwiek własnością?

Są też problemy prawne, nad którymi wprawdzie zwykle w dyskusjach biznesowych przechodzi się do porządku dziennego, ale nie można zamykać oczu na to, że „Traktat o przestrzeni kosmicznej” z 1967 oraz inne umowy międzynarodowe zawarte później, w tym Układ Księżycowy z 1979 i konwencji o międzynarodowej odpowiedzialności za szkody wyrządzone przez obiekty kosmiczne z 1972 r., właściwie w ogóle zasadniczo nie zezwalają na prywatną czy państwową działalność górnictwa w celu osiągnięcia partykularnych korzyści (3).

W aktach tych, których jest jeszcze kilka znajdujemy m.in. takie stwierdzenia, że przestrzeń kosmiczna, łącznie z Księżycem i innymi ciałami niebieskimi, nie podlega zawłaszczeniu przez państwa ani poprzez ogłoszenie suwerenności, ani w drodze użytkowania lub okupacji, ani w jakiegokolwiek inny sposób. W odniesieniu do górnictwa kosmicznego szczególne znaczenie ma art. 11 Układu Księżycowego, zgodnie z którym Księżyc i jego zasoby są wspólnym dziedzictwem ludzkości. Na mocy tego traktatu ustanowiono zakaz pozyskiwania zasobów mineralnych znajdujących się na Księżycu i innych ciałach niebieskich, przy

czym zakaz ten powinien obowiązywać do czasu stworzenia przez wspólnotę międzynarodową reżimu regulującego eksploatację i wykorzystanie tych zasobów.

Te międzynarodowe konwencje są jednak dość bezceremonialnie traktowane przez państwa, które widzą w górnictwie kosmicznym szansę na realne zysku. Dobitnie o tym świadczą uchwalone w ostatnich latach akty prawne w USA i w Luksemburgu. W 2015 r. prezydent USA, Barack Obama, podpisał ustawę o zapewnieniu konkurencji w obszarze usług wynoszenia obiektów w przestrzeń kosmiczną. Przepisy powstałe w Luksemburgu mają podobny charakter. Ustawa amerykańska, a następnie Luksemburska, jako sprzeczna z postanowieniami międzynarodowego prawa kosmicznego, w tym w szczególności zasadą traktowania przestrzeni kosmicznej jako wspólnego dziedzictwa ludzkości i oczywiście są szeroko krytykowane, zwłaszcza przez Rosję i Chiny. Głównym powodem kontrowersji stało się prawo własności do surowców mineralnych wydobytych na Księżycu i innych ciałach niebieskich. Sama dopuszczalność prowadzenia działalności wydobywczej nie budzi bowiem kontrowersji. Podstawową zasadą uregulowaną w Układzie kosmicznym jest zasada wolności badania i użytkowania przestrzeni kosmicznej, łącznie z Księżycem i innymi ciałami niebieskimi, przez wszystkie państwa bez jakiegokolwiek dyskryminacji, na zasadzie równości i zgodnie z prawem międzynarodowym. Dostęp do wszystkich obszarów ciał niebieskich jest zatem wolny, pod warunkiem że badanie i użytkowanie przestrzeni kosmicznej jest prowadzone lub wykonywane dla dobra i w interesie wszystkich państw, niezależnie

od stopnia ich rozwoju gospodarczego czy naukowego.

Sprawa jednak nie jest taka oczywista, gdyż np. stany Zjednoczone nie podpisały Układu Księżycowego a z Traktatu z 1967 roku zakaz prywatnego posiadania wydobywanych bogactw jasno nie wynika. Stąd też dające się czasem zauważyć lekceważenie przepisów prawa międzynarodowego, jeśli chodzi o kosmiczne górnictwo. ■

Mirosław Usidus

3. Koparka w kosmosie





1. Plac budowy wielkiego tokamaka ITER we Francji

Dlaczego jest tak wiele zapowiedzi budowy reaktorów termojądrowych?

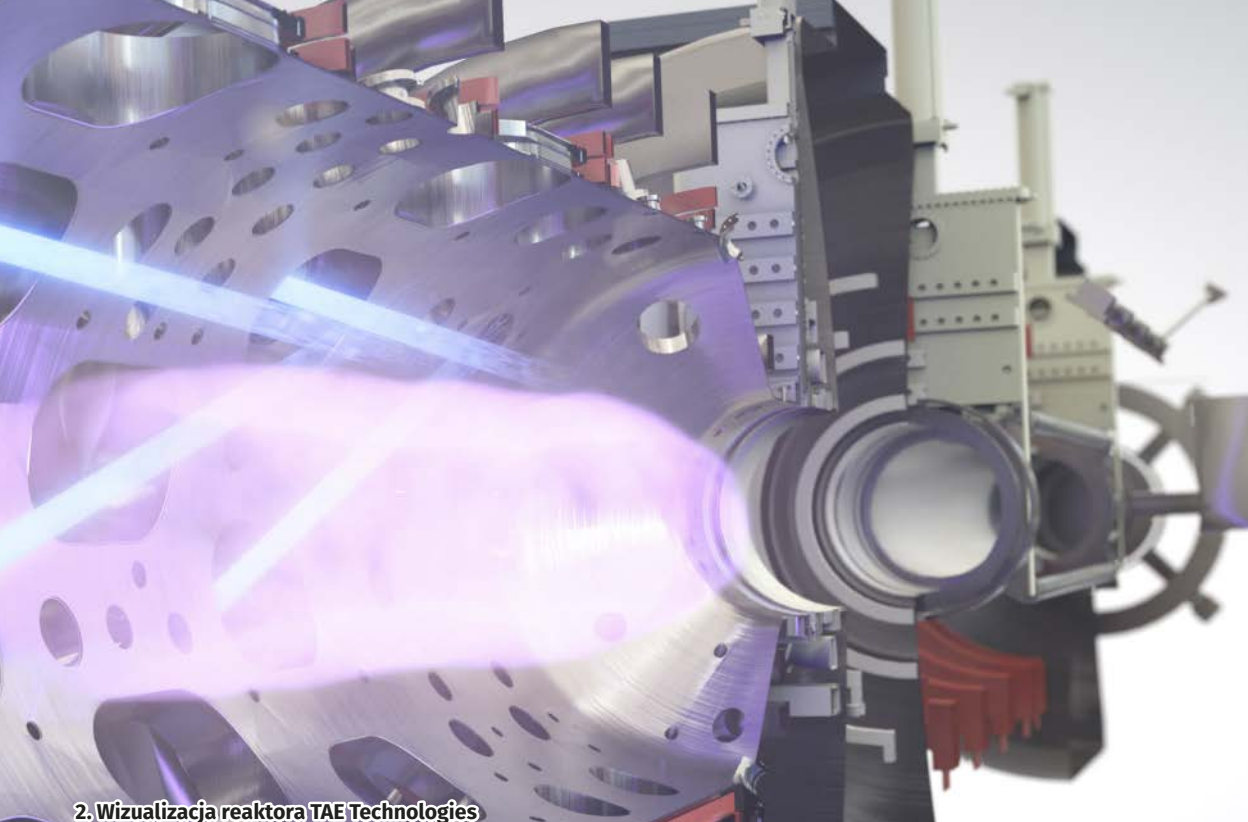
Fuzja obietnic

Co mamy myśleć o mnożących się szumnych zapowiedziach budowy reaktorów syntezy termojądrowej? Przecież oficjalnie i naukowo problem opanowania jej jest wciąż nierozwiązany. Buduje się wielkie tokamaki (1) i stellaratory, które nawet jeszcze nie mają produkować energii, tylko raczej dowieść, że da się ją produkować. A tu raz po raz ktoś wyskakuje z nowym projektem...

Jednym z takich dość zaskakujących termojądrowych komunikatów jest zapowiedź firmy TAE Technologies z kwietnia 2021 roku, że jej reaktor, nazwany Norman, może wytworzyć stabilną plazmę o temperaturze ponad 50 milionów stopni Celsjusza w perspektywie około dekady (2). Tak podawał serwis TechCrunch. Tak, nie przesłyszeliśmy się. Ich konstrukcja ma pozwolić na utrzymanie trwałych reakcji fuzji jądrowej do normalnej produkcji energii elektrycznej. Według informacji

TechCrunch, plazma w reaktorze Norman osiągnęła temperaturę ponad 50 milionów stopni Celsjusza w kilkuset cyklach testowych w ciągu półtora roku poprzedzającego publikację.

TAE Technologies twierdzi, że pozyskała od inwestorów niemal 900 mln dolarów. Wśród nich jest macierzysta spółka Google Alphabet. Pieniądze zostaną przeznaczone na rozwój reaktora demonstracyjnego o nazwie Copernicus, który ma generować temperatury



2. Wizualizacja reaktora TAE Technologies

jeszcze wyższe niż Norman, przekraczające 100 milionów stopni Celsjusza.

Reaktor Norman jest już piątą generacją technologii reaktorowej firmy TAE. Do jego budowy wykorzystano uczenie maszynowe, w tym algorytm Optometrist firmy Alphabet oraz moc obliczeniową dostarczoną przez program INCITE Departamentu Energii USA. Nazwa reaktora pochodzi od nieżyjącego już Normana Rostokera, który był założycielem TAE Technologies w 1998 roku. Był autorem najważniejszych rozwiązań technicznych stojących za projektem, w tym oryginalnych mechanizmów kontroli mocy, podawania paliwa oraz kontroli próżni. TAE Technologies opracowuje również w pełni zintegrowany system zarządzania energią, który łączy kontroler, konwerter i komponenty baterii w jeden skalowalny moduł mocy.

Firma twierdzi, że jej projekt reaktora linowego wytwarza gorętszą i bardziej stabilną plazmę niż tokamaki, choćby takie jak budowany we Francji, gigantyczny ITER. Ponadto TAE ma w planach budowę reaktor wodorowo-borowy. Jak na razie jednak, że reaktor Norman wykorzystuje izotopy wodoru tryt i deuter.

Synteza na ciężarówce?

Jeśli chodzi o reaktory oparte na borze, to jak najbardziej można, wręcz trzeba, je zaliczyć do tej

zastanawiającej serii alternatywnych wobec głównego nurtu badań inicjatyw. W MT pisaliśmy kilka miesięcy temu o tym, że ogłoszonym przez uczonych z Australii przełom w pracach nad reaktorem opartym na syntezie wodorowo-borowej.

Prace zespołu z Uniwersytetu Nowej Południowej Walii, który proponuje wykorzystanie laserów dużej mocy zamiast podgrzewania paliwa deuterowo-trytowego w ekstremalnie wysokich temperaturach, trwają już od dawna. W 2017 roku cytowaliśmy na łamach MT szefa zespołu Heinricha Horę: „Z punktu widzenia inżynierii nasze podejście będzie o wiele prostsze, a ponieważ paliwa i odpady są bezpieczne, więc reaktor nie będzie potrzebował wymiennika ciepła i generatora z turbinami parowymi, zaś niezbędne lasery można kupić na rynku”.

Wodorowo-borowa reakcja fuzji H-11B jest aneutroniczna, czyli nie wymaga działania neutronów, ani ich wytwarzania. Innymi słowy, nie jest radioaktywna. Nie wymaga też radioaktywnego paliwa i nie wytwarza odpadów radioaktywnych. W odróżnieniu od większości innych metod generowania energii, nie są potrzebne wymienniki ciepła, ani też turbiny parowe, fuzja wodorowo-borowa uwalnia energię prawie bezpośrednio w postaci elektryczności. Nie potrzeba skomplikowanych konstrukcji reaktorów, takich jak tokamaki lub stellatory. Produktem „odpadowym” jest nieszkodliwy



hel. Zarówno produkty wejściowe, jak i wyjściowe są nietoksyczne. Zdaniem prof. Heinricha Hory, „reakcja 12 mg paliwa w postaci boru może wytworzyć ponad 1 GJ energii”.

W „zwykłym” procesie syntezy używa się silnie ściśniętych jąder deuteru i trytu. Aby wytworzyć wysokie ciśnienie potrzeba bardzo dużo energii. Zastosowanie wodoru i boru powoduje, że nie trzeba paliwa tak mocno „ściskać”. Dzięki temu proces potrzebuje znacznie mniej energii do uruchomienia. Jak twierdzi Heinrich Hora, konstrukcja reaktora to przede wszystkim pusta metalowa kula, wyposażona w dwa lasery dużej mocy, w której będzie można umieścić niewielką pastylkę paliwową HB11. Jeden laser będzie wytwarzał pole magnetyczne powstrzymujące plazmę, a drugi będzie wyzwał „lawinową” reakcję łańcuchową syntezy. Projekt reaktora nowego typu został przez prof. Horę opatentowany.

Brak zaawansowanych laserów utrzymywał koncepcje profesora Hory i innych uczonych proponujących fuzję wodorowo-borową przez długi czas w sferze rozważań czysto teoretycznych. Sytuacja zmieniła się diametralnie po wynalezieniu w 1985 roku technologii laserowej o nazwie „Chirped Pulse Amplification” (CPA). CPA spełniło wymagania określone przez Horę dla osiągnięcia reakcji HB11. Moc dostępnych laserów od kilkunastu lat wzrasta wykładniczo, zatem osiągnięcie fuzji wodorowo-borowej wydaje się być już blisko, co nie znaczy, że reaktor tego typu powstanie w ciągu roku czy dwóch. Skale czasowe, o których się mówi, to raczej minimum kilkanaście lat.

Warto dodać, że poza Australijczykami od kilku lata dwie firmy, Tri Alpha Energy i LPPFusion, pracują nad mini-reaktorami termojądrowymi, wykorzystującymi jako paliwo mieszaninę wodorowo-borową (PB), która nie wytwarza neutronów.

Fuzja jądrowa na mieszanke wodorowo-borowej (PB) jest trudniejsza do osiągnięcia niż fuzja deuteru-trytu (DT) w tokamakach, bo wymaga o wiele wyższej temperatury. Firma Tri Alpha Energy spodziewa się, że temperatura plazmy w jej reaktorze osiągnie około trzy miliardy stopni Kelvina, czyli ponad dwieście razy więcej niż temperatura w środku Słońca. Firma Tri Alpha Energy i Google do opracowania reaktora fuzyjnego użyje superkomputera o mocy przetwarzania rzędu eksaflopsów i wspomnianego już algorytmu Optometrist, służącego do obliczania procesów stabilizacji plazmy.

Wzrost zainteresowania techniką kontrolowanej syntezy w sektorze prywatnym zaobserwować

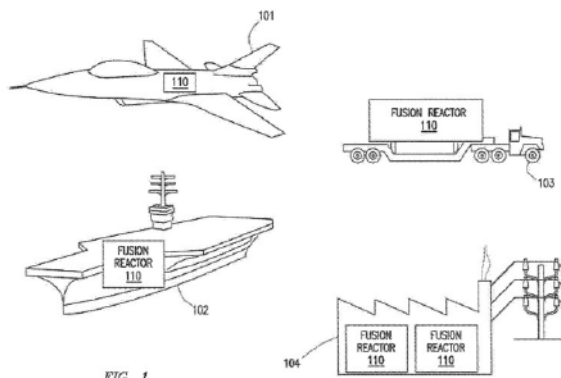


FIG. 1

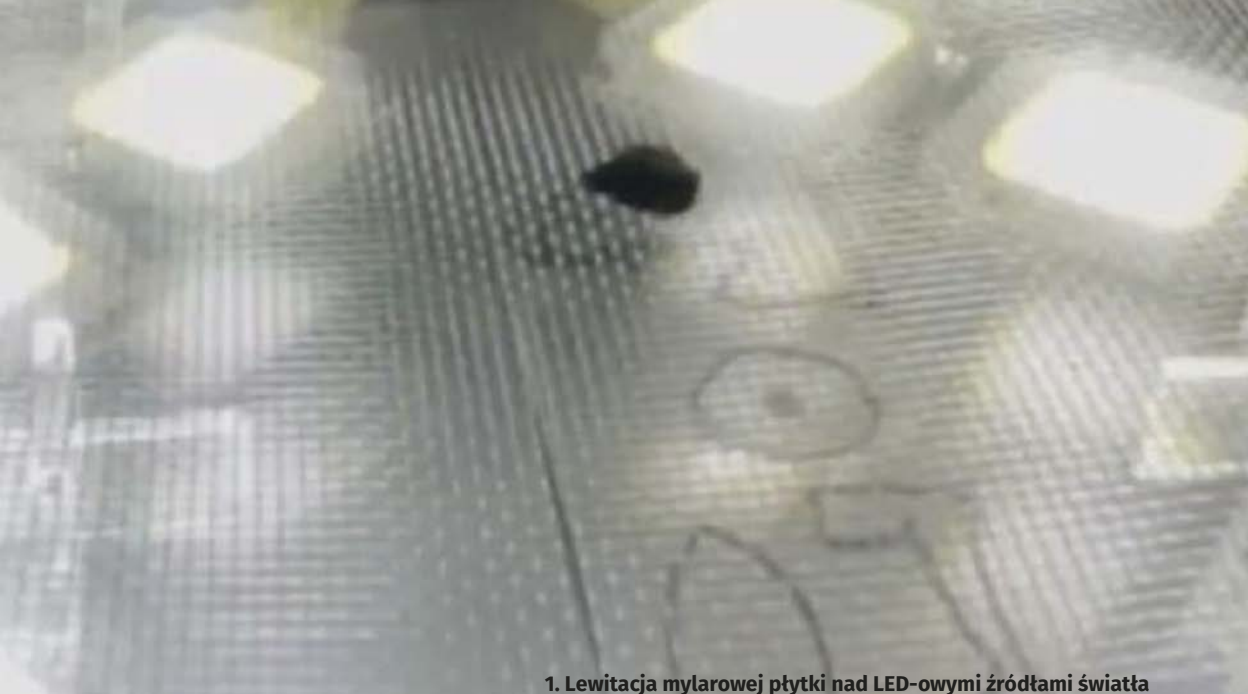
3. Ilustracja z patentu firmy Lockheed Martin pokazująca zastosowania niewielkiego reaktora fuzyjnego

daje się od około dekady. Za przełom i silny sygnał, że coś ciekawego (i w pewnym sensie dziwnego) dzieje się w tej branży uznać można ogłoszenie w październiku 2014 roku przez koncern Lockheed Martin planu opracowania prototypu kompaktowego reaktora termojądrowego (CFR) w ciągu dekady (czyli do 2024 – już blisko). Lockheed Martin pracuje rozwiązaniem kompaktowym, wielkości ciężarówki, które, jeśli zadziała, mogłoby zapewnić wystarczającą ilość energii elektrycznej do zaspokojenia potrzeb liczącego 100 tys. mieszkańców miasta (3).

Do rywalizacji o to, kto zbuduje pierwszy praktyczny reaktor syntezy, dołączały kolejne firmy i ośrodki badawcze, m.in. wspomniany TAE Technologies, ale także Massachusetts Institute of Technology. Nawet Jeff Bezos z Amazona i Bill Gates z Microsoftu zaangażowali się w projekty termojądrowe. NBC News naliczył kilka lat temu w USA siedemnaście mniejszych firm zajmujących się wyłącznie syntezą termojądrową. Startupy takie jak General Fusion czy Commonwealth Fusion Systems stawiają na mniejsze reaktory, oparte na innowacyjnych nadprzewodnikach.

Co więc począć z tym zaskakującymi projektami małych a czasem dużych reaktorów termojądrowych, które trochę, lub nawet czasem mocno kłócą się z obowiązującą wiedzą na temat prac nad kontrolowaną syntezą jąder? Chyba najlepsza odpowiedź to – „zaczekać”. Za kilka rat sporo powinno się wyjaśnić, bo autorzy pierwszych z tych szumnych zapowiedzi powinni zacząć demonstrować praktyczne efekty swoich prac. ■

Miroslaw Usidus



1. Lewitacja mylarowej płytki nad LED-owymi źródłami światła

Naukowe zabawy światłem

Tajemnice kwantowej walizki

Światło to rzecz nauce dobrze znana i fizycznie dokładnie opisana. Naukowcy potrafią jednak, niejako bawiąc się różnymi, zaskakującymi czasem jego właściwościami, odkrywać jego nieznane oblicza. Niektóre z niespodziewanych wyników eksperymentów każą zastanowić się, czy naprawdę tak dobrze rozumiemy światło.

Przykładem takiego ciekawego eksperymentu jest wprawienie przez badaczy z Uniwersytetu Pensylwanii w stan lewitacji wykonanych w tworzywa sztucznego płytek przy użyciu samego tylko światła. Dokładniej mówiąc – dzięki energii pochodzącej z zestawu jasnych diod LED umieszczonych w komorze próżniowej, udało im się wprawić w ruch dwie małe płytki z Mylaru (1). Jak twierdził magazyn „Wired”, eksperymenty te to prawdziwy przełom, ponieważ naukowcy nigdy wcześniej nie byli w stanie sprawić, by tak duży obiekt unosił się w powietrzu przy użyciu samego światła.

Uczeni chcieliby wykorzystać tę technologię do badania mezosfery, wysoko położonego regionu atmosfery, znanego z niedostępności. Nawet NASA jest zainteresowana potencjalnymi zastosowaniami

tej technologii w eksploracji Marsa, zwłaszcza, że ciśnienie w atmosferze czerwonej planety jest podobne jak w ziemskiej mezosferze. Napędzane światłem lewitatory mogłyby pomóc w zbieraniu danych o temperaturze i składzie marsjańskiej atmosfery.

W październiku 2020 roku media podały, że fizycy z powodzeniem przeprowadzili kontrolowany transport zmagazynowanego światła. Zespół fizyków kierowany przez profesora Patricka Windpassingera z Uniwersytetu Johannesena Gutenberga w Moguncji (JGU) przetransportował światło przechowywane w pamięci kwantowej na odległość 1,2 milimetra. Wykazali oni, że proces ma niewielki wpływ na właściwości przechowywanego światła. Jako medium do przechowywania



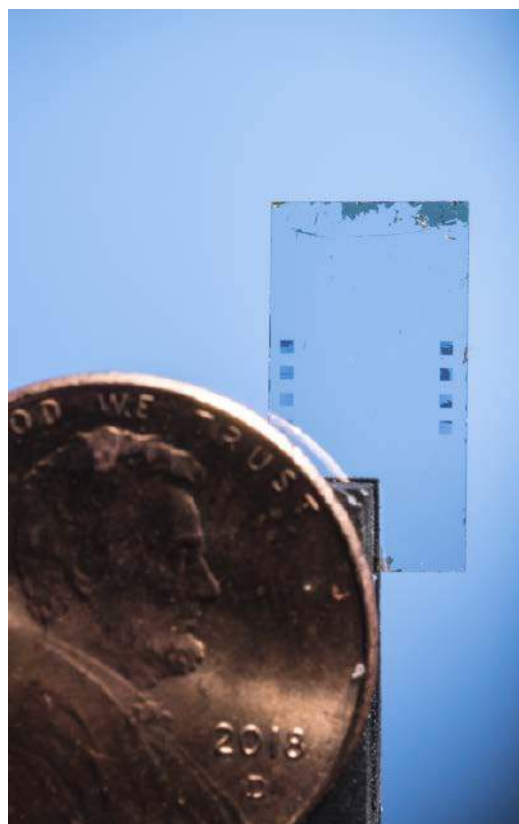
światła naukowcy wykorzystali ultrazimne atomy rubidu-87.

„Przechowywaliśmy światło, wkładając je niejako do walizki, tyle że w naszym przypadku walizka była wykonana z chmury zimnych atomów. Przenieśliśmy tę walizkę na niewielką odległość, a następnie wyciągnęliśmy światło z powrotem. Jest to bardzo interesujące, choćby dla dziedziny komunikacji kwantowej, ponieważ światło nie jest zbyt łatwe do ‘przechwycenia’, a jeśli chcesz je przetransportować gdzie indziej w kontrolowany sposób, zwykle kończy się to jego utratą”, powiedział „Science” profesor Patrick Windpassinger, wyjaśniając ten skomplikowany proces.

Kontrolowana manipulacja i przechowywanie informacji kwantowej, jak również możliwość jej odzyskania, są niezbędnymi warunkami osiągnięcia postępu w komunikacji kwantowej oraz wykonywania odpowiednich operacji komputerowych w świecie kwantowym. Optyczne pamięci kwantowe, które pozwalają na przechowywanie i pobieranie na żądanie informacji kwantowych przenoszonych przez światło pozwalają skalować sieci komunikacji kwantowej. Mogą one na przykład stanowić ważne elementy konstrukcyjne powtarzaczy kwantowych lub narzędzia w liniowych obliczeniach kwantowych. Przy wykorzystaniu techniki znanej jako elektromagnetycznie indukowana przezroczystość (EIT), impulsy światła mogą być uwięzione i spójnie mapowane w celu wywołania kolektywnego wzbudzenia atomów. Ponieważ proces ten jest w dużej mierze odwracalny, światło może być ponownie odzyskane z wysoką wydajnością.

Miesiąc wcześniej naukowcy z Uniwersytetu Stanforda poinformowali o udanym spowolnieniu i sterowaniu światłem za pomocą rezonansowych nanoanten. W artykule opublikowanym w „Nature Nanotechnology”, badacze z laboratorium Jennifer Dionne, opisują wykorzystanie ultracienkich chipów krzemowych do konstrukcji nanoskalowych prętów, które rezonansowo więżą światło, a następnie uwalniają je lub przekierowują. Rezonatory takie mogą być stosowane w przyszłości do obliczeń kwantowych, wirtualnej rzeczywistości i rzeczywistości rozszerzonej, WiFi opartego na świetle (Li-Fi), a nawet do wykrywania wirusów takich jak SARS-CoV-2. Urządzenia wykazały tzw. współczynniki jakości do 2500, o dwa rzędy wielkości wyższe niż jakiegokolwiek podobne urządzenia osiągnęły wcześniej.

Dwa lata wcześniej zespołowi Roberto Merlina i Meredith Henstridge udało się za pomocą mikroskopijnych urządzeń wygiąć światło wewnątrz



2. Mikroskopijne urządzenie zaginające światło

kryształu, aby wygenerować promieniowanie synchrotronowe w laboratorium (2). Naukowcy użyli urządzenia do zginania światła widzialnego, aby uzyskać światło o długości fali w zakresie terahercowym. Duże obiekty generujące promieniowanie synchrotronowe są zazwyczaj wielkości kilku stadionów piłkarskich. Zamiast tego, zespół opracował sposób wytwarzania promieniowania synchrotronowego poprzez drukowanie wzoru mikroskopijnych złotych anten na wypolerowanej powierzchni kryształu tantalenu litu. Użyto lasera do impulsowania światła przez wzór anten, które ugięły światło i wytworzyły promieniowanie synchrotronowe, przydatne do obrazowania w naukach biomedycznych, wykorzystywane np. do rozróżniania tkanki nowotworowej od zdrowej.

Z opisanych wyżej laboratoryjnych zabaw wynika, jak widać, nie tylko przyjemność odkrywania nowych zjawisk i możliwości manipulacji światłem. Wyłaniają się z nich praktyczne i czasem bardzo obiecujące zastosowania w nauce, medycynie, w technice. ■

Miroslaw Usidus



TAJEMNICA ZIEMI

Nasza obca planeta



1. Cyfrowa mapa Europy

W dzisiejszych czasach, gdy niemal wszystko jest zdigitalizowane, tworzenie i korzystanie z cyfrowych map również stało się codziennością **(1)**. Ręczne rysowanie dróg, linii kolejowych, a nawet całych miast, nie jest już konieczne.

Kartografia, jakiej nigdy dotąd nie mieliśmy

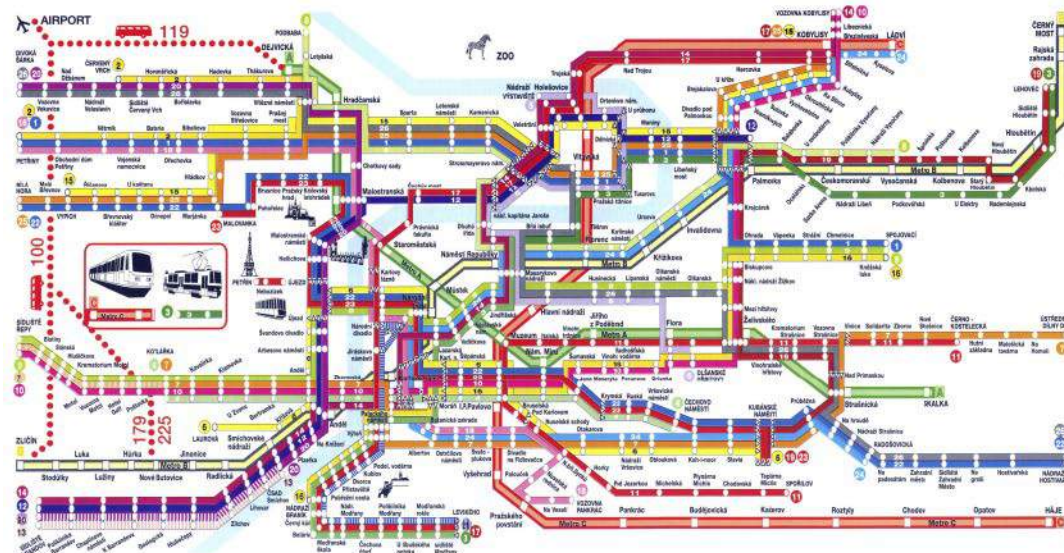
W CENTRUM KAŻDEJ MAPY JEST CZŁOWIEK

Odkąd ludzie zaczęli tworzyć mapy, używano ich do wyjaśniania i nawigowania po świecie. Dziś, przede wszystkim cyfrowe, mapy odgrywają ogromną rolę w wizualizacji danych. Patrząc

na mapę można dotrzeć do nie tylko celu, ale także do informacji w szybszy sposób niż czytając tekst. Na podstawie danych lokalizowanych na mapach rozwija się geoanalitika, wykorzystywana chociażby w dziedzinie bezpieczeństwa. Wiele osób korzysta z map codziennie, w telefonie, poruszając się po świecie, jeżdżąc samochodem, rowerem i biegając, używają ich w wiadomościach i publikacjach społecznościowych. Wystarczy wejść np. na Google Maps, wpisać lokalizację i już. Można przez wprowadzanie danych, np. miejsc, zdjęć a nawet filmików wideo, samemu współtworzyć internetowe mapy.

Ręczne rysowanie to przeszłość

Kartografia, jako połączenie praktyki, nauki i sztuki, kartografia określa zasady i praktyczne



2. Mapa sieci transportowej stolicy Czech – Pragi

standardy tworzenia map. Jest dziedziną interdyscyplinarną, w której nakładają się na siebie geografia, nauki o ziemi, topologia, a nawet polityka. Jej centralną koncepcją jest oparcie się na lokalizacji, od której wychodząc kartografia pomaga nam zrozumieć nasze miejsce w świecie, analizować relacje przestrzenne. Zbieranie danych do tworzenia map jeszcze ok. trzy dekady temu obejmowało analizę zdjęć satelitarnych i lotniczych, jak również pomiary w terenie. Proces mapowania dopiero zaczął się wtedy digitalizować. Dziś kartografia opiera się na mapach generowanych cyfrowo i bazach danych, a nie na ręcznej pracy.

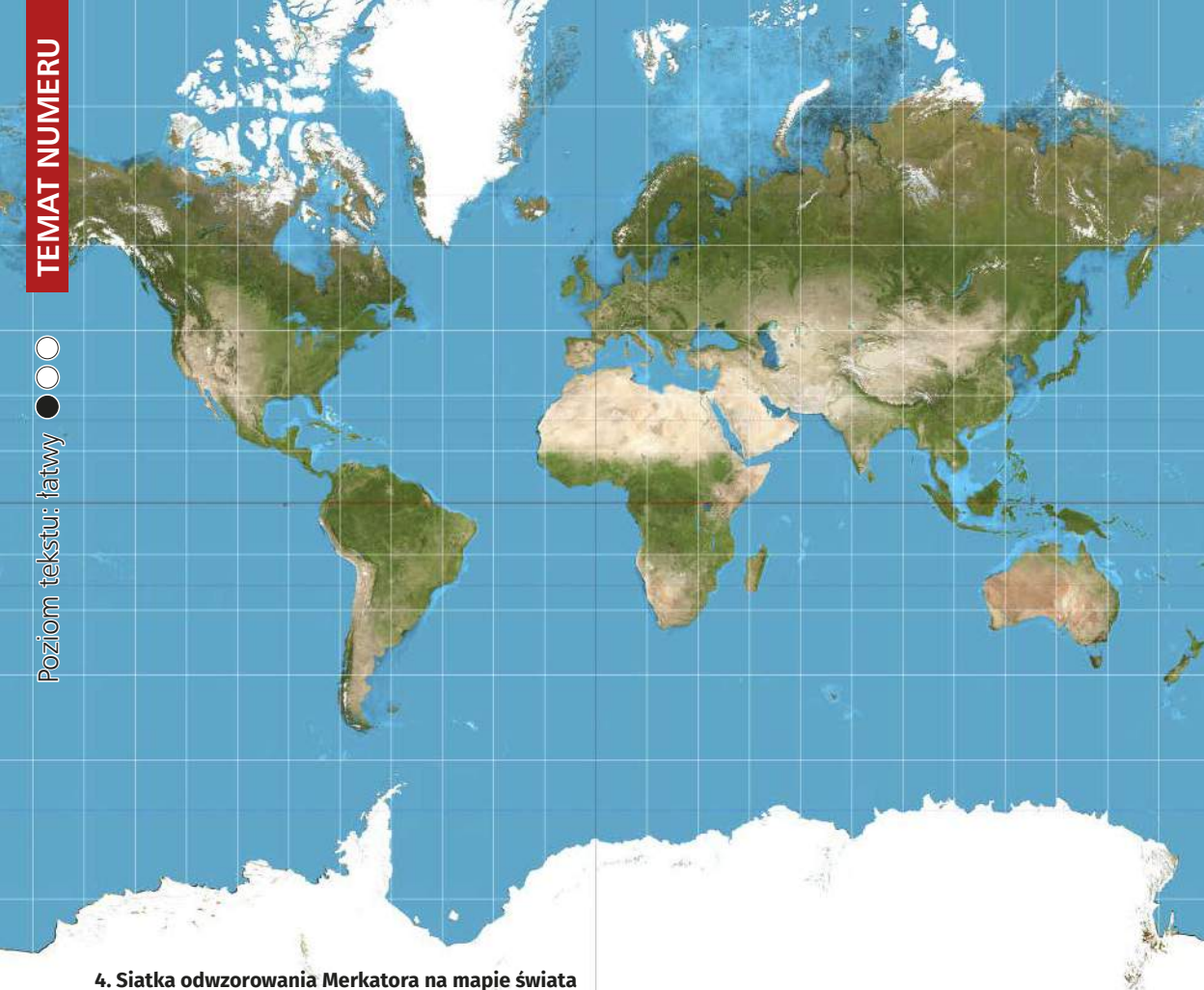
Ważne jest, aby pamiętać, że kartografia zajmuje się reprezentacjami świata, ukształtowanymi przez cel mapy i intencje jej twórcy/twórców. Idealna wierność nie zawsze jest celem kartografa, czego przykładem są plany mapy sieci komunikacji miejskiej (2). Na przestrzeni dziejów mapy były wykorzystywane m.in. do pokazywania bieżących wydarzeń, zajmowały się polityką i handlem, a nawet miały cele wyznaniowe.

Najstarsza znana mapa na świecie pochodzi z siódmego tysiąclecia p.n.e. Uważa się, że przedstawia prawdopodobnie plan i lokalizację anatolijskiego miasta Çatalhöyük. Została namalowana na ścianie jaskini. Należy zauważyć, że istnieją różne interpretacje tych rysunków. Nawet jeśli jest to mapa, ogólna wierność jest wyraźnie niska. Bo najwcześniejsze znane mapy, ze względu na ograniczoną wiedzę ich twórców, były raczej

ekspresją artystyczną niż dokładnym odwzorowaniem. Babilońska „Mapa Świata”, gliniana tabliczka stworzona w Mezopotamii około 700–500 r. p.n.e. (3), przedstawia Babilon jako centrum znanego świata otoczone okrągłym szlakiem wodnym nazwanym „słonym morzem” i otoczony ośmioma trójkątnymi regionami.

3. Babilońska mapa świata





4. Siatka odwzorowania Merkatora na mapie świata

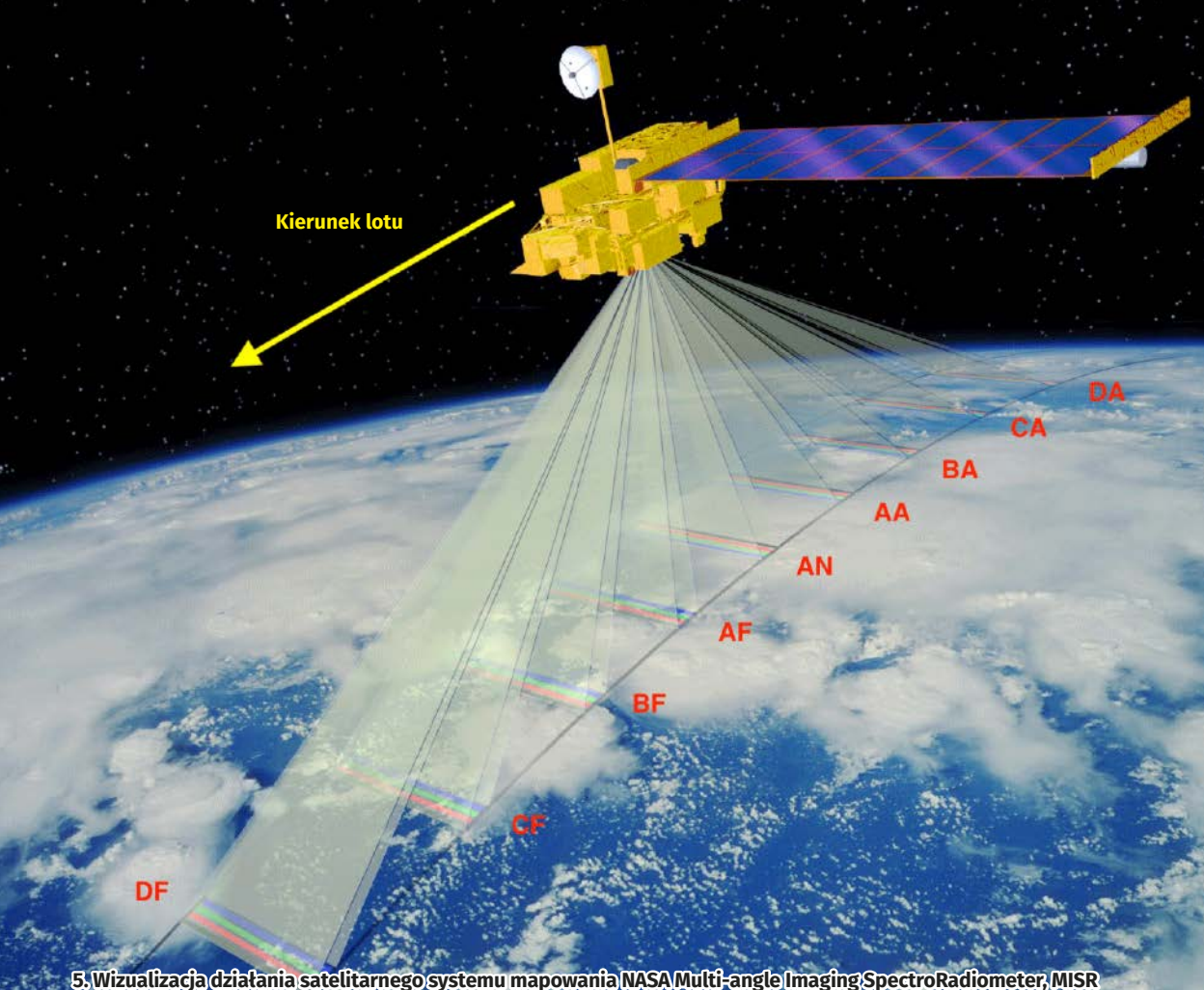
Tworzenie map nabrało bardziej realistycznych kształtów w II wieku, kiedy grecki uczyony Klaudiusz Ptolemeusz opierając się na metodach matematycznych i geometrycznych, opracował system szerokości i długości geograficznej dla prawie 10 tys. miejsc w Europie, Azji i Afryce Północnej. Uważa się, że grecki filozof Anaksymander stworzył pierwszą opublikowaną mapę świata. Mapa ta zaginęła dawno temu, ale na podstawie opisów ze źródeł wtórnych powstały jej reprodukcje. Oczywiście jest, że dokładność tej mapy była ograniczona. Mimo to była to jedna z pierwszych znanych prób dokładnego przedstawienia świata jako całości.

Choć z czasem kartografia stawiała się coraz bardziej precyzyjną nauką, mapy często były bardziej związane z tradycjami kulturowymi, opowiadaniami i wierzeniami. XII-wieczna mapa autorstwa islamskiego uczonego Al-Sharif al-Idrisi przedstawiała świat, którego szczytem był południowy krańiec, ponieważ tam właśnie znajdowała się Mekka,

i z półkulą północną podzieloną na siedem stref klimatycznych. W średniowieczu tak zwane „mappa mundi” często porządkowały świat według czasu biblijnego, z Chrystusem górującym nad Jerozolimą w ich centrum.

W połowie XVI wieku rozwój żeglugi i handlu wywołał zapotrzebowanie na mapy, które pozwoliłyby nawigatorom morskim wyznaczać dużych odległości za pomocą linii prostej, z uwzględnieniem krzywizny powierzchni Ziemi. Zapotrzebowanie to skłoniło europejskich kartografów do przeprowadzenia szeroko zakrojonych badań terenu, eksploracji niezbadanych obszarów i stworzenia tak szczegółowych map, jakie do tej pory powstały. To właśnie w tym okresie powstało odwzorowanie walcowe równokątne Gerarda Merkatora. Siatka, na której oparła się jego mapa świata z 1569 r., wciąż jest używana w mapach tworzonych na całym świecie (4).

Technika tworzenia map była w kolejnych latach doskonalona. Z 1860 r. pochodzi



5. Wizualizacja działania satelitarnego systemu mapowania NASA Multi-angle Imaging SpectroRadiometer, MISR

najstarsze zachowane zdjęcie lotnicze zostało wykonane przez Jamesa Wallace'a Blacka, lecącego w balonie na ogrzane powietrze nad Bostonem. Fotografia lotnicza, kontynuowana później z samolotów a w dzisiejszych czasach dronów, była kolejnym wielkim krokiem naprzód w dziedzinie kartografii. Nie trzeba już było wysyłać w teren dużej liczby geodetów i kartografów, aby przygotować podstawowe mapy. Wykorzystanie zdjęć lotniczych w procesie tworzenia map znacznie się rozszerzyło w czasie I i II wojny światowej, stanowiąc podstawę dla satelitów mapujących NASA, wystrzelonych po raz pierwszy w 1984 roku a dziś tworzących konstelację GPS, nie będącą jedynym systemem nawigacji globalnej. Dziś jej kontynuacją jest obrazowanie z satelitów (5) lub z dronów, wykorzystujące zresztą czujniki i detektory wychodzące poza wyłącznie zakresy światła widzialnego.

Nowe instrumenty umożliwiły pozyskanie wcześniej nieosiągalnych danych mapowych. Na przykład dane do mapy świata trzęsień ziemi

z 1981 roku, stworzonej przez geologów Alvarę Espinosa i Wilbura Rineharta, zostały pozyskane ze stacji monitorowania sejsmicznego. Bazowa mapa dna oceanicznego autorstwa oceanograf Marie Tharp, która potwierdziła teorię tektoniki płyt, powstała z danych uzyskanych dzięki echosondom opracowanym dla okrętów podwodnych w czasie II wojny światowej.

I w końcu przyszła cyfryzacja map, era, która trwa obecnie. Miesięcznie za pośrednictwem aplikacji i platform Google i Apple Maps przeglądanych są miliardy map. Stały się one ważnym interfejsem do korzystania z zasobów danych, wymiany informacji, komunikacji i do wielu innych zadań. Cyfryzacja map przyniosła rozwój i masowe przyjęcie systemów informacji geograficznej (GIS), platform oprogramowania używana do przechwytywania, zarządzania, analizowania, przechowywania i prezentowania warstw danych geoprzestrzennych, która pozwala na lepsze zrozumienie wzorców i relacji geograficznych, na przykład danych dotyczące

nierówności w wynagrodzeniach kobiet i mężczyzn w poszczególnych regionach.

System GPS dostarcza geodetom i autorom map precyzyjne współrzędne geograficzne elementów powierzchni Ziemi za pośrednictwem ogólnoswiatowej sieci orbitujących satelitów i jednostek odbiorczych. Satelity zwiększyły szybkość gromadzenia danych i radykalnie poszerzyły zakres możliwych do odwzorowania informacji. To, co kiedyś zajmowało miesiące lub lata, teraz może być zrobione w ciągu godzin lub minut. Powierzchnia Ziemi jest obecnie nieustannie mapowana przez liczne satelity teledetekcyjne, dzięki czemu powstają ogromne archiwa danych mapowych, które są odbierane, analizowane i przechowywane przez kartografów, naukowców i techników. Tylko pięć czujników mapowania środowiska należącego do NASA satelity Terra gromadzi kwartalnie prawie 620 terabajtów danych. Przez ciągle rejestrowanie obrazu powierzchni Ziemi, satelity umożliwiły stworzenie tysięcy, jeśli nie milionów map, wykorzystywanych w rolnictwie, gospodarce komunalnej, leśnictwie, naukach o ziemi, zmianach globalnych i planowaniu regionalnym.

Wizualizacje wielkich zasobów danych

Wspomniane systemy informacji geograficznej (GIS) pozwalają nam mapować nasz świat jak nigdy dotąd, od renderowania trójwymiarowych map dna oceanów po znalezienie najbliższej pralni chemicznej. GIS służy do wprowadzania, gromadzenia, przetwarzania oraz wizualizacji danych geograficznych, składa się z bazy danych geograficznych, sprzętu komputerowego, oprogramowania oraz twórców i użytkowników GIS. Gdy gromadzi dane opracowane w formie mapy wielkoskalowej (tj. w skalach 1:5000 i większych), może być nazywany systemem informacji o terenie (ang. land information system, LIS). Powstanie GIS jest wynikiem połączenia prac prowadzonych w różnych dziedzinach, geografii, kartografii, geodezji, informatyce, elektronice.

Opisy obiektów geograficznych zasadniczo składają się z dwóch części, zawierających dwa różne rodzaje danych: dane przestrzenne – mogące zawierać informacje zarówno o kształcie i lokalizacji bezwzględnej poszczególnych obiektów w wybranym układzie odniesienia, jak również o ich rozmieszczeniu wzajemnym względem innych obiektów (topologia) oraz dane opisowe (zwane także danymi nieprzestrzennymi lub atrybutowymi), opisujące cechy ilościowe lub jakościowe

obiektów geograficznych nie związane z ich umiejscowieniem w przestrzeni. Uzupełnieniem informacji o obiektach świata rzeczywistego reprezentowanych w bazie danych jest symbolika, tj. graficzny opis postaci, w jakiej obiekty te mają być przedstawiane użytkownikowi. Istotnym składnikiem GIS jest cyfrowa geograficzna baza danych, zawierająca opisy poszczególnych obiektów geograficznych. Baza danych przestrzennych jest zazwyczaj ściśle zintegrowana z pozostałymi modułami funkcjonalnymi GIS, tzn. dostęp do niej jest możliwy tylko poprzez GIS. Alternatywnym rozwiązaniem jest usytuowanie jej na zewnątrz systemu. Wówczas stanowi ona odrębny system, komunikujący się z GIS poprzez dostęp do wspólnych zbiorów danych.

Szeroką grupę zastosowań GIS stanowi wszelkiego typu ewidencja – gruntów, budynków, a ogólnie rzecz biorąc wszelkiego rodzaju zasobów. Szczegółowe informacje tego typu wykorzystują urbaniści, geodeci, konstruktorzy. Zastosowanie warstwowej organizacji map umożliwia łatwą modyfikację jedynie wybranych obiektów, bez konieczności przerysowywania całej mapy. Komputerowa ewidencja własności gruntów z powodzeniem może zastąpić tradycyjną, prowadzoną za pomocą rejestrów i map geodezyjnych. GIS są też bardzo wygodnym narzędziem w rejestracji poziomów emisji wszelkiego rodzaju zanieczyszczeń. Dla potrzeb monitoringu środowiska naturalnego akwizycja danych dla GIS może być prowadzona z wykorzystaniem zdalnych czujników i urządzeń pomiarowych sterowanych komputerowo. W tej grupie zastosowań mieści się również wykorzystanie GIS do analizy i obrazowania danych o charakterze statystycznym, takich jak np. zagrożenie przestępczością, występowanie chorób, struktura użytkowania gruntów. GIS mogą również być bardzo wygodnym narzędziem do przetwarzania danych o infrastrukturze technicznej terenu, tj. o sieciach wodociągowych, gazowniczych, energetycznych, liniach komunikacyjnych.

Współczesna kartografia wykracza daleko poza zwykłe odnajdywanie lokalizacji na mapie. Techniki nazywane ogólnie inteligencją lokalizacyjną połączone z modelowaniem 3D i tworzeniem map w czasie rzeczywistym to obecnie główne kierunki rozwoju cyfrowej kartografii. Inteligencja lokalizacyjna, znana również jako inteligencja przestrzenna, pomaga użytkownikom uzyskać wgląd w dane geoprzestrzenne i odkryć znaczące relacje między nimi. Do tworzenia map trójwymiarowych dziś wykorzystuje się najczęściej

skanowanie LiDAR-em, urządzeniem, które opiera się na świetle laserowym do pomiaru odległości, działając trochę podobnie do opartego na falach dźwiękowych sonaru. Skany LiDAR-owe tworzą chmury punktów, które są bardzo szczegółowymi mapami 3D wszystkiego, od tętniącej życiem metropolii po Wielki Kanion. Tworzenie map w czasie rzeczywistym umożliwiają chmury obliczeniowe. W przeciwieństwie do lokalnie zainstalowanego oprogramowania, dostęp do platform GIS opartych na chmurze można uzyskać za pomocą dowolnej przeglądarki internetowej.

Mapy zapachowe

Mapy umożliwiają odnajdywanie drogi, eksplorację otoczenia, planowanie wycieczek, kierowanie ruchem, lokalizowanie ścieżek rowerowych i wiele innych. Mapy mogą stanowić podstawę do objaśniania historii, jak czynili to Aztekowie w swoich kodeksach, barwnie pokazując migracje swoich przodków w czasie i przestrzeni. Mogą opowiadać historii wojen, tak jak to robiły gazety podczas II Wojny Światowej, pokazując dzień po dniu ruchy i zmiany sytuacji wojsk w Europie. Mogą śledzić rozprzestrzenianie się chorób. Wykorzystał je już lekarz John Snow podczas epidemii cholery w Londynie w 1854 roku, kiedy poprosił, aby każdy przypadek był zapisywany na mapie centrum Londynu. Zauważył dzięki temu, że wiele przypadków cholery łączy się geograficznie z pompą na Broad Street i nakazał usunięcie tego ujęcia. W ten sposób fala epidemiczna opadła. Mapy mogą dziś nadawać sens rozkładowi głosów wyborczych, śledzeniu klęsk głodu i powodzi oraz oczywiście danych demograficznych, zmianą liczby ludności i różnic ekonomicznych. Pomagają wyjaśnić zmiany społeczne, religijne, polityczne, językowe, genetyczne i technologiczne oraz ich konsekwencje.

Popularne są mapy tras, które prowadzą użytkowników z punktu A do punktu B, z jednej lokalizacji do drugiej. Zawierają tylko te informacje, które są przydatne do wykonania zadania, bez zbędnych informacji, które mogą rozpraszać uwagę lub dezorientować. Wyolbrzymiają, a nawet zniekształcają przydatne informacje, aby ułatwić ich odnalezienie i śledzenie. Dodawaj słowa i symbole tam, gdzie jest to przydatne.

Eric Gunderson, dyrektor generalny Mapbox, platformy, która dostarcza narzędzia do mapowania firmom takim jak Foursquare, Evernote i Snapchat mawia, że mapy dziś „odzwierciedlają ludzkie zachowania i doświadczenia” a nie jedynie geograficzne ramy, w których ludzie żyją. Ich twórcy coraz bardziej



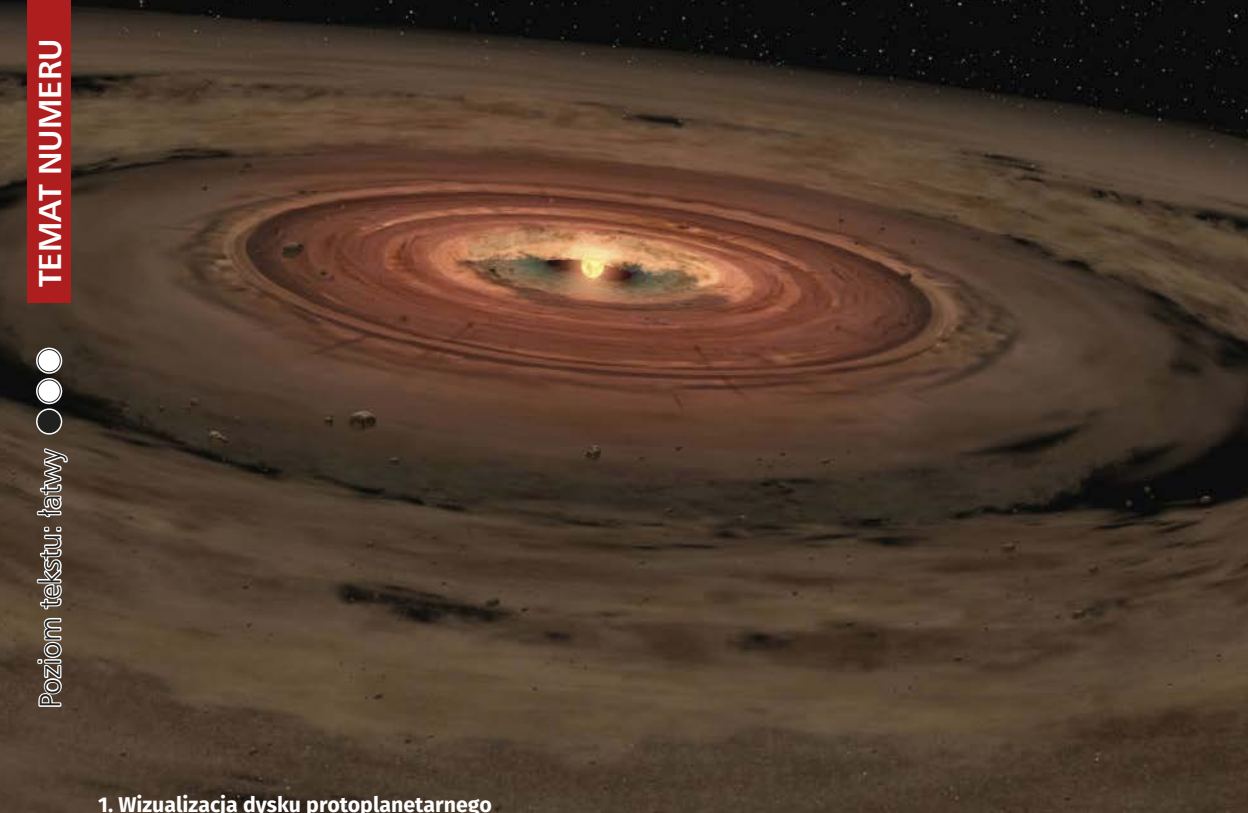
6. Fragment mapy zapachowej Nowego Jorku

odchodzą od stosowanych przez wieki schematów. Przykładem odchodzenia od utartych schematów kartografii są projekty Good City Life, zespołu badaczy bez wykształcenia kartograficznego, który tworzy mapy na samopoczucie mieszkańców miast. Biorą pod uwagę takie czynniki jak dane sensoryczne, emocje, odczucia piękna i przyjemność wizualną.

Jednym z powstałych w ten sposób projektów jest Happy Maps, seria map online, które wykorzystują algorytmy do sortowania geotagowanych zdjęć i wyznaczania najbardziej malowniczych tras. Inną jest Chatty Maps, projekt, który dokumentuje to, co ludzie słyszą na ulicach i rozważa, w jaki sposób pejzaż dźwiękowy wpływa na ich postrzeganie otoczenia. Każda mapa jest oznaczona kolorem, a dźwięki mają różne odcienie. Smelly Maps, trzeci eksperyment Good City Life, działa w oparciu o podobne założenie, ale skupia się na śledzeniu zapachów (6), na bazie danych pochodzących z mediów społecznościowych.

Jeśli zastanowić się nad głębiej, to ten najnowszy nurt nie jest czymś wcześniej niespotykanym. Mapy, choć doczekały się naukowych podstaw i wyrafinowanych narzędzi, zawsze były antropocentryczne, służąc nie komu innemu, lecz tylko ludziom jako pomoce i narzędzia do realizacji celów, czy były to podróże, czy chęć przekazania i zdobycia informacji. ■

Miroslaw Usidus



1. Wizualizacja dysku protoplanetarnego

To co stało się przed miliardami lat doprowadziło do stanu, który widzimy obecnie, badamy i staramy się zrozumieć. Czyli – aby zrozumieć współczesność naszej planety, musimy poznać dokładnie przebieg wydarzeń w czasach, gdy dopiero powstawała. Niestety, nie wszystko jest w dziejach Ziemi jasne i oczywiste.

Seria bombardowań w sprzyjających warunkach

NIEZNANA HISTORIA PLANETY ZIEMIA

Pięć miliardów lat temu lub trochę dawniej gdzieś w naszej kosmicznej okolicy doszło do wybuchu supernowej. Powstała w ten sposób fala uderzeniowa (strumień materii), która, przechodząc przez mgławicę, z której później powstał nasz Układ

Słoneczny, zainicjowała zagęszczanie się materii, wprawiając ją w lub zwiększając jednocześnie jej ruch obrotowy. W wyniku przyciągania grawitacyjnego zagęszczenie przyspieszało. Większość masy (ponad 99 proc.) formującego się dysku skoncentrowała się w jego centralnej części. Zapadanie grawitacyjne materiału mgławicy przekształcało energię grawitacyjną obłoku w energię ciepłą. W centrum mgławicy szybkość przemiany energii grawitacyjnej w ciepłą przewyższała szybkość przenoszenia tej energii na zewnątrz, co prowadziło do znacznego rozgrzania się centralnej części dysku. Dalsze zapadanie wywołało reakcję termojądrową przemieniającą atomy wodoru w hel. Zajaśniało nasze Słońce (1).

Po zapłonie Słońca jego promieniowanie czyli wiatr słoneczny wydmuchnął wszystkie drobniejsze składniki mgławicy, zostawiając Układ



2. Hipotetyczna synestia powstała hipotetycznej kolizji proto-Ziemi z Theia

Słoneczny w takiej postaci, w jakiej go znamy. Powstające w dysku niejednorodności narastały i powiększały się, różnice w prędkości obrotowej sprawiały, że zagęszczenia przyjmowały najpierw formę pierścieni, później, gdy wystąpiły w nich większe gęstości, pod wpływem grawitacji trwał lokalny proces dalszego ich zagęszczania. Sukcesywnie dochodziło do kolizji różnych obiektów, co prowadziło do powiększania ich masy. W zderzających się z dużą prędkością drobnych ciałach dominowało kruszenie, takie, jakie obserwuje się obecnie w pierścieniach planetarnych. Większość istniejących wówczas drobnych obiektów w późniejszych okresach spadła na planety. Tylko niewielka część tych okruchów pozostała do dziś w Układzie Słonecznym i są one klasyfikowane jako drobne ciała niebieskie. Ważną rolę odegrały w tym gazy, które wyhamowywały obiekty i umożliwiały im zlepianie się. W ten sposób powstały protoplanety. Jedną z nich, oddaloną od Słońca o około 150 milionów kilometrów, była Ziemia. Według obecnego datowania Ziemia uformowała się $4,54 \pm 0,05$ mld lat temu, czyli zaledwie ok. 100 mln lat po powstaniu Słońca. Nie wiemy jednak dokładnie, jak to się stało.

Nowe, dokładniejsze, pomiary laboratoryjne wskazują, że głównym czynnikiem pobudzającym tworzenie Układu Słonecznego były magnetyczne

fale uderzeniowe w chmurze pyłów i gazów otaczających nowopowstałe Słońce. Takie wnioski z badań najstarszych i najbardziej prymitywnych meteorytów zwanych chondrytami, opublikowali pod koniec ubiegłego roku w „Science” naukowcy z Massachusetts Institute of Technology oraz Uniwersytetu Stanowego w Arizonie.

Naukowcy wciąż nie mają dokładnej teorii, jak formują się planety (albo czy w ogóle jest tylko jeden sposób). Obecnie walczą (a może się uzupełniają?) dwa modele. Pierwszy i najszerzej akceptowany mówiący o akrecji jądra, sprząda się w przypadku formowania się planet skalistych takich jak Ziemia, ale ma kłopoty z planetami dużymi. Drugi, oparty na niestabilności dysku, może odpowiadać za powstanie największych planet. Dochodzą teorie „komplukujące”, jak np. opublikowane kilkanaście miesięcy temu w czasopiśmie „Science Advances American Association for the Advancement of Science”, badania sugerujące, że wiele oryginalnych planetarnych elementów naszego Układu Słonecznego mogło w rzeczywistości zacząć funkcjonować nie jako obiekty skaliste, ale jako gigantyczne kule ciepłego błota.

Chociaż populacja komet i asteroid przechodzących przez wewnętrzny Układ Słoneczny jest dziś niewielka, były one bardziej obfite, gdy

planety i Słońce były młode. Zderzenia z tymi lodowymi ciałami prawdopodobnie zdeponowały na powierzchni Ziemi znaczną część jej wody. Planeta nasza znajduje się w strefie Złotowłosej, czyli w regionie, w którym woda może pozostać w stanie ciekłym, co zdaniem wielu naukowców odgrywa kluczową rolę w rozwoju życia.

Ziemia się rozpada – Jowisz wymiata

Z jednej strony obowiązuje coś w rodzaju głównego nurtu przekonań naukowych na temat powstania i wczesnego okresu w historii Ziemi. Z drugiej mnóstwo jest zagadek, konkurujących hipotez i nowych koncepcji, dawniej nie branych pod uwagę. Należy do nich popularna dziś

7 teorii na temat pochodzenia życia

1. Zaczęło się od iskry elektrycznej

Jak pokazał słynny eksperyment Millera-Ureya z 1953 roku, iskry elektryczne mogą tworzyć aminokwasy i cukry z atmosfery wypełnionej wodą, metanem, amoniakiem i wodorem. Sugeruje to, że pomoc w stworzeniu kluczowych elementów składowych życia na Ziemi w jego początkach mogły wyładowania atmosferyczne. Chociaż badania przeprowadzone od tego czasu wykazały, że wczesna atmosfera Ziemi była w rzeczywistości uboga w wodór, naukowcy zasugerowali, że chmury wulkaniczne we wczesnej atmosferze mogły zawierać metan, amoniak i wodór, a także oczywiście pełne błyskawice.

2. Z gliny powstałeś

Według koncepcji chemika organicznego Alexandra Grahama Cairns-Smitha z Uniwersytetu w Glasgow w Szkocji, występująca na Ziemi glina mogła nie tylko koncentrować organiczne związki, ale także pomagać w ich organizowaniu we wzory odpowiednie dla struktur DNA. Uczony sugeruje, że kryształy mineralne w glinie mogły ułożyć cząsteczki organiczne w zorganizowane wzory. Po pewnym czasie cząsteczki organiczne przejęłyby to zadanie i zaczęły organizować się tańczuchy same.

3. Życie zaczęło się w głębinowych otworach hydrotermalnych

Podmorskie otwory hydrotermalne wyrzucają kluczowe, bogate w wodór cząsteczki. Skaliste zakamarki w ich otoczeniu mogły skupiać te cząsteczki i dostarczać mineralnych katalizatorów dla krytycznych w powstawaniu życia reakcji. Kominy te, bogate w energię chemiczną i ciepłą, podtrzymują dziś tętniące życiem ekosystemy.

4. Początek na zimno

Naukowcy twierdzą, że 3 miliardy lat temu, kiedy Słońce świeciło o jedną trzecią słabiej niż teraz, ziemskie oceany były pokryte lodem. Ta warstwa lodu, być może gruba na setki metrów. Mogła chronić delikatne związki organiczne w wodzie przed promieniowaniem ultrafioletowym i zniszczeniem w wyniku uderzeń kosmicznych. Także niska temperatura miała swoją rolę w ochronie wczesnych form organicznych.

5. „Świat RNA”

W dzisiejszych czasach DNA potrzebuje białek, a białka do powstania wymagają DNA, więc jak mogło powstać życie? Odpowiedzią może być RNA, które może przechowywać informacje jak DNA, służyć jako enzym jak białka i pomagać w tworzeniu zarówno DNA jak i białek. Późniejsze DNA i białka zastąpiły wczesny „świat RNA”, ponieważ są bardziej wydajne. RNA pełni różne funkcje w organizmach, działa jako włącznik i wyłącznik dla niektórych genów. Pozostaje pytanie, jak powstało lub znalazło się na Ziemi RNA. Niektórzy naukowcy uważają, że cząsteczka mogła spontanicznie powstać na Ziemi. Inni twierdzą, że jest to bardzo mało prawdopodobne.

6. Proste złożone początki

Zamiast rozwijać się ze złożonych cząsteczek, takich jak RNA, życie mogło zostać zainicjowane jako mniejsze cząsteczki oddziałujące ze sobą w cyklach reakcji. Mogły one być zamknięte w prostych kapsułkach przypominających błony komórkowe, a z czasem mogły wyewoluować bardziej złożone cząsteczki, które wykonywały te reakcje lepiej niż te mniejsze. Są to scenariusze określane jako modele „po pierwsze metabolizm”, w przeciwieństwie do modelu „po pierwsze gen”.

7. Życie przybyło z kosmosu

Być może życie wcale nie zaczęło się na Ziemi, ale zostało tu przyniesione z innego miejsca w kosmosie, co jest znane jako hipoteza panspermii. Na przykład, skały regularnie są wyrzucane z Marsa przez kosmiczne uderzenia a na Ziemi znaleziono wiele marsjańskich meteoratów, które, jak sugerują niektórzy badacze, przyniosły tu mikroby, potencjalnie czyniąc z nas potomków Marsjan. Inni naukowcy sugerują nawet, że życie mogło podróżować na kometach z innych systemów gwiazdnych. Być może „świat RNA” przybył na Ziemię gotowy z zewnątrz. Teoria ta jednak nie tyle rozwiązuje zagadkę powstania życia, ile przesuwając je rozwiązanie ze skały ziemskiej do kosmicznej.

hipoteza zakładająca, że miliardy lat temu w wewnętrznym Układzie Słonecznym istniał obiekt wielkości Marsa, nazywany Theia, umiejscowiony tuż przy orbitalnym szlaku protoplanety, z której ukształtowała się ostatecznie Ziemia. Gdy doszło do kolizji Thei z obiektem, którego być może nie należy jeszcze nazywać Ziemią, w przestrzeń kosmiczną wyrzucone zostało mnóstwo materiału, z którego później powstała Ziemia i Księżyc. Co zaś stało się z Theią, tego nie określa się dokładnie.

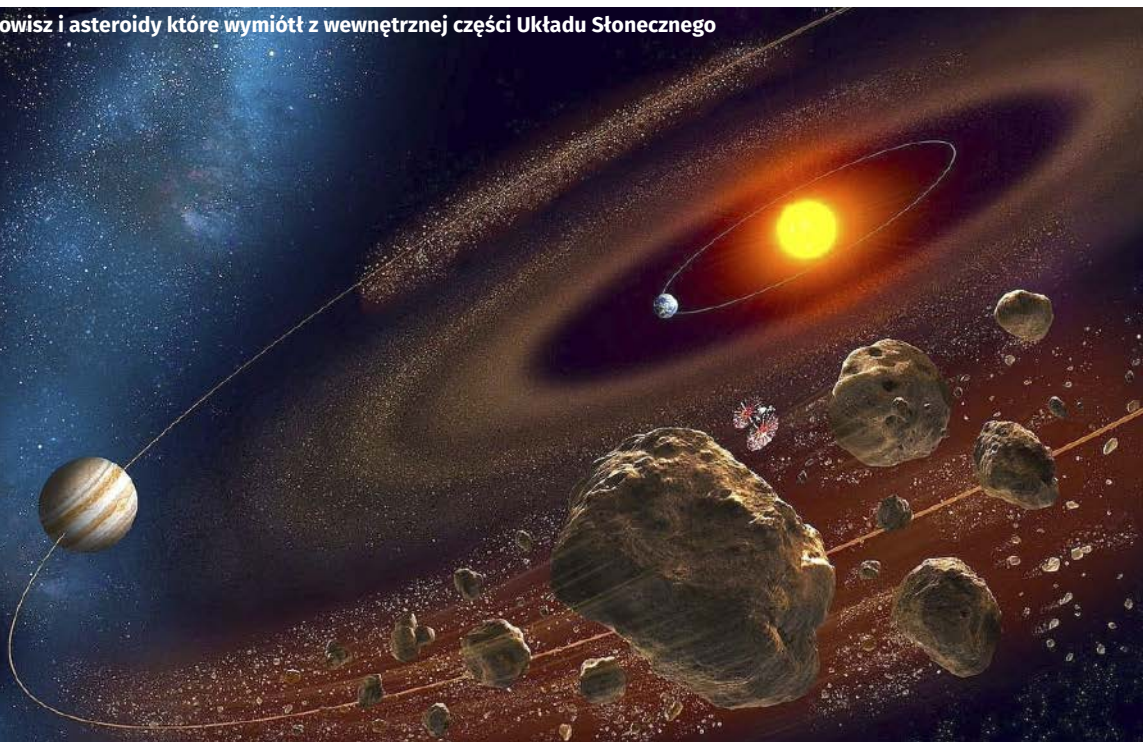
Obecnie astronomowie z Harvardu przypuszczają, że takie kosmiczne katastrofy skutkują bardziej widowiskowymi efektami. Naukowcy stworzyli model komputerowy obrazujący zderzenie dwóch dużych i gorących obiektów o podobnej masie obrotowej. Stwierdzili, że najbardziej gwałtowne z takich kolizji wytworzyłyby synestię – osobliwe halo gazowe powstałe z rozpuszczonych skał, wirujących wokół płynnego rdzenia (2). Nowo powstały obiekt, żyjący jedynie przez paręset lat, miałby znacznie większą objętość niż obiekty, które się zderzyły. Nie byłoby na nim ani cieczy, ani obiektów stałych. Zespół uważa również, że po około stu-leciu trwania w tej formie synestia traci tyle ciepła, że z powrotem zastyga do skalistej postaci. Część skał synestii znalazłaby się na orbicie zestalono-ponownie obiektu, dając początek Księżycowi.

Autorzy przypuszczają, że większość obecnych planet była początkowo synestiami. Model ten wywołał mieszane reakcje w szeregach planetologów, z których wielu zachowuje sceptycyzm wobec tej teorii.

Innym, bardziej akceptowanym scenariuszem, który miał i ma nadal duże znaczenie dla losów Ziemi jest teoria migracji Jowisza. Według wyliczeń Francesco DeMeo, naukowiec z amerykańskiego Harvard-Smithsonian Center for Astrophysics, Jowisz w przeszłości znajdował się tak blisko Słońca jak obecnie Mars. Następnie, podczas migracji ku swojej obecnej orbicie, Jowisz zniszczył prawie w całości Pas Planetoid – pozostało w nim do dziś jedynie 0,1 proc. populacji asteroid. Z drugiej strony, ta migracja wyekspediowała mniejsze obiekty do krańców Układu Słonecznego (3), co miało potem duże znaczenie dla „spokoju” naszej planety, jednak na wczesnym jednak etapie mogło to doprowadzić do silnego bombardowania dostarczającego wczesnej Ziemi dużo ważnych dla dalszego rozwoju wypadków substancji, przede wszystkim wodę.

Do roli Jowisza, ale nie tylko jego nawiązuje zdobywający w ostatnich latach popularność dwu-etapowy model formowania się Ziemi nazywany modelem pojawienia się biopierwiastków (ABEL),

3. Jowisz i asteroidy które wymiotti z wewnętrznej części Układu Słonecznego



który przedstawia nieco inną opowieść o wczesnej Ziemi niż ta, do której przez dekady nas przyzwyczajano. Co nie znaczy, że różni się od innych teorii radykalnie, bo nie tyle wprowadza nowe elementy co reinterpretuje istniejące.

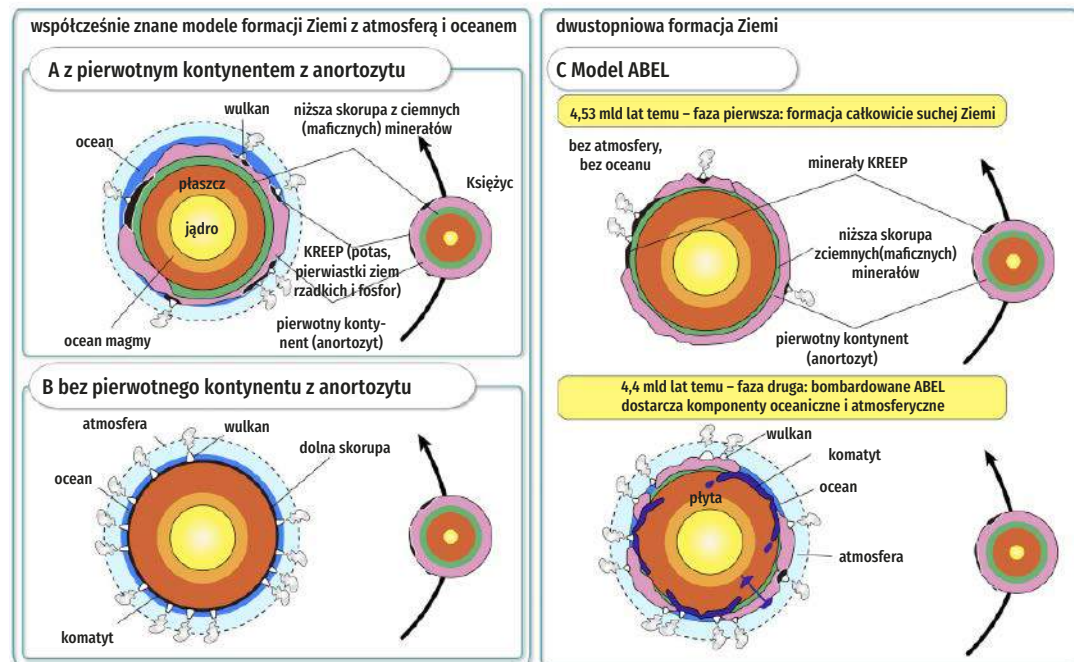
Dwie wcześniejsze teorie powstania Ziemi naukowcy zakładały, że oceany powstały w wyniku ucieczki pary wodnej i innych gazów z atmosfery. Kiedy powierzchnia Ziemi ochłodziła się do temperatury poniżej punktu wrzenia wody, w wyniku opadów deszczu woda spływała do wielkich zagłębień w powierzchni Ziemi, tworzyły się oceany. Grawitacja uniemożliwiła wodzie opuszczenie Ziemi. Model ABEL zakłada, że Ziemia narodziła się jako sucha planeta bez atmosfery i składników oceanicznych 4,56 mld lata temu, zaś potem dopiero nastąpiła wtórna akrecja biopierwiastków, takich jak węgiel (C), wodór (H), tlen (O) i azot (N), która osiągnęła szczyt w 4,37–4,20 mld lat temu. Przyjmuje się zarazem założenie, że ziemska woda pochodzi głównie z węglowego materiału chondrytowego, co wynika z proporcji izotopów wodoru (4).

Gdyby bombardowanie ABEL nie miało miejsca, życie nigdy nie powstałoby na Ziemi, twierdzą zwolennicy tej hipotezy. Ponadto z powodu wstrzykiwania lotnych substancji do początkowej suchej Ziemi

wyzwolilo przejście Ziemi od litej pokrywy do tektoniki płyt co również miało przełomowe znaczenie. Dlatego też jest jednym z najważniejszych dla naszej planety wydarzeń. Uważa się, że bombardowanie ABEL miało miejsce w okresie 4,37–4,20 mld lat temu.

Przez długi czas uważano, że formowanie się oceanu i atmosfery jest równoczesne z pierwotną akrecją Ziemi. Jednak rachunek ilości wody, także przy założeniu zderzenia z Theią i formowaniem Księżyca, nie bardzo się zgadzał. Model ABEL może przezwyciężyć te trudności, starając się wyjaśnić dostawę wody i lotnych substancji potrzebnych do zapoczątkowania prebiotycznej ewolucji chemicznej, która w dalszej konsekwencji doprowadziła do powstania życia na Ziemi, dzięki mieszaninowi się materiału redukcyjnego i utleniającego podczas teoretycznego wczesnego bombardowania. Nastąpić miało ponad 200 milionów lat później niż etap T-tauri (dopiero co powstałego słońca), dlatego uważa się, że to wydarzenie nie jest związane z główną częścią procesu formowania się planet. Jeśli tak, to powinien istnieć inny czynnik wywołujący bombardowanie ABEL. Jeden z możliwych scenariuszy jest następujący: trzy gazowe giganty – Jowisz, Saturn

Modele formowania się Ziemi



4. Jak formowała się planeta Ziemia – modele z uwzględnieniem ABEL

i hipotetyczna planeta „czarna owca” zostały uformowane jako pierwsze. Grawitacyjne interakcje wyrzuciły „czarną owcę”, która być może znajduje się obecnie w Pasie Kuipera a może została całkiem wypchnięta z Układu Słonecznego.

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

Dostawa wody z zewnątrz? Niekoniecznie

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

Są inne wyniki, wykazujące, że przynajmniej połowa wody, która jest obecnie na planecie Ziemia powstała zanim powstała nasza gwiazda. Dla naukowców, którzy napisali o tym w magazynie „Science”, nie jest jasne jednak, czy woda pochodzi z chmury będącej fazą załączkową Układu Słonecznego, czy w ogóle przybyła z zewnątrz. Skąd uczeni wiedzą, jak stara jest woda? Zbadano jej wiek za pomocą analizy stosunku izotopu deuteru w atomach wodoru. Woda powstająca w określonych warunkach różni się jego zawartością od wody formującej się w innych. np. zawartość deuteru w wodzie powstałej w warunkach międzygwiazdowych jest relatywnie wysoka.

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

Pochodzenie wody na Ziemi nie jest więc wcale jasne. A co z atmosferą? Uważa się, że nieostrygła jeszcze kula magmy wydzielala gazy takie jak azot, wodór, węgiel i tlen, tworząc najstarszą znaną wersję atmosfery. „Stwierdziliśmy, że atmosfera, którą według naszych obliczeń była obecna na Ziemi miliardy lat temu, miała podobny skład do tego, co znajdujemy dziś na Wenus i Marsie”, pisze Paolo Sossi z ETH w Zurichu, współautor pracy na ten temat opublikowanej w czasopiśmie „Science Advances” w 2021 r. Aby poznać dokładny skład wczesnej atmosfery Ziemi, naukowcy musieli stworzyć w laboratorium miniaturową wersję wczesnej Ziemi (i jej atmosfery). Zebrali elementarne składniki wczesnego płaszcza Ziemi (znanego geologom jako perydotyt), podgrzali go laserem do momentu, aż stał się stopioną lawą, a następnie lewitowali tę kulę stopionej lawy w strumieniu gazu, który miał reprezentować najwcześniejszą atmosferę Ziemi.

Mówi się, że Ziemia miała w swojej historii trzy atmosfery. Pierwsza atmosfera, przechwycona być może z mgławicy słonecznej, składała się z lekkich pierwiastków z mgławicy słonecznej, głównie wodoru i helu. Działanie wiatru słonecznego i ciepła Ziemi wyparło tę atmosferę. Po zderzeniu, które stworzyło Księżyc, stopiona Ziemia uwolniła lotne gazy; a później więcej gazów zostało uwolnionych przez wulkany, uzupełniając drugą atmosferę bogatą w gazy cieplarniane, ale ubogą w tlen. Wreszcie trzecia atmosfera, bogata w tlen, pojawiła się, gdy bakterie zaczęły produkować tlen około 2,8 mld lat temu.

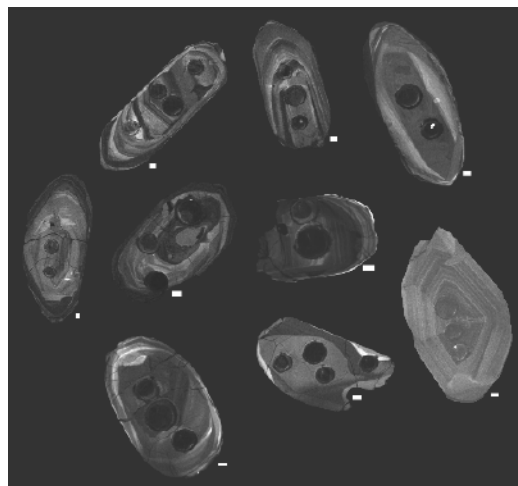
Tajemnice pradawnych skał

W rejonie wzgórz Jack Hills w Australii już wiele lat temu znajdowano ślady cyrkonu (5) świadczące o tym, że Ziemia, przynajmniej jej część, musiała się schłodzić i skryształizować w formę stałą już 4,4 mld lat temu. W dodatku poziom tlenu w tamtejszych skałach sugeruje, że klimat był wówczas na tyle łagodny, że pozwalał na istnienie wody w stanie ciekłym. Jednak według naszej obecnej

wiedzy (publikacje w „Nature” w ostatnich latach), najwcześniejsze formy życia pojawiły się znacznie później, co najmniej 3,95 miliarda lat temu, w czasie, gdy trwał intensywny nalot komet i asteroid. Kiedy japońscy badacze w tym badaniu, pod kierunkiem geologa Tsuyoshi Komiya z Uniwersytetu w Tokio, zbadałi próbki skał z Kanady, zauważyli ślady życia w postaci biogenicznie modyfikowanej skały grafitowej. Te wyniki zresztą są dla wielu zbyt kontrowersyjne – inni eksperci wątpią w ich datowanie. Ogólnie jednak w ostatnich latach rośnie grono badaczy spychających początki życia do czasów kiedy naukowcy kiedyś uważali Ziemię za niezdatną do zamieszkania.

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

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



5. Starożytne cyrkonie

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

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

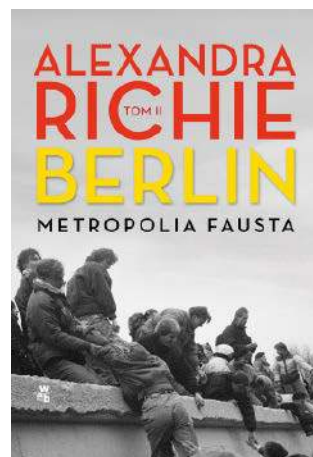
Mirosław Usidus

Metropolia Fausta. Berlin. Tom 2

Alexandra Richie

Wydawnictwo W.A.B., liczba stron: 816, cena: 79,99 zł

Drugi tom niesamowitej historii Berlina! Od jego średniowiecznych początków, przez zdobycze oświeconych władców okresu nowożytnego – przede wszystkim Fryderyka Wielkiego – z dynastii Hohenzollernów, po epokę Bismarcka. Najwięcej miejsca w tej imponującej pracy zajmują rzecz jasna XX-wieczne, tragiczne dzieje miasta, pełne nagłych zwrotów: od dekadenskiej stolicy Republiki Weimarskiej, przez nazistowską, ponurą metropolię, po miejsce rozdarte na pół „żelazną kurtyną”, żyjące w cieniu Muru Berlińskiego, które zjednoczyło się dopiero u progu XXI stulecia i nadal próbuje przezwyciężyć zimnowojenne podziały.



W XVII wieku astronom Edmond Halley, odkrywca wielkiej komety, ogłosił, że „Ziemia jest pusta w środku”. Dziś wiemy na pewno, że tak nie jest, ale to jedna z niewielu rzeczy, które na temat wnętrza Ziemi wiemy na pewno.

Podróż do wnętrza Ziemi

NIEZNANY ŚWIAT POD SKORUPĄ

Wewnątrz Ziemi spoczywają nie tylko zaginione kontynenty, jak ten pod miejscem, w którym płyty tektoniczne Europy i Afryki zachodzą na siebie, lub inny, pogrzebany pod Kanadą lub ten pod oceanem w pobliżu Nowej Zelandii (1), ale również zapewne fragmenty pradawnej planety Theia, która miliardy lat temu zderzyła się z Ziemią, doprowadzając do powstania Księżyca.

Mimo że mamy za sobą setki lat innowacji technologicznych i odkryć naukowych, które umożliwiły nam tworzenie map gwiazd i umieszczenie człowieka na Księżycu, to, co leży bezpośrednio pod naszymi stopami, jest nadal w dużej mierze wiedzą tajemną. Nasza planeta trochę jak cebula (i ogry) składa się z wielu warstw (2). Ludzie, zwierzęta oraz wszystkie nadające się do zamieszkania lądy i oceany leżą na najbardziej zewnętrznej warstwie, znanej jako skorupa ziemiska. Pod nią znajduje się sztywna zewnętrzna warstwa płaszczu zwana litosferą, a poza nią bardziej miękka część płaszczu zawierająca magmę. W tej chwili nie możemy wwiercić się głębiej niż do litosfery, w dodatku na głębokość zaledwie jej „naskórka”.

Mokry płaszcz

Promień Ziemi to 6371 km, czyli podróż do środka naszego globu jest daleką drogą. Ze względu na ogromne ciśnienie i temperaturę sięgającą około 5 tys. stopni Celsjusza trudno wyobrazić sobie dotarcie do samego środka za pomocą jakiegokolwiek urządzenia stworzonego ludzką ręką. Jak zatem dowiedzieliśmy się tego co wiemy



1. Mapa zaginionego kontynentu pod Nową Zelandią

o budowie głębokiego wnętrza Ziemi? Informacji dostarczają nam głównie fale sejsmiczne generowane przez trzęsienia ziemi. W ośrodku sprężystym (skalnym) mogą się rozchodzić dwa rodzaje



2. Warstwy wewnętrzne Ziemi, według współczesnej wiedzy: 1) Skorupa, 2) Nieciągłość Mohorovičića (Moho), 3) Górny płaszcz, 4) Dolny płaszcz, 5) Strefa D, 6) Jądro zewnętrzne, 7) Granica pomiędzy materiałem ciekłym a stałym, 8) Jądro wewnętrzne

objętościowych fal sprężystych – szybsze fale podłużne i wolniejsze fale poprzeczne. Fale podłużne są drganiami ośrodka zachodzącymi wzdłuż kierunku propagacji fali, podczas gdy w falach poprzecznych drgania ośrodka są prostopadłe do kierunku rozchodzenia się fali. Fale podłużne rejestrowane są jako pierwsze (łac. *primae*), a poprzeczne jako drugie (łac. *secundae*). Stąd też ich tradycyjne oznaczenia w sejsmologii – fale podłużne p i poprzeczne s.

Informacje, których dostarczają nam fale sejsmiczne, pozwalają zbudować model wnętrza Ziemi na podstawie własności sprężystych. Inne własności fizyczne możemy określić na podstawie pola siły ciężkości (gęstość, ciśnienie), obserwacji prądów magnetotellurycznych generowanych w płaszczu Ziemi (rozkład przewodnictwa elektrycznego), czy rozkładu ziemskiego strumienia ciepłego. Skład wnętrza Ziemi możemy też określać na podstawie porównań z laboratoryjnymi badaniami własności minerałów i skał w warunkach wysokich ciśnień i temperatur. Na podstawie badań wiemy, że Skorupa Ziemi składa się z tlenu – 46,6 proc. wagowo, krzemu – 27,7 proc., glinu – 8,1 proc. żelaza – 5 proc., wapnia – 3,6 proc., sodu – 2,8 proc., potasu – 2,6 proc. i magnezu, 2,1 proc. Skorupa dzieli się na ogromne płyty, które unoszą się na płaszczu, kolejnej warstwie. Są w ciągłym ruchu. Wstrząsy mają miejsce, kiedy płyty te „trą” o siebie.

Poniżej stosunkowo cienkiej skorupy ziemskiej (10–80 km) i za granicą Mohorovičića (Moho), występuje stanowiący 84 proc. objętości Ziemi skalny płaszcz, zbudowany głównie ze skał krzemianowych bogatych w magnez i żelazo. Silne ciepło powoduje unoszenie się skał, które następnie schładzają się i opadają z powrotem na jądro. Konwekcja tej substancji o konsystencji, jak się uważa, karmelu, powoduje ruch płyt tektonicznych i wybuchy wulkanów. Płaszcz dzieli się na cztery odrębne warstwy. Górna warstwa płaszczu obejmuje litosferę i astenosferę. Strefa Przejściowa (do 660 km, gdzie wykrywane są podziemne zasoby wody). Dolny płaszcz (do głębokości 2891 km) i w końcu granica płaszczu z jądrem.

Ostatnie lata przyniosły serię zaskakujących odkryć dotyczących wody w płaszczu ziemskim. Według wyników badań geologicznych i sejsmologicznych opublikowanych w magazynie „Science”, na głębokościach od 400 do 660 km pod powierzchnią Ziemi znajdują się wielkie jej zasoby wody związane w minerały o nazwie ringwoodyt



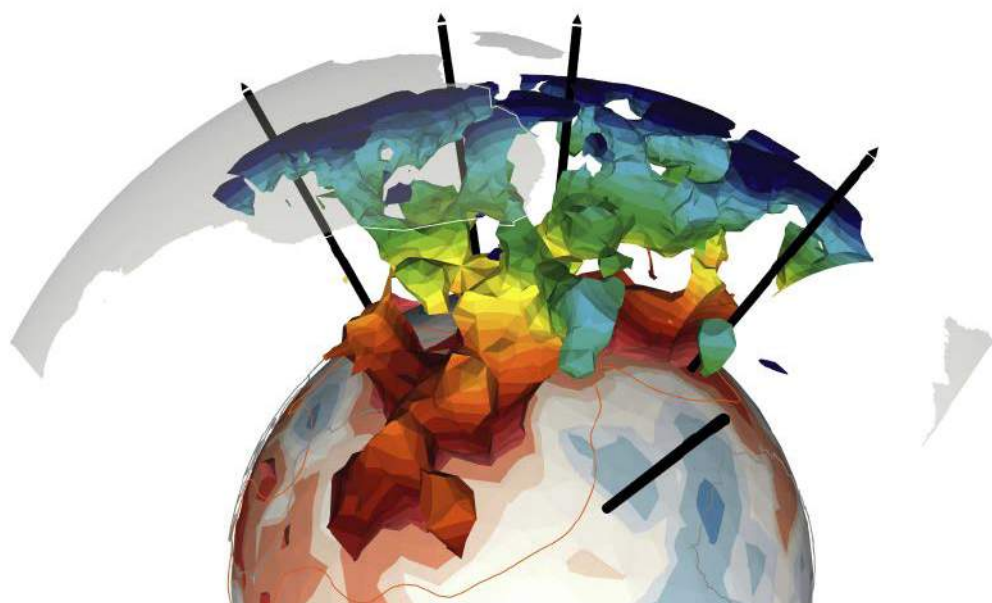
3. Kryształ ringwoodytu

(3), formie oliwiny. „W tej strefie wody może być tyle, ile we wszystkich ziemskich oceanach”, ocenił w artykule Brandon Schmandt, sejsmolog z Uniwersytetu Nowego Meksyku. Także znaleziony w rumowisku rzeczonym w Brazylii, „poobijany” diament, stanowił dla naukowców dowód, że ok. 410 kilometrów pod ziemią, znajduje się „mokra strefa”. Znalezione pięciomilimetrový diament był składnikiem bogatego w wodę minerału. Naukowcy twierdzą, że został wyrzutycony przez erupcję wulkaniczną z głębokości 500 km pod powierzchnią.

Czasopismo „Nature” informowało potem o odkryciu ogromnych rezerwarów wody słodkiej pod dnem oceanów. Według autorów publikacji, z australijskiego Uniwersytetu Flinders, w pobliżu wybrzeży Australii, Ameryki Północnej, Chin i RPA, znaleziono zasoby szacowane na 500 tysięcy kilometrów kwadratowych. Odkrywcę określają tę wodę, jako „wodę o niskim zasoleniu”, co oznacza, że dałaby się wykorzystać bezpośrednio do potrzeb ludzkich. Skąd ta „dobra” woda się wzięła. Według naukowców, gromadziła się pod ziemią w wyniku naturalnej cyrkulacji, przez ostatnich kilkaset tysięcy lat, gdy poziom morza był znacznie niższy.

Czyżby wewnątrz Ziemi spoczywały szczątki Thei?

W ostatnich latach naukowcy zaglądający do wnętrza Ziemi znaleźli dwie struktury wielkości kontynentu, które podobnie jak podkłady wodne burzą dotychczasowy obraz płaszczu ziemskiego. Struktury te mają długość



4. Tomograficzny obraz struktury w płaszczu ziemskim pod Afryką

kontynentów i są do stu razy wyższe niż Mount Everest. Znajdują się na dnie skalistego płaszczu Ziemi, ponad jądrem zewnętrznym. Są zbudowane ze skał, tak jak reszta płaszczu, ale zdają się być gorętsze i cięższe. Naukowcy po raz pierwszy zauważyli te „plamy” w późnych latach 70-tych XX wieku. Zespół geofizyków z Uniwersytetu w Maryland w 2020 r., analizując tysiące danych o trzęsieniach ziemi za pomocą algorytmu zwanego Sequencer, znalazł dowody na istnienie gorących, gęstych struktur i opublikował swoje odkrycia w czasopiśmie „Science”. Według uczonych, struktury mają swoje podstawy prawie 2900 km pod powierzchnią Ziemi.

Dane wykazały istnienie wcześniej nieznanej strefy ultraniskiej prędkości, LLSVP, (inne określenia tych struktur) pod pacyficznym Archipelagiem Markizy. Okazało się również, że tego samego rodzaju strefa pod Hawajami jest znacznie większa niż wcześniej sądzono. Występują, według wyników badań, także pod Afryką (4) i częścią Oceanu Atlantyckiego. Jednak na podstawie odczytów danych sejsmicznych nie można na razie określić gęstości materiału, z którego są zbudowane. Naukowcy spekulują, że LLSVP mogą zasilać wulkany w takich aktywnych rejonach jak Hawaje.

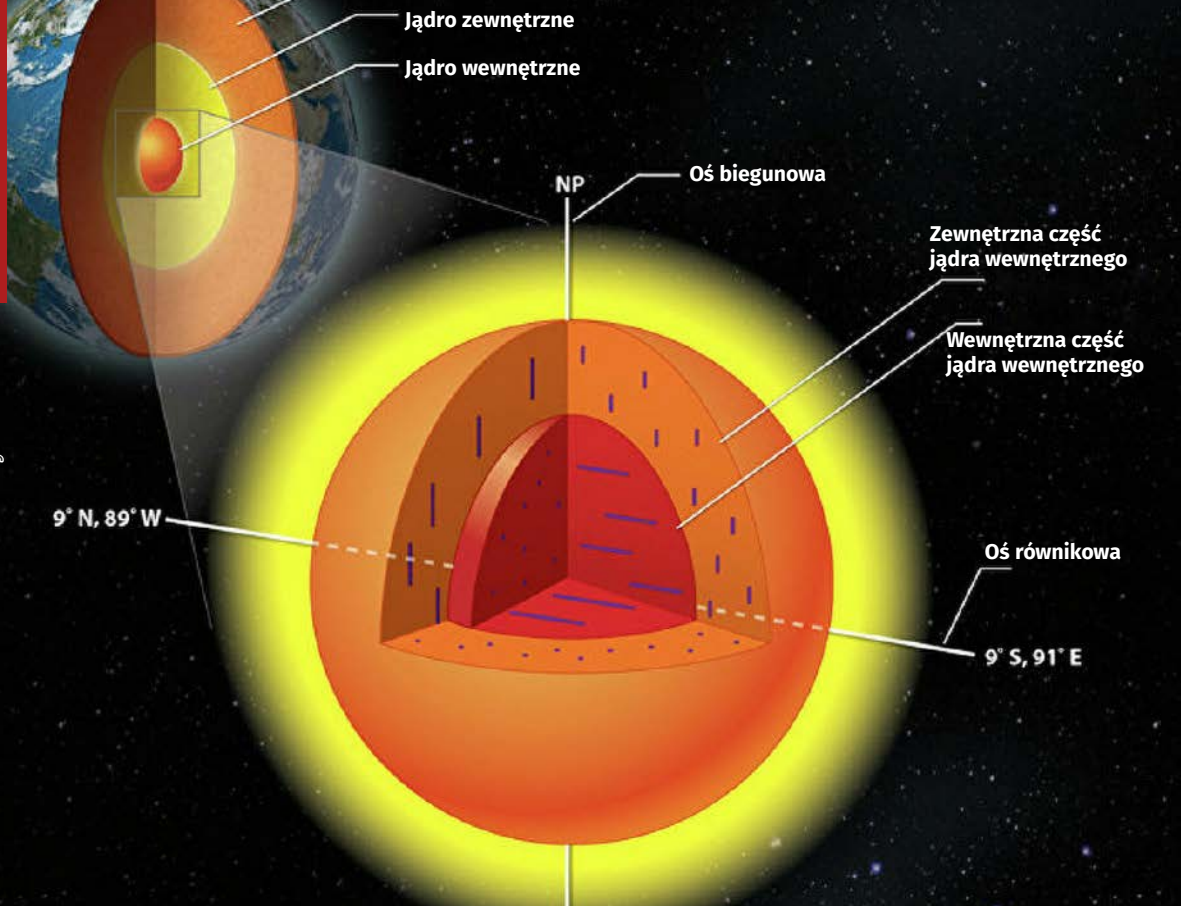
Są sugestie, że te zakopane głęboko w płaszczu Ziemi struktury mogą być pozostałością obiektu wielkości Marsa, nazywanego Theia, który miał zderzyć się z Ziemią u zarania jej dziejów. Podczas 52. konferencji Lunar and Planetary Science, Qian Yuan z Uniwersytetu w Arizonie (ASU) powiedział, że nowe dowody izotopowe i modelowanie wskazują „pozaziemskie” pochodzenie

tych struktur. Wcześniej sądzono, że LLSVP mogły po prostu wykrystalizować się z głębi pierwotnego oceanu magmy ziemskiej lub mogą być skoncentrowanymi relikami pierwotnych skał płaszczu, które przetrwały potężne uderzenie, które doprowadziło do powstania Księżyca.

Schodząc do wewnętrznej części wewnętrznego jądra

Granica Gutenberga-Wiecherta na głębokości 2900 km oddziela skalny płaszcz od ciekłego jądra zewnętrznego, które otacza wewnętrzne, zestalone. Fale s całkowicie w tym miejscu zanikają a fale p zmniejszają swoją prędkość, co wyraźnie wskazuje na ciekłą, stopioną naturę zewnętrznego jądra. Na podstawie badań sejsmologicznych w obrębie jądra wyróżniono trzy strefy – jądro zewnętrzne, jądro wewnętrzne i położoną między nimi strefę przejściową (tzw. nieciągłość Lehmann-Bullena). Zgodnie z powszechnie przyjętymi poglądami jądro Ziemi jest kulą o promieniu według równych oszacowań od 3300 do niemal 3500 km, masie $1,85 \cdot 10^{24}$ kg i gęstości 9,6–18,5 g/cm³. Jądra zewnętrzne i wewnętrzne mają podobny skład żelazowo-niklowy (około 85 proc. żelaza, około 6 proc. niklu, 5 proc. krzemu, reszta to domieszki siarki, chromu, fosforu i węgla).

Według NASA, stałe, wewnętrzne jądro żelazne ma promień około 1220 km (inne szacunki mówią o 1250 km). Na tej głębokości ciśnienie jest tak ogromne, że żelazo i nikiel obecne w jądrze mogą istnieć w stanie stałym pomimo wysokiej temperatury. Jest to aż 5700 Kelwinów, czyli temperatura powierzchni Słońca.



5. Anomalia ziemskiego jądra

Punkt Curie, czyli maksymalna temperatura, do której dany materiał może zachować swoje stałe właściwości magnetyczne dla żelaza wynosi około 1043 Kelwinów. Oczywiście jest więc, że stałe wewnętrzne żelazne jądro nie jest powodem istnienia pola magnetycznego. Pochodzenie pola magnetycznego jest wyjaśnione za pomocą bardziej skomplikowanej teorii geodynamiki, którego funkcjonowanie wyjaśnia magnetoohydrodynamika. Łączy ona zachodzącą w jądrze zewnętrznym konwekcję, unoszącą silniej ogrzaną materię znajdującą się bliżej środka Ziemi wyżej, z działaniem prawa Ampere'a. Efekt zmiany kierunku ruchu unoszącej się i opadającej substancji można opisać jako efekt siły Coriolisa. A unosząca się i opadająca substancja tworzy wiry w większości kręcące się w płaszczyźnie obrotu Ziemi i w tę samą stronę. Pętla prądu elektrycznego może generować pole magnetyczne, a zmieniające się pole magnetyczne może w zamian generować prąd elektryczny. Pola te wywierają siłę Lorentza na naładowane cząstki.

Na skutek wyżej opisanego efektu oraz zmiennej prędkości obrotowej płaszczki ziemskiej, jądro wewnętrzne obraca się z inną prędkością

niż płaszcz. Obecnie jądro wewnętrzne obraca się szybciej niż płaszcz Ziemi o około 0,3–0,5 stopnia rocznie. Wewnętrzne jądro jest, jak wspomniano, zbyt gorące, aby utrzymać stałe pole magnetyczne, ale prawdopodobnie działa stabilizująco na pole magnetyczne wytwarzane przez zewnętrzne jądro ciekłe. Średnie natężenie pola magnetycznego w wewnętrznym jądrze Ziemi szacuje się na 25 gaussów, czyli jest pięćdziesiąt razy silniejsze niż pole magnetyczne na powierzchni Ziemi.

Tradycyjnie przyjmowano, że wewnętrzne jądro Ziemi uformowało się około miliarda lat temu, kiedy to supergorąca bryła żelaza spontanicznie zaczęła krystalizować się w centrum naszej planety w formę metalicznej kuli. Z tą teorią jest jednak taki problem, że to po prostu niemożliwe – ogłosili kilka lat temu uczeni z Uniwersytetu Case Western Reserve. Zespół badawczy w opublikowanym w „Earth and Planetary Science Letters” artykule przypomina o tym, o czym dobrze wiadomo – że materiał musi mieć temperaturę poniżej temperatury zamarzania, aby stał się ciałem stałym. Okazuje się, że krystalizacja cieczy wymaga dodatkowej energii do pokonania bariery nukleacyjnej. Aby pokonać tę barierę i rozpocząć krzepnięcie,

ciecz musi być schłodzona znacznie poniżej temperatury krzepnięcia, co naukowcy nazywają „przechłodzeniem”. Powstanie załączkowego kryształu z cieczy wymaga zatem dodatkowej energii. Nic nie wiadomo o jakimkolwiek dodatkowym bodźcu, który by to wspierał. „Musiało zdarzyć się coś, co obniżyło barierę nukleacyjną, pozwalając na krystalizację w wyższej temperaturze. Naukowcy dokonują takiej sztuki w laboratorium, dodając kawałek zestalonego metalu do lekko schłodzonego ciekłego metalu, co tworzy zarodek krystalizacji a niejednorodny materiał szybko się zestala. Trudno jednak określić w skali planetarnej, jak mogłoby dojść do czegoś podobnego, w jaki sposób fragment ciała stałego wzmacniającego nukleację mógłby znaleźć drogę do środka Ziemi i umożliwić utwardzenie (i rozbudowę) wewnętrznego jądra”, czytamy w artykule. Naukowcy mają przypuszczenia. Jedno z nich to hipoteza, że ze skalnego płaszczka Ziemi do środka stopniowo opadały zestalone bryły metalu, tworząc potrzebne zarodki krystalizacyjne. Musiałyby być jednak ogromne, co znów wzbudza wątpliwości, czy takie zjawisko w ogóle jest możliwe.

To nie koniec zagadek związanych ze środkiem Ziemi. Na pytanie – dlaczego zewnętrzne jądro jest płynne – znamy mniej więcej odpowiedź. Jak w przypadku wszystkich pierwiastków, to czy żelazo jest ciałem stałym, cieczą, gazem zależy zarówno od ciśnienia jak i temperatury żelaza. Jądro Ziemi jest płynne, ponieważ jest wystarczająco gorące, aby stopić żelazo, ale tylko w miejscach, gdzie ciśnienie jest wystarczająco niskie. W miarę jak Ziemia się starzeje i ochładza, coraz większa część jądra zestala się w tempie szacowanym na milimetr na rok, a kiedy tak się dzieje, Ziemia trochę się kurczy. Więc ostatecznie, Ziemia ma płynne jądro ponieważ nie skończyła jeszcze stygnąć. W dodatku proces ten nie jest równomierny. Według doniesień naukowych z lata 2021 żelazne jądro naszej planety stygnie szybciej po jednej stronie niż po drugiej i nikt nie wie dlaczego. „Wewnętrzne jądro jest asymetryczne”, głosi opublikowany w czerwcu komunikat prasowy badaczy z Uniwersytetu Kalifornijskiego w Berkeley. Stygnie szybciej po jednej stronie, pod Indonezją, w porównaniu do drugiej strony, która znajduje się pod Brazylią.

Dawniej uważano, że jądro wewnętrzne może tworzyć żelazny monokryształ. Pojawiła się również koncepcja georeaktora, zaproponowana przez J. Marvinę Herdona, mówiąca, że w centrum wewnętrznego jądra Ziemi znajduje się uran tworzący naturalny reaktor jądrowy, w którym uran ulega rozszczepieniu w reakcji łańcuchowej. Hipoteza ta budzi

jednak wiele kontrowersji i nie jest uznawana przez geofizykę, choć niewątpliwym jest, że uran i inne radioaktywne jądra rozpadają się we wnętrzu Ziemi, wydzielając ciepło i dając tym samym istotny wkład do bilansu cieplnego.

Najnowsze badania pokazują, że samo jądro wewnętrzne ma jeszcze jedną warstwę zwaną wewnętrznym jądrem wewnętrznym, zaś najbardziej wewnętrzne jądro w rzeczywistości uformowało się miliardy lat wcześniej niż wcześniej sądzono, krótko po uformowaniu się naszej planety. Zaskakującym faktem dotyczącym tego regionu jest to, że kryształy żelaza są tu zorientowane w osi wschód-zachód, w przeciwieństwie do innych części, gdzie są one zorientowane w osi północ-południe (5). Nie wiemy dlaczego.

A nuż się dowiercimy...

Badacze z Japońskiej Agencji do spraw Nauki i Technologii Morskiej (JAMSTEC). Chcą oni jako pierwsi w historii świata przewiercić się przez ziemską skorupę aż do płaszczka i pobrać z niego próbki do badań. Próbował tego dokonać międzynarodowy zespół naukowców pracujący w ramach projektu Slow Spreading Ridge Moho Project (SloMo Project). Wybrali rejon Oceanu Indyjskiego, w którym grubość skorupy szacuje się na tylko ok. 5 kilometrów. Odwiertu dokonywano dzięki statkowi wiertniczemu JOIDES Resolution (6). Niestety, ich wysiłki na razie speliły na niczym, gdyż zdołano wykonać odwiert tylko na niecałe 800 metrów.

Za pomocą odwiertów badacze mają nadzieję zbadać tzw. nieciągłość Mohorovičića. Jest to kilkusetmetrowa granica między skorupą i płaszczem Ziemi, ułożona ok. 5–8 km pod oceanami, 35 km pod kontynentami i aż 80 km pod wysokimi górami (np. Himalajami). W tej strefie fale sejsmiczne

6. Statek JOIDES Resolution



gwałtownie zmieniają prędkość, a typowe skały skorupy (granit, bazalt) ustępują miejsca głębinowym perydotytom. Niektórzy naukowcy przypuszczają, że w nieciągłości może pojawiać się woda, która tworzy odmienny rodzaj skał, tzw. serpentynit.

Nie są to oczywiście pierwsze ambitne projekty naukowych odwiertów pod dnem morskim. Próby dotarcia pod skorupę datują się od lat 60. We wrześniu 2012 roku jednostka „Chikyu” wiercąca w pobliżu wysp japońskich ustanowiła nowy rekord 2133,6 metra. Pomimo, że warstwa do przewiercenia jest znacznie cieńsza niż na lądzie, to w oceanicznym przedsięwzięciu wcale nie jest łatwo. Po pierwsze, trzeba bezpiecznie opuścić na głębokość prawie dwóch kilometrów sprzęt. Po drugie, mimo zastosowania wiertła wykonanych z super-twardego węgla wolframowego i tak muszą być one wymieniane co 50–60 godzin, skały bowiem z spod dna morza charakteryzują się bardzo wysoką twardością. Po trzecie, wyzwaniem jest relatywnie bardzo mała średnica otworu, ok. 28 cm.

Niedawno naukowcy z Kraju Kwitnącej Wiśni postanowili wykorzystać ponownie statek Chikyu. Nowy projekt wierceń będzie prowadzony na Hawajach. Plan zakłada zanurzenie wiertła ok. 4 kilometry poniżej powierzchni oceanu, następnie wwiercenie się w skorupę ziemską, pokonanie 6 kilometrów jej grubości, i w końcu pobranie próbek z płaszczu. Termin rozpoczęcia właściwych odwiertów wyznaczono na koniec lat 20.

Przypomnijmy, że najgłębszy wykonany dotąd przez człowieka odwiert to Supergłęboki Odwiert Kolski oznaczany SG-3. Choć był to projekt owiany tajemnicami, to kilka informacji w końcu dotarło do opinii publicznej. Prace nad wierceniami rozpoczęto 24 maja 1970 roku. SG-3 przestał być drażniony w latach 90. ubiegłego stulecia. Do zakończenia badań naukowych doprowadziła niemożność

dalszego drażenia skał w wysokich temperaturach, rozpad ZSRR, niedostateczne finansowanie. Według pierwotnych planów miał dotrzeć na głębokość 15 km. Okazało się, że już na długości 12 262 metrów panują warunki, które nie pozwalają na dalsze wiercenie. Nieco wcześniej, jeszcze na głębokości 4 km znaleziono materiał biologiczny w postaci jednokomórkowych organizmów. Łącznie wydobyto ponad 20 ich gatunków. Na siódmym kilometrze natrafiono na mineralny materiał płynny. Stwierdzono, że woda utknęła tam z powodu przemieszczania się płyt tektonicznych.

Jak widać wiercenie jest dużym wyzwaniem. Dlatego wielu specjalistów, np. Mark Russell, szef firmy Hypersciences, chciałoby radykalnie zmienić technikę wykonywania otworów w skorupie ziemskiej. Zamiast powolnego, żmudnego, drogiego i niejednokrotnie niebezpiecznego wiercenia za pomocą śwідrów, proponuje on wstrzeliwanie w głąb ziemi pocisków. Specjalnie skonstruowane pociski przebijające otwór w skale, w wyniku zapłonu gazów w komorze uzyskiwałyby przyspieszenie ponad 7000 km/h (prawie 2 km/s). Russell ma pomysł, aby w głowicy tych pocisków umieszczać dodatkowo materiały wybuchowe, aby ich eksplozja pod ziemią zwiększała siłę przebicia. Działa strumieniowe strzelać miałyby seriami, pocisk za pociskiem. Sceptycy zauważają jednak, że nie wiadomo, jak owo „strzelanie w ziemię” wpłynie na podziemne środowisko, wody gruntowe i geologię terenu.

Wątpliwe więc, by pomysł strzelania w głąb Ziemi pocisków znalazł uznanie. Może w końcu uda się opracować dokładniejsze i nieinwazyjne techniki wglądu i skanowania wnętrza planety. Uчени bardzo liczą w tym względzie na okiełznanie neutron, które z łatwością przenikają wszystko. Niestety chyba nam wciąż daleko do panowania nad tymi efemerycznymi cząstkami. ■

Mirosław Usidus

Odłot. Elon Musk i szalone początki SpaceX

Eric Berger

Wydawnictwo Kompania Mediowa, cena: 44,90 zł

To książka o podboju kosmosu, spełnianiu marzeń, wielkich ambicjach i ogromnym sukcesie. SpaceX powstał 20 lat temu. W ciągu dwóch dekad stał się liderem posiadającym największą liczbę satelit, tworzy pionierskie rakiety kosmiczne i jako pierwsza prywatna firma wystrzelił człowieka na orbitę. Pół wieku po wyścigu kosmicznym to SpaceX stał się najważniejszą obok NASA firmą zajmującą się eksploracją kosmosu.



Po powierzchni Księżyca chodziło dwunastu ludzi, a tylko cztery osoby na własne oczy widziały dno Rowu Mariańskiego. Ponad pół wieku temu – Don Walsh i Jacques Piccard na pokładzie batyskafu „Trieste”. Potem reżyser James Cameron, który dotarł w 2012 r. na głębokość 10 989 metrów w pojeździe „Deepsea Challenger”, a ostatnio milioner Victor Vescovo.

Przygoda w otchłani dopiero się zaczyna

W GŁĘBINACH

Mamy lepsze mapy powierzchni Marsa (rozdzielczość 4 m na piksel) niż naszych oceanów (rozdzielczość jest niższa niż 1 km na piksel w większości miejsc). Mniej niż 20 proc. ich powierzchni zostało zmapowane w rozdzielczości schodzącej poniżej 100 metrów. Istnieją ogromne obszary oceanu, które sondowano technikami batymetrii pod kątem głębokości co najmniej raz na 100 km kwadratowych powierzchni. Jednocześnie to, co wiemy z tych fragmentów dna oceanicznego, które zbadaliśmy,

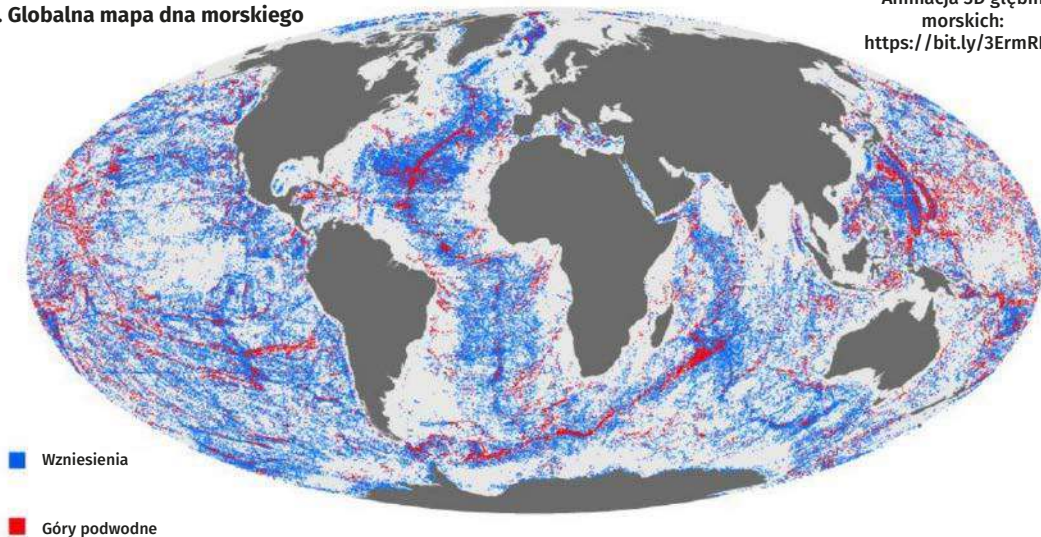
to fakt, że topografia jest tak samo fascynująca i złożona jak na lądzie, wypełniona łańcuchami górskimi, dolinami, kanionami, wulkanami, rowami i rozległymi równinami otchłani (1). Ricardo Aguilar, kierownik wielu ekspedycji oceanograficznych, podkreśla, że badacze tacy jak on nie mogą polegać na informacjach batymetrycznych. „Mamy dobre informacje na temat mniej niż 5 proc. światowych oceanów i bardzo skąpe na temat kolejnych 10 proc.”, mówił w mediach.

Ograniczenia na eksplorację oceanicznych głębi nakłada fizyka. Na dużych głębokościach mamy do czynienia z zerową widocznością, ekstremalnie niskimi temperaturami i miazdzącym ciśnieniem. „Pod pewnymi względami dużo łatwiej jest wysłać ludzi w kosmos niż na dno oceanu”, mówi Gene Carl Feldman, oceanograf z Goddard Space Flight Center NASA w wywiadzie dla agencji Oceana. „Podczas nurkowania na dno Rowu Mariańskiego, który ma prawie 7 mil głębokości, mamy do czynienia z ciśnieniem ponad 1000 razy



Animacja 3D głębin morskich:
<https://bit.ly/3ErRN6>

1. Globalna mapa dna morskiego





2. Robot podwodny – wizualizacja

większym niż na powierzchni. To odpowiednik ciężaru 50 jumbo jetów naciskających na twoje ciało”.

Zatem eksploracja z udziałem ludzi jest trudna. W ostatnich latach rozwija się technika bezzałogowa (2), zarówno nawodna jak i podwodna. Mogą to być pojazdy kierowane przez człowieka (HOV), zdalnie sterowane (ROV) oraz autonomiczne, a także hybrydowe. Takie właśnie autonomiczne aparaty dokonały w 2020 roku odkrycia ogromnych rzek zimnej, słonej wody, które wypływają z australijskiego wybrzeża na głębiny oceanu w chłodniejszych miesiącach z powodu opadania silnie zasolonej wody.

W 2019 roku Victor Vescovo, w ramach własnego projektu zdobycia pięciu najgłębszych miejsc pod wodą („Five Deeps”), zszedł na głębokość 10927 metrów do dna Głębi Challenger, południowego krańca Rowu Mariańskiego na Oceanie Spokojnym. Jak poinformował jego zespół, Vescovo ustanowił rekord najgłębszego samodzielnego zanurzenia w historii. Poprzedni rekord należał do reżysera „Titanica”, Jamesa Camerona. Badacz odkrył cztery nowe gatunki stworzeń. Widział też tam niestety plastikową torbę, puszkę po napoju i opakowania po cukierkach (3).

Innego rodzaju zanieczyszczenia znalazł głębi Challenger oceanograf amerykańskiej Narodowej Administracji Oceanów i Atmosfery, Robert Dziak i zespół badawczy, który w 2015 roku spuścił specjalistyczny sprzęt akustyczny na dno Rowu Mariańskiego. Jak powiedział Dziak w komunikacie prasowym, „można by pomyśleć, że najgłębsza część oceanu będzie jednym z najczystszych

miejsz na Ziemi”. Ku zaskoczeniu zespołu, obszar ten okazał się być miejscem pełnym hałasu. W ciągu dwudziestu czterech dni ich sprzęt zarejestrował kakofonię dźwięków, w tym dudnienie trzęsień ziemi, gaworzenie wielorybów i odgłosy tajfunu szalejącego nad powierzchnią morza. Podwodny mikrofon uchwycił także wiele głosów nienaturalnych, w tym ryk śmigieł łodzi, a także wojskowy sonar i sejsmiczne działa powietrzne. Te ostatnie są kontrowersyjną metodą poszukiwania złóż ropy i gazu ukrytych pod dnem oceanu.

3. Zdjęcie plastikowego śmiecia znalezione na dnie Rowu Mariańskiego przez ekspedycję Victora Vescovo



Dno oceanów badane jest nie tylko pod powierzchnią wody, z jej powierzchni i z wybrzeża, ale także w dużej wysokości, co, jak się okazuje, przynosi doskonałe wyniki. np. w 2014 roku misja agencji ESA o nazwie CryoSat, której głównym zadaniem jest wykonywanie precyzyjnych pomiarów wysokości globalnej pokrywy lodowej za pomocą radaru, posłużyła zarazem do odkrycia tysięcy wcześniej nieznanych morskich grzbietów, rowów oraz innych struktur oceanicznych. Nowa mapa, opisana w piśmie „Science”, w połączeniu z dotychczasowymi danymi, ukazała szczegóły tysięcy podwodnych gór wznoszących się na kilometr lub więcej ponad dno oceanu.

Istnieje ambitny projekt mapowania całego dna oceanu światowego do 2030 roku o nazwie Seabed 2030. Został uruchomiony w 2017 roku przez japońską fundację Nippon Foundation i organizację non-profit GEBCO. Naukowcy ogłosili w ubiegłym roku, że skompletowali nowe dane batymetryczne 14,5 miliona kilometrów kwadratowych dna.

Budowanie podwodnej sieci badawczej

Z roku na rok rośnie liczba projektów badawczych, których celem jest eksploracja dna oceanów. Coraz częściej tworzy się stałe podwodne laboratoria badawcze, które zbierają wiele frapujących danych. Tylko na stacji Aquarius, zlokalizowanej u wybrzeży Florydy, naukowcy odkrywają siedem nowych gatunków na godzinę. Aquarius znajduje

się w pobliżu wyspy Key Largo na Florydzie. Przytwierdzona jest do dna oceanicznego na głębokości 20 m i odległości 6,5 km od brzegu stanowi bazę wypadową dla zamieszkujących ją oceanografów, biologów i inżynierów, którzy badają parametry wód, temperaturę, zasolenie, prądy morskie czy poziom tlenu.

Aquarius jest jedyną stacją na świecie, gdzie naukowcy przebywają bez przerwy (oczywiście nie ci sami). Są też inne stacje, i to znacznie bardziej zaawansowane, choć dłuższe przebywanie w nich często nie jest możliwe. Wspomnijmy choćby o kanadyjskim projekcie NEPTUN, North-East Pacific Time-series Undersea Experiments, który tworzy sieć jedenastu bezzałogowych obserwatoriów głębinowych (4). Za ponad 300 mln dolarów wybudowano szereg baz u zachodnich wybrzeży Kanady. Grube przewody elektryczne i światłowodowe łączą głębię oceanu z wyspą Vancouver.

Sieć ma w sumie 800 km długości i składa się na nią ponad 2 tys. kilometrów przewodów. Jest to największe tego typu laboratorium na świecie. Posiada m.in. kamery wysokiej rozdzielczości, sejsmografy, hydrofony, sondy, prądomiery i roboty podwodne. Najgłębiej położona baza znajduje się prawie 3000 m pod wodą. Dane z głębin przekazywane są do bazy Port Alberni na wyspie Vancouver, a następnie do Uniwersytetu Victorii, gdzie przez internet mogą z nich korzystać naukowcy z całego świata. Działające 24 godziny na dobę kamery, przesyłają dane wprost do internetu. Oczywiście sam obraz to nie wszystko. Zainstalowane na dnie czujniki (w sumie ok. sześć set) zbierają informacje m.in. o zasoleniu i temperaturze wody, nasyceniu jej planktonem i dwutlenkiem węgla, a nawet o drganiach wody powodowanych ruchami tektonicznymi dna czy podwodnymi wybuchami wulkanów.

Kolejny podwodny projekt badawczy, tym razem znów amerykański, o nazwie MARS (Monterey Accelerated Research System) realizowany jest w zatoce Monterey, 100 km na południe od San Francisco. Wybudowano tam stację węzłową, usytuowaną na głębokości 900 m. Z lądem połączona jest za pomocą długiego na 52 km światłowodu. Jedną z głównych części MARS-a jest niewielki robot-łazik, który poruszając się po oceanicznym dnie, mierzy aktywność życiową mieszkańców głębin. Na głębokości 160 m unosi się platforma połączona z obserwatorium na dnie pionową liną. Po linie niczym winda porusza się specjalna kapsuła z czujnikami mierzącymi właściwości wody





4. Sieć bezzałogowych obserwatoriów NEPTUN u wybrzeży Kanady

m.in. temperaturę, zawartość tlenu, przezroczystość, ilość planktonu etc.

Podobne rozwiązanie zastosowali francuscy i brytyjscy naukowcy, którzy rozpoczęli testy automatycznego, podwodnego laboratorium u wybrzeży Bretanii. Projekt o nazwie MeDON (Marine e-Data Observatory Network), na rzecz którego składa się podwodne laboratorium umieszczone w rezerwacie biosfery Iroise, zakłada badanie wpływu człowieka na morskie ekosystemy. Baza znajduje się 2 km od brzegu, na głębokości 20 metrów. Uzyskane informacje, przekazywane są kablem do stacji lądowej. W niedalekiej przyszłości tego typu laboratoria mają powstać na dnie całego Morza Śródziemnego.

Podwodne czarne dziury

Największa głębokość osiągnięta przez wciąż jeszcze niewielkie roboty podwodne to 6–7 tysięcy metrów, co teoretycznie pozwala na badanie niemal 99 proc. podmorskich głębin. Najgłębiej żyjącym stworzeniem, jakie kiedykolwiek dostrzeżono, była ryba żyjąca w Rowie Mariańskim, która majestatycznie przepłynęła obok batyskafu „Triest”, wprawiając w zdumienie jego pasażerów.

Żałogowe jak i bezzałogowe aparaty dokonują niezwykłych odkryć, takich jak np. odkrycie

w 1977 roku kominów hydrotermalnych (5) otoczonych wianuszkami kolorowych, żyjących stworzeń. Pierwszy natrafił na nie „Alvin”, należący do The Woods Hole Oceanographic Institution (WHOI) żałogowy pojazd podwodny. Poszukując wyłotów hydrotermalnych na Krawędzi Galapagos na głębokości przeszło 2000 m, naukowcy ujrzeli mięsiste rurki (nazwane później robakami rylcowymi), otoczone przez wielkie małże i pływające między nimi białe kraby i homary. Odnalezienie tego pierwszego ekosystemu opartego na odmiennym od znanego nam sposobie odżywiania uruchomiło lawinę podobnych odkryć. Dziś liczba rozpoznanych gatunków w takich ekosystemach stworzeń wynosi kilka tysięcy.

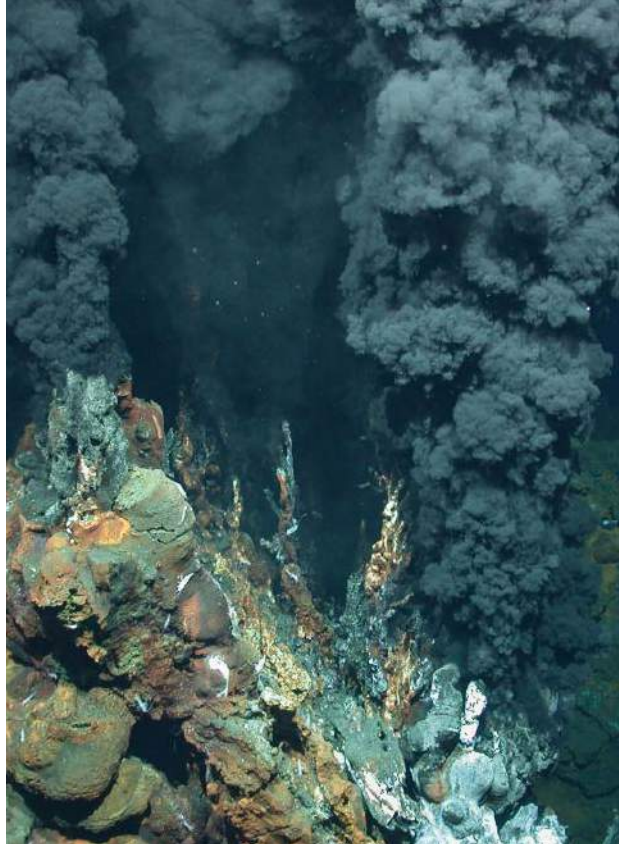
Otwory hydrotermalne są związane z częściami dna oceanicznego, które wykazują wysoki poziom aktywności tektonicznej, jak np. grzbiety śródoceaniczne. W tych regionach, gorące komory magmowe pod dnem morskim podgrzewają wodę, która przeniknęła do dna oceanu. Woda, która wsiąka w szczeliny skorupy najpierw traci minerały (tlen, potas, magnez), ale później jest ładowana nowymi związkami chemicznymi (miedź, cynk, żelazo, siarkowodor), gdy osiągnie najwyższą temperaturę (350–400°C). Czarny „dym”

w kominów jest tworzony przez czarne minerały metalowo-siarczkowe, gdy gorący płyn wznosi się i jest uwalniany do zimnej, bogatej w tlen wody morskiej. Istnieją również chłodniejsze otwory, nazywane białymi kominami.

Inne zidentyfikowane niezwykle zjawiska na dnie oceanów to także wulkany błotne. Nazywa się tak formacje powstałe w wyniku geowymuchiwania płynów i gazów, w największym stopniu metanu. Są też góry morskie, wznoszące się niekiedy ponad tysiąc metrów od dna oceanu, nie sięgające jednak powierzchni wody. Zazwyczaj powstają one z wygasłych wulkanów. Szacuje się, że w oceanie znajduje się ponad 30 tys. gór podwodnych, ale tylko kilka z nich zostało zbadanych. Są to miejsca pełne życia. Mogą generować lokalne wiry i strefy upwellingu (wynoszenia wody w górę).

Specyficzną atrakcją dna morskiego są zwłoki dużych zwierząt, które opadają na dno. Są atrakcją w sensie dosłownym dla żyjących głębiej stworzeń, które uczują na opadłych ciałach wielkich wielorybów lub mątw.

Od niedawna wiemy, że w otchłaniach oceanów występują twory wypisz wymaluj przypominające kosmiczne czarne dziury. Zdaniem naukowców są to ogromne wiry oceaniczne, gdy coś do nich wpadnie, to już się stamtąd nie wydostanie. Niektóre z wirów są ogromne – mogą osiągać średnicę równą nawet 150 km. Ten specyficzny ruch wody może trwać nawet kilka miesięcy lub lat. Co więcej, wiry jako całość również się poruszają. Naukowcy zauważyli, że w miejscu, które zgodnie z ich wyliczeniami stanowiło granicę wiru, cząsteczki wody zachowują się tak samo jak fotony wpadające do czarnych dziur. Co więcej, podobnie jak w przypadku tych kosmicznych obiektów, w rotującej masie wody działają tak duże siły, że nic nie jest w stanie się wydostać z matni. Analogia dotyczy również samej wody tworzącej



5. Kominy hydrotermalne

wir. Jej skład chemiczny (np. poziom zasolenia) może znacznie różnić się od otaczającego go środowiska. Badacze zaobserwowali siedem wirów tego typu, które przez prawie rok nie wypuściły poza „horyzont zdarzeń” ani kropli wody, z której powstały, przy jednoczesnym ciągłym przemieszczaniu się.

Co jeszcze niezwykle i fascynującego czeka na nas w głębinach? Naukowcy i inni eksperci są zwykle zgodni, że jesteśmy dopiero na początku naszej przygody z eksploracją oceanicznych otchłani, więc prawie wszystko jeszcze przed nami. ■

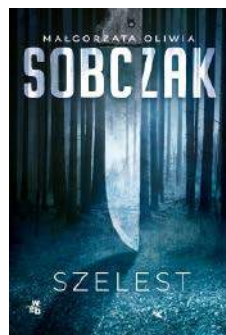
Mirosław Usidus

Szelest

Małgorzata Oliwia Sobczak

Wydawnictwo W.A.B., liczba stron: 352, cena: 42,99 zł

Nowy thriller kryminalny Małgorzaty Oliwii Sobczak, autorki bestsellerowej serii „Kolory zła” („Czerwień”, „Czerń”, „Biel”), dzięki której została okrzyknięta gwiazdą polskiego kryminatu. Dziennikarka Alicja Grabska dostaje zlecenie przetestowania mobilnej aplikacji „Place to Rest”, prowadzącej do zagadkowych miejsc na mapie Trójmiasta. W czasie jednego z „wypadów” natyka się na ciało zamordowanej kobiety. Śledztwo prowadzone przez detektywa Oskara Kordę wiedzie do mrocznego, pełnego tajemnic świata, w którym przeszłość nigdy nie daje o sobie zapomnieć.





Wielki Teleskop Jamesa Webba

Kosmiczne oko na rzeczy dotąd
nie widziane



Według planu, który obowiązywał, gdy przygotowywaliśmy to wydanie „Młodego Technika”, start misji Kosmicznego Teleskopu Jamesa Webba (1), miał nastąpić w grudniu 2021 r. Przygotowujemy więc niezbędnik informacyjny na temat przedsięwzięcia, które porównuje się skalą i znaczeniem z zasłużonym dla badań kosmosu teleskopem kosmicznym Hubble’a.

1. Teleskop Kosmiczny Jamesa Webba – wizualizacja rozłożonej konstrukcji w przestrzeni kosmicznej

W maju JWST po raz ostatni na Ziemi otworzył swoje złote zwierciadło. W okolicach sierpnia-września tego roku zakończyły się wreszcie testy i skomplikowana konstrukcja JWST (skrót od angielskojęzycznej nazwy James Webb Space Telescope) gotowa była do spakowania i wysłania na miejsce startu. „Liczne testy i kontrole Webba miały na celu upewnienie się, że to najbardziej złożone obserwatorium kosmiczne w historii będzie działało zgodnie z założeniami, gdy znajdzie się w przestrzeni kosmicznej”, pisała NASA w komunikacie o zakończeniu testów.

Teleskop jest wspólnym projektem NASA, Europejskiej Agencji Kosmicznej i Kanadyjskiej Agencji Kosmicznej. Jego koncepcja została zbudowana na spuściznie i doświadczeniach wcześniejszych misji, takich jak teleskopy Hubble'a i Spitzera, a sam JWST będzie stanowił fundament, na którym będą mogły powstać przyszłe duże astronomiczne obserwatoria kosmiczne. Kosmiczny Teleskop Jamesa Webba będzie najważniejszym na świecie kosmicznym obserwatorium naukowym. Jego zadaniem jest zgłębianie tajemnic początków naszego Wszechświata, badania, jak powstawały gwiazdy, planety i galaktyki, eksploracja Układu Słonecznego, a także, co chyba najbardziej wszystkich ekscytuje, poszukiwanie chemicznych śladów życia lub warunków mu sprzyjających na egzoplanetach.

Wszystko sprawdzone – można wysłać

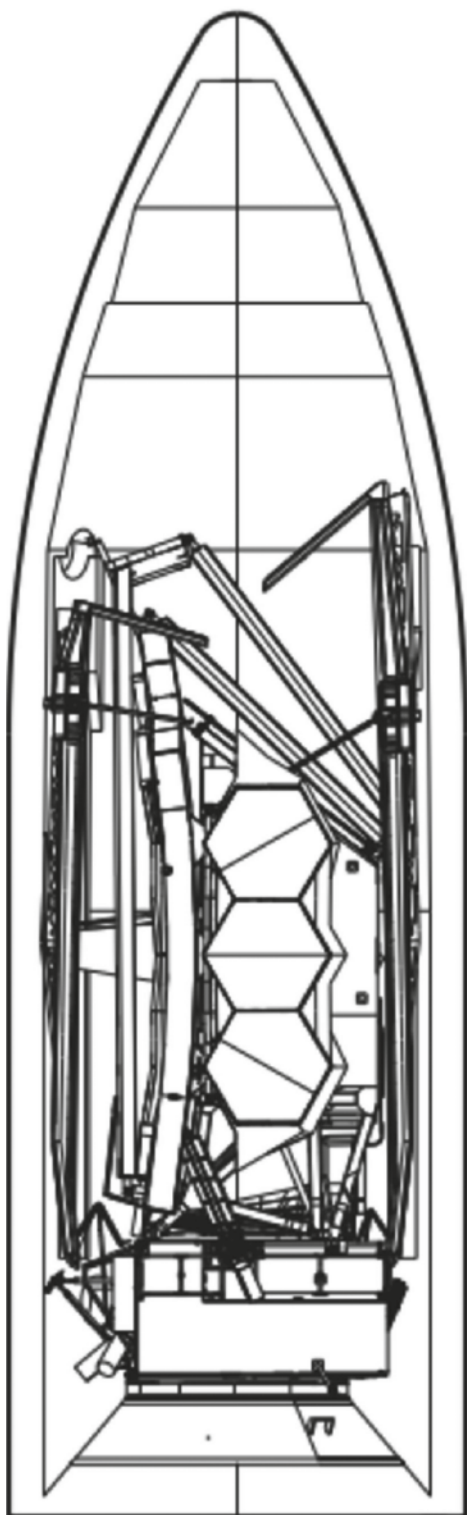
Przygotowania do wysyłki obserwatorium zakończyły się we wrześniu. Potem był czas na drogę z siedziby Northrop Grumman w Kalifornii, przez Kanał Panamski, aż do Gujany Francuskiej w Ameryce Południowej. James Webb ma zostać wyrzuty na rakiety Ariane 5. Podawano dzień 18 grudnia, ale dokładna data może się jeszcze zmienić. W czasie gdy są w toku przygotowania do startu, zespoły znajdujące się w Mission Operations Center (MOC) Webba w Space Telescope Science Institute (STScI) w Baltimore przeprowadzają intensywne testy sieci komunikacyjnej, z której Webb będzie korzystał w przestrzeni kosmicznej.

Teleskop złożony jak origami (2) rozłoży się w przestrzeni kosmicznej podczas podróży na swoją orbitę, prawie półtora miliona kilometrów od Ziemi. Jeśli wszystko pójdzie zgodnie z planem, to programy badań naukowych rozpoczną się około sześć miesięcy po starcie.

Gdy Webb dotrze do Gujany Francuskiej, zespoły zajmujące się obsługą startu skonfigurują obserwatorium do lotu. Będzie to kolejna seria kontroli, aby upewnić się, że obserwatorium nie zostało uszkodzone podczas transportu i usuwanie zabezpieczeń transportowych. Następnie zespoły inżynierów połączą obserwatorium z rakiętą Ariane 5 (3),

2. Skompletowany, przetestowany i w pełni zmontowany Kosmiczny Teleskop Jamesa Webba





3. Konfiguracja startowa JWST wewnątrz rakiety Ariane 5

dostarczoną przez ESA (Europejską Agencję Kosmiczną), zanim wyruszy ono na stanowisko startowe. Webb jest międzynarodowym programem prowadzonym przez NASA wraz z jej partnerami, ESA (Europejską Agencją Kosmiczną) i Kanadyjską Agencją Kosmiczną.

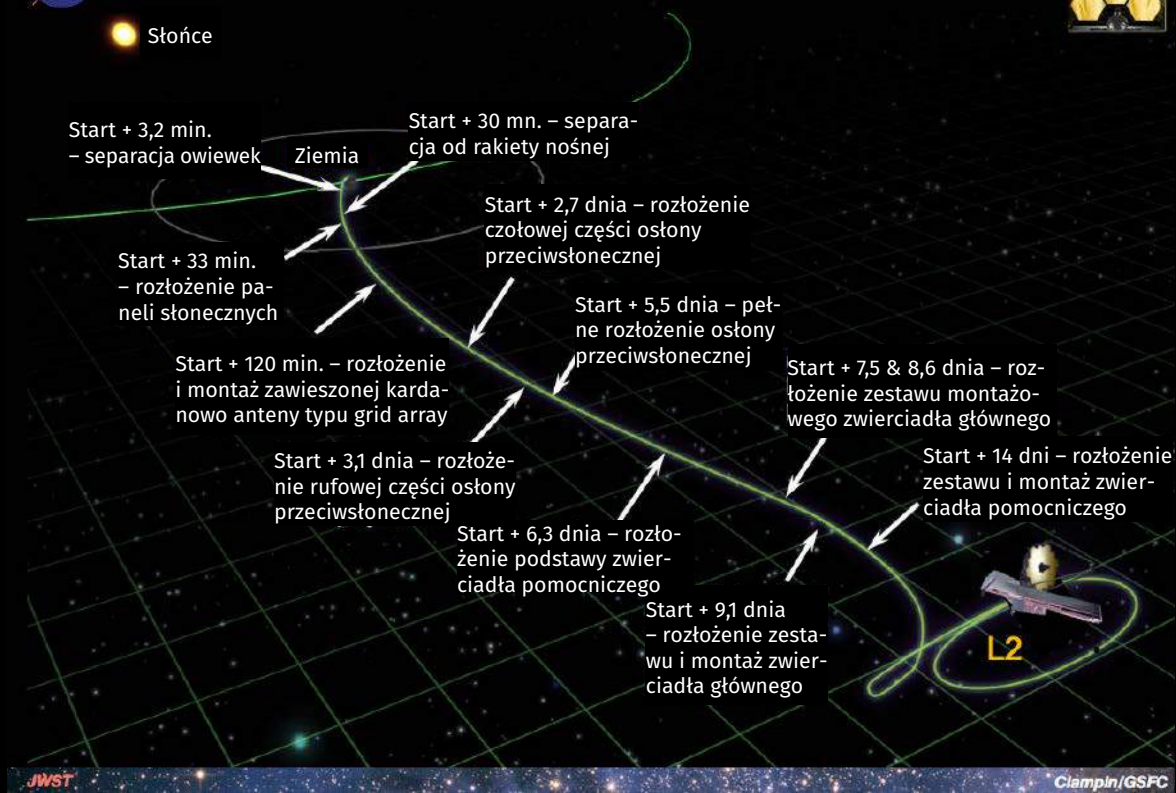
Po wystrzeleniu obserwatorium przejdzie pełen krytycznych wydarzeń sześciomiesięczny okres rozruchu. Chwilę po zakończeniu 26-minutowego lotu na pokładzie rakiety Ariane 5, statek kosmiczny odłączy się od rakiety, a jego panele słoneczne zostaną automatycznie rozłożone. Wszystkie kolejne operacje (4) w ciągu następnych kilku tygodni będą inicjowane z kontroli naziemnej znajdującej się w STScI. JWST będzie potrzebował miesiąca, aby dolecieć do planowanej pozycji orbitalnej w przestrzeni kosmicznej, ok. półtora miliona km od Ziemi, powoli rozwijając kolejne układy w trakcie lotu.

Po ustawieniu JWST przejdzie przez proces rozmieszczania osłony przeciwsłonecznej, zwierciadła i ramienia, co zajmie około trzech tygodni. Gdy osłona przeciwsłoneczna zacznie się rozsuwać, teleskop i instrumenty znajdą się w cieniu i zaczną się ochładzać. Przez następne tygodnie zespół misji będzie ściśle monitorował proces stygnięcia obserwatorium, kontrolując naprężenia na instrumentach i strukturach. W tym samym czasie rozłożony zostanie statyw z lustrem wtórnym, nastąpi również rozłożenie lustra pierwotnego, instrumenty Webba będą powoli włączane, a silniki napędowe wprowadzą obserwatorium na wyznaczoną orbitę. Prawie miesiąc po starcie, zostanie zainicjowana korekta trajektorii, aby umieścić JWST na orbicie Halo w punkcie Lagrange'a L2. Gdy obserwatorium ostygnie i ustabilizuje się w niskiej temperaturze roboczej, przez kilka miesięcy będzie trwało dostrajanie jego optyki i kalibrowanie instrumentów naukowych.

„Teleskop, który zjadł astronomię”

Nazwa Kosmicznego Teleskopu Jamesa Webba jest uhonorowaniem drugiego w historii administratora NASA, który kierował agencją w latach 1961–1968, gdy ta pracowała nad lądowaniem ludzi na Księżycu. Krytycy Webba twierdzą, że był on współwinny dyskryminacji gejów i lesbijek wśród pracowników NASA podczas swojej kadencji. Pojawiły się nawet ostatnio postulaty by zmienić nazwę misji, ale NASA uznała te oskarżenia i pomysły za bezzasadne.

Wczesne prace nad następcą Hubble'a, wyniesionego na orbitę w 1990 roku, prowadzone w latach



4. Harmonogram umieszczenia i rozłożenia teleskopu JWST na orbicie

1989–1994, doprowadziły do powstania koncepcji teleskopu Hi-Z, teleskopu na podczerwień o czterometrowej aperturze, który miałby zostać umieszczony na orbicie w odległości 3 jednostek astronomicznych. Ta odległa orbita korzystać miała z niskiego szumu świetlnego. Inne wcześnie plany zakładały misję teleskopu prekursorskiego NEXUS.

W erze „szybciej, lepiej, taniej” w połowie lat 90., liderzy NASA naciskali na stworzenie taniego teleskopu kosmicznego. Wynikiem tego była koncepcja NGST, z 8-metrową aperturą, zlokalizowanego na L2, którego koszt szacowano na 500 milionów dolarów. W 1999 roku NASA wybrała Lockheed Martin i TRW do przeprowadzenia wstępnych badań koncepcyjnych. Start planowano wówczas na 2007 rok, ale data startu była później wielokrotnie przesuwana. W 2003 roku NASA przyznała firmie TRW kontrakt na wykonanie NGST, obecnie nazwanego Kosmicznym Teleskopem Jamesa Webba, o wartości 824,8 mln USD. Projekt przewidywał zwierciadło główne o średnicy 6,1 m i datę wyniesienia na orbitę w 2010 r. Jeszcze w tym samym roku TRW zostało przejęte przez Northrop Grumman

i przekształciło się w Northrop Grumman Space Technology.

Wzrost kosztów ujawniony wiosną 2005 r. doprowadził do ponownego planowania w sierpniu 2005 r. Głównymi technicznymi rezultatami ponownego planowania były znaczące zmiany w planach integracji i testów, 22-miesięczne opóźnienie startu (od 2011 do 2013 r.) oraz wyeliminowanie testów na poziomie systemu dla trybów pracy obserwatorium przy długości fali krótszej niż 1,7 mikrometra. Inne główne cechy obserwatorium pozostały niezmienione. W ponownym planie z 2005 roku, koszt cyklu życia projektu został oszacowany na około 4,5 mld USD. Składało się na to około 3,5 miliarda dolarów na projekt, rozwój, uruchomienie i oddanie do użytku oraz około miliarda dolarów na dziesięć lat eksploatacji. ESA wносиła około 300 milionów euro, w tym na uruchomienie. Kanadyjska Agencja Kosmiczna zobowiązała się do przekazania 39 milionów dolarów kanadyjskich w 2007 roku, a w 2012 roku dostarczyła swój wkład w postaci sprzętu do wycelowania teleskopu i wykrywania warunków atmosferycznych na egzoplanetach.



W kwietniu 2010 r. teleskop przeszedł pomyślnie techniczną część przeglądu Mission Critical Design Review (MCDR). Pozytywny wynik MCDR oznaczał, że zintegrowane obserwatorium może spełnić wszystkie naukowe i inżynierskie wymagania dla swojej misji. Harmonogram projektu został poddany przeglądowi, co doprowadziło do zmiany planu misji. Wówczas założono start w 2015 roku, ale nie później niż w 2018 roku.

W 2011 roku projekt JWST był już w końcowej fazie projektowania i produkcji (Faza C). Montaż sześciokątnych segmentów zwierciadła pierwotnego, który odbywał się za pomocą robotycznego ramienia, rozpoczął się w listopadzie 2015 r. i został zakończony w lutym 2016 r. Budowa teleskopu Webba została zakończona w listopadzie 2016 r., po czym rozpoczęły się szeroko zakrojone procedury testowe.

W marcu 2018 r., NASA opóźniła start JWST o dodatkowy rok, do maja 2020 r., po tym jak osłona słoneczna teleskopu rozerwała się podczas praktycznego rozkładania – linki osłony słonecznej nie napięły się wystarczająco mocno. W czerwcu 2018 roku NASA opóźniła start JWST o dodatkowe dziesięć miesięcy, do marca 2021 roku, opierając się na ocenie niezależnej komisji przeglądowej zwołanej po nieudanych testowym rozmieszczeniu z marca 2018 roku. Przegląd wykazał także, że JWST ma 344 potencjalne awarie w pojedynczych punktach, z których każda może zakończyć projekt. W sierpniu 2019 roku zakończono mechaniczną integrację teleskopu, co było opóźnione o 12 lat w stosunku do pierwotnego planu. Do października 2019 r. szacowany koszt projektu osiągnął 10 mld USD.

Oprócz licznych problemów technicznych pojawiały się też polityczne. W lipcu 2011 r. komisja ds. handlu, sprawiedliwości i nauki Izby Reprezentantów Stanów Zjednoczonych podjęła działania w celu anulowania projektu Teleskopu Kosmicznego Jamesa Webba, usuwając przeznaczone nań pieniądze z budżetu NASA. W tamtym momencie JWST kosztował już 3 mld USD, a trzy czwarte potrzebnego w obserwatorium sprzętu było w produkcji. Komisja zarzuciła projektowi „przekroczenie budżetu o miliardy dolarów i złe zarządzanie”. W odpowiedzi Amerykańskie Towarzystwo Astronomiczne wydało oświadczenie popierające JWST. Także w prasie międzynarodowej również pojawiło się wiele artykułów popierających JWST. W listopadzie 2011 roku Kongres wycofał się z planów anulowania JWST i zamiast

tego ograniczył dodatkowe fundusze na dokończenie projektu do 8 mld USD. W lutym 2019 roku, pomimo krytyki wzrostu kosztów, Kongres zwiększył limit kosztów misji o 800 mln USD.

Pomimo wsparcia projektu przez środowiska naukowe niektórzy naukowcy wyrażali zaniepokojenie rosnącymi kosztami i opóźnieniami teleskopu Webba, który tym samym zagraża finansowaniu innych programów nauki o kosmosie. Artykuł opublikowany w „Nature” w 2010 roku opisywał JWST jako „teleskop, który zjadł astronomię”.

Do powstania JWST przyczyniło się kilka tysięcy naukowców, inżynierów i techników z piętnastu krajów. W projekcie uczestniczy 258 firm, agencji rządowych i instytucji akademickich – 142 z USA, 104 z krajów europejskich i 12 z Kanady.

Obserwatorium jest przymocowane do rakiety Ariane 5 za pomocą pierścienia adaptacyjnego, który może być użyty przez przyszły statek kosmiczny do uchwycenia obserwatorium w celu usunięcia poważnych problemów jakie mogłyby się pojawić w trakcie jego rozkładania. Jednak sam teleskop nie jest serwisowalny i astronauta nie mogliby wykonywać takich zadań jak wymiana instrumentów, jak to miało miejsce w przypadku Teleskopu Hubble’a. Zatem, jeśli ktoś wie, ile to kosztowało, rozumie poziom napięcia odpowiedzialnych za misję.

Nominalny czas misji wynosi pięć lat, docelowo ma wynosić dziesięć lat. JWST musi korzystać z paliwa do utrzymania swojej orbity halo wokół L2. Zapas paliwa może wystarczyć na dziesięć lat i to jest górna granica czasowa misji.

W podczerwieni widać więcej, ale tylko gdy jest zimno

Kosmiczny Teleskop Jamesa Webba jest zoptymalizowany do oglądania kosmosu w zakresie światła podczerwonego. Główne zwierciadło JWST składa się z osiemnastu sześciokątnych segmentów lustrzanych wykonanych z połączanego berylu. Razem tworzą one zwierciadło o średnicy 6,5 metra, prawie trzy razy szersze niż Hubble’a z sześciokrotnie większą powierzchnią zbierającą, 25,4 m² (5). Teleskop Hubble’a miał niesamowitą passę odkryć, ale nie pozostało mu wiele czasu. Ostatnia misja serwisowa Hubble’a miała miejsce w 2009 roku, a problemy stają się coraz częstsze, jak na przykład kilka miesięcy temu, kiedy obserwatorium przestało działać z powodu awarii zasilania. Po uruchomieniu Teleskop Webba będzie mógł kontynuować pracę tam, gdzie zakończył ją Hubble, obserwując



5. Porównanie teleskopów Hubble i Webba – rozmiary ich zwierciadeł

bardziej odległe i ciemniejsze obiekty większemu zwierciadłu. Zapewni lepszą rozdzielczość i czułość w podczerwieni niż Hubble.

Hubble prowadzi obserwacje w bliskim ultrafiolecie, świetle widzialnym i bliskiej podczerwieni (0,1 do 1 μm). JWST będzie prowadził obserwacje w niższym zakresie częstotliwości, od światła widzialnego w zakresie fal dłuższych, przez bliską podczerwień, do średniej podczerwieni (0,6 do 28,3 μm), co pozwoli mu na obserwacje obiektów o dużym przesunięciu ku czerwieni, które są zbyt stare i zbyt odległe, by mógł je zaobserwować Hubble. Teleskop musi być utrzymywany w bardzo niskiej temperaturze, aby móc prowadzić obserwacje w podczerwieni bez zakłóceń, dlatego zostanie umieszczony w przestrzeni kosmicznej w pobliżu punktu Lagrange'a Słońce-Ziemia L2. Obiekty znajdujące się w pobliżu tego punktu mogą okrążać Słońce w synchronizacji z Ziemią, co pozwala teleskopowi pozostać w mniej więcej stałej odległości i używać pojedynczej osłony słonecznej do blokowania ciepła i światła ze Słońca i Ziemi. Ponieważ L2 jest tylko punktem równowagi bez przyciągania grawitacyjnego, orbita halo nie jest orbitą w zwykłym sensie. Statek kosmiczny

jest w rzeczywistości na orbicie wokół Słońca, a orbitę halo można traktować jako kontrolowane dryfowanie w celu pozostania w pobliżu punktu L2. Wymaga to pewnych korekcyj położenia stacji, do czego służyć będą pędniki.

Duża osłona słoneczna JWST wykonana z kaptonu, folii poliimidowej pokrytej krzemem i aluminium, utrzyma jego zwierciadło i instrumenty w temperaturze poniżej 50 K (-223°C). Przypadkowe rozdarcia delikatnej struktury folii podczas testów w 2018 roku były jednym z czynników opóźniających projekt. Osłona przeciwsłoneczna została zaprojektowana tak, aby można ją było złożyć dwanaście razy, dzięki czemu zmieści się w owiewce ładunku użytecznego rakiety Ariane 5 (4,57×16,19 m). Po umieszczeniu w punkcie L2 osłona rozłoży się do rozmiarów 14,162×21,197 m. Osłonę przeciwsłoneczną zmontowano ręcznie w firmie ManTech (NeXolve) w Huntsville w Alabamie, a potem dostarczono ją do Northrop Grumman w Redondo Beach w Kalifornii w celu przeprowadzenia testów.

Ze szkodliwym nagrzewaniem radzono sobie różnie, choćby przez umieszczanie teleskopu kosmicznego w tzw. naczyniu Dewara, pojemniku



izotermicznym z ekstremalnie zimną substancją, np. ciekłym helem. Oznacza to, że większość teleskopów na podczerwień ma żywotność ograniczoną przez dostępność chłodziwa, kilka miesięcy, a maksymalnie kilka lat. W niektórych przypadkach udało się utrzymać temperaturę na tyle niską dzięki konstrukcji statku kosmicznego, by umożliwić obserwacje w bliskiej podczerwieni bez konieczności dostarczania chłodziwa, jak w przypadku przedłużonych misji Spitzer i Wide-field Infrared Survey Explorer. Innym przykładem jest instrument Hubble'a Near Infrared Camera and Multi-Object Spectrometer (NICMOS), który początkowo korzystał z bloku lodu azotowego, który wyczerpał się po kilku latach, ale następnie został zamieniony na kriokomórkę, która pracowała bez przerwy. Teleskop Kosmiczny Jamesa Webba został zaprojektowany tak, aby chłodził się bez naczynia Dewara, wykorzystując kombinację osłon słonecznych i promienników, przy czym instrument pracujący w średniej podczerwieni korzysta z dodatkowego chłodzącego układu kriokomórkowego.

Projekt kładzie nacisk na bliską i średnią podczerwień z trzech głównych powodów:

- obiekty o dużym przesunięciu ku czerwieni mają swoje emisje widzialne przesunięte w podczerwieni
- zimne obiekty, takie jak dyski i planety emitują najsilniej w podczerwieni
- pasmo to jest trudne do zbadania z ziemi lub przez istniejące teleskopy kosmiczne, takie jak Hubble

Promieniowanie podczerwone może przechodzić przez obszary pyłu kosmicznego, które rozpraszają światło widzialne. Obserwacje w podczerwieni pozwalają na badanie obiektów i regionów przestrzeni, które byłyby przesłonięte przez gaz i pył w widmie widzialnym, takich jak obłoki molekularne, w których rodzą się gwiazdy, dyski okołobiegunowe, z których powstają planety, oraz rdzenie aktywnych galaktyk. Także stosunkowo chłodne obiekty (o temperaturze poniżej kilku tysięcy stopni) emitują swoje promieniowanie głównie w podczerwieni. W rezultacie, większość obiektów chłodniejszych od gwiazd lepiej badać w tym zakresie. Do planów obserwacyjnych dołączają więc brązowe karły, planety zarówno w naszym jak i innych układach słonecznych, komety oraz obiekty pasa Kuipera, które będą obserwowane za pomocą instrumentu MIRI (Mid-Infrared Instrument).

Pojedyncze tak duże zwierciadło jak to w JWST byłoby zbyt duże dla istniejących rakiet nośnych,

stąd segmentowa konstrukcja. Do ustawienia segmentów zwierciadła we właściwym położeniu zostanie wykorzystana detekcja czoła fali w płaszczyźnie obrazu poprzez odzyskiwanie fazy za pomocą bardzo precyzyjnych mikrosiłników w liczbie 126. Po tej wstępnej konfiguracji będą one wymagały jedynie sporadycznych aktualizacji co kilka dni, aby zachować optymalne skupienie. Jest to rozwiązanie inne niż stosowane w teleskopach naziemnych, na przykład w obserwatorium Kecka, które nieustannie dostosowują segmenty zwierciadła za pomocą optyki aktywnej, aby przezwyciężyć efekty grawitacji i wiatru. Konstrukcja optyczna JWST to trójlustrowy anastygmat, który wykorzystuje zakrzywione lustra wtórne i trzeciorzędowe, aby dostarczyć obrazy wolne od aberracji optycznych w szerokim polu. Dodatkowo, istnieje szybkie zwierciadło sterujące, które może zmieniać swoją pozycję wiele razy na sekundę, aby zapewnić stabilizację obrazu.

Osiemnaście segmentów zwierciadła głównego, zwierciadła wtórne, trzeciorzędowe i lustra precyzyjnego sterowania oraz części zamienne zostały wyprodukowane i wypolerowane przez Ball Aerospace & Technologies na podstawie berylowych półfabrykatów segmentów wyprodukowanych przez kilka firm, w tym Axsys, Brush Wellman i Tinsley Laboratories.

Super-aparatura dla super-teleskopu

JWST ma na pokładzie szereg instrumentów naukowych, z których najważniejsze umieszczone zostały w Zintegrowanym Module Instrumentów Naukowych (Integrated Science Instrument Module – ISIM). Jest to konstrukcja, która zapewnia teleskopowi Webba zasilanie elektryczne, zasoby obliczeniowe, możliwość chłodzenia, a także stabilność strukturalną. Jest wykonana z kompozytu grafitowo-epoksydowego przymocowanego do spodniej części struktury teleskopu. W ISIM znajdują się cztery instrumenty naukowe i kamera prowadząca.

Pierwszy z tych instrumentów to NIRCam (Near InfraRed Camera), kamera obrazująca w podczerwieni, która będzie miała pokrycie spektralne od krawędzi światła widzialnego (0,6 mikrometra) do bliskiej podczerwieni (5 mikrometrów). NIRCam będzie również służyć jako czujnik czoła fali obserwatorium, który jest wymagany do działań kontrolnych.

Z kolei MIRI (Mid-InfraRed Instrument) będzie pracował w zakresie fal od średniej do długiej

podczerwieni, czyli od 5 do 27 mikrometrów. Zawiera zarówno kamerę na średnią podczerwień, jak i spektrometr obrazowy. MIRI został opracowany w ramach współpracy NASA z konsorcjum krajów europejskich. MIRI ma podobne mechanizmy koła jak NIRSpec, które również zostały opracowane i zbudowane przez Carl Zeiss Optonics GmbH. Temperatura MIRI nie może przekraczać 6 kelwinów (K). Jego chłodzenie to zapewnia mechaniczna chłodnica helowa umieszczona po ciepłej stronie osłony środowiskowej. NIRCам i MIRI posiadają korony blokujące światło gwiazdowe do obserwacji słabych obiektów, takich jak planety pozasłoneczne i dyski okołobiegunowe bardzo blisko jasnych gwiazd.

NIRSpec (Near InfraRed Spectrograph) będzie wykonywał spektroskopię w tym samym zakresie długości fal. Został on zbudowany przez Europejską Agencję Kosmiczną w ESTEC w Noordwijk w Holandii.

Kanadyjskim wkładem jest FGS/NIRISS (Fine Guidance Sensor and Near Infrared Imager and Slitless Spectrograph), używany do stabilizacji linii widzenia obserwatorium podczas obserwacji naukowych. Pomiary dokonywane przez FGS są wykorzystywane zarówno do kontroli ogólnej orientacji statku kosmicznego, jak i do operowania precyzyjnym zwierciadłem sterującym w celu stabilizacji obrazu. Kanadyjska Agencja Kosmiczna dostarcza również moduł Near Infrared Imager and Slitless Spectrograph (NIRISS) do obrazowania astronomicznego i spektroskopii w zakresie długości fali od 0,8 do 5 mikrometrów.

Zespół inżynierów Zintegrowanego Modułu Instrumentów Naukowych (ISIM) i Obsługi Dowodzenia i Danych (ICDH) Teleskopu Kosmicznego Jamesa Webba (JWST) do przesyłania danych pomiędzy instrumentami naukowymi a urządzeniami do obsługi danych używa SpaceWire, sieci komunikacyjnej statków kosmicznych, koordynowanej przez Europejską Agencję Kosmiczną.

Głównym elementem nośnym Teleskopu Kosmicznego Jamesa Webba jest „magistrala” (ang. „Spacecraft Bus”). Region 3 modułu ISIM, który zawiera podsystem dowodzenia i obsługi danych ISIM oraz kriokomórkę MIRI, znajduje się wewnątrz magistrali. Konstrukcja Spacecraft Bus waży 350 kg i musi utrzymać 6,5-tonowy teleskop kosmiczny. Wykonana jest głównie z grafitowego materiału kompozytowego. Magistrala znajduje się po „ciepłej” stronie zwróconej ku Słońcu i działa w temperaturze około 300 K.

Wszystko, co znajduje się po stronie zwróconej ku Słońcu, musi być w stanie wytrzymać warunki termiczne orbity halo JWST, która ma jedną stronę w ciągłym świetle słonecznym, a drugą w cieniu osłony słonecznej statku kosmicznego. Ważnym modułem magistrali kosmicznej są centralne urządzenia obliczeniowe, pamięciowe i komunikacyjne. Procesor i oprogramowanie kierują dane do i z instrumentów, do rdzenia pamięci półprzewodnikowej oraz do systemu radiowego, który może wysyłać dane na Ziemię i odbierać polecenia. Komputer kontroluje również orientację i moment obrotowy statku kosmicznego, pobierając dane z czujników żyroskopów i trackera gwiazdowego oraz wysyłając niezbędne komendy do kół reakcyjnych lub pędników.

Szerokość pasma i przepustowość cyfrowa satelity została zaprojektowana na 458 gigabitów danych dziennie przez cały okres trwania misji. Większość przetwarzania danych w teleskopie jest wykonywana przez konwencjonalne komputery. Konwersja analogowych danych naukowych do postaci cyfrowej jest wykonywana przez specjalnie zbudowany układ SIDECAR ASIC (System for Image Digitization, Enhancement, Control And Retrieval Application Specific Integrated Circuit).

Chcesz obserwować za pomocą JWST? Złóż wniosek

Do funkcji centrum naukowo-operacyjnego dla JWST wybrano Space Telescope Science Institute (STScI), zlokalizowany w Baltimore, stanie Maryland, na kampusie Homewood Uniwersytetu Johns Hopkinsa. Instytut będzie odpowiedzialny za naukową obsługę teleskopu i dostarczanie danych społeczności astronomicznej. Dane będą przesyłane z JWST na ziemię przez NASA Deep Space Network, przetwarzane i kalibrowane w STScI, a następnie dystrybuowane online do astronomów na całym świecie. Podobnie jak w przypadku Hubble'a, każdy, z dowolnego miejsca na świecie, będzie mógł składać propozycje obserwacji. Każdego roku kilka komitetów naukowych złożonych z wybitnych astronomów będzie recenzować nadesłane propozycje, aby wybrać projekty obserwacyjne na nadchodzący rok. Autorzy wybranych propozycji będą mieli zazwyczaj przez rok wyłączny dostęp do nowych obserwacji, po czym dane staną się publicznie dostępne do pobrania przez każdego z internetowego archiwum STScI.

Wybór programów obserwacyjnych w Cyklu 1 został ogłoszony już 30 marca 2021 roku.



6. Gwiazda Beta Pictoris

Zatwierdzono 266 propozycji, w tym 13 dużych programów.

W sensie ogólnym już dość dawno wyznaczono Kosmicznemu Teleskopowi Jamesa Webba cztery główne cele badawcze. Należą do nich

- poszukiwanie światła z pierwszych gwiazd i galaktyk, które uformowały się we Wszechświecie po Wielkim Wybuchu;
- badanie formowania się i ewolucji galaktyk;
- pogłębianie wiedzy procesach powstawania gwiazd i układów planetarnych;
- badanie systemów planetarnych pod kątem możliwości powstania i utrzymania się w nich życia.

Gdyby zaś mówić o bardziej szczegółowych planach i projektach naukowych to w ostatnim czasie pojawił się szereg publikacji zapowiadających tego rodzaju obserwacje, gdy JWST zostanie uruchomiony i udostępniony.

Naukowcy chcą go na przykład wykorzystać do zbadania gwiazdy Beta Pictoris (β Pic) w gwiazdozbiórze Malarza (6), która znajduje się w odległości ok. 63,4 lat świetlnych od Słońca i znana jest z tego, że ma intrygujący młody układ planetarny z najmniej dwiema planetami, mieszaniną mniejszych, skalistych, obiektów, egzokometami oraz dyskiem pyłowym. Kosmiczny Teleskop Jamesa Webba NASA będzie obserwował Beta Pictoris w podczerwieni, zarówno za pomocą swoich koronagrafów, jak i rejestrując widma,

co pozwoli naukowcom dowiedzieć się znacznie więcej o gazie i pyłe w dysku.

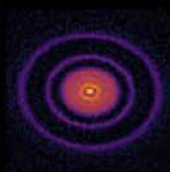
Beta Pictoris jest celem kilku planowanych programów obserwacyjnych Webba, w tym jednego prowadzonego przez Chrisa Starka z Goddard Space Flight Center NASA i dwóch prowadzonych przez Christine Chen z Space Telescope Science Institute w Baltimore, Maryland. Program Starka będzie bezpośrednio obrazował system po zablokowaniu światła gwiazdy, aby zebrać dane na temat jej pyłu. Program Chen zbierze widma, które wskażą, jakie pierwiastki są w tam obecne. Naukowcy mają nadzieję, że uda się ocenić, jak bardzo system ten jest podobny do naszego Układu Słonecznego, co pomoże nam ocenić wyjątkowość naszego Układu.

Kolejny program badawczy, który wykorzysta możliwości teleskopu Webba to planowane obserwacje siedemnastu znajdujących się obecnie w fazie formowania układów planetarnych. Wybrane układy były wcześniej badane przez obserwatorium Atacama Large Millimeter/submillimeter Array (ALMA) w Chile w 2018, największy radioteleskop na świecie. Do badań za pomocą JWST wybrano 17 z 20 dysków protoplanetarnych obserwowanych przez ALMA (7). Dyski protoplanetarne, które będą obserwowane w tym programie są bardzo jasne i znajdują się stosunkowo blisko Ziemi, co czyni je doskonałymi obiektami do badań.

Webb zmierzy widma cząsteczek w wewnętrznych regionach tych protoplanetarnych dysków.



AS 205



AS 209



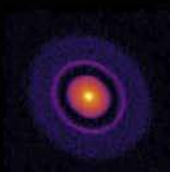
DoAr 25



DoAr 33



Elias 20



Elias 24



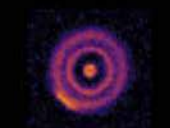
Elias 27



GW Lup



HD 142666



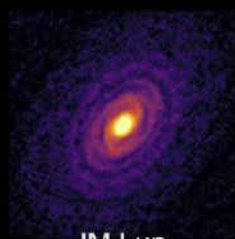
HD 143006



HD 163296



HT Lup



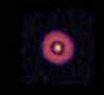
IM Lup



MY Lup



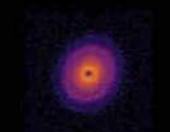
RU Lup



SR 4



Sz 114



Sz 129



WaOph 6



WSB 52

7. Dyski protoplanetarne, które obserwować będzie teleskop Jamesa Webba

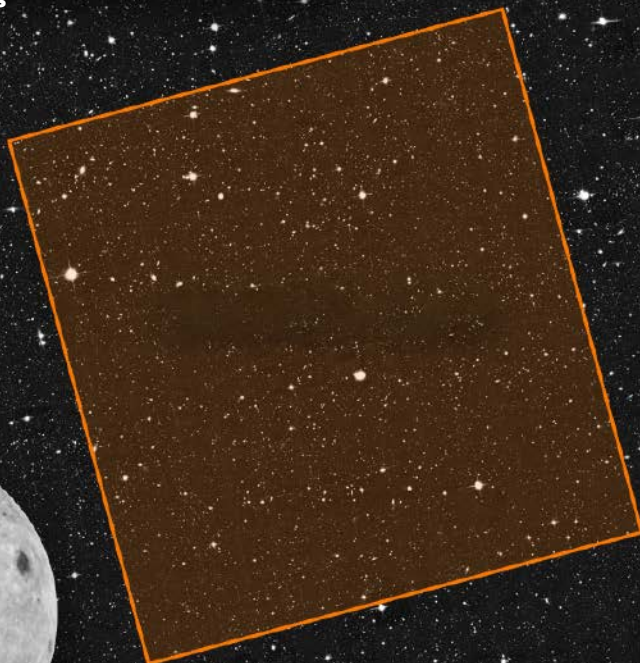
Te regiony to miejsca, gdzie mogą zacząć formować się skaliste, podobne do Ziemi planety. Zespół badawczy kierowany przez Colette Salyk z Vassar College w Poughkeepsie, w Nowym Jorku, i Klausa Pontoppidana z Space Telescope Science Institute w Baltimore, Maryland, chce ocenić ilość wody, tlenku węgla, dwutlenku węgla, metanu i amoniaku, oraz wielu innych cząsteczek, w każdym z tych dysków. Co ważne, będą

możliwe oszacowanie zawartości cząsteczek, które zawierają pierwiastki niezbędne do życia, jakie znamy, w tym tlen, węgiel i azot. Za pomocą spektroskopii uchwycić będzie można całe światło emitowane w centrum każdego dysku protoplanetarnego jako widmo, które tworzy szczegółowy wzór barwny na bazie długości fal emitowanego światła. Ponieważ każda molekula odścisła unikatowy wzór na widmie, naukowcy mogą



Tarcza Księżyca dla porównania

**COSMOS Hubble – na podstawie obserwacji
za pomocą instrumentu Advanced Camera
for Surveys**



8. COSMOS-Webb

stworzyć spis zawartości każdego dysku protoplanetarnego. Siła tych wzorów niesie ze sobą informacje o temperaturze i ilości każdej molekuly.

Międzynarodowy zespół korzystający z Teleskopu Kosmicznego Jamesa Webba będzie również badał Mglavicę Oriona, a konkretnie jej część zwaną Pasem Oriona, aby dowiedzieć się więcej o wpływie masywnych gwiazd na ich otoczenie. „Fakt, że masywne gwiazdy kształtują strukturę galaktyk przez eksplozje jako supernowe jest znany od dawna. Jednak ostatnio odkryto, że masywne gwiazdy wpływają na swoje otoczenie nie tylko jako supernowe, ale także poprzez wiatry i promieniowanie w trakcie swojego życia”, powiedział w komunikacie prasowym jeden z głównych badaczy zespołu, Olivier Berné, naukowiec z francuskiego Narodowego Centrum Badań Naukowych w Tuluzie.

Gdy Kosmiczny Teleskop Jamesa Webba NASA rozpocznie operacje naukowe w 2022 roku, jednym z jego pierwszych zadań będzie także ambitny program mapowania najwcześniejszych struktur we Wszechświecie, nazwany COSMOS-Webb (8). Ten szeroki i głęboki przegląd pół miliona galaktyk jest największym projektem, jaki Webb podejmie w pierwszym roku swojego istnienia. Mając ponad dwieście godzin czasu obserwacyjnego przegląd COSMOS-Webb zmapuje 0,6 stopnia kwadratowego nieba za pomocą kamery bliskiej podczerwieni (NIRCam) – to mniej więcej obszar trzech Księżyców w pełni, jednocześnie mapując



**Operacja montażu
teleskopu Webba
w przestrzeni
kosmicznej:**
<https://bit.ly/3xZ2qET>



9. Wizualizacja kwazara

0,2 stopnia kwadratowego za pomocą instrumentu średniej podczerwieni (MIRI).

Badanie COSMOS rozpoczęło się w 2002 roku jako program Hubble'a mający na celu zobrazowanie znacznie większego skrawka nieba, o powierzchni około dziesięciu księżyców w pełni. Od tego momentu prace w ramach programu nabrały tempa. Zaangażował się w nie większość najważniejszych teleskopów na Ziemi i w kosmosie. Obecnie COSMOS jest przeglądem o wielu długościach fali, który obejmuje całe spektrum od promieniowania rentgenowskiego do radiowego. COSMOS-Webb jest kolejną odsłoną programu, w której używa się możliwości JWST do rozszerzenia (i w pewnym sensie – pogłębienia) horyzontu spojrzenia na Wszechświat.

COSMOS-Webb jest programem klasy „Treasury”, który z definicji ma na celu stworzenie zbiorów danych o trwałej wartości naukowej. Programy Treasury mają na celu rozwiązanie wielu problemów naukowych za pomocą jednego, spójnego zbioru danych. Dane uzyskane w ramach programu Treasury zazwyczaj nie mają okresu wyłącznego dostępu, co umożliwia natychmiastową analizę przez innych badaczy.

W ramach kolejnego ważnego programu badawczego zespół badaczy po wystrzeleniu teleskopu zespół naukowców skieruje Kosmiczny Teleskop Jamesa Webba na sześć najbardziej odległych i jasnych kwazarów, by badać ich właściwości a także zbierać dane na temat goszczących je galaktyk. Zespół wykozysta również kwazary do zbadania gazu w przestrzeni pomiędzy galaktykami, szczególnie w okresie kosmicznej rejonizacji, która zakończyła się, gdy Wszechświat był bardzo młody. Wszystkie te kwazary, które badamy, istniały bardzo wcześnie, kiedy Wszechświat miał mniej niż 800 milionów lat. Tak

więc te obserwacje dają nam możliwość studiowania ewolucji galaktyk oraz formowania się i ewolucji supermasywnych czarnych dziur w bardzo wczesnym okresie istnienia.

Zespół naukowy użyje kwazarów jako źródeł światła tła, aby zbadać gaz znajdujący się pomiędzy nami a kwazarem. Gaz ten pochłania światło kwazara na określonych długościach fal. Za pomocą techniki zwanej spektroskopią obrazową, uczeni będą szukali linii absorpcyjnych w gazie między kwazarami. Im jaśniejszy jest kwazar, tym silniejsze będą te linie absorpcyjne w widmie. Określając, czy gaz jest neutralny czy zjonizowany, naukowcy liczą na zdobycie wiedzy o wczesnych procesach zjonizacji i rejonizacji.

Zespół przeanalizuje światło pochodzące z kwazarów za pomocą NIRSpec w poszukiwaniu pierwiastków cięższych od wodoru i helu. Pierwiastki te powstały w pierwszych gwiazdach i pierwszych galaktykach, a następnie zostały wyrzucone. Gaz przemieszcza się z galaktyk, w których pierwotnie się znajdował, do ośrodka międzygalaktycznego. Zespół planuje zmierzyć wytwarzanie tych pierwszych „metali”, jak również zbadać sposób, w jaki są one wypychane do ośrodka międzygalaktycznego przez te wczesne procesy.

Jak widać, kalendarz badawczy JWST ma już mocno napięty. A to tylko pierwszy rok i zapewne w kolejce czeka wiele innych projektów i programów obserwacji. Miejmy nadzieję, że za jakieś pięć lat będziemy dzięki tej wspaniałemu zaawansowanemu obserwatorium kosmicznemu dużo mądrzej niż jesteśmy obecnie. Oby tylko wszystko poszło jak trzeba w skomplikowanej procedurze umieszczania, rozmieszczania, rozkładania i uruchomienia tej skomplikowanej maszynierii. Trzymamy kciuki. ■

Miroslaw Usidus

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



**do 6* wydań
gratis!**

* Cena prenumeraty rocznej wynosi 130,90 zł.
Przy zamówieniu prenumeraty dwuletniej w cenie 214,20 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 ca98e133b3

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.



1. Almon Brown Strowger

CV: Almon Brown Strowger

Data i miejsce urodzenia: 11.02.1839 (zm. 26.05.1902)

Adres zamieszkania: nie żyje

Obywatelstwo: amerykańskie

Stan cywilny: nie żyje

Majątek: milioner, dokładna wartość majątku nie daje się oszacować

Kontakt: nie żyje

Edukacja: niewiele wiadomo o jego wykształceniu. Musiał je mieć, skoro został nauczycielem

Doświadczenie zawodowe: kawalerzysta, nauczyciel wiejski i przedsiębiorca pogrzebowy, 1891 – założyciel firmy Strowger Automatic Telephone Exchange Company, która posiadała patent na centralę telefoniczną Strowgera

Zainteresowania: kwestie społeczne, lokalna polityka

Nie oszukuj, bo wynajdę coś, co ci się nie spodoba – **Almon Brown Strowger**

Autorem jednej z ważniejszych innowacji w telekomunikacji był wkurzony przedsiębiorca pogrzebowy z Kansas City. Almon Brown Strowger (1) uznał, że traci klientów przez telefonistki obsługujące lokalną centralę, przełączające telefony do konkurencji. I wymyślił mechanizm umożliwiający automatyczne łączenie abonentów. Telefonistki straciły pracę, a przedsiębiorczy wynalazca zwiększył zyski.

Przyszły wynalazca wychował się w małym domu przy Whalen Road w Penfield w stanie Nowy Jork. Rodzice, potomkowie pierwszych osadników w mieście mieli już trzynastkę pociech, gdy Almon pojawił się na świecie 19 października 1839 roku.

Przez lata chłopiec nie wyróżniał się specjalnie ani w zaciszu domowym, ani w grupie rówieśników. Jak na wnuka pierwszego młynarza w mieście przystało miał sporo do czynienia z drobnymi naprawami sprzętu mechanicznego. Ponoć chętniej niż domowymi zadaniami zajmował się rozmyślaniami o maszynach, które uwolniłyby dzieci od uciążliwych obowiązków. Jak na wynalazcę, początek obiecujący, ale do wielkich odkryć było jeszcze daleko.

Lata miały, Almon Brown Strowger zdążył dorosnąć i wybrał karierę wojskową. W wieku 22 lat wstąpił do armii, gdzie dosłużył się stopnia podporucznika. Brał udział w wojnie secesyjnej, m.in. w drugiej bitwie pod Bull Run w pobliżu Manassas w Wirginii.

Skończyć z telefonistkami

Pod koniec 1864 r. otrzymał honorowe zwolnienie z armii i wrócił w rodzinne strony. Weteran wojenny miał się rozmaitych zajęć. Pracował m.in. jako nauczyciel, a następnie był dyrektorem szkoły w Penfield. Po śmierci pierwszej żony Rosalthy porzucił dotychczasowe życie. Opuścił rodzinne strony, kupił przedsiębiorstwo pogrzebowe w Kansas i tam poślubił kolejną żonę – Annę Condon.

Być może nic szczególnego do historii telekomunikacji nie wniósłby i wiódłby stateczne i dostatnie życie, gdyby nie przebiegła konkurencja. W czasach, gdy ludzie dość często umierali z powodu chorób płuc i innych infekcji, firmy funeralne nie narzekały na brak klientów. Jednocześnie konkurencja w tej branży była niemała. W Kansas City, gdzie operował Almon, tak się złożyło, że żona właściciela konkurencyjnego zakładu pogrzebowego pracowała jako telefonistka w lokalnej centrali. W czasach, gdy operator ręcznie łączył rozmowy telefoniczne, przedsiębiorcza dama zyskała w ten sposób możliwość świadczenia pomocy małżonkowi w budowaniu monopolu i przełączała zrozpaczone wdowy i wdowców wprost do firmy męża.

Dla rodzin w żałobie nie miało zapewne większego znaczenia, która firma zajmuje się organizacją pogrzebu. Ale dla Almona stało się to sporym wyzwaniem biznesowym. Najpierw jednak



2. Łącznica Strowgera – rekonstrukcja

przedsiębiorca musiał ustalić przyczyny braku klientów. Koronnym dowodem nieuczciwych praktyk telefonistki stało się przełączenie do konkurencji rodziny zmarłego przyjaciela Almona. Tym razem działanie telefonistki nie mogło zostać przypisane przypadkowi.

Almon Brown Strowger przygotował również perfidny plan zemsty, ale również odzyskania utraconych dochodów. Uznał, że telefonistki ręcznie przekładające wtyczki powinny zniknąć. Postanowił zautomatyzować centralę telefoniczną. Z opracowaniem nowego systemu przełączania rozmów nie szło mu tak gładko. Poszukiwał pomocy wśród członków rodziny. Jeden z braci Almona miał doświadczenie z urządzeniami elektrycznymi, kuzyni dysponowali odpowiednimi środkami finansowymi.

Razem stworzyli w 1892 r. Strowger Automatic Telephone Company. Zanim jednak ruszyli z produkcją Almon Brown Strowger przedstawił w 1888 r. swój pomysł na automatyczną łącznicę (2). Było to okrągłe drewniane pudełko z metalowymi złączkami wykonanymi ze szpilek. Trzy lata później, w 1891 roku, uzyskał patent na swój wynalazek (3), a w kolejnym roku rozpoczęto produkcję łącznicy w warsztatach Strowger Automatic Telephone Company.

Już w listopadzie 1892 r. w mieście La Porte w stanie Indiana została zainstalowana pierwsza komercyjna centrala firmy Strowger, o pojemności dziewięćdziesięciu dziewięciu numerów. Nowy automatyczny system łączenia abonentów opierał się na systemie dziesiętnym. Wysoce innowacyjnym elementem w urządzeniu był tzw. wybierak Strowgera z napędem krokowym oraz trzema

oddzielnymi elektromagnesami -podnoszącym, ob-
racającym oraz zwalniającym. W celu uzyskania
połączenia abonent za pomocą naciskania przyci-
sków w swoim telefonie kontrolował odpowiednią
liczbę impulsów, aby krok po kroku uzyskać po-
łączenie z żądanym innym abonentem. Wybierak
z ramieniem unoszącym się do góry ustawiał się
najpierw w rzędzie, wskazanym przez użytkow-
nika za pomocą 1 przycisku w telefonie, a następ-
nie obracał się zgodnie z ruchem wskazówek zegara
i wskakiwał w odpowiednią szczelinę (wskazaną
przez drugi przycisk telefonu, zaś po zakończeniu
rozmowy wracał do pierwotnej pozycji. Jeśli połą-
czenie nie mogło być zrealizowane, to dzwoniący
słyszał sygnał zajętości. Dwustopniowy mecha-
nizm, podnosząco-obrotowy wybieraka, służył
jednocześnie do zliczania impulsów oraz do ze-
stawienia połączeń. W ten sposób centrala łączyła
bezpośrednio dwóch użytkowników, którzy mo-
gli swobodnie rozmawiać bez obawy o podsłuchi-
wanie przez telefonistki. Wcześniej rozmówcom
usłyszeć nawet przekleństwa i komentarze opera-
torów, kontrolujących przy okazji czas rozmowy
do naliczenia opłaty.

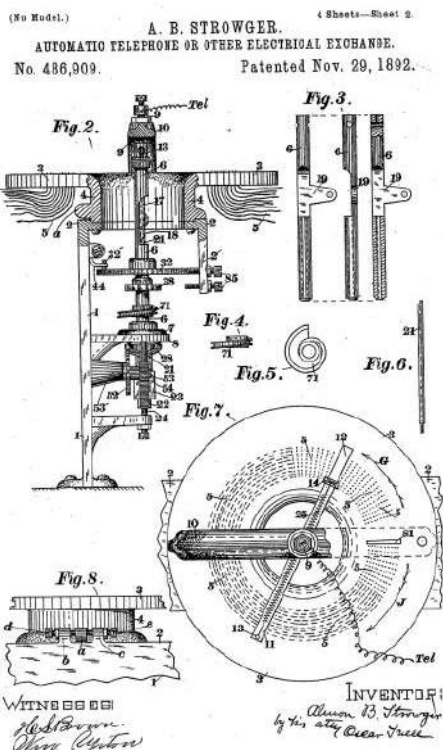
Strowger nie był pierwszym wynalazcą, który
próbował stworzyć automatyczny przełącznik dla
centrali telefonicznej. Ale to jego wybierak okazało
się udanym projektem, a cały system miał dodat-
kową zaletę – umożliwiał rozwijanie bazy klientów.

Innowacyjne rozwiązanie zrewolucjonizowało
telekomunikację. Strowger reklamował łącznicę
jako „nie wymagające telefonistek, wolne od prze-
kleństw, niepsujące się i nie wymagające oczę-
kiwania”. Rozmowy telefoniczne rzeczywiście
realizowano szybko i bez zakłóceń dla połączeń in-
nych użytkowników, ale początkowo mogły obsłu-
giwać tylko rozmowy lokalne. Nawiasem mówiąc
jeszcze przez dłuższy czas telefonistki były po-
trzebne, gdyż wciąż niezbędne do łączenia np. roz-
mów międzymiastowych.

Centrale telefoniczne z wybierakiem Strowgera
zyskiwały popularność, zakładano je w kolejnych
miastach amerykańskich i europejskich. Biznes roz-
wijał się znakomicie, firma wprowadzała kolejne
innowacje m.in. tak charakterystyczną dla telefo-
nów retro tarczę obrotową. Mechanizm, który za-
stąpił przyciski, wymyślił i skonstruował w 1896 r.
współpracownik Strowgera.

Powrót do branży pochówkowej

Choć wszystko układało się pomyślnie,
to Almon Brown Strowger postanowił wycofać się



3. Rysunek z patentu Strowgera

z przedsięwzięcia. W 1896 sprzedał swój patent
współpracownikom za tysiąc osiemset dolarów,
a dwa lata później odstąpił wspólnikom udziały
w Automatic Electric Company za 10 tysięcy.
Po skapitalizowaniu patentów i udziałów osie-
dlił się na Florydzie, gdzie wrócił do prowadze-
nia firmy pogrzebowej. Zyski inwestował głównie
w nieruchomości, z czasem uchodził za człowieka,
do którego należy „pół miasta” St. Petersburg
na Florydzie.

Gdy zmarł 26 maja 1902 r., jego trzecia
żona Susan oskarżyła wspólników o zaniże-
nie wartości udziałów męża w Automatic Electric
Company. W nekrologu umieściła informację, że po-
wstrzymywała męża przed ujawnieniem innych
rewolucyjnych wynalazków, ponieważ tylko osoby
postronne na nich zarabiał. Gdy w 1916 roku Bell
Systems chciał odkupić patenty Strowgera, musiał
zapłacić 2,5 miliona USD.

Almon B. Strowger został przyjęty w 1965 r.
do galerii sław Independent Telephone Association.
Centrale Strowgera były w powszechnym użyciu
jeszcze w latach 70-tych XX wieku. ■

Miroslaw Usidus

Elon Musk zapowiadał kilka miesięcy temu, że satelitarna usługa internetowa Starlink opuści fazę testów beta w październiku 2021. Cokolwiek to oznacza, opinie i kontrowersje wokół tego pełnego rozmachu projektu nie gasną. Obecnie nie dotyczą już tylko problemu przesłaniania nieba przez setki nowych satelitów (1), co krytykowali obserwatorzy gwiazd.

Krzywdza Niebiańskiego Emu

Starlink – nowe oblicze internetu czy utrapienie?

SpaceX dążyła do pełnego pokrycia zasięgiem internetowym całej kuli ziemskiej do września tego roku, więc październik jako koniec fazy beta brzmi logicznie. Informowała też pod koniec sierpnia o wynikach sprzedaży terminali do odbioru sygnału Starlink. Miało być ich 100 tysięcy, głównie w Ameryce Północnej i części Europy, z dodatkami niektórych innych krajów, takich jak Chile, Australia i Nowa Zelandia. Według ogłoszonych publicznie planów, w pierwszym kwartale 2022 liczba użytkowników tej usługi na świecie miałyby sięgnąć pół miliona a docelowo miałyby dotrzeć ona do pięciu procent światowej populacji, czyli prawie 400 mln ludzi.

Sam Musk twierdził podczas prezentacji na Mobile World Congress, że sieć Starlink jest na dobrej drodze do objęcia usługą szerokopasmowego internetu całego świata z wyjątkiem regionów polarnych do sierpnia 2021 roku. Przyznał jednak, że jest niezwykle kosztowna inwestycja. Całkowite nakłady SpaceX zawrą się, według przewidywań, w dość szerokich widełkach od pięciu do dziesięciu miliardów dolarów, zanim bilans sprzedaży zacznie wychodzić na plus.

Musk powiedział, że podpisał dwie umowy z „głównymi krajami” i ich narodowymi operatorami telekomunikacyjnymi, ale nie mógł ich jeszcze wymienić. Ponadto jest w trakcie rozmów z kolejnymi. Sieć satelitarna obecnie przynosi około 30 terabitów danych na sekundę, a Musk



1. Elon Musk i satelity Starlink na niebie

podaje, że jego celem jest opóźnienie w czasie reakcji sieci na poziomie mniej niż 20 milisekund i oferowanie usługi łączności o przepustowości 300 Mb/s do końca 2021 roku. Sierpniowy raport firmy Ookla po serii testów prędkości transmisji danych ujawnił, że rzeczywista prędkość łączności Starlink, wynosząca prawie 14 Mb/s, zbliżyła się do średniej prędkości oferowanej przez dostawców internetu przewodowego szerokopasmowego.

Ciekawą, choć dla interesujących się tematem wcale nie zaskakującą, informacją jest ta o uruchomieniu nowej wersji satelitów Starlink, wyposażonej w inter-satelitarne łącza laserowe które minimalizują opóźnienia czasowe związane z tradycyjnym przekazywaniem danych przez stacje

naziemne. Istnieje plan, by tym typem pokryć rejon polarne Ziemi. Inżynierowie SpaceX pracują także nad nowym, tańszym w sprzedaży detalicznej (obecnie zestaw kosztuje równowartość tysiąca dolarów) terminalem naziemnym (2).

Pierwsze 51 satelitów internetowych Starlink na orbity polarne wystrzelono już we wrześniu 2021 roku, co podbiło łączną liczbę satelitów systemu Starlink do 1 791 jednostek, przy czym ponad 1400 z nich było uważanych za w pełni sprawnych.

Operatorzy telekomunikacyjni protestują

W Stanach Zjednoczonych, gdzie firma Muska działa w komfortowych warunkach. Otrzymała środki finansowe w ramach programów Federalnej Komisji Komunikacji (FCC) mających na celu zapewnienie łączności sieciowej na obszarach wiejskich. W 2018 roku Komisja zatwierdziła plan Starlink dotyczący wysłania na orbitę 12 000 satelitów Starlink. A w 2020 roku firma otrzymała prawie 900 milionów dolarów w ramach Rural Digital Opportunity Fund, federalnego programu wdrożonego podczas pandemii, którego celem jest podłączenie terenów wiejskich i rzadziej zaludnionych do sieci. W Stanach Zjednoczonych, zaludnionych bardzo nierównomiernie, kwestia ta była od dawna problemem, zaś pandemią pokazała tylko, jak palącym.

Z takim dofinansowaniem oferta SpaceX stała się w USA bardzo konkurencyjna. Do tego stopnia, że wielu dostawców usług internetowych złożyło protest do FCC, twierdząc, że technologia stosowana przez firmę ma charakter eksperymentalny, jest niewystarczająco przetestowana i będzie powodować problemy w przyszłości. Zdaniem wielu komentatorów, oznacza to, że tradycyjni gracze sektora telekomunikacyjnego boją się Starlinka i konkurencji, nie tylko na terenach słabo zaludnionych ale również w aglomeracjach.

Kwestia zapewnienia łączności słabo zaludnionym obszarom wiejskim, gdzie rozmieszczenie infrastruktury nie jest ekonomicznie opłacalne, to jeden z największych i najdawniejszych problemów telekomunikacyjnych, nie tylko w USA, lecz na całym świecie. Tereny te, zwłaszcza, gdy pomyśli się o sytuacjach klęsk żywiołowych, akcji ratowniczych, stanowią wyzwanie dla świetnie połączonych, ale tylko w gęsto zaludnionych okolicach, świata.

Ekspansja Starlinka porusza jednak dość delikatne sfery. np. pojawia się pytanie o opłacalność



2. Satelitarny zestaw antenowy do odbioru Starlink

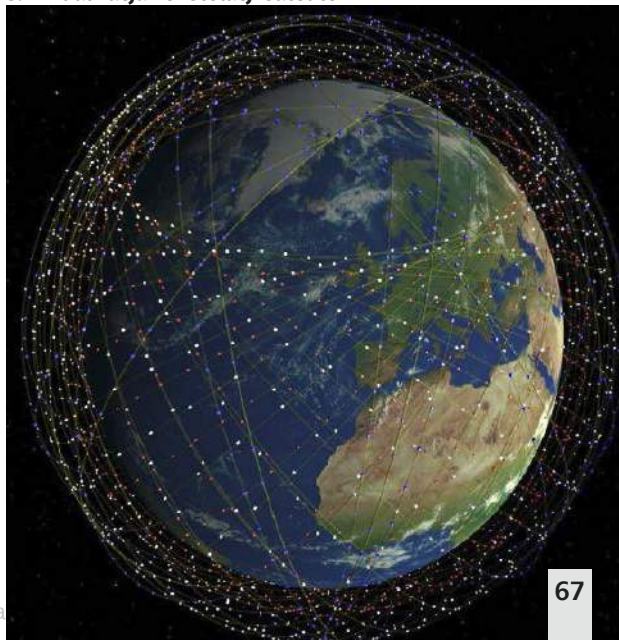
i zwrot inwestycji w zakup licencji na częstotliwości, gdy z góry spada konkurencja internetu satelitarnego, z konkurencyjnymi ofertami na każdym praktycznie rynku. Jest też problem kontroli i cenzury w wielu państwach. Czy np. użytkownicy satelitarnego internetu Starlink mogą być przesładowani za korzystanie z tego sprzętu i usług?

Plemiona protestują

Kontrowersje związane z utrudnianiem przez sieci satelitów Starlink obserwacji astronomicznych opisywaliśmy już w „Młodym Techniku” dość dokładnie. Od tego czasu pojawiło się wiele innych zastrzeżeń z różnych punktów widzenia. O tym, że mogą narobić sporo zamieszania na rynku usług telekomunikacyjnych już wspominaliśmy. Dochodzi do tego krytyka z innych powodów, czasem dość zaskakujących. Megakonstelacje satelitarne (3) są np. od niedawna krytykowane przez przedstawicieli plemion indiańskich lub aborygeńskich za coś co określa się jako „astrokolonializm”.

W kosmologii kanadyjskich Indian Cree, Gwiezdna Kobieta wpadła przez dziurę w niebie zwaną Pakone-Kisik, która wciąż może być

3. Wizualizacja konstelacji satelitów



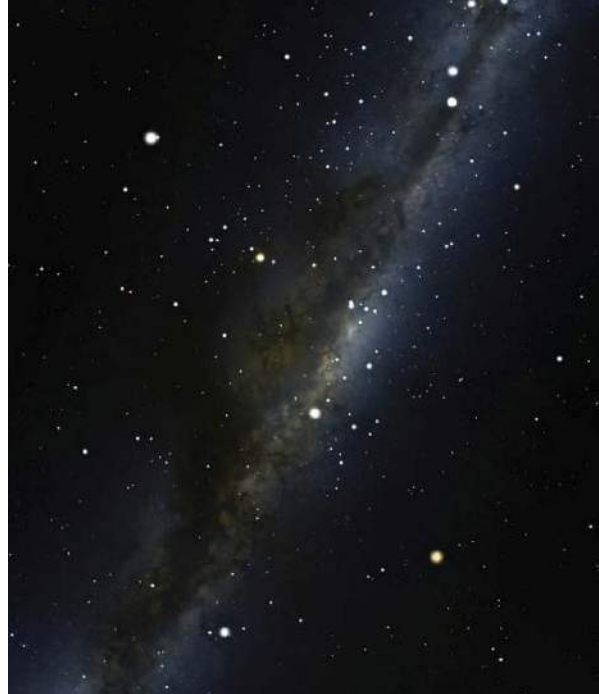
widziana jako skupisko jasnych punktów w przestrzeni. Dla Tainui Māori gwiazdzisty wzór na południowym niebie tworzy Te Punga, kotwicę podniebnego kajaka. W inuickiej legendzie trzej myśliwi są uosobieniem promienistych światła na arktycznym niebie zwanych Ullaktut. My nazywamy te konstelacje odpowiednio Plejady, Krzyż Południa i Pas Oriona, co przez wielu przedstawicieli rdzennych społeczności było już krytykowane jako narzucanie wzorów kulturowych i wymazywanie plemiennej tradycji.

Ogromne sieci orbitalnych statków kosmicznych są postrzegane przez niektórych przedstawicieli społeczności rdzennych jako kontynuacja tego wymazywania. Choć obserwatorzy nieba dostrzegają sztuczne obiekty od zarania ery kosmicznej, to jednak sam natłok nowych satelitów dramatycznie zmienia obraz Wszechświata, którego doświadczały niezliczone pokolenia.

Konstelacja Starlink może, w ocenie tych krytyków, burzyć te pierwotne zakorzenione w plemiennych kulturach pojęcia i obrazy w zaskakujące czasem sposoby. Może np. wymazać pojęcie całkiem ciemnego nieba. Tymczasem np. w aborygeńskiej Australii, konstelacje takie jak Niebiański Emu są tworzone z ciemnych plam w przestrzeni (4). W zanieczyszczonym światłem niebie nie ma miejsca dla tej i wielu innych „ciemnych” konstelacji rdzennych Australijczyków.

Z drugiej strony Starlink program testów beta na obszarach zamieszkałych przez te ludy, zostało z zadowoleniem przyjęte przez wielu ich mieszkańców, którzy przez lata starali się o doprowadzenie szybkich łączy szerokopasmowych w swoich regionach. Było tak m.in. w Pikangikum First Nation w kanadyjskiej prowincji Ontario, Kanada, lub na terenach zamieszkałych przez plemię Hoh w zachodnim stanie Waszyngton, w USA. Gdy na terenach Hoh uruchomiono Starlink we wrześniu 2020 r., prędkości połączeń w społeczności skoczyły z poziomu 0,3 i 0,7 megabitów na sekundę do średnio 45–65 Mb/s, a w niektórych miejscach nawet powyżej 100 Mb/s. Nagle, uczniowie z plemienia Hoh mogli wszyscy jednocześnie uzyskać dostęp do zdalnych platform edukacyjnych. Pojawiły się wcześniej tam nieznane lub ekstremalnie trudno dostępne usługi strumieniowego audio i wideo. Wzrost szybkości łączy ułatwił również życie niektórym starszym członkom plemienia, zaczęły bez problemów działać telefony komórkowe i wiele innych usług.

Zresztą zarówno SpaceX jak i też inne firmy pracujące nad konstelacjami satelitów zaczęły dość



4. Niebiański Emu australijskich Aborygenów

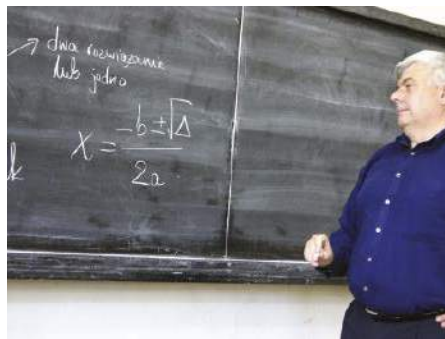
szybko pracować nad rozwiązaniami zmniejszającymi odbłaskowość urządzeń, czyli zanieczyszczanie nieba światłem. W przypadku SpaceX i Amazona są to specjalne pokrycia, zaś OneWeb, który umieszcza swoje satelity na orbicie o wysokości ponad tysiąca kilometrów nie tworzy zasadniczo tego problemu.

Bo nie należy zapominać, że SpaceX i jej Starlink to nie jedyny projekt budowy megakonstelacji satelitarnej. Wspomniana OneWeb, firma z siedzibą w Wielkiej Brytanii, wystrzeliła około połowy planowanej sieci 648 satelitów, zaś Amazon przygotowuje się do uruchomienia własnego „Project Kuiper”, która ma się składać z trzech tysięcy satelitów. W dodatku Chiny rozwijają państwową konstelację o nazwie GW, która może składać się z około 13 tysięcy obiektów.

Każdemu projektowi oczywiście oficjalnie przyświecają szczytne hasła „zanieśienia internetu na cały glob, wszędzie tam gdzie go nie ma a ludzie potrzebują”. Jednak za tymi projektami kryją się równe inne cele, np. Starlink już jest postrzegany w USA jako strategiczna alternatywa na wypadek wyłączenia GPS przez wroga, zaś Chiny, wiadomo, nie zamierzają wpuścić niecenzurowanego internetu, więc zapewne chcą zbudować kosmiczny internet dla swoich obywateli za „wielkim sieciowym firawallem”. I to jest wszystko ta mniej idealistyczna twarz satelitarnych projektów. ■

Miroslaw Usidus

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



Klocki, monety, kasztany i Encyklopedia Ciągów Liczbowych

Obecny odcinek moich Rozmaitości Matematycznych jest streszczeniem zajęć, jakie miałem z uczniami klas ósmych w ramach Krajowego Funduszu na rzecz Dzieci oraz (po kilkunastu dniach) z uczniami Liceum im. Adama Mickiewicza w podwarszawskim Piastowie. Wśród licealistów przeważali uczniowie pierwszej klasy, różnica między nimi a ośmioklasistami była zatem niewielka. I tu i tam chodziło o pokazanie – cokolwiek to znaczy – kolorów matematyki. Pomagała jesienna pogoda, drzewa za oknami mieniły się wszystkimi odcieniami zieleni, żółci i czerwieni.

Uczniowie zdziwili się od razu pierwszym zadaniem. Oto ono:

Zadanie 1. Narysuj $1+2+3+4+5$. Ułóż z klocków.

Po chwili konsternacji (jak można narysować dodawanie? A czy może chodzi o narysowanie liczby 15?) i w jednej i w drugiej grupie uczniów odważny uczeń podszedł do tablicy (ciekawe, że przyszło im to samo do głowy!)

„No, tak, dobrze”, powiedziałem, „ale może ładnie, z kółeczkami?”. Domyślili się, o co mi chodzi. Powstały dwa rysunki. „O świetnie”, ale ułóżmy to z monet. Szybko powstały dwie układanki,



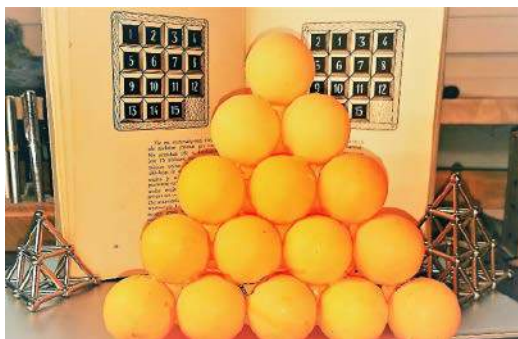
1. Jeden dodać dwa, dodać trzy, dodać cztery, dodać pięć

A ja pokazałem trzecią – z piłeczek pingpongowych, ułożonych na tle książki Hugona Steinhausa „Kalejdoskop matematyczny” i czworościanów ułożonych z pałeczek i kuleczek magnetycznych.

„No, a teraz ułóżmy to jeszcze raz z klocków, ale tak, żeby zrozumieć ogólny wzór na sumę kolejnych liczb: $1+2+3+\dots+n$. Wyjaśniłem, że matematyka jest tam, gdzie odkrywamy ogólną prawidłowość,



2. Piętnaście groszy ułożone w trójkąt równoboczny i trójkąt prostokątny



3. Piętnaście piątek sklejonych w trójkąt równoboczny. Ile sklepień musiałem wykonać?

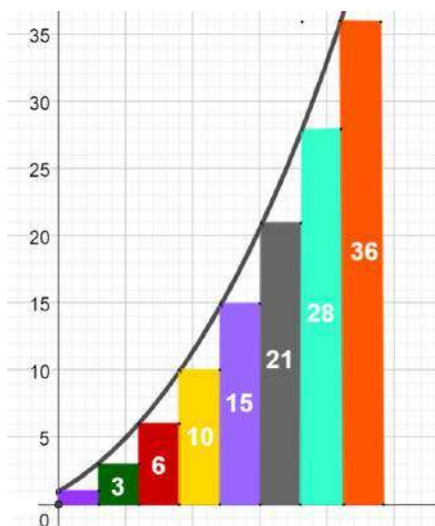
wzór, formułę. Jednostkowe łamigłówki są oczywiście ciekawe i mogą być bardzo trudne – ale właśnie tym różnią się od matematyki, że nie odkrywają ogólnych zależności.

Możemy się domyślić z rysunku, że liczby takie jak sumy $1+2$, $1+2+3$, $1+2+3+4$, a ogólnie $1+2+3+\dots+n$ nazywamy liczbami trójkątnymi. Napiszę trzydzieści początkowych liczb trójkątnych. Jak można się domyślać, obliczył mi to komputer. No, nie sam, tylko stosowny program, który ja (człowiek!) uruchomiłem.

Result:

{1, 3, 6, 10, 15, 21, 28, 36, 45, 55, 66, 78, 91, 105, 120, 136, 153, 171, 190, 210, 231, 253, 276, 300, 325, 351, 378, 406, 435, 465}

Zobaczmy to na wykresie.



4. Wysokości słupków odpowiadają kolejnym liczbom trójkątnym 1, 3, 6, 10, 15, 21, 28, 36. Wysokości te zwiększają się kolejno o 2, 3, 4, 5, 6, 7 i 8. Wierchołki prostokątów leżą na paraboli

Do wyprowadzenia ogólnego wzoru na liczby trójkątne już należało uczniów podprowadzić, by odkryli coś samodzielnie, ale w końcu ułożyliśmy taki kwadrat (rys. 5). Trzeba się w niego wpatrzeć, żeby zrozumieć, o co chodzi. No, to wpatrujemy się. Jest to kwadrat 6 na 6, ale ogólnie widzimy n na n . Ale niech będzie nie n na n , tylko $(n+1)$ na $(n+1)$. To przecież wszystko jedno, a potem końcowy wzór będzie prostszy. Na przekątnej, tej z północnego zachodu na południowy wschód (często mówię na nią: od Szczecina do Przemyśla!) są monety ułożone orłami do góry. Jest ich sześć – zgodziliśmy się, że to jest dla nas $n+1$.



5. Trzydzieści sześć groszy ułożone w kwadrat

Nad i pod przekątną złożoną z orłów mamy dwie piramidki jak z rysunku 2, czyli dwa razy po $1+2+3+4+5$, a ogólnie $1+2+3+\dots+n$. A więc $2 \cdot (1+2+3+\dots+n) + (n+1) = (n+1)^2$.

Sumę kolejnych liczb nazwaliśmy liczbą trójkątną, T_n . Mamy więc $2 \cdot T_n + (n+1) = (n+1)^2$.

Nietrudno stąd wyprowadzić końcowy wzór na liczby trójkątne:

$$T_n = \frac{n(n+1)}{2}$$

Wzór prosty, ale ciekawy. Dałem następne zadanie.

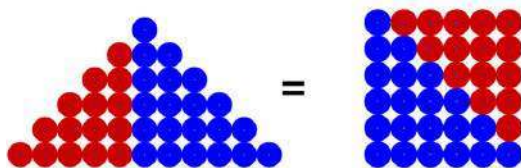
Zadanie 2. Narysuj $1+3+5+7$. Oblicz $1+3+5+7+9+11+13+15+17+19+21+23+25+27+29+31$. Zapisz wzorem. Oblicz.

Tu poszło gorzej. Najpierw obliczaliśmy $1=1$, $1+3=4$, $1+3+5=9$, $1+3+5+7=16$, $1+3+5+7+9=25$, $1+3+5+7+9+11=36$, ale uczniom ciąg 1, 4, 9, 16, 25, 36, ... z niczym się nie kojarzył. Nawet mnie to zdziwiło, ale gdy im powiedziałem, co to za ciąg, to z kolei oni zdziwili się,

że tego nie dostrzegli ..., przecież to kolejne kwadraty:

$1^2, 2^2, 3^2, 4^2, 5^2, 6^2, \dots$

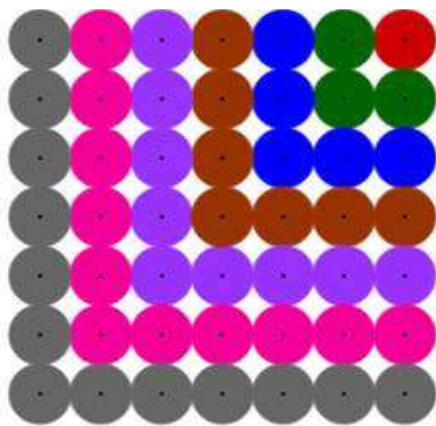
Ale jak to zobaczyć na monetach, albo klockach? Można na kilka sposobów. Na przykład tak:



6. Trzydzieści sześć jako liczba trójkątna i jako liczba kwadratowa

Spójrzmy na rys. 6. Po lewej stronie mamy (w poziomych rzędach) kolejne liczby nieparzyste: 1, 3, 5, 7, 9, 11. Układamy kółka w kwadrat i widzimy, że suma jest równa 36. Suma kolejnych liczb nieparzystych jest kwadratem pewnej liczby. Pewną sztuką jest zgadnąć, jakiej. Oto i stosowny wzór: $1+3+5+\dots+(2n-1)=n^2$.

I jeszcze ułożymy inaczej sumę kolejnych liczb nieparzystych (rys. 7). Kółka tego samego koloru wypełniają „korytarzyk” w kształcie litery L (poza pojedynczym kółkiem w prawym górnym narożniku). Ile jest kółek poszczególnych kolorów? Liczby od góry: 1, 3, 5, 7, 9, 11, 13 a łącznie wypełniają kwadrat 7 na 7. Wszystko się zgadza. Suma liczb nieparzystych $1+3+5+7+9+11+13$ to 7^2 .



7. Suma kolejnych liczb nieparzystych

Ponieważ, jak wspomniałem, pora roku była kasztanowa, więc uczniowie chętnie stworzyli układankę jak na rys. 8. Trzydzieści sześć kasztanów można ułożyć w kwadrat albo w trójkąt równoboczny.



8. Kasztany jako pomoc w nauce arytmetyki

A czy są inne liczby trójkątno-kwadratowe, większe od 36? Są. Jest ich nieskończenie wiele, ale najmniejsza po 36 to dopiero

$$1225=35^2=1+2+3+\dots+49,$$

a potem idą

41616, 1413721, 48024900, 1631432881, 55420693056, 1882672131025, 63955431761796, 2172602007770041, 73804512832419600, 2507180834294496361, 85170343853180456676, 2893284510173841030625, ...

Sam zdziwiłem się (wiele lat temu), gdy dowiedziałem się o istnieniu OEIS, The On-line Encyclopedia of Integer Sequences. Są w niej najrozmaitsze ciągi liczbowe, które są interesujące z matematycznego punktu widzenia. Liczby trójkątne mają kod A000217 (można powiedzieć, że to ich PESEL), a liczby trójkątno-kwadratowe kod A001110. Można się z tej encyklopedii dowiedzieć, jak wiele interesujących własności one mają i że nawet w XXI wieku ich własności nie są do końca odkryte. Ale już w 1778 roku Leonhard Euler odkrył ogólny wzór na ten liczby. Oznaczone są one symbolem N_k , a zatem $N_1=1$, $N_2=36$, $N_3=1225$. Wzór ten wygląda dość zawiłe:

$$N_k = \frac{1}{32} \left((1+\sqrt{2})^{2k} - (1-\sqrt{2})^{2k} \right)^2$$

Jeden z uczniów zadał kłopotliwe pytanie: czy to wszystko ma jakieś zastosowanie, czy to tylko tak sobie? Pytanie stawiane często matematykom. Co komu ze znajomości własności ciągów liczbowych? Można to pytanie postawić w jednym rzędzie z innymi: po co nam wiadomości o dawno wymarłych gadach, o budowie dalekich słońc, o Marsie sprzed miliona lat i Ziemi za milion lat, z czego zasłynął król Rurytanii Burburyk Pierwszy, co chciał powiedzieć Adam Mickiewicz przez swoje 44, dlaczego gramy Chopina (gdy to piszę, kończy się XVIII Konkurs Chopinowski) tak dalej i tak dalej. Jaki sens ma piłka nożna i zjeżdżanie na sankach po specjalnie spreparowanym lodowym torze?

No cóż, to wszystko jest dla matematyków po prostu interesujące, jak dla fanów muzyki kolejne wykonania klasycznych utworów, wciąż pięknych – to kolejna aluzja do Fryderyka Chopina. Ale matematyka jest w lepszej sytuacji. Chociaż nie wszystkie odkrycia matematyczne znajdują zastosowanie, to „jakoś tak jest”, że bardzo wiele z tych odkryć biorą fizycy, inżynierowie, ekonomiści, informatycy, a nawet językoznawcy i literaturoznawcy. Trudno tylko powiedzieć, co i kiedy zostanie wykorzystane. Wielkie odkrycia biorą się niekiedy z przypadku. Jednym z przykładów jest odkrycie promieniotwórczości. Gdyby Roentgen nawet dostał „grant” na znalezienie metody prześwietlania ciała ludzkiego, pewnie by nic nie odkrył. A tymczasem zauważył, że przez nieuwagę naświetliła się pewna płyta fotograficzna... i stąd mamy aparaty rentgenowskie. To długi temat, za długi na ten artykuł. Jako swego rodzaju perełkę zamieszczę fragment pierwszej strony pracy dwóch Chińczyków. Rzecz jest bardzo aktualna, 2020 rok. Spostrzegli oni, że ciąg liczb trójkątno-kwadratowych jest szczególnym przypadkiem tak zwanych ciągów Horadama (Alwyn Horadam, australijski matematyk, 1923–2016, jego podobizna poniżej). Praca została opublikowana w uznanym czasopiśmie naukowym. To dowód, że wszędzie jeszcze jest coś do odkrycia. Czy to odkrycie Chena i Pana znajdzie kiedyś zastosowanie? Nie wiadomo. Być może. Jest pogląd, że cała nauka jest w ogóle „na wszelki wypadek”. Żeby „jak co”, to uczeni – tak jak wojsko – byli pod ręką. Ale daleko odszedłem od prostej lekcji o matematyce dla młodzieży... ■

Michał Szurek



Journal of Integer Sequences, Vol. 23 (2020),
Article 20.5.8

Greatest Common Divisors of Shifted Horadam Sequences

Kwang-Wu Chen and Yu-Ren Pan
Department of Mathematics
University of Taipei

9. Tytuł pracy Chena i Pana



10. Alwyn Horadam (1923–2016)

Źródło: <https://binged.it/31EG0I6>

Jak dogadać się z każdym w każdej sytuacji

Laurence i Emily Alison

Wydawnictwo MUZA S.A., liczba stron: 400, cena: 49,90 zł

Zdroworozsądkowe i proste do opanowania techniki, jak budować dobry kontakt z innymi ludźmi, z którymi niekoniecznie musimy się zgadzać, ale z którymi lepiej współpracować niż rywalizować. Umiejętność budowania więzi może okazać się decydująca, gdy musimy poprosić szefa o podwyżkę, zwrócić uwagę nauczycielce na epizod przemocy w klasie dziecka, poprosić sąsiada, aby ściszył muzykę czy wpoić nastoletniemu dziecku zasady powrotu do domu o określonej porze. Ułatwia nawiązywanie nowych przyjaźni, poprawienie już istniejących relacji rodzinnych i zawodowych oraz otwarcie się na drugiego człowieka.



1. Biometria behawioralna

Hasło

Pożegnanie z 123456?

Na początku 2018 roku analitycy firmy Gartner przewidywali, że hasła do 2022 roku zastąpi behawioralna biometria (1). „Smartfony będą rozszerzeniem użytkownika, zdolnym rozpoznać i przewidzieć jego następny ruch,” czytamy w raporcie. Mają rozumieć, kim jesteś, czego chcesz, kiedy chcesz, jak chcesz, co zrobić i wykonywać zadania. A więc wpisywanie haseł stanie się zbędne.

Biometria behawioralna to nowa gałąź technik cyfrowych, która analizuje zachowania użytkowników, w tym dynamikę klawiszy, wzorce chodu, identyfikację głosu, cechy korzystania z myszy, rozpoznawanie podpisu oraz biometrię kognitywną, tworząc i zapisując unikatowy szablon biometryczny na urządzeniu. Gdy rejestrowane przez urządzenie zachowanie nie jest zgodne z szablonem, domniemanemu włamywaczowi/oszustowi blokowany jest dostęp, lub też włącza się procedura uwierzytelnienia dodatkowego (to na wypadek gdyby autoryzowany użytkownik popełnił błąd).

Pojawiają się ciekawe praktyczne propozycje rozwiązań tego rodzaju autentykacji. Trzy lata temu australijscy naukowcy opracowali prototyp ubrania, które wykorzystuje sposób chodzenia użytkownika jako token uwierzytelniający. Kolejną tego typu technikę ogłosiła pierwszych miesiącach 2021 roku firma BioCatch. Technologia autentykacji behawioralnej polega w nim na porównywaniu zapisanych

w bazie zachowań użytkownika w określonych aplikacjach z bieżącymi zachowaniami, co ma ze stuprocentową dokładnością wykrywać intruzów. Analiza wzorców zachowań schodzi tu bardzo głęboko. BioCatch np. potrafi zarejestrować szybkość poruszania się kursora i ścieżkę jego przesuwania podczas wizyty w serwisie internetowym. Mechanizm potrafi też rozpoznać parametry takie jak szybkość pisanie na klawiaturze a nawet

2. Rozpoznawanie wzorców pisanie na klawiaturze



siłę uderzenia w jej klawisze (2). W przypadku ekranów dotykowych rozpoznaje typowe, powtarzające się, ścieżki po których wędrują końcówki palców, jak również siłę dotyku. System „nauczony” kojarzyć określony zespół parametrów z naszymi danymi, rozpoznaje inne ścieżki i sposoby pisanja. Żąda wówczas dodatkowych potwierdzeń, co powinno zniechęcić ewentualnego włamywacza.

Gartner zapowiadał w swoim raporcie sprzed trzech lat wdrożenie w 2022 roku wielu innych systemów podobnego rodzaju w urządzeniach przenośnych. Miałoby być to m.in. rozpoznawanie emocji, pełne maszynowe rozumienie języka naturalnego, analizę dźwięków i wiele innych. Rok 2022 za progiem, więc... przekonamy się.

Czasy hasła mijają? Niekoniecznie i nie do końca

Wydawca Cybersecurity Ventures szacuje, że do 2020 roku stworzonych zostało przez ludzi co najmniej 100 miliardów hasel. Dość często spotyka się jednak opinię, że to szczyt epoki hasła i stopniowo będzie się ona kończyć. Bill Gates powiedział już w 2004 roku, że hasła odchodzą do lamusa. Trzeba jednak przyznać, że w chwili obecnej, pod koniec 2021 roku, świat techniki cyfrowej, z której korzystamy, jest w ogromnym stopniu uzależniony od wszelkiego rodzaju hasel. Są one potrzebne, aby uzyskać dostęp do naszych kont e-mail, sieci społecznościowych, telefonów komórkowych i do dziesiątków innych systemów oraz urządzeń.

Co więcej, eksperci wciąż potrafią przekonywać, że potrzebujemy znacznie większej liczby hasel niż te, z których korzystamy. Według raportu LastPass z 2019 roku, pracownicy dużych firm powinni mieć 25 unikalnych loginów, zaś w przypadku pracowników mniejszych organizacji liczba ta wynosi aż 85. Zatem nie widać tendencji do ograniczania roli hasła a raczej do czegoś przeciwnego – rozbudowy i komplikowania opartych na hasłach rozwiązań autentykacyjnych.

Problem w tym, że kiedy wymaga się od ludzi posiadania zbyt wielu unikalnych hasel, znacznie większej liczby, niż są w stanie zapamiętać, zaczynają chodzić na skróty. Hasła są ponownie wykorzystywane (recykling hasel) lub stosuje się zbyt łatwe do odgadnięcia dane uwierzytelniające. Według brytyjskiego Narodowego Centrum Cyberbezpieczeństwa, najczęściej łamanym hasłem w 2019 r. było „123456” co oznacza także, że było nader często stosowane przez użytkowników. Z innych badań, przeprowadzonych przez

TraceSecurity, wynika, że aż 81 proc. naruszeń danych firmowych jest spowodowanych złymi protokołami hasel. Firmy starają się znaleźć alternatywne metody uwierzytelniania swoich pracowników. Jednak, chociaż wiele osób zniechęciło się do hasel, są one również do nich przyzwyczajone – przekonanie ich do zmiany na coś innego może nie być łatwe. Era hasel być może dobiega końca, ale panuje opinia, że będzie to wolne odejście.

Rick McElroy, główny strateg ds. bezpieczeństwa w VMware Carbon Black w wypowiedzi dla magazynu „Wired” zwracał uwagę, że dane zebrane w ciągu ostatnich dwudziestu lat pokazują, że im bardziej wymaga się złożoności i rotacji hasel, tym bardziej prawdopodobne jest, że użytkownicy będą zapisywać hasła i zwiększy się liczba zgłoszeń do działu help desk w organizacji. Wpłynie to oczywiście na wydajność pracowników. Jak dodał, „używanie hasła jest tak przestarzałe, jak używanie standardowego klucza do drzwi wejściowych. Jasne jest, że są zamknięte, ale ktoś może skopiować klucz lub wybrać zamek i uzyskać dostęp”.

„Coś co wiesz + coś co masz + coś czym jesteś”

Firmy nie rezygnują z hasel, ale często przyjmują tak zwane „uwierzytelnianie dwuskładnikowe” lub „wieloskładnikowe”. Zazwyczaj podejście to polega na tym, że osoba musi posiadać różne typy informacji, aby uzyskać dostęp, jedna jest tym co wie, druga często opiera się na tym, co ma lub czym jest. Pierwszy składnik, to co użytkownik wie, to wciąż najczęściej hasło. Natomiast to coś, co posiada, może być urządzeniem, takim jak smartfon albo specjalnym rodzajem USB, takim jak YubiKey (3), natomiast to coś, czym jest, może opierać się na danych biometrycznych.

3. YubiKey



Do zestawu – „coś co wiesz + coś co masz + coś czym jesteś” dołączany bywa składnik „coś co robisz”. Tradycyjnie rozumiana autentykacja behawioralna może opierać się na powtarzalnych w określonym czasie czynnościach, np. wybieraniu pieniędzy z bankomatu w piątek po południu. Może mieć jednak charakter bardziej złożony i opierać się na analizach algorytmicznych, co zostało opisane na początku artykułu.

Dodatkowe warstwy zabezpieczeń znacznie utrudniają hakerom dostęp do aktywów osobistych lub biznesowych. Ponadto, wiele firm korzysta z programów typu menedżer haseł w celu wzmocnienia poziomu bezpieczeństwa. Jednocześnie ponieważ wielu pracowników jest obecnie proszonych o zapamiętanie kilku haseł do różnych programów a jednocześnie organizacje szukają sposobu na zwalczenie tendencji do ponownego wykorzystywania haseł. Menedżer haseł jest metodą, która to umożliwia, bez wielkich uciążliwości. Menedżerami haseł nazywa się aplikacje, które przechowują informacje dla wielu urządzeń i usług cyfrowych, automatycznie logując użytkowników. Przechowują bazę danych haseł, która jest zaszyfrowana i dostępna wyłącznie za pomocą hasła głównego, dzięki czemu użytkownicy mają znacznie mniej do zapamiętania i są bardziej skłonni do wybierania silniejszych haseł.

Coraz częściej w funkcje typu menedżer haseł wyposażane są przeglądarki internetowe. Najpopularniejsze systemy operacyjne, takie jak iOS, Android, MacOS i Windows, również mają tego typu funkcje.

Biometria nie bez zastrzeżeń, ale jednak czyni postępy

Identyfikacja biometryczna, np. odcisk palca czy skan tęczówki, jest trudniejsza do oszukania, ale konsekwencje wprowadzania takich rozwiązań sięgają, wiążąc się wieloma problemami i zastrzeżeniami. Firmy musiałyby przechowywać kopie tych informacji wewnętrznie w celu weryfikacji każdej osoby, co czyni systemy jeszcze atrakcyjniejszym celem dla hakerów chcących ukraść osobiste dane uwierzytelniające.

Uwierzytelnianie biometryczne stosowane w izolacji ma poważną wadę. Jeśli dane dotyczące rozpoznawania twarzy lub odcisku palca zostaną skradzione, ponownie pojawią się te same zagrożenia, co w przypadku ponownego użycia hasła, czyli atakujący może użyć tych danych, aby uzyskać dostęp do innych kont, które używają biometrii do uwierzytelniania. Jeszcze bardziej niepokojące jest, że nie



4. Windows Hello

można po prostu zresetować swoich cech biometrycznych tak jak hasła – są one trwałe. Hasła, które zostały złamane, można zmienić – odcisku palca nie.

Pomimo rozlicznych wątpliwości i obaw, biometria wkracza w świat techniki cyfrowej nasilającą się falą. We wrześniu firma Microsoft ogłosiła, że pozwala ludziom pozbyć się swoich haseł na rzecz używania odcisków palców, rozpoznawania twarzy a także aplikacji uwierzytelniających. Jej zdaniem, że ten ruch sprawi, że użytkownicy będą bardziej bezpieczni, ponieważ hasła mogą być łatwym celem dla włamywaczy. Vasu Jakkał, wiceprezes Microsoftu ds. bezpieczeństwa, zgodności i tożsamości, napisał na blogu firmowym, że obecnie „mierzymy się z liczbą 579 ataków na hasła co sekundę, czyli 18 miliardami każdego roku”. Jak zauważa, „ludzie tworzą hasła, które są łatwe do zapamiętania, a zatem mogą być łatwiejsze do odgadnięcia dla hakra”.

Microsoft podał także, że do końca 2020 roku liczba konsumentów korzystających z Windows Hello (4) – biometrycznej funkcji logowania – wzrosła z 69,4 proc. do 84,7 proc. Zapewnia, że oferowana od września opcja „całkowitego pozbycia się haseł” przez użytkowników Windows, została starannie przetestowana przez firmę.

W przyszłości być może zapomnimy nie tylko o hasłach, ale również o opisywanych wyżej wieloskładnikowych i biometrycznych rozwiązaniach autentykacyjnych, a do uzyskiwania dostępu używać będziemy własnego DNA, czyli czegoś czym jesteśmy w najgłębszej warstwie naszego istnienia. Może to to nie wystarczy i potrzebne będą dodatkowe składniki, takie jak bicie naszego serca i wzorce fal mózgowych. Brzmi jak fantastyka, tymczasem techniki te są już eksperymentalnie testowane... w wojsku. Dreszczem tylko przeszywa myśl o możliwości wykradzenia takich „haseł”. ■

Miroslaw Usidus

Eksperymenty z ubiegłego miesiąca dowiodły chemicznego pokrewieństwa kobaltu i niklu z żelazem. Jednak, jak to w rodzinie bywa, podobieństwo nie jest całkowite i każdy z kuzynów ma indywidualne cechy charakteru. W dzisiejszym odcinku wykonasz próby, które pozwolą ci zauważyć najważniejsze różnice w zachowaniu metali z triady żelazowców.

Na cześć złośliwych duchów (2)

Nie tylko stopy

Nikiel i kobalt, jak wiele innych pierwiastków, stanowią dodatki do różnych rodzajów stali. Metale te jednak tworzą również własne stopy. O zastosowaniu niklu do wytwarzania monet była już mowa w poprzednim odcinku. Z udziałem tego metalu produkuje się także stopy podobne do srebra, np. nowe srebro, używane jako ozdoby, sztucce, elementy sprzętu medycznego i aparatury naukowej. Stosowane w urządzeniach grzejnych oraz opornikach stopy oporowe (nikielin, konstantan, chromo-nikielina) także zawierają nikiel. Nawet niewielki dodatek tego metalu zwiększa odporność chemiczną i dlatego stanowi on składnik stopów wykorzystywanych w konstrukcji aparatury dla przemysłu chemicznego. Prawie cała produkcja kobaltu zużywana jest w specjalnych stopach, np. bardzo twardych stellitach (zawierających także chrom i wolfram), stalach szybkotnących i odpornych na ścieranie. W jeszcze twardszej widii ziarna węglików spojęne są za pomocą metalicznego kobaltu. Metal ten zwiększa własności magnetyczne stali i sam służy do produkcji magnesów (kobalt to ferromagnetyk, podobnie jak żelazo i nikiel).

Nikiel stosowany jest do galwanicznego pokrywania przedmiotów stalowych oraz metalizowania specjalnie przygotowanych tworzyw sztucznych. Ponieważ powłoka niklu najlepiej wiąże się z miedzią, zarówno stal, jak i plastiki najpierw pokrywa się warstwą tego metalu. Nikiel jest odporny na korozję i dzięki jego powłoce przedmioty długo zachowują ceniony przez użytkowników metaliczny połysk. Jednak uszkodzenie niklowej powłoki powoduje przyspieszone niszczenie stali (patrz odcinki o korozji).

Nikiel i kobalt to również katalizatory dla przemysłu chemicznego. Kobaltu używa się w roli katalizatora utleniania, natomiast nikiel, podobnie jak pierwiastki grupy 10 – pallad i platyna – przyspiesza reakcje z udziałem wodoru, dlatego często stosuje się go jako zamiennik tych drogich metali, np. podczas otrzymywania margaryny z ciekłych tłuszczów roślinnych.

Związki niklu używane są do produkcji akumulatorów, w tym i do sprzętu domowego (niklowo-kadmowe NiCd i niklowo-wodorkowe NiH), (1) natomiast połączenia kobaltu to wspomniane uprzednio pigmenty. Organizmy żywe również wykorzystywały właściwości niklu i kobaltu (patrz: W świecie przyrody), a ten ostatni metal posłużył



1. Akumulatory niklowe to ważne zastosowanie tego metalu

do skonstruowania specjalnej bomby (patrz: Bomba kobaltowa).

Problemy z amoniakiem

Próby wykonane w ubiegłym miesiącu pokazały, że kobalt i nikiel reagują tak, jak żelazo: pod wpływem zasad jony tych metali tworzą nierozpuszczalne w wodzie wodorotlenki, w których kationy o ładunku +2 utleniają się do połączeń trójwartościowych (kobaltu i żelaza łatwo, niklu trudniej). Ponadto wodorotlenki te nie są amfoteryczne i nie rozpuszczają się w nadmiarze dodanej zasady.

Roztwór amoniaku w wodzie ma odczyn zasadowy i również służy do wytrącania osadu wodorotlenków. Sporządź roztwory soli kobaltu i niklu. Do każdej z probówek dodawaj kroplami wodę amoniakalną (roztwór NH_3). Zaobserwujesz powstawanie osadów: niebieskiego w przypadku jonów kobaltu(II) i zielonego w naczyniu z jonami niklu(II). Tego mogłeś się spodziewać – reakcje przebiegają w taki sam sposób, jak w ubiegłym miesiącu pod wpływem roztworu NaOH . Zatem nic nowego, potwierdziłeś tylko wnioski z poprzedniego odcinka (jony żelaza zareagują w ten sam sposób: kationy Fe^{2+} dadzą osad zielonkawy, a Fe^{3+} – brunatny; sam wykonaj odpowiednie próby). Jednak nie kończ jeszcze doświadczenia, nadal dodawaj roztwór amoniaku i mieszaj zawartościami probówek.

Reakcja przebiega dalej: osady znikają, a powstające roztwory barwią się na kolor szafirowy w przypadku niklu (2) i różowy dla kobaltu (ale barwa ta wkrótce staje się brunatna). Jeżeli równolegle przeprowadzisz próby dla żelaza, nie zauważysz rozpuszczania osadów. Czyżby wnioski z ubiegłego miesiąca były fałszywe i osady wodorotlenków niklu i kobaltu są jednak amfoteryczne (rozpuściły się przecież w nadmiarze dodanego odczynnika), a samo pokrewieństwo z żelazem niepewne (skoro jego jony tak nie zareagowały)?



2. Komplex niklu(II) z amoniakiem



3. Z lewej – kompleks kobaltu(II) z amoniakiem, z prawej – kompleks szybko się utlenia

Nie, wnioski były prawidłowe, a ty byłeś świadkiem powstawania połączeń kompleksowych niklu i kobaltu. Są one rozpuszczalne w wodzie, dlatego osady zniknęły (podobnie reaguje miedź – jej wodorotlenek również rozpuszcza się w roztworze amoniaku). Brunatnienie roztworu zawierającego jony kobaltu to efekt utleniania do połączeń trójwartościowych (3). Dlaczego żelazo nie daje takich wyników? Co prawda jony Fe^{2+} i Fe^{3+} tworzą liczne kompleksy, lecz te z amoniakiem są akurat bardzo mało trwałe.

Powtórz eksperyment, ale tym razem do roztworów soli dodaj najpierw roztwór chlorku amonu NH_4Cl . Gdy teraz wlejesz wodę amoniakalną, nie zauważysz wytrącania osadów – roztwory od razu przybiorą barwy charakterystyczne dla kompleksów. Chlorek amonu cofną dysocjację amoniaku tak, że roztwór nie wykazuje odczynu zasadowego i nie dochodzi do wytrącania osadów, natomiast amoniak bezpośrednio tworzy związki kompleksowe. Przedstawiony sposób jest „chemiczną sztuczką” pozwalającą otrzymać kompleksy z amoniakiem z pominięciem przeszkadzającego czasem wytrącania osadu. W przypadku jonów Fe^{3+} wytrąci się brunatny osad, ale to efekt jego ekstremalnie niskiej rozpuszczalności – tworzy się on nawet w słabo zakwaszonych roztworach.



4. Po ogrzaniu kartki ujawnia się napis wykonany roztworem chlorku kobaltu(II)

Różowy i niebieski

Wodne roztwory soli kobaltu(II) mają różowe zabarwienie, ale związki tego pierwiastka często przyjmują również niebieski kolor (co zauważyłeś wytrącając osad przy pomocy zasady). Jaki jest powód zmiany barwy? Sporządź dość stężony roztwór chlorku kobaltu(II) o ciemnoróżowej barwie i wlej go do probówek (wystarczy po 0,5 cm³). Do pierwszej z nich dodaj kroplami stężony kwas solny, do drugiej kilka kryształów bezwodnego (koniecznie!) chlorku wapnia CaCl₂ lub magnezu MgCl₂, a do trzeciej nieco acetonu. W każdym przypadku roztwór zmienił barwę na niebieską lub niebieskofioletową. Dodane substancje „wyciągnęły” cząsteczki wody z najbliższego otoczenia jonu Co²⁺, powodując zmianę zabarwienia. Zapamiętaj, że uwodnione sole kobaltu mają różowe zabarwienie, natomiast bezwodne (tak naprawdę miejsce cząsteczek wody w pobliżu jonu zajmują inne drobiny) – niebieskie.

5. Niebieskie zabarwienie kryształu tiosiarczanu sodu potwierdza obecność jonów kobaltu(II)



Opisane zmiany barwy znalazły zastosowanie do sporządzenia **atramentu sympatycznego**. Napis wykonany rozcieńczonym roztworem chlorku kobaltu(II) jest po wyschnięciu praktycznie niewidoczny. Po ogrzaniu do 30–40°C na kartce ukazują się niebieskie litery (4). Związek pochłania wodę z otoczenia i po ochłodzeniu napis ponownie znika. Również żel krzemionkowy używany do osuszania nasycy się roztworem CoCl₂. Zmiana zabarwienia żelu z niebieskiego na różowy sygnalizuje pochłonięcie znacznej ilości wody przez absorbent. Ogrzanie powoduje odwodnienie (powraca niebieski kolor) i regenerację środka suszącego.

I jeszcze bardzo prosty test analityczny. Odrobinę tiosiarczanu sodu Na₂S₂O₃ zwilż kroplą analizowanego roztworu. Jeżeli zawierał on jony Co²⁺, kryształ tiosiarczanu zabarwi się na niebieski kolor (5).

Żelazo przeszkadza

Jednym z odczynników do identyfikacji kationów Co²⁺ jest rodnik potasu KNCS lub amonu NH₄NCS. Wykonaj próbę mieszając roztwór rodanku

W świecie przyrody

Wiele roślin, bakterii i grzybów wykorzystano nikiel w postaci kationu jako element enzymów umożliwiających im reakcje z udziałem wodoru. Przykładem jest ureaza – enzym katalizujący rozkład mocznika. Natomiast organizmy wyższe unikają niklu jako biopierwiastka. U ludzi często spotykana jest alergia na ten metal, objawiająca się zapaleniem skóry w miejscach kontaktu z przedmiotami zawierającymi nikiel, np. oprawkami okularów czy biżuterią. Ponieważ nikiel wchodzi w skład praktycznie wszystkich przedmiotów stalowych, trudno uniknąć kontaktu z alergenem. Naukowcy ostrzegają również przed niklem zawartym w produktach żywnościowych, np. utwardzonych tłuszczach roślinnych (pochodzi z użytego katalizatora reakcji) czy też przygotowywanych w stalowych naczyniach. Kobalt to ważny mikroelement w świecie organizmów żywych. Roślinom motylkowym ułatwia wiązanie azotu atmosferycznego, co ma kluczowe znaczenie dla życia na Ziemi. W przypadku zwierząt niedobór powoduje zaburzenia krzepnięcia krwi. Dla nas najważniejszym związkiem kobaltu jest **witamina B₁₂** (kobalamina) odpowiadająca za prawidłowy skład krwi, oraz działanie układu nerwowego i przemiany metabolicznej. W przewodzie pokarmowym jest wytwarzana przez bakterie, a ponieważ część produkcji ulega wydaleniu, musi być także dostarczana z pokarmem. Źródłem witaminy B₁₂ są produkty pochodzenia zwierzęcego, natomiast w przypadku zaburzeń wchłaniania czy też diety wegańskiej zalecana jest suplementacja jej preparatami.

oraz soli kobaltu – niebieska barwa pozwala potwierdzić obecność jonów tego metalu. [reakcja_z_rodankiem]

Żelazo i kobalt często występują obok siebie, a fakt ten sprawia trudności w ich wykrywaniu. Klasyczna analiza opiera się na wytrącaniu osadów i obserwacji zmian zabarwienia. W mieszaninie trudno jednak zauważyć kolory poszczególnych związków. Dla przykładu: wykonanie opisanej wyżej próby w roztworze zawierającym jony Fe^{3+} i Co^{2+} nie powiedzie się, ponieważ krwista barwa związku żelaza (znana jako „chemiczna krew”) zamaskuje kolor połączenia kobaltu. Rozwiązaniem jest dodatek fluorku sodu NaF przed waniem do probówki roztworu rodanku. Jony fluorkowe tworzą bardzo trwałe kompleksy z żelazem i czerwona barwa nie wystąpi, co umożliwia obserwację niebieskiego koloru połączenia kobaltu. Czynność ta – częsta w praktyce laboratoryjnej – to **maskowanie**. Opisane wyżej próby pozostawiam ci do samodzielnego wykonania. ■

Krzysztof Orlński



6. Niebieska barwa rodankowego kompleksu kobaltu(II)



7. Radioterapia przy użyciu bomby kobaltowej (1951, ze zbiorów National Cancer Institute, USA)

Bomba kobaltowa

Sztucznie otrzymywany izotop **kobalt-60** emituje silne promieniowanie gamma, które jest używane do radioterapii w leczeniu nowotworów, defektoskopii („prześwietlanie” pomagające zlokalizować ukryte wady materiału) czy też sterylizacji. Bomba kobaltowa to urządzenie zawierające izotop Co-60 oraz osłony i mechanizm sterujący. Dlaczego bomba? Pierwsze egzemplarze miały kształt bomby lotniczej (7).

Bomba kobaltowa ma również drugie znaczenie. To ładunek nuklearny, którego pancerz wykonano ze stopu zawierającego kobalt. W czasie wybuchu jądrowego powstaje promieniotwórczy Co-60, silnie skażający teren.

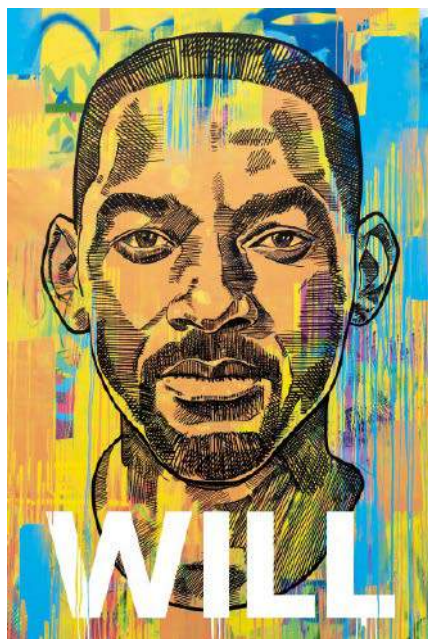
Will

Will Smith, Mark Manson

Wydawnictwo Insignis, cena: 49,99 zł

Odważna i inspirująca książka, w której Will Smith, jedna z najbardziej wyrazistych i rozpoznawalnych postaci współczesnego przemysłu rozrywkowego, opowiada szczerze o swoim życiu i dzieli się tym wszystkim, co pozwoliło mu zrozumieć, że pielęgnowanie więzi międzyludzkich jest równie ważne, co poczucie szczęścia płynącego z własnego wnętrza i satysfakcja z sukcesów zawodowych. To także książka, która pozwala poznać ze szczegółami przebieg jednej z najbardziej niesamowitych karier w świecie muzyki i filmu.

To nie tylko pasjonująca opowieść o metamorfozie Willa Smitha z dziecka dorastającego w zachodniej Filadelfii w czołową gwiazdę rapu swoich czasów, a następnie w jednego z najpopularniejszych aktorów w historii Hollywood. To o wiele więcej. Will Smith miał wszelkie powody, by uważać, że osiągnął w życiu sukces: nie tylko zrobił oszałamiającą karierę, wspinając się na sam szczyt świata rozrywki – wprowadził tam także całą swoją rodzinę. Tyle tylko że jego bliscy wcale nie byli tym zachwyceni: czuli się jak gwiazdy w cyrku Willa, odgrywając przez siedem dni w tygodniu rolę, o które wcale nie prosili. To pokazało, że Will musi się jeszcze wiele nauczyć.





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

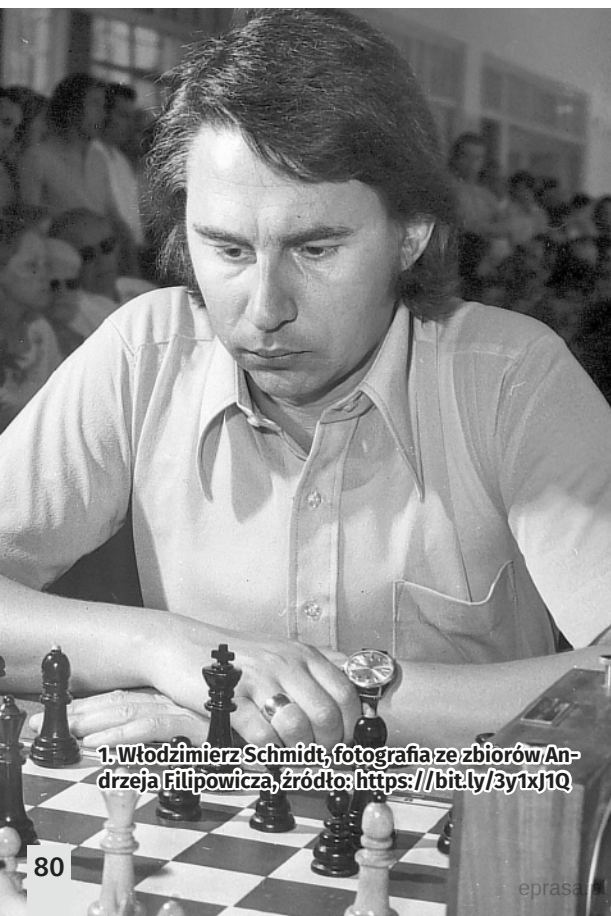
Włodzimierz Schmidt – legenda polskich szachów

Włodzimierz Schmidt zaczął grać w szachy w wieku 9 lat, a swoje umiejętności doskonalił w poznańskim klubie „Ogniwo”. Od początku związany był ze swoim rodzinnym Poznaniem, któremu jest wierny do dziś (1). Pierwszy szachowy sukces

zanimował w 1955 roku, zajmując drugie miejsce w turnieju młodzików poznańskiego klubu. Rok później zdobył mistrzostwo szkół podstawowych Poznania. Pod koniec lat pięćdziesiątych czterokrotnie występował w finałach Mistrzostw Polski juniorów (do 18 lat) zdobywając dwa tytuły mistrzowskie i jeden wicemistrzowski.

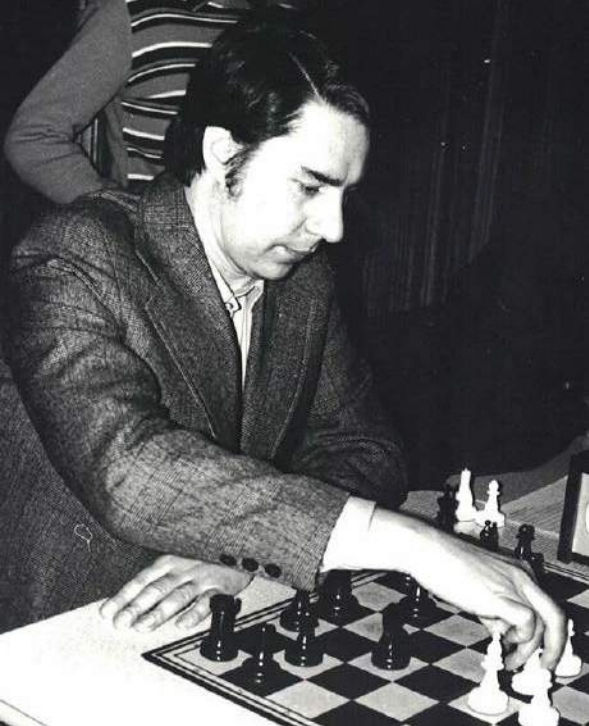
W 1962 r. zadebiutował w finale Mistrzostw Polski seniorów zdobywając brązowy medal. Po tym sukcesie na kilka lat ograniczył udział w rozgrywkach szachowych, więcej uwagi poświęcając studiom na wydziale elektrycznym Politechniki Poznańskiej. W finałach indywidualnych Mistrzostw Polski występował aż 28 razy (2). W swoim dorobku ma 15 medali mistrzostw kraju. Siedmiokrotnie zdobywał tytuł Mistrza Polski (w latach 1971, 1974, 1975, 1981, 1988, 1990 i 1994).

2. Włodzimierz Schmidt (po prawej) w partii z Andrzejem Filipowiczem, Łódź 1974, źródło: <https://bit.ly/3rFwsMM>



1. Włodzimierz Schmidt, fotografia ze zbiorów Andrzeja Filipowicza, źródło: <https://bit.ly/3y1xj1Q>





3. Włodzimierz Schmidt – Mistrz Polski w szachach błyskawicznych, Kalisz 1979, źródło: archiwum Macieja Sroczyńskiego

W drużynowych Mistrzostwach Polski zdobył 17 medali, w tym siedem złotych (Bydgoszcz 1973, Poznań 1974, Augustów 1975, Ciechocinek 1976, Katowice 1977, Lublin 1979 i Lubniewice 1981). Włodzimierz Schmidt znany jest z gruntownego przygotowania teoretycznego i dobrej techniki, a także wyjątkowego refleksu. W Indywidualnych Mistrzostwach Polski w szachach błyskawicznych zdobył 16 złotych medali, a w Drużynowych Mistrzostwach Polski w szachach błyskawicznych 8 złotych (3).

W 1968 roku Międzynarodowa Federacja Szachowa (fr. Fédération Internationale des Échecs, FIDE) nadała Schmidtowi – za wyniki na XV Olimpiadzie Szachowej w Złotych Piaskach (1962) oraz na VI Memoriale Akiby Rubinsteina (1968) w Polanicy Zdroju tytuł mistrza międzynarodowego. Pięć lat później (1973) Schmidt zrezygnował z pracy na Uniwersytecie im. Adama Mickiewicza w Poznaniu i poświęcił się całkowicie szachom.

Podczas turnieju rozgrywanego w Lipsku w 1973 roku wypełnił pierwszą normę arcymistrzowską a Międzynarodowa Federacja Szachowa przyznała mu tytuł arcymistrza po wypełnieniu drugiej normy na XIV Memoriale Akiby Rubinsteina w Polanicy Zdrój w 1976 roku. Włodzimierz Schmidt został pierwszym polskim

szachistą (nie licząc nominowanych w 1950 r. Akiby Rubinsteina, Ksawerego Tartakowera i Mieczysława Najdorfa), któremu Międzynarodowa Federacja Szachowa przyznała tytuł arcymistrza. Tytuły arcymistrzowskie FIDE zaczęła przyznawać od 1950 roku. Wówczas otrzymało je trzech polskich szachistów – Akiba Rubinstein, Ksawery Tartakower i Mieczysław Najdorf. Polska była bowiem przed II wojną światową prawdziwą szachową potęgą. Reprezentacja sześciokrotnie stawała na podium, wygrywając olimpiadę w 1930 roku w Hamburgu. Na kolejnego arcymistrza Polska czekała aż 26 lat, gdyż wywalczenie tego tytułu było obwarowane bardzo ostrymi kryteriami. W następnych latach FIDE doszła do wniosku, że trzeba obniżyć kryteria zdobywania tytułów mistrzów międzynarodowych i arcymistrzów, żeby zwiększyć popularność szachów na świecie.

Najwyższy ranking FIDE (2505) osiągnął Włodzimierz Schmidt 1 stycznia 1978 i do dzisiaj legitymuje się najwyższym rankingiem (2190) wśród wszystkich aktywnych (uczestniczących nadal w turniejach) polskich szachistów w swojej kategorii wiekowej.

Włodzimierz Schmidt nadal uczestniczy w wielu wydarzeniach szachowych, szkoleniach młodzieży i symultanach szachowych (4).



4. Symultana szachowa Włodzimierza Schmidta, 25 października 2019, Uniwersytet im. Adama Mickiewicza w Poznaniu, źródło: Życie Uniwersyteckie, UAM w Poznaniu



5. Włodzimierz Schmidt, Festiwal Szachowy, Poznań 2021, fot. Krzysztof Szeląg

78-letni obecnie Włodzimierz Schmidt (5) jest trenerem (w 2004 r. otrzymał jako pierwszy Polak trenerski tytuł *FIDE Senior Trainer*) i działaczem szachowym (był m.in. wiceprezesem do spraw sportowych Polskiego Związku Szachowego w latach 1999–2004 i 2009–2017). Ukończył Podyplomowe Studium Trenerskie na AWF Warszawa i był m.in. trenerem reprezentacji Polski na Olimpiadach Szachowych: 1992 – Manila (Filipiny), 2000 – Stambuł (Turcja), 2002 – Bled (Słowenia) i 2004 – Calviá (Wyspy Kanaryjskie – Hiszpania), trenerem reprezentacji Polski na Drużynowych Mistrzostwach Europy: 2001 – Leon (Hiszpania), 2003 – Płowdiw (Bułgaria) oraz



6. Włodzimierz Schmidt ze statuetką „Hetmana 2016” przyznaną przez Polski Związek Szachowy, źródło: <https://bit.ly/3pxLYr4>

trenerem reprezentantów Polski na Indywidualnych Mistrzostwach Świata i Europy Juniorów.

Jest Honorowym Członkiem Polskiego Związku Szachowego, laureatem „Hetmana 2016” w kategorii Całokształt Działalności, przyznanego przez Polski Związek Szachowy (6). Za swoją działalność na polu szachowym odznaczony został m.in. Złotym Krzyżem Zasługi, Krzyżem Kawalerskim Orderu Odrodzenia Polski i Krzyżem Oficerskim Orderu Odrodzenia Polski.

Wśród wielu wspaniałych partii szachowych rozegranych przez Włodzimierza Schmidta na szczególną uwagę zasługuje wygrana w eliminacjach drużynowych mistrzostw Europy w Aarhus w 1971



7. Bent Larsen – Włodzimierz Schmidt, pozycja po 21. Gg5



8. Bent Larsen – Włodzimierz Schmidt, pozycja po 27...Gd5



9. Bent Larsen – Włodzimierz Schmidt, pozycja po 31. Wd3

roku z czołowym wówczas szachistą świata – duńskim arcymistrzem Bentem Larsenem.

Białe: Bent Larsen, **Czarne:** Włodzimierz Schmidt, **Eliminacje Drużynowych Mistrzostw Europy, Aarhus (Dania), 8.09.1971**

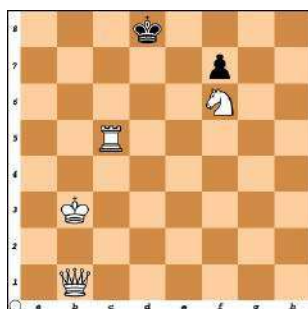


10. Bent Larsen – Włodzimierz Schmidt, pozycja w której Larsen poddał partię

1. f4 Tak grał wielokrotnie w swojej karierze Larsen 1...Sf6 2. Sf3 g6 3. d3 d5 4. e3 Gg7 5. Ge2 O-O 6. O-O c5 7. He1 Sc6 8. Hh4 Larsen wykonał już dwa posunięcia hetmanem w debiucie 8...b6 9. Sbd2 Ga6 10.

Wf2 Se8 11. c3 e5 12. H:d8 W:d8 13. f:e5 S:e5 14. S:e5 G:e5 15. a4 Gb7 16. d4 Gg7 17. d:c5 b:c5 18. Sb3 Wc8 19. e4 d4! Jeżeli 19...d:e4? to 20. Ge3 c4 21. Sc5 Gc6 22. G:c4 z dobrą pozycją białych 20. c:d4 c:d4 21. Gg5 (diagram 7) 21...Sd6! Czarne poświęcając wieżę za gońca zostają z parą silnych gońców i wolnym pionem d4 22. Ge7 S:e4 23. G:f8 G:f8 24. Wff1 Gh6 25. Wfd1 Ge3+ 26. Kf1 Wd8 27. Sa5 Gd5 (diagram 8) 28. Gb5? Prowadzi do wyraźnej przewagi czarnych, należało grać 28. Sc4! G:c4 29. G:c4 i różnobarwne gońce dawały białym duże szanse na remis 28...Wc8! 29. Wa3 Sd2+ 30. Ke1 Wc2 31. Wd3 (diagram 9) 31...Sf1! 32. W:d4 Skoczek czarnych jest nietykalny, po 32. K:f1?? nastąpiłoby 32...G:g2+ 33. Ke1 Gf2# 32...G:g2 33. Wd8+ Kg7 34. Ge2 S:h2 35. W1d3 Jeżeli 35. W8d3?? To 35...Sf3+ 36. G:f3 Gf2# 35...Gb6 36. b4 G:d8 37. W:d8 Gf3 38. Gc4 Sg4 39. We8 Kf6 40. Gb3 Wb2 41. b5 h5 i białe poddały partię (diagram 10). ■

Zadania do samodzielnego rozwiązania



Zadanie 1
11. I. Bakajew, W. Iwanow, 1995
Mat w 3 posunięciach



Zadanie 2
12. Samuel Loyd, 1857
Mat w 3 posunięciach

Rozwiązanie zadań z MT 11/2021

Zadanie 1
W. Czeliznyj 1986
Mat w 3 posunięciach
Rozwiązanie: 1.Kd5 d6 2.Gc5 d:c5 3.Sf2#

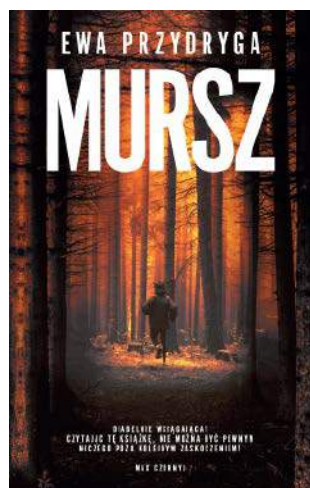
Zadanie 2
W. Simonow 1992
Mat w 3 posunięciach
Rozwiązanie: 1.Ga1 b2 2.Wb1 b:a1H 3.Wb8#

Mursz

Ewa Przydryga

Wydawnictwo MUZA S.A., cena: 39,90 zł

Las, który odcisnął krwawe piętno na życiu okolicznych mieszkańców, przypieczętuje także jej los... Pola z trudem usiłuje wrócić do nowego życia bez męża i syna. Po zakończeniu terapii regularnie koresponduje z Kasandrą, z którą zaprzyjaźniła się podczas pobytu w szpitalu. Wkrótce w ręce Poli trafia list, zupełnie inny niż poprzednie, będący zapowiedzią samobójstwa przyjaciółki. Ujawniony w liście sekret prowadzi Polę do owianego złą stawą Murszu, wymarłej części kaszubskiego lasu. To tam rok wcześniej została zamordowana Lara, koleżanka Kasandry. Wyznania Kas rzucają jednak zupełnie nowe światło na makabryczne zdarzenia tamtej nocy. Każda z leśnych dróg prowadzi Polę pod drzwi tajemniczej spalonej leśniczówki. Strawiona ogniem chata skrywa mroczne sekrety lasu i mieszkających w nim ludzi. Tych żyjących i tych, którzy na przestrzeni lat w tragicznych okolicznościach stracili tu życie.





Pewien Amerykanin w latach 90-tych ubiegłego stulecia został porażony piorunem. Nie byłoby w tym nic nadzwyczajnego gdyby nie to, że w wyniku porażenia doszło do przedziwnej zmiany w jego mózgu. Mężczyzna przed wypadkiem był chirurgiem, którego zainteresowania związane z muzyką kończyły się na słuchaniu rocka. Trzy miesiące po wypadku poczuł nieodpartą potrzebę słuchania muzyki klasycznej z naciskiem na utwory Chopina. „Polonez wojskowy”, „Zimowy wiatr”, „Czarny klucz” nie dawały mu spokoju. Obsesja narastała do tego stopnia, że nauczył się grać je na pianinie, a następnie sam zaczął komponować. W wyniku rażenia piorunem Tony Cioria, bo tak nazywał się ów lekarz, został niemalże „nawiedzony” przez dźwięki, opętany przez muzykę. Fascynacja przerodziła się w nowe hobby. Zabawy dźwiękiem to hobby dla wybranych, ale nie trzeba igrać z piorunami by studiować akustykę. Zapraszamy.

Akustyka

Akustyka to nie tylko kierunek dla audiofilów. Wciąż starzejące się społeczeństwo wymaga coraz większego wsparcia. Rozwój cywilizacji to także pojawianie się nowych chorób i niepełnosprawności. Złe nawyki, przebywanie w ciągłym hałasie, stałe stymulowanie się dźwiękami powodują uszkodzenia słuchu. Już małe dzieci mają kłopoty z przetwarzaniem słuchowym, a problem stale narasta. Akustyka w tym wymiarze stanowi może olbrzymie wsparcie w zakresie profilaktyki i kompensowania powstałych ubytków. Akustyka to także prace związane z ochroną środowiska przed hałasem. Jest to zatem kierunek interdyscyplinarny, który łączy w sobie wiedzę z zakresu: fizyki, medycyny, ochrony środowiska, sztuki czy nawet psychologii.

Akustykę można studiować zarówno na uczelniach prywatnych jak i państwowych. Kierunek ten dostępny jest na uniwersytetach, akademiach i politechnikach. Edukacja na wyżej wymienionych uczelniach różnić się będzie ze względu na charakter szkoły. Akustyka oferowana na uniwersytetach ma bardziej charakter interdyscyplinarnych. Uczelnie techniczne większą wagę skupią na inżynierii akustycznej. Dla porównania, na Uniwersytecie Warszawskim można spodziewać się takich przedmiotów jak na przykład: audiometria tonalna, audiometria mowy, dopasowanie aparatów słuchowych, hałas komunikacyjny, natomiast na Akademii Górniczo Hutniczej:

elektrotechniki, elektroniki, teorii drgań, matematyki, fizyki. Wszyscy absolwenci tego kierunku są ponadto przygotowani do profesjonalnego rejestrowania, przetwarzania, archiwizowania, transmitowania i odtwarzania dźwięku. Każdy absolwent nabywa kompetencji do wspomagania kultury i sztuki oraz do twórczego realizowania się w tym obszarze. Tym samym jest to doskonałe przygotowanie do realizowania pasji związanych z dźwiękiem. Dostępność tego kierunku na uczelniach wyższych nie jest wysoka. Na mapie Polski znajdziemy niewielką ilość uczelni oferujących akustykę. Szkoły oddają skromną ilość miejsc dla studentów, dlatego konkurencja w procesie rekrutacji jest stosunkowo duża. Zdając maturę należy szczególną uwagę zwrócić na: biologię, fizykę, chemię, informatykę. Odpowiednio wysoki wynik z tych przedmiotów może zagwarantować miejsce na akustyce. Po ukończeniu procesu rekrutacji przed studentami 5 lat edukacji. Wybierając uniwersytet, po pierwszym etapie student podchodzi do obrony pracy dyplomowej w celu uzyskania tytułu licencjata. Decydując się na uczelnię techniczną, pierwszy etap kończy się również obroną pracy, natomiast absolwent uzyskuje tytuł inżyniera. Kolejnym etapem są studia uzupełniające, których ukończenie nagradzane jest przyznaniem tytułu magistra. Studenci w trakcie nauki mogą dokonać wyboru specjalizacji. W zależności od uczelni wybór będzie różnorodny. I tak



na przykład Uniwersytet Warszawski oferuje studentom: Protetykę słuchu i ochronę przed hałasem, natomiast Akademia Górniczo Hutnicza specjalizuje studentów w: Drganiach i hałasie w technice i środowisku oraz w Inżynierii dźwięku w mediach i kulturze. Nie tylko wybór uczelni, ale i specjalizacji może w znacznym stopniu wpłynąć na przyszłą pracę zawodową absolwenta.

Poziom trudności na tym kierunku zależy od umiejętności studentów. Jest to niewątpliwie wymagająca dziedzina nauki, gdyż łączy ze sobą różne zagadnienia. Nie sposób być tutaj jedynie audiofilem, ciężko będzie osobom, których jedyną mocną stroną jest analityczne myślenie. Osiągnięcie zadowolenia na tym kierunku, a co za tym idzie powodzenia odzwierciedlającego się w kolejnych zaliczonych semestrach, jest możliwe gdy łączy się i wykorzystuje umiejętności z różnych dziedzin. Dla osób, które potrafią płynnie przemieszczać się pomiędzy różnymi zagadnieniami nauki, akustyka może okazać się kierunkiem prostym, łatwym i przyjemnym. Jednak gdy łączenie treści z różnych dziedzin dawać będzie efekt w muzyce zwany kafeoną, studia te mogą okazać się wyjątkowo trudne. Należy natomiast zwrócić uwagę na fakt, że jest to kierunek, na którym można spotkać dużo ciekawych ludzi. Akustyka przyciąga wspomnianych już audiofilów, czyli fascynatów i miłośników dźwięku. Tu spotkać można przyszłych artystów, „zwarowanych” techników i inżynierów dźwięków.

To miejsce, które wytwarza dobre vibracje, a większość studentów potrafi słuchać... przynajmniej dźwięków. Warto zwrócić uwagę, że jest to miejsce dla osób, którym „słoi na ucho nie nadeptał”.

Akustyka to kierunek interdyscyplinarny, a co za tym idzie stwarza ogromne możliwości dla jej absolwentów. Praca w zawodzie dostępna jest w medycynie, szeroko pojętym przemyśle, sztuce i mediach. Studia te otwierają także możliwości przed osobami przedsiębiorczymi i dobrze czującymi się w handlu. Rozwój zawodowy to nie tylko klasycznie rozumiane zatrudnienie. Akustyka daje możliwość rozwoju na polu naukowym. Opracowanie nowych patentów może stanowić doskonale źródło dochodu, kształcenie specjalistów, konsulting w zakresie ochrony przed hałasem to świetny sposób na łączenie rozwoju wiedzy z pracą zarobkową. Dla chętnych znajdują się także miejsca pracy w administracji państwowej lub samorządowej. Jeśli natomiast absolwent ma duszę artystyczną, to z pewnością odnajdzie się w studiu nagraniowym, będzie mógł próbować swoich sił w reżyserii dźwięku, a także może liczyć na wdzięczność wszelkich gwiazd estrady za doskonale przygotowanie sceny pod kątem dźwięku. Wynagrodzenia w tej branży są tak różnorodne jak możliwości zatrudnienia. I tak na przykład realizator dźwięku i protektyk słuchu mogą liczyć na pensję w okolicach 5000 zł. W administracji można spodziewać się około 4000 zł. Inżynier akustyk zatrudniony w przemyśle zarobi ok. 7000 zł. W handlu zarobki przeważnie uzależnione są od realizowanych celów sprzedażowych, ale swobodnie można mówić o kwotach rzędu 5000 zł i więcej. Znalezienie pracy nie powinno stanowić dużego problemu ze względu na szeroki zakres kompetencji inżyniera akustyka, jednak absolwent powinien być przygotowany na poświęcenie czasu na poszukiwania.

Akustyka to kierunek skierowany jedynie dla fascynatów i hobbystów. Studiowanie wymaga przede wszystkim predyspozycji takich jak na przykład dobry słuch. To także umiejętność łączenia i wykorzystywania informacji z różnych dziedzin, gdyż jest to kierunek interdyscyplinarny. Akustyka to niewątpliwie kierunek ciekawy, stwarzający możliwości zawodowe. Inżynierowie akustyki nie są najbardziej poszukiwanymi pracownikami w Polsce, ale z pewnością znajdzie się dla nich zatrudnienie. Należy pamiętać, że nie jest to miejsce dla osób przypadkowych, ale jedynie dla tych, którzy wiedzą czego chcą od życia. Zapraszamy na akustykę. ■

Michał Pacholski



Szkoła Wynalazców

dozwolone do lat 15

Jak będzie szybciej: mieszać gorącą czarną kawę z mlekiem i poczekać aż mieszanina ostygnie do odpowiedniej temperatury, czy najpierw poczekać aż czarna kawa ostygnie, po czym mieszać ją z mlekiem i pić?

Nie oczekiwaliśmy analizy z zastosowaniami prawa Stefana – Boltzmanna, różnych dróg „wędrówek” ciepła: przez przewodzenie, promieniowanie, konwekcję, parowanie, topnienie, itd. Problem jest prosty i wymagał jedynie elementarnej orientacji w procesach ogrzewania i stygnięcia różnych ciał. Najpierw trzeba sobie jasno i ściśle odpowiedzieć na pytanie: czym jest studzenie czegośkolwiek gorącego? Z punktu widzenia elementarnej fizyki studzenie, to odprowadzenie pewnej ilości ciepła, albo do ośrodka chłodzącego – jak w chłodnicy samochodowej, albo rozproszenie do otoczenia. Najprostsze wzory na ilość ciepła odprowadzanego do otoczenia niezależnie od innych czynników podają zależność od różnicy temperatur ciała chłodzonego i temperatury otoczenia lub ciała chłodzącego (np. wody). Jasne jest, że gorąca, świeżo zaparzona kawa odda więcej ciepła w jednostce czasu, niż kawa częściowo schłodzona. Zadanie byłoby trudniejsze, gdybyśmy założyli, że garnek z kawą jest chromowany i błyszczący, a garnuszek na mieszkankę kawy z mlekiem matowy i czarny. Nie pogłębiając tej analizy łatwo zauważyć, że gorąca, świeżo zaparzona kawa, w jednostce czasu odda do otoczenia więcej ciepła niż kawa z mlekiem o temperaturze znacznie niższej niż kawa. W sumie: krócej będzie poczekać

aż kawa trochę ostygnie i dopiero wtedy mieszać ją z mlekiem. A co proponowali nasi czytelnicy?

Jacek Zieliński (4 pkt.) uważa, że proces schładzania zbyt gorącej, czarnej kawy można radykalnie przyspieszyć, wlewając ją do płytkiego talerza. Duża powierzchnia parowania przyspieszyłaby w znacznym stopniu schłodzenie kawy.

Bardzo dobry pomysł, chociaż wybiega „w bok” od zasadniczego pytania. Chodziło wyłącznie o to, która procedura da szybszy efekt: poczekanie aż kawa ostygnie i dopiero wtedy dolanie mleka, czy mieszanie kawy z mlekiem i czekanie na ostygnięcie mieszaniny.

Robert Orzechowski (4 pkt.) pisze: najlepiej byłoby wlać kawę do metalowego garnka, który jak wiadomo dobrze przewodzi ciepło i pozwala kawie stygnąć dzięki parowaniu i promieniowaniu ciepła przez ścianki garnka. Taka trochę ochłodzona kawa zmieszana z mlekiem, nadawałaby się już do picia.

Też dobry pomysł, ale również „w bok” od zasadniczego problemu.

Sebastian Makuch (5 pkt.) podał w pełni prawidłową odpowiedź: „oczywiście szybciej wytraci ciepło gorąca kawa, trzeba tylko czekać do momentu gdy dolanie mleka da mieszaninę zdatną do picia.

I o to dokładnie chodziło.

Dziękuję za udział w konkursie i zapraszam do następnych.

Wszystkim kolegom życzę sukcesów w kolejnych zadaniach i sławy „wynalazcy” w rodzinie, gronie przyjaciół i w szkole!

Nowe zadanie

W domu Jurka w pokoju „jadalnym” jest stół z gładkim, błyszczącym blatem. Mama Jurka- dla ochrony ładnego blatu- do codziennego użytku nakrywa go foliowym obrusem. Ten obrus kiedyś przesunął się mocno w jedną stronę i mama poleciła Jurkowi go wyrównać. Jurek energicznie pociągnął za jedną krawędź obrusa i... urwał spory kawałek. Okazało się, że obrus stawia duży opór

Ranking Szkoły Wynalazców

1. Mateusz Konarski.....(13 pkt.)
2. Marek Miernik.....(13 pkt.)
3. Mateusz Gawlikowski.....(14 pkt.)
4. Jacek Zieliński.....(15 pkt.)
5. Zbigniew Borek.....(5 pkt.)
6. Sebastian Makuch.....(5 pkt.)
7. Małgorzata Mickiewicz.....(5 pkt.)
8. Robert Pająk.....(4 pkt.)

Ranking Klubu Wynalazców

1. Ryszard Andruszaniec.....(14 pkt.)
2. Zbigniew Przygodzki.....(14 pkt.)
3. Rajmund Kosiński.....(11 pkt.)
4. Robert Pająk.....(11 pkt.)
5. Grzegorz Siwkowski.....(10 pkt.)
6. Mateusz Winkler.....(10 pkt.)
7. Mirosław Lubicz.....(6 pkt.)
8. Leszek Wawrzekiewicz.....(5 pkt.)
9. Krzysztof Wojnar.....(5 pkt.)

i jego przesunięcie wymaga jakiejś metody. Jakiej? To właśnie wasz problem:

„W jaki sposób należy postąpić przy próbie przesunięcia foliowego obrusa leżącego na gładkim blacie stołu”.

Zasadnicze pytanie: dlaczego przesuwaniu foliowego obrusa po gładkiej powierzchni stołu towarzyszy spory opór? Wchodzi tu w grę kilka zjawisk: na pewno niewielka lepkość folii; wiemy jak kłopotliwe jest otwieranie torebek foliowych,

obecność niewielkiej wilgoci, efekt elektrostatyczny; zwłaszcza gdy obrus leży na stole poli-turowanym i oczywiście jest często wycierany „na sucho” co powoduje naładowanie folii ładunkami elektrycznymi. Te zjawiska należy zgrabnie pokonać, żeby bezawaryjnie zdjąć obrus ze stołu lub przesunąć.

Wszystkim życzymy wnikliwości i spojrzenia na sprawę okiem fizyka. Termin nadsyłania propozycji: do końca stycznia 2022 roku.

Klub Wynalazców

bez ograniczeń wieku

Zadanie „stare jak świat”, ale ciągle zaskakuje. Konkretnie zadanie waszym było: Zaproponować sposób uniknięcia szkód, do których może dojść w wyniku zjawiska „uderzenia wodnego.

Zjawisko występuje wtedy, gdy gdzieś w jakimś miejscu sieci wodociągowej dochodzi do gwałtownego zatrzymania wypływu. Było to dość częste, gdy pojawiły się baterie jednodźwigniowe, umożliwiające zamknięcie wypływu wody jednym szybkim ruchem. Pozytywnym przykładem wykorzystania uderzenia wodnego jest tzw. „taran wodny”, który w Polsce można zobaczyć we wsi Kajny na Mazurach. Pompa, wykorzystująca efekt uderzenia wodnego działa tam nieprzerwanie od ponad 124 lat, bez zużywania energii. Wymaga jedynie konserwacji okresowej. Jak więc widać, uderzenie wodne może zdziałać wiele dobrego, ale też i złego. W uproszczonym ujęciu zasada efektu tarana wodnego, polega na przerwaniu dynamicznego procesu wypływania wody i jeśli rurociąg jest sztywny i woda nie ma możliwości łagodnie wytracić ciśnienia dynamicznego, następuje impulsowy skok ciśnienia w sieci.

W takiej sytuacji w rurociągach słychać dość głośny stuk, a najsłabsze elementy pękają. A co na to wszystko nasi czytelnicy? Zobaczmy.

Zbigniew Przygodzki (5 pkt.) pisze: nie wnikając zbyt głęboko w szczegóły, myślę, że wystarczy do rurociągu dołożyć element buforowy. Może to być zbiornik wodno – powietrzny, może być zastąpienie odcinka rurociągu rurą gumową, może być także zawór otwierający się na skutek uderzenia i wypuszczający chwilowy nadmiar wody do otoczenia.

Bardzo dobre pomysły. Oczywiście należy wziąć pod uwagę realia: a więc jeśli rurociąg jest

dostatecznie długi, to wtedy elastyczność długiej rury może być wystarczająca. W mało rozległych sieciach np. domowych, rzeczywiście wspomniane przez kolegę sposoby będą na pewno skuteczne.

Robert Pająk (3 pkt.) oglądał kiedyś film z serii „Da Vinci”, w którym bohaterowie filmu – żeby uniknąć rozmrożenia butelki wsypywali do niej kilkanaście kulek styropianowych. Gdy woda w butelce zamarza, lód, który zyskuje nieco większą objętość, ścisną kulki, a butelka pozostaje cała. W przypadku rurociągów, a raczej instalacji domowych powinna wystarczyć tuleja gumowa lub styropianowa włożona do odcinka rury.

Oczywiście, jest to ta sama idea jaką przedstawił Zbigniew, ale tu Robert zaproponował nieco inne rozwiązanie, w każdym razie skuteczne.

Kolegom gratuluję i zapraszam do następnych zadań.

Nowe zadanie:

Tym razem problem „z wysokiej półki”, no ale zadanie jest przeznaczone dla starszej grupy naszych doświadczonych czytelników. Jest to problem, z którym kiedyś pewna pani minister załatwiła się prosto, mówiąc: „sorry, taki mamy klimat!”

Oczywiście chodziło o lód, namarżający na przewodach elektrycznych, przewodach trakcyjnych sieci kolejowych itp. Najprościej powtórzyć wypowiedź pani minister i oczywiście wysłać ekipy ludzi z drągami, żeby ten lód strącali. No tak, ale



linie przesyłowe wysokiego napięcia mają tysiące kilometrów długości! No to może przepuszczać prąd o dużym natężeniu i topić lód. Też niedobrze, bo wtedy należałoby wyłączyć wszystkich odbiorców energii elektrycznej, obsługiwanych przez nagrzewaną linię. No to co można zrobić? Ktoś powie: przelecieć śmigłowcem tuż nad linią i lód spadnie od poddmuchu wywołanego jego wirnikiem. Nad linią przesyłową, a nawet kolejową może tak, ale czy spadnie cały lód? I to jest wasz problem:

„Zaproponować sposób walki z oblodzeniami linii elektrycznych wszystkich rodzajów: przesyłowych, trakcyjnych, lokalnego zasilania obiektów przemysłowych i domów mieszkalnych.”

Wiadomo, że najbardziej zawodnym elementem każdego systemu jest człowiek. W nocy, np. człowiek śpi, a lód... robi swoje! Rankiem nie mamy prądu, pociąg nie odjedzie itp., bo linie zostały uszkodzone przez nawisy lodowe. A więc musi to być system i sposób, działający autonomicznie. Oczywiście nie mamy tu na myśli robotów, pełzających wzdłuż linii i omiatających miotełką, a może nawet laserem, powierzony sobie odcinek linii. To musi być coś całkiem innego. Sięgnijcie po podręczniki fizyki z ostatnich klas liceum i po metody TRIZ.

Wszystkim życzymy powodzenia i przypominamy o terminie: do końca stycznia 2020 roku.

Vademecum Młodego Wynalazcy

Przechodzimy od tego odcinka VMW do „Systemu standardów rozwiązań zadań wynalazczych”

Wydawałoby się, że samo zestawienie pojęć: standardy i wynalazek – to antynomie! Wynalazek to przecież ma być i na ogół jest – nowość – coś czego wcześniej nie było, a standardy to ramy, przepisy i reguły postępowania. W miarę analizowania tysięcy (później milionów) wynalazków okazało się, że istnieją grupy wynalazków pokrewne w sensie podstawowej idei. Analiza tych podobieństw idei wynalazczych i ich ujęcie w system trwała wiele lat i jest w zasadzie do dziś niedokończonym procesem.

Od samego początku opracowywania TRIZ było jasne – *trzeba mieć potężny bank informacji, zawierający przede wszystkim typowe metody usuwania technicznych sprzeczności. Prace nad jego tworzeniem trwały wiele lat: przeanalizowano początkowo ponad 40 000 wynalazków, ujawniono 40 typowych metod (razem z podmetodami – ponad 100). Ustalono, że wewnątrz technicznych sprzeczności – tkwią sprzeczności fizyczne.*

W samej swej istocie fizyczne sprzeczności (FS) prowadzą do podwójnych wymagań w stosunku do obiektu: być jednocześnie ruchomym i nieruchomym, gorącym i chłodnym itp.

Nic dziwnego, że badanie metod usuwania FS doprowadziło do wniosku, że powinny istnieć parzyste metody, bardziej skuteczne, niż pojedyncze. Bank informacyjny TRIZ uzupełniono spisem parzystych metod (rozdrobienie – łączenie itd.). Później okazało się, że rozwiązywanie złożonych zadań zazwyczaj związane jest z zastosowaniem

metod kompleksowych, zawierających kilka zwyczajnych (w tej liczbie i parzystych) metod i fizyczne efekty. Ostatecznie więc wydzielono osobno silne połączenia metod i efektów fizycznych – i zestawiono pierwszą, jeszcze nieliczną grupę standardów rozwiązywania zadań wynalazczych. Pierwsze standardy ustalono empirycznie: pewne połączenia metod i efektów fizycznych spotykało się w praktyce tak często i dawały rozwiązania tak skuteczne, że narzucała się myśl o przekształceniu ich w standardy. W rezultacie więc:

Standardy – to przepisy syntezy i przekształcania technicznych systemów, bezpośrednio wynikające z praw rozwoju tych systemów.

Początkowo standardy nie były uporządkowane: były włączane do bazy w miarę ujawniania. Liczba ich szybko powiększała się: 5, 9, 11, 18... W 1979 roku zestawiono pierwszy system, zawierający 28 standardów. Systematykę prowadzono z pozycji **analizy wępolowej**. Określono podstawowe klasy standardów:

1. standardy na zmianę systemów (i zmiany w systemach);
2. standardy na zdefiniowanie i ilościowe ujęcie systemów (i zmiany w systemach);
3. standardy na zastosowanie standardów.

Do końca 1984 roku w większości szkół TRIZ stosowano systemy, zawierające 54, 59 i 69 standardów. Praktyka pokazała, że standardy to bardzo skuteczne narzędzie TRIZ.

Zaznaczyła się perspektywa: podstawowa część zadań powinna być rozwiązywana z wykorzy-

staniem standardów, podczas gdy ARIZ (Algorytm Rozwiązywania Innowacyjnych Zadań) należy wykorzystać przede wszystkim dla analizy nietypowych zadań i gromadzenia informacji, pomagającej formułować nowe standardy. Prócz tego, pojawiła się nadzieja, że przy dalszym udoskonalaniu systemu standardów przekształci się – w odróżnieniu od ARIZ – w narzędzie prognozowania rozwoju technicznych systemów.

W 1983–1986 latach prowadzono intensywne prace badawcze praw rozwoju technicznych systemów. Wg nowoczesnych pojęć rozwój systemów przebiega przez kolejne etapy: niepełne wopolowe systemy – pełne wepola – złożone wepola – wzmacniane wepola i kompleksowo wzmacniane wepola. W każdym ogniwie tego łańcucha możliwe jak przejście wwyż, na kolejny systemowy poziom, jak i przejście w dół, na niższy systemowy poziom.

Udało się odkryć pewne mechanizmy, realizujące ten ogólny schemat: przejście do bi i polisystemów, operacje zwijania, przejście na mikropoziom itd. Nowa wiedza o prawach rozwoju systemów technicznych pozwoliła wnieść korekty w strukturę systemu standardów i uzupełnić ją o nowe, skuteczne standardy. Innowacje były wypróbowane na seminariach w 1984-1986 latach. Okazało się możliwym przejście do systemu, liczącego 76 standardów.

Różnice nowego systemu:

1. Klasyfikacja standardów doprowadzona do zgodności z ogólnym schematem rozwoju technicznych systemów: proste wepola – złożone wepola – silne wepola – kompleksowo – silne wepola, przejście w nadsystem i do podsystemów.
2. Wprowadzono szereg nowych standardów. Wskazanie niektórych z nich uwarunkowane było pogłębieniem wiedzy o prawach rozwoju technicznych systemów i odpowiedziane przez logikę samego systemu standardów (zapełnienie „pustych” klatek).
3. Znacznie powiększono liczbę typowych przykładów na standardy. Przykłady dopełniają ogólną formułę standardu ważnymi w praktyce szczegółami i niuansami. W tym celu w bazę standardów włączono 15 szkolnych zadań.

Standardy – to narzędzia usuwania technicznych i fizycznych sprzeczności. Ich cel – pokonanie sprzeczności, a w ostateczności – ich obejście. Zwyciężyć sprzeczność, połączyć sprzeczne, urzeczywistnić niemożliwe – w tym leży sens standardów. Wierzmy, że znajomość systemu 76 standardów da innowatorowi skuteczne

narzędzia twórczego rozwiązania praktycznych produkcyjnych zadań.

Klasa 1. Synteza i analiza systemów wopolowych

1.1. Synteza wepola

1.1.1. Synteza wepola

1.1.2. Przejście do wewnętrznego kompleksowego wepola

1.1.3. Przejście do zewnętrznego kompleksowego wepola

1.1.4. Przejście do wepola na obszarze zewnętrznego otoczenia

1.1.5. Przejście do wepola na obszarze zewnętrznego otoczenia z dodatkami

1.1.6. Minimalne parametry oddziaływania na substancję

1.1.7. Maksymalne parametry oddziaływania na substancję

1.1.8. Selektynie maksymalne parametry

1.2. Burzenie wepoli

1.2.1. Usuwanie szkodliwych oddziaływań przez wprowadzenie obcej substancji

1.2.2. Usuwanie szkodliwych oddziaływań metodą modyfikacji posiadanych substancji

1.2.3. Likwidacja szkodliwego działania pola

1.2.4. Przeciwdziałanie szkodliwym oddziaływaniom za pomocą pola

1.2.5. Wyłączenie oddziaływań magnetycznych

Cały ten system – początkowo niezrozumiały, stanie się jasny po zaprezentowaniu zadań rozwiązywanych z pomocą tego „niezrozumiałego” systemu.

Synteza wepola

Główna idea tej klasy wyraża się w standardzie

1.1.1: dla syntezy zdolnego do pracy technicznego systemu trzeba - w najprostszym wypadku – przejść od niewepola do wepola. Nierzadko zbudowanie wepola napotyka na trudności, wynikające z różnych ograniczeń swobody we wprowadzaniu substancji i pól. Standardy 1.1.2 – 1.1.8 pokazują typowe drogi obejścia w takich wypadkach.

1.1.1. Synteza wepola

Jeżeli dany obiekt, źle poddaje się potrzebnym zmianom, a warunki nie nakładają ograniczeń na wprowadzanie substancji i pól, zadanie rozwiązuje synteza wepola i wprowadzenie brakujących elementów.

Przykłady:

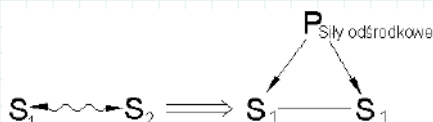
Patent Nr 283885.

Sposób deaeracji (odgazowanie, odpowietrzanie) proszkowych substancji, znamieny tym, że w celu



intensyfikacji procesu, deaerację przeprowadza się pod działaniem sił odśrodkowych.

Dane dwie substancje – proszek i gaz – same ze sobą nie współdziałające. Wprowadzono pole, (sił odśrodkowych) i utworzyło się wepole:



Odczytujemy to następująco: substancje S_1 i S_2 niewspółpracujące ze sobą, pod wpływem sił odśrodkowych zaczynają współpracować i substancja S_1 pozbywa się gazu np. powietrza.

Patent Nr 305363

Sposób nieprzerwanego dozowania sypkich materiałów przy zachowaniu stałej masy w jednostce objętości, na przykład ścierniwa, przy przyspieszonych badaniach trwałościowych silnika spalinyowego, znamienny tym, że w celu podniesienia dokładności, ścierniwo wstępnie nanosi się równomierną warstwą na powierzchnię giętkiej taśmy z łatwopalnej substancji, i podaje się ją z zadaną prędkością w strefę spalania, gdzie taśma się spala, a ścierniwo wchodzi do strefy pracy badanej pary kinematycznej.

Analogicznie przeprowadza się mikrodozowanie wg patentu **Nr 421327**: roztwór preparatów biochemicznych nanosi się na papier, a pobranie niezbędnej mikrodozy wykonuje się przez odcięcie odpowiedniego kawałka papieru.

Okazuje się, że idea wprowadzania dodatkowego elementu do wepola pomaga rozwiązać bardzo wiele różnych zadań, np. takie:

Przy walcowaniu stalowych blach na gorąco trzeba podawać płynny smar w strefę stykania się metalu z walcami. Istnieje mnóstwo systemów podawania smaru: grawitacyjnie, za pomocą różnego typu „szczotek” i „pędzli”, pod ciśnieniem (tj. stróżkami) itd. Wszystkie te systemy są bardzo złe: smar rozbryzguje się, dostaje się w pożądaną strefę nierównomiernie i w niedostatecznej ilości, duża część smaru gubi się, zanieczyszcza powietrze.

Potrzebny jest sposób smarowania, który zapewni dostarczenie w pożądaną strefę niezbędnej ilości smaru – bez jego strat i bez istotnej komplikacji urządzenia. Stosując ideę taką jak w powyższych przykładach, łatwo sprecyzować nowy sposób smarowania walców:

Patent Nr 456643

Sposób podawania płynnego smaru w strefę zgniatania przy walcowaniu na gorąco, znamienny tym, że w celu uniknięcia zanieczyszczenia środowiska i zmniejszenia zużycia, płynnym smarem nasącza się nośnik, który podaje się w strefę zgniatania razem z walcowanym metalem. W charakterze nośnika wykorzystuje się materiał, likwidujący się w temperaturze walcowania, na przykład, papierową taśmę.

Podkreślić należy, że to dopiero pierwszy standard: 1.1.1.1.

Procedura stosowania standardów wydaje się złożona i zbyt sformalizowana. Stąd liczne próby zredukowania ilości standardów, np. Jurij Bielski (wykładowca TRIZ na politechnice w Melbourne) proponuje 25 standardów, inni również próbują coś tu zmienić. Wydaje się, że komputeryzacja systemu TRIZ może w znacznym stopniu uprościć procedury stosowania standardów i w ogóle stosowanie TRIZ w praktyce. Prezentacja tych programów np. Tech Optimizer jest w MT niemożliwa z uwagi na prawa autorskie i wysoki koszt uzyskania licencji, zwłaszcza, że Tech Optimizer wykupuje się na konkretną liczbę stanowisk komputerowych. Dlatego w ramach VMW pokażemy jeszcze kilka standardów i bardzo ważny „Standard na korzystanie ze standardów”.

Warto zaznaczyć, że np. Koreańczycy – pracownicy koncernu Samsunga (koncern ten posiada obecnie ponad 38 tysięcy osób wyszkolonych w TRIZ) tylko w pierwszych ok. 10 latach od wdrożenia tej metodyki uzyskali ponad dziesięciokrotny wzrost ilości patentów zgłoszonych w ciągu jednego roku na międzynarodowym rynku. Wtedy nie mieli dostępu do Tech Optimizera, jedynie własną pracowitość. Może warto ich naśladować?

Prezes Klubu Wynalazców
Instruktor TRIZ
Jan Boratyński

KITY AVT PRZEDSTAWIA

AVTEDU

Zostań mistrzem lutownicy!

Przejdź całą serię

#AVTEDU #NaukaLutowania #KityAVT

AR

**bierz udział w konkursie
Active Reader i zgarniaj
nagrody!**

**Nieustannie czekamy na Wasze pomysły ulepszeń,
innowacji, zmian.**

Swoje propozycje nadsyłajcie na adres redakcji z dopiskiem „Pomysły” lub na e-mail: activerreader@mt.com.pl.

Zachęcamy Was również do głosowania na „Pomysł miesiąca”. Jeżeli spośród prezentowanych pomysłów jeden spodoba Wam się szczególnie, możecie na niego oddać głos, wysyłając e-mail na wyżej podany adres.

Wystarczy podać numer wybranego pomysłu.

Ten, który zbierze najwięcej głosów, zdobywa tytuł „Pomysłu miesiąca” i będzie dodatkowo nagrodzony oraz przypomniany w kolejnym numerze.

Nagrodą za pomysł miesiąca jest książka wybrana z listy nagród w konkursie Active Reader (www.mt.co.pl/ActiveReaderNagrody)

Monika Rachwalska Mieszka w bloku mieszkalnym na IV piętrze. Bloki o tej wysokości zazwyczaj nie mają wind, Wychodzić na to IV piętro bez obciążenia, to można wytrzymać, ale co ma zrobić gospodyni domowa, która niesie dwie torby z zakupami, co razem może ważyć nawet 15 kg? Monika zauważyła w internecie mini wciągarki elektryczne, do których można dorobić lub dokupić pilota. Wciągarka taka mogłaby być zamocowana na poręczy balkonu. Gospodyni wraca z zakupami, z pilota włącza wciągarkę, która opuszcza linkę z hakiem, na którym można powiesić torby. Pilotem włączamy ruch w górę i mamy „towa” już na wysokości balkonu. Idziemy bez obciążenia do mieszkania i odbieramy zakupy. Wciągarka powinna być tak zamocowana, żeby torby znalazły się nieco powyżej poręczy balkonu, żeby łatwo można było je zdjąć.

Problem jest i znają go wszystkie panie, które mieszkają w kamienicach bez wind. Wydaje się jednak, że rozsądniej byłoby zainstalować taką wciągarkę przy oknie klatki schodowej, żeby mogła obsłużyć wszystkich mieszkańców tej klatki. Małe wciągarki rzeczywiście są dostępne i są dość tanie. Rzecz na pewno warta przemyślenia.

Stanisław Wótek pozostający pod wrażeniem konkursu im. Fryderyka Chopina dostrzegł możliwość ułatwienia pracy jurorom. Uważa, że miernikiem sprawności technicznej kandydatów może być elektroniczny system zliczania ilości zagranych nut. Dziś da się to zrobić i system, zastosowany zwłaszcza na wstępnym etapie – eliminacyjnym i może jeszcze na drugim, pozwoliłby obiektywnie ocenić techniczne przygotowanie. Ponieważ konkurs ma jednak cechy rywalizacji sportowej, więc tak jak w piłce nożnej i w siatkówce wprowadzono system VAR, tak i w pianistyce coś podobnego mogłoby się przydać. Oczywiście system można by rozszerzyć, wprowadzając też zliczanie błędnie zagranych nut. Stanisław nie wierzy w możliwość weryfikacji uchem poprawności bardzo szybkich fragmentów, występujących np. w walcach Chopina, w etiudach itp. Oczywiście, od pewnego etapu konkursu przewagę w ocenie gry miałyby wyłącznie wartości artystyczne, której wymierzyć się już obiektywnie nie da.

Już słyszymy „uszami duszy” protesty wszelkich ty-pów artystów. Propozycja jest jednak interesująca i być może warto byłoby ją sprawdzić.

Rafał Kowalik przyglądając się kiedyś zabiegowi kucia konia, czyli zaopatrywania konia w podkowy, zauważył, że konie,

Pomysł miesiąca 12/2021

Wyróżniamy pomysł stworzenia „kalkulatorów do treningu ścisłego matematycznego myślenia”. Wprawdzie w dzisiejszych czasach nie musi to być oddzielne urządzenie, ale np. ambitnie i pomysłowo napisana aplikacja na smartfona, ale idea nam się bardzo podoba.

Autorem pomysłu jest Marek Wypych

„Pomysły” nie są wołaniem na puszczy! Komentujemy, oceniamy i staramy się wyrazić nasz szczerzy podziw i uznanie dla pomysłowości Czytelników. Gorąco zachęcamy wszystkich do prezentowania swoich koncepcji, również tych najbardziej zwariowanych! Wszystkie mają wartość, nawet te z pozoru niedorzeczne, bo ich krytyka może stać się twórczym zaczątkiem czegoś ciekawego!

A oto plon ostatniego miesiąca:

zwłaszcza młode, boją się tego zabiegu i proponuje zaopatrzenie kowala – podkuwacza w pistolet pneumatyczny do wbijania hufnali, który skróciłby niemiły dla konia zabieg do minimum.

Rafał najwyraźniej nie widział czynności kowala poprzedzających wbijanie hufnali. Tych czynności jest sporo: najpierw zdjęcie starej podkowy i wykonanie „pedicure” podkowy. Do tego przymierzanie podkowy i korekta jej kształtu. Koń nie ma wyjścia: musi się przyzwyczaić!

Marek Wypych rozmawiał z bratem – asystentem na wyższej uczelni technicznej i brat opowiadał mu jakie spustoszenie w umiejętnościach matematycznych czynią kalkulatory! Dziś student, mając ułamek np. skraca go sobie beztrudno na: Marek zaproponował opracowanie i wyprodukowanie kalkulatorów do treningu ścisłego matematycznego myślenia. Kalkulatory te pisałyby „odjechane” wzory i równania, a zadaniem studenta byłoby je poprawić. Oszczędziłoby to pracy asystentom, a studenci nauczyliby się prawidłowych działań.

Zjawisko coraz gorszej sprawności matematycznej obserwuje się od dawna. Na uniwersytecie im. Piotra i Marii Curie, w Paryżu, studentowi III roku fizyki, na kolokwium z grawitacji „wyszedł” promień kuli ziemskiej 10 mm! Na pytanie: czyś ty myślał, pisząc ten wynik? Student odparł: ale mnie tak wyszło z kalkulatora! Być może takie kalkulatory do nauki uważnego myślenia miałyby sens!

Tadeusz Majkowski proponuje wprowadzić drobną, ale ważną korektę do zwyczaju obdarowywania się prezentami na św. Mikołaja i pod choinkę. W rezultacie okazania sobie wzajemnej sympatii, w mieszkaniach gromadzą się góry gadżetów, niechcianych książek, nie mówiąc już o pluśzakach – jeśli są dzieci. Tadeusz oczywiście wie, że problem można w znacznym stopniu załatwić, wręczając tzw. „kartę podarunkową” ale to rzecz niewielka i nieefektywna. Proponuje wdrożyć produkcję „paczek świątecznych” – sporych prostopadłościaków ze styropianu, efektywnie zapakowanych, ze świątecznymi akcesoriami, ale mające wewnątrz kartę podarunkową. Wtedy „oko będzie zadowolone” i ręka, która znajdzie kartę też!

Pomysł istotnie niezły z wyjątkiem dzieci – które nie będą zachwycone styropianowym klockiem, a karty mogą nawet nie zauważyć. Ale jako idea sprzyjająca odgracaniu mieszkań, można pomysł polecać.



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

Klaskacz 12 V na jedno lub dwa kłaśnięcia

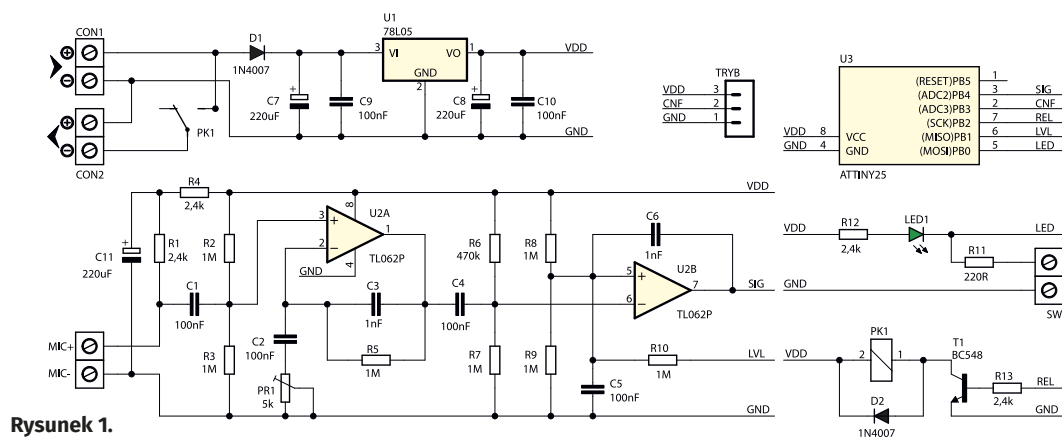
AVT3144

<http://sklep.avt.pl>

Dwufunkcyjny włącznik akustyczny wyróżnia się na tle innych tego typu urządzeń możliwością wyboru sposobu sterowania. Możliwe jest sterowanie poprzez pojedyncze lub podwójne kłaśnięcie. Układ zasilany jest bezpiecznym napięciem 12 V, a do wyjścia można dołączyć bezpośrednio taśmy LED lub żarówki 12 V. Urządzenie doskonale sprawdzi się jako zdalny włącznik oświetlenia lub efektowny sterownik urządzeń.

Schemat ideowy włącznika pokazano na **rysunku 1**. Pracą układu steruje mikrokontroler ATtiny25 takto: wewnątrz sygnałem zegarowym. Włącznik powinien być zasilany napięciem stałym o wartości 12 V. Może to być dowolny zasilacz o wydajności prądowej odpowiadającej dołączonemu obciążeniu. Dioda D1 zabezpiecza układ przed niewłaściwą polaryzacją napięcia wejściowego. Stabilizator U1 dostarcza napięcia 5 V a elementy C7...C10 zapewniają odpowiednią

filtrację tego napięcia. Sygnał z mikrofonu trafia do przedwzmacniacza, do budowy którego wykorzystany został układ TL072. Pasma przenoszenia wzmacniacza zostało ograniczone do dolnej części pasma słyszalnego. Potencjometr PR1 służy do regulacji czułości układu, natomiast zworka TRYB pozwala wybrać sposób sterowania włącznikiem. Ustawiona w pozycji „I” skonfiguruje układ do pracy w trybie pojedynczego kłaśnięcia, natomiast w pozycji „II” umożliwi sterowanie poprzez podwójne kłaśnięcie. Inne zbliżone dźwięki np. głośne puknięcie czy nawet szczekanie psa mogą również zostać zinterpretowane przez układ i spowodować zadziałanie przekaźnika. Choć w proponowanym rozwiązaniu udało się w znacznym stopniu, zredukować podatność układu na dźwięki otoczenia, a tym samym przypadkowe

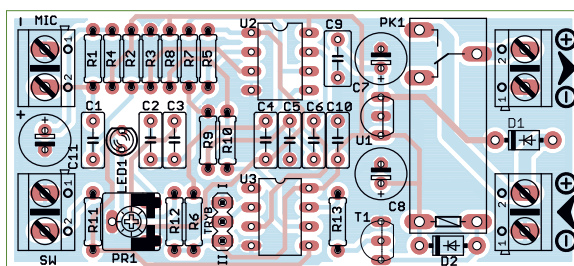


Rysunek 1.

jego zadziałanie przypadkowego zadziałania włącznika nie można wykluczyć. Jako układ wykonawczy zastosowano przełącznik o obciążalności styków 8 A/230 VAC. Pomimo znacznej obciążalności przełącznika przy sterowaniu dużymi mocami należy zwrócić uwagę na obciążenie ścieżek płytki drukowanej. Aby poprawić ich obciążalność można pocynować ścieżki lub ułożyć na nich i przylutować drut miedziany. Dioda LED1 pełni rolę sygnalizatora stanu urządzenia. Natomiast złącze SW umożliwia dołączenie dodatkowego przycisku, dzięki któremu możliwe będzie bezpośrednie przełączanie przełącznika bez konieczności kłaskania.

Montaż układu rozpoczynamy od wlutowania w płytkę rezystorów i innych elementów o niewielkich rozmiarach, a kończymy montując podstawki, kondensatory elektrolityczne, złącza śrubowe oraz przełącznik. Po zmontowaniu należy wstępnie ustawić potencjometr PR1 w środkowym położeniu. Do złącza CON2 można dołączyć dowolny odbiornik 12 V, natomiast do złącza MIC mikrofon elektretowy z zachowaniem właściwej polaryzacji. Na koniec należy dołączyć zasilanie do złącza CON1. Prawidłowo zmontowany układ działa od razu, należy tylko doświadczać wyregulować jego czułość oraz dobrać najbardziej optymalne ukierunkowanie mikrofonu.

Obsługa włącznika w trybie „I” nie wymaga specjalnego komentarza, układ reaguje na pojedyncze kłasknięcia, gdzie każde kolejne wyzwolenie zmienia stan przełącznika na przeciwny. W trybie „II” układ reaguje tylko na podwójne, następujące po sobie, w określonych odstępach czasu kłasknięcia. Drugie kłasknięcie musi nastąpić w czasie od 1 s do 2 s po pierwszym kłasknięciu. Pierwsze kłasknięcie powoduje mignięcie diody LED, która po upływie około 1 s zaświeci się sygnalizując tym samym, że to właściwy moment na kolejne kłasknięcie. Aby zapobiec przypadkowemu zadziałaniu w momencie

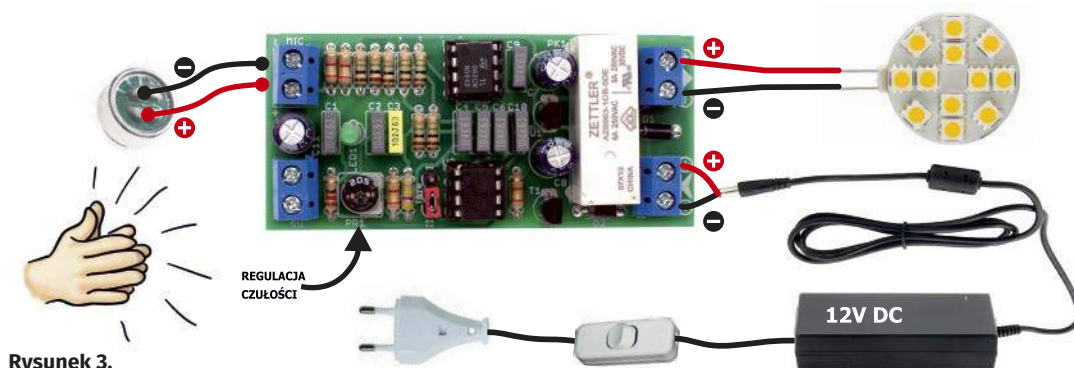


Rysunek 2.

wykrycia przez układ niewłaściwej sekwencji dźwięków jego działanie zostanie zablokowane na czas kilku sekund. Po kilku próbach sterowanie włącznikiem stanie się intuicyjne, a wyłapanie odpowiedniego momentu na kłasknięcie nie będzie stanowiło żadnego problemu.

Włącznik posiada dodatkową funkcjonalność polegającą na załączeniu przełącznika po dołączeniu zasilania.

Dzięki temu, jeśli układ zostanie włączony do istniejącej instalacji oświetleniowej za główny włącznik, to po jego włączeniu oświetlenie załączy się od razu, a po wyłączeniu głównego włącznika wyłączy się. Cecha ta nie wpływa w żaden sposób na instalację, a daje dodatkową możliwość włączania i wyłączania oświetlenia za pomocą kłasknięcia w dłoń. Przykład takiego wykorzystania włącznika przedstawia rysunek 3.



Rysunek 3.


AVT3144
<http://sklep.avt.pl>



Wszystkie niezbędne części do tego projektu zawiera kit AVT3144, w cenie 35 zł, dostępny pod adresem: <http://sklep.avt.pl/avt3144.html>



Efektowna skarbonka

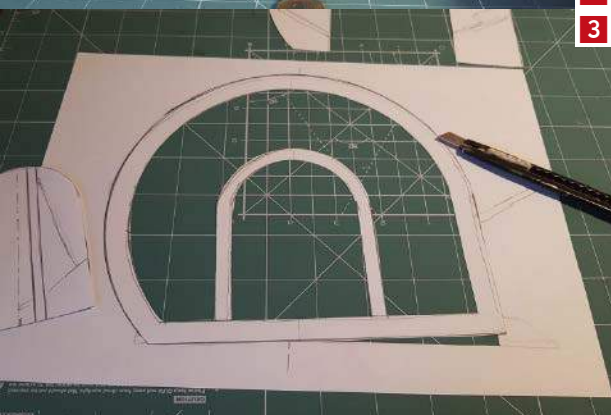
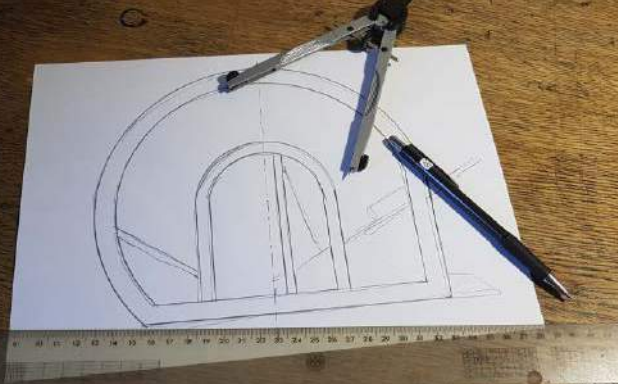
Tym razem na warsztacie efektowna skarbonka. Łatwy w budowie model będzie przypominać jednorękiego bandytę ale takiego dobrotliwego, bo wrzucone monety będą się gromadzić i można je w każdej chwili ze skarbonki wyjąć i wydać na obrany cel.

Do budowy użyjemy eko materiałów w postaci niepotrzebnego opakowania z tekstury falistej, patyczków od lodów, gumy recepturki czyli wszystkiego co pewnie już jest w domu. Jeśli chodzi o teksturę falistą to jest ich kilka rodzajów. Wykorzystajmy to. Najgrubsza jaką spotkamy ma grubość 7 milimetrów z niej można zrobić podstawę oraz tylną ściankę modelu. Z tekstury cieńszej o grubości 2 milimetrów z łatwością wytniemy przednią ściankę. Taka grubość bardzo ułatwi wycinanie trochę skomplikowanego kształtu tej ścianki, natomiast z grubszej 3 milimetrowej wewnętrzne elementy mechanizmu modelu. Płat przezroczystego celuloиду wzmocni przednią ściankę

i zapewni widok wnętrza skarbonki. Celuloid pozyskamy z jakiegoś opakowania. Narzędzia do budowy też pewnie już mamy, bo najważniejsze z nich to nóż i glutownica z klejem na gorąco. Model ma postać pudełka o przezroczystej ścianie bocznej i zaokrąglonym wierzchu. Z bocznej ścianki wystaje drewniana rączka dźwigni. To dlatego model przypomina znany z kasyn automat, nazywany jednoręki bandytą. Do szczeliny w bocznej ścianie, tuż powyżej wystającej rączki wkładamy monetę. Moneta turla się i zatrzymuje się w uchwycie na końcu dźwigni i jest widoczna przez szybkę. Opuszczamy dźwignię i puszczaamy uchwyt a ten gwałtownie wraca na swoje miejsce



1 2
3 4



5 6



1. Gotowa skarbonka, 2. Przygotowanie szablonu przedniej ścianki, 3. Wycinanie szablonu, 4. Tylna ścianka skarbonki, 5. Przednia ścianka skarbonki, 6. Gotowe ścianki skarbonki

i moneta wyrzucona w górę efektownie szybuje. Opada po drugiej stronie na pochylnie, przewraca się na płask. Wreszcie wsuwa się przez szczelinę do środka czyli skarbczyka naszej skarbonki. Aby dokładniej zaobserwować działanie naszego urządzenia, będziemy sięgać do kieszeni po monety raz po raz. Tak to działa nasza skarbonka. Aby tego doświadczyć i wszystko zobaczyć na własne oczy proponuje niezwłocznie zabrać się do pracy czyli budowy naszego modelu.

Materiały: tektura falista, kilka patyczków od lodów lub lepiej szpatulek lekarskich, patyczki

od szaszłyków lub wykałaczki, gumka recepturka, przezroczysty plastik na przednią ściankę skarbonki.

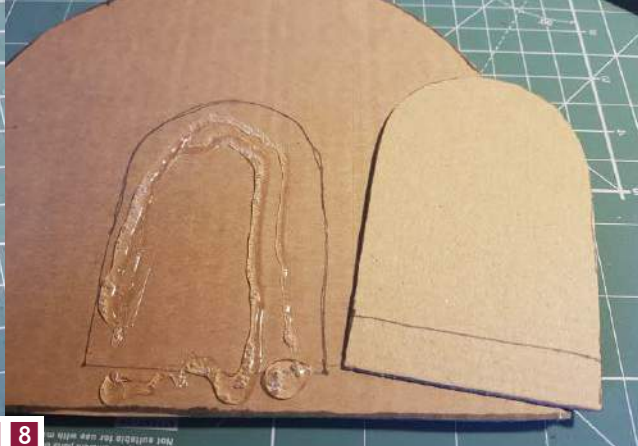
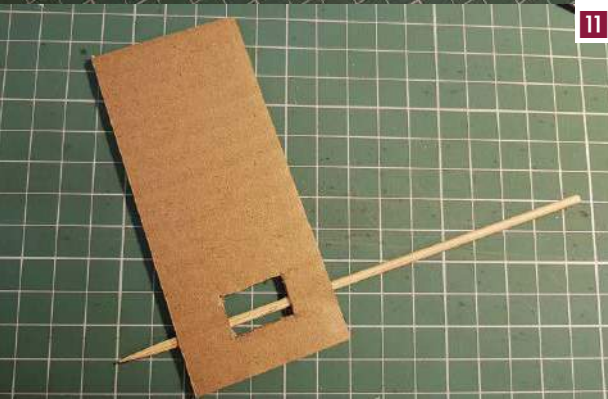
Narzędzia: ołówek, cyrkiel, nożyczki, papier ścierny, nóż z łamanymi ostrzami do tapet, glutownica serwująca klej na gorąco.

Przednia ścianka: najpierw na papierze projektujemy kształt o wymiarach 200×175 milimetrów. Promień większego koła ma 100 milimetrów a mniejszego tworzącego łuk 40 milimetrów. Prawa strona kończy się prostym odcinkiem. Do łuków dorysowujemy wewnątrz czyli linie w odległości około

Na warsztacie



7 8

9 10
11 12

7. Suport ścianki wewnętrznej, 8. Suport przyklejamy do tylnej ścianki, 9. Ścianka wewnętrzna, 10. Klejenie ścianki wewnętrznej, 11. Montowanie osi dźwigni, 12. Uchwyt gumki

12 milimetrów. Końcowy kształt przedniej ścianki widzimy na fotografii. Z tektury falistej według zaprojektowanego wzoru odrysowujemy i wycinamy ten element.

Tylna ścianka: jest zrobiona z grubej tektury falistej o zewnętrznym kształcie takim samym jak ścianka przednia. Ołówkiem zaznaczamy położenie wzmacniającego suportu wewnętrznej ścianki zakończonego łukiem. Z papierowego szablonu obrysowujemy zaprojektowany kształt i wycinamy z tektury nożem.

Suport wewnętrznej ścianki. Wycinamy z tektury kształt o wymiarach 75×110 milimetrów w górnej części zakończonej łukiem. Suport przyklejamy do tylnej ścianki w wyznaczonym szablonem miejscu. Element ten bardzo ułatwi nam prawidłowe uformowanie i przyklejenie wewnętrznej ścianki skarbonki.

Podstawa: to prostokąt o wymiarach 210×90 milimetrów wycięty z tektury. W dnie wykonujemy otwór zamykany klapką. Przez ten otwór



13 14
15 16



13. Wahliwa dźwignia zamontowana, 14. Testowanie pochylni, 15. Mechanizm skarbonki gotowy, 16. Chwytał monety

w przyszłości opróżnimy naszą skarbonkę. Jego położenie wyznaczamy doświadczalnie. Robimy tylko trzy nacięcia a z czwartej strony teksturę bigujemy. To pozwoli nam na otwieranie klapki. Widzimy to na zdjęciu. Dla zamaskowania otworu a także dla tego, żeby pieniądze się przypadkowo nie wysypały, oklejamy papierem spód podstawy. Gdy przyjdzie czas na otwarcie tej klapki wystarczy nożem przeciąć papier w wyznaczonym miejscu i wtedy odyskamy zgromadzone pieniądze.

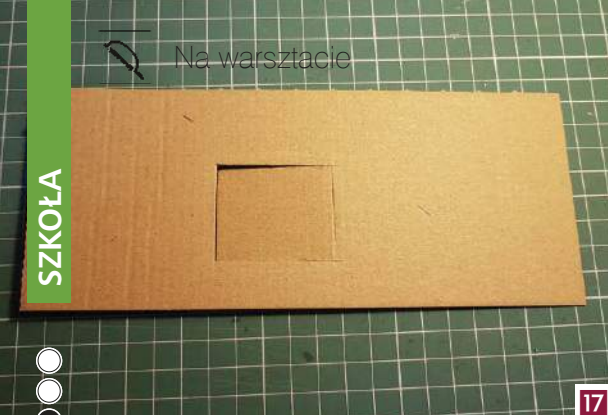
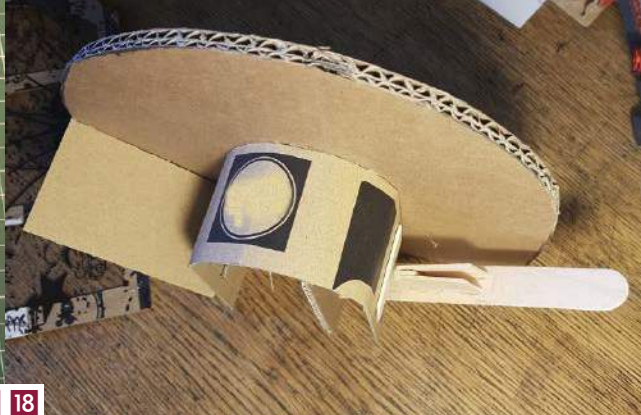
Ramię dźwigni: zbudujemy ze szpatulek laryngologicznych, które kupimy za kilka złotych w najbliższej aptece. W ich zastępstwie możemy użyć patyczków od lodów.

Wewnętrzna ścianka: z tekstury wycinamy pasek o długości 270 mm i szerokości 50 mm. Pionowe odcinki mają wysokość po 80 milimetrów. Środek paska bigujemy gęsto, raz koło razu tak jak gęsto jak są fale środkowej warstwy. Bigowanie to po prostu zgniecenie tekstury tępa stroną noża. Trzeba uważać, żeby jej przy tym nie przeciąć. Dzięki bigowaniu teksturę z łatwością uformujemy w kształt łuku. Z jednego końca tego paska w odległości

35 milimetrów od krawędzi wycinamy otwór szczelinę o wymiarach 30 na 8 milimetrów. Tym otworem moneta zostanie się do skarbczyka. Na drugim końcu wytniemy prostokątny otwór o wymiarach 30 milimetrów wysokości 18 szerokości w odległości 12 milimetrów od dolnej krawędzi. W tym otworze będzie poruszać się drewniana dźwignia zrobiona ze szpatułki. Ściankę wewnętrzną formujemy w łuk i przyklejamy klejem na gorąco do formy ścianki tylnej. Otwór na dźwignię powinien znaleźć się po prawej stronie. W precyzyjnym ustawieniu tego kluczowego elementu bardzo pomoże nam suport uprzednio przyklejony do tylnej ścianki. Widzimy to na zdjęciu.

Support dźwigni: jest wycięty z tekstury o grubości 3 milimetrów. Ma postać prostokąta o wymiarach 110×50 milimetrów. W tej części wycinamy prostokątny otwór o wymiarach 20×17 milimetrów w odległości 10 milimetrów od dolnej krawędzi. Do tego otworu wsuwamy patyczek od szaszłyka i skracamy równo z szerokością tej ścianki. Patyczek tkwiący w grubości tekstury ma możliwość osiowego obracania się w ścianie tekstury i będzie

Na warsztacie


17 18
19 20


17. Dno skarbonki, 18. Chwytek monety przymocowany do dźwigni, 19. Wzmocnienie nad otworem ścianki, 20. Przyziarniki do przyklejenia dna

osią obrotu dźwigni. Drugi odcinek patyczka przyklejamy do powierzchni ścianki pionowo, ale tak by lekko odchylał się od powierzchni a górna część nie była pokryta klejem. O ten patyczek będzie zaczepiona gumka trzymająca dźwignię w górze. Gotowy suport wkładamy od środka wewnętrznej ścianki, przyklejając go równocześnie do zaokrąglonego prostokąta przyklejonego uprzednio do tylnej ściany skarbonki.

Dźwignia: jest zrobiona ze szpatułki laryngologicznej. Jeden jej koniec powlekamy klejem na gorąco i przyklejamy do osi obrotu tkwiącej w suporcie dźwigni. Widzimy to na fonografii. Do szpatułki, za pomocą kleju na gorąco, przyklejamy od dołu, prostopadle do jej powierzchni odcinek patyczka nieco szerszy od szpatułki. Będzie on zapobiegał zsunięciu się gumki, którą połączymy dźwignię z pionowo przyklejonym do suportu patyczkiem. Gumkę odpowiedniej długości włożymy otworem w ścianie wewnętrznej i zaczepimy o pionowy patyczek pomagając sobie pęsetą. Dla próby palcem naciskamy naszą dźwignię szpatułką a gdy ją puścimy, powinna pod wpływem kurczącej się gumki gwałtownie wrócić na swoją górną pozycję.

Do powierzchni dźwigni przyklejamy jeszcze chwytak monety.

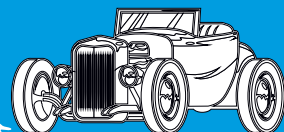
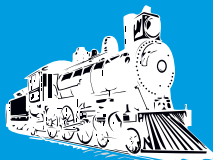
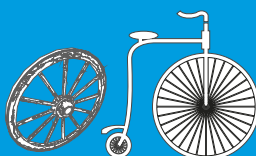
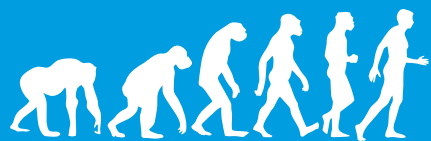
Chwytek monety: zrobimy go z trzech odcinków szpatułki. Dwa z nich mają długość 35 milimetrów i trzeci 10. Wkładając najkrótszy pomiędzy dwa dłuższe skleamy je na kształt litery „U” Aby ramiona były lekko rozwarłe włożymy podczas klejenia dwie monety. Chwytek obrabiamy papierem ściernym i przyklejamy do powierzchni dźwigni w takiej odległości, by nie mógł zawadzić o powierzchnię ścianki zewnętrznej. Opuszczając dźwignie możemy z łatwością wyznaczyć to miejsce.

Pochylnia: prostokątny kawałek tektury długości 90 mm i szerokości 50 mm. Na pochylni ląduje moneta po wystrzeleniu i po niej zsuwa się do skarbyczyka. Pochylnię przyklejamy klejem na gorąco tuż pod szczeliną w wewnętrznej ścianie i do tylnej ściany skarbonki. Powinna być przyklejona pod kątem około 30 stopni. Warto położyć monetę i wypróbować czy ta będzie zsuwać się bez problemu i czy nie utknie na jakiejś nierówności w szczelinie.

Obudowa czyli ścianka zewnętrzna: dla ułatwienia montażu składa się z dwóch części. Jedna pionowa o wymiarach 120 milimetrów ma wycięty otwór na dźwignię i szczelinę na monetę.



99

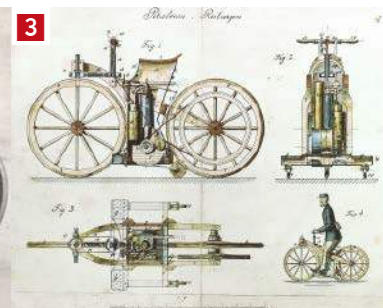
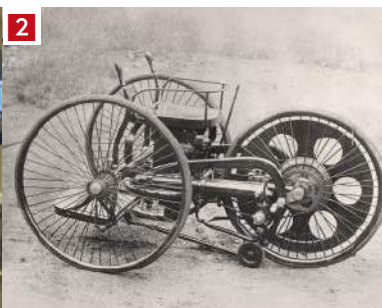


MOTOCYKLE

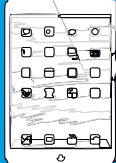
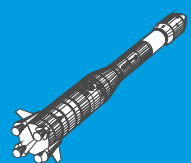
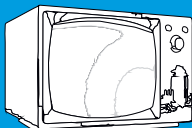
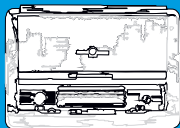
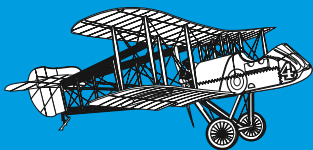
W Paryżu hrabia Made de Sivrac demonstruje wielce oryginalny na owe czasy pojazd. Dwa koła pożyczone z wozu konnego umieszczone jedno za drugim połączono w nim drewnianą ramą.

Amerikanin Sylvester Howard Roper zostaje wynalazcą napędzanego parą, dwucylindrowego velocypedu. Była to w zasadzie wczesna forma roweru, ale, co ważniejsze, był to pierwszy rower napędzany parą. Choć można kwestionować, czy projekt Ropera powinien być uważany za „motocykl”, to niewątpliwie miał dwa koła i silnik parowy opalany węglem, co było bardzo nowoczesne jak na swoje czasy. Konstruktor aż do śmierci, która miała miejsce w 1896 r., rozwijał swój pomysł. Powstały udoskonalone modele, z których najbardziej rozpoznawalna jest konstrukcja z 1894 r.

Pierwszym samojezdnym jednośladem był motocykl parowy Michaux- Perreaux, zbudowany we Francji. Konstrukcja właściciela fabryki rowerów Pierre'a Michaux oraz jego syna Ernesta stanowiła połączenie nowoczesnego, jak na tamte czasy, stalowego bicykla z lekką, jednocylindrową maszyną parową Louis-Guillaume'a Perreaux (1). Velocyped mógł być napędzany zarówno za pomocą pedałów umieszczonych w przednim kole, jak również za pomocą niewielkiego silnika parowego, który umieszczony był na ramie pod siodełkiem i przenosił napęd za pomocą paska na tylne koło. Taka konstrukcja umożliwiała też użycie obydwu napędów równocześnie. Z tego względu opisywany velocyped można uznać za pierwszy motorower w historii (rower z silnikiem pomocniczym), z drugiej strony, pod pewnymi względami, jest to również jeden z prekursorów późniejszych motocykli. Para wytworzona w niskociśnieniowym kotle podgrzewanym palnikiem naftowym wprawiała w ruch maszynę parową dwustronnego działania, która za pośrednictwem dwóch przekładni pasowych, umieszczonych po obu stronach piasty tylnego koła, napędzała pojazd. Mógł on osiągnąć prędkość 16 km/godz. Pomysł był rozwijany w kolejnych latach – w 1884 r. doczekał się wersji trójkołowej, jednakże pojazdy te nie odniosły komercyjnego sukcesu. Historycy motoryzacji mają dziś problem z ustaleniem, który z velocypedów był pierwszy: Perreaux-Michaux czy Ropera. Wiele wskazuje na to, że powstawały równolegle, nie czerpiąc wzorów z siebie. W każdym razie dały one początek późniejszym motocyklom i motorowerom.



1. Rekonstrukcja pojazdu Michaux- Perreaux, 2. Trójkołowy benzynowy pojazd Butlera, 3. Reitwagen Daimlera



lata 80–90. XIX wieku

Angielska firma The Humber Company buduje pierwszy motocykl napędzany silnikiem elektrycznym. Po latach prób i testowania prototypów w 1898 powstaje model Pennindton (tandem z napędem elektrycznym), a zaraz po nim Ladies Motor Safety z silnikiem umieszczonym za siedzeniem kierowcy.

1881–1888

Amerikanin Luis Copeland produkuje dwukotowe welocypedy napędzane silnikiem parowym. W 1887 r. doczekały się również wersji trójkotowej. Nie miał on pedałów podobnie jak wynalazek Ropera, co bardziej zbliżało go do późniejszych motocykli. Siodełko umiejscowione było na dużym, tylnym kole, a układ kierowniczy obejmował mniejsze, przednie koło. Silnik umieszczony był pod kierownicą i za pomocą paska napędzał koło tylne. Copeland nie tylko sprawił, że kocioł parowy był mniejszy, ale mógł jeździć z prędkością 12 km/h, co było sporym osiągnięciem jak na jego czasy.

1884

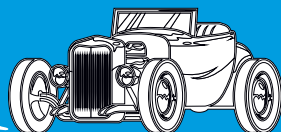
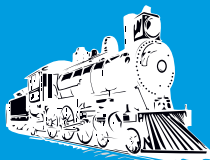
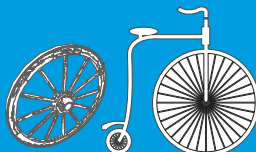
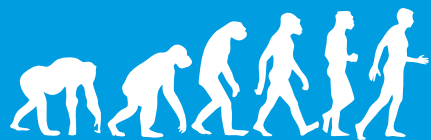
Brytyjski wynalazca Edward Butler prezentuje trójkotowy pojazd z silnikiem benzynowym na wystawie Stanley Cycle Show. Konstrukcja Butlera miała dwa symetrycznie rozmieszczone koła z przodu oraz jedno koło z tyłu (2), napędzane łańcuchem od silnika gaźnikowego chłodzonego cieczą. Brak pedałów zbliżało go również do późniejszych motocykli, choć z drugiej strony można dopatrzeć się analogii do wczesnych (choć późniejszych niż Butlera) konstrukcji samochodowych np. Benz Motorwagen 1 z 1886 r. W annałach archiwalnych historyk brytyjskiej motoryzacji George Nicholas Georgano opisuje ten wynalazek Butlera jako wynalazek, który powstał jako wczesny trójkotowy samochód benzynowy o nazwie Butler Petrol Cycle. Z powodu braku zainteresowania, wynikającego m.in. z ograniczeń prędkości w Anglii i obowiązku tzw. czerwonej flagi, Butler porzucił prace nad swoim projektem, pojazd oddał na złom i sprzedał prawa patentowe do niego, a sam skierował swe zainteresowania w stronę silników stacjonarnych i morskich.

1885

W rok przed zbudowaniem pierwszego praktycznego samochodu, niemiecki konstruktor i przemysłowiec Gottlieb Daimler zbudował tzw. Reitwagen mit Petroleumotor (3) – pierwszy motocykl napędzany silnikiem spalinowym, zbudowanym przez Daimlera specjalnie do tego celu. Był to jednocylindrowy silnik chłodzony powietrzem, o pojemności skokowej 264 cm sześciennych i mocy 0,5 KM. Rama i koła tego motocykla były drewniane, zaś koła opasane żelaznymi obręczami, bez amortyzacji zawieszenia. Napęd z silnika przenoszony był na tylne koło przekładnią pasową dwubiegową. Jako kierownicę zastosowano pojedynczy drążek, tzw. rumpel, którym kierowało się podobnie jak np. we współczesnych łodziach żaglowych. Oryginalną cechą tego pojazdu były dwa boczne kółka podpierające, zabezpieczające przed wywróceniem. Konstrukcja pojazdu została odpowiednio wzmocniona, a testów nowego jednoślada, dokonał Paul Daimler, syn Daimlera. Nad projektem pracowano jeszcze w 1884 roku, ale oficjalnie Daimler i Maybach, z którym Daimler współpracował nad projektem, otrzymali patent na tę konstrukcję 29 sierpnia 1885 roku. Pojazd nigdy nie wyszedł poza stadium prototypu, a pierwsza przejażdżka zakończyła się awarią, gdy siedzisko zajęło się ogniem.

1892

Pojazd, który można uznać za prototyp współczesnego motocykla, zbudował Francuz Felix Theodor Millet. Jego pojazd miał ramę wykonaną z rur stalowych, koła podobne do współczesnych, składały się ze stalowych szprych i stalowych obręczy z nałożonymi oponami pneumatycznymi. Napędzany był silnikiem spalinowym o pięciu cylindrach.



1894

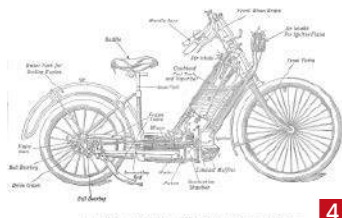
1895

1903

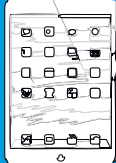
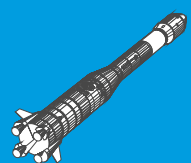
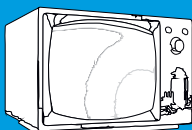
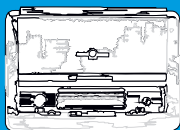
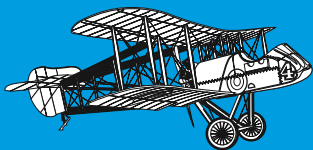
Prawdziwy przełom i rozwój epoki motocykli datuje się na 20 stycznia 1894 r., gdy Heinrich i Wilhelm Hildebrand, wspólnie z Aloisem Wolfmüllerem, opatentowali pojazd (4) pod nazwą „Motorrad” (pol. „motocykl”) i od 1896 r. rozpoczęli produkcję seryjną. Zastosowano w nim dwucylindrowy silnik czterosuwowy o pojemności 1489 cm³, mocy 2,5 KM przy 240 obr./min., chłodzony cieczą umieszczoną w zbiorniku pod tylnym błotnikiem. Jednostka ta pozwalała motocyklowi o masie ok. 50 kg na osiągnięcie prędkości ok. 45 km/h. Ciekawostką jest fakt, iż pojazd nie posiadał typowego sprzęgła, a wałem korbowym, podobnie jak w lokomotywach, było bezpośrednio tylne koło. Sama nazwa pojazdu weszła na trwałe do słownika motoryzacji i funkcjonuje do dziś. Pojazd ten uznawany jest zarazem za pierwszy motocykl produkowany seryjnie. W ciągu trzech lat warsztat braci Hildebrandów i Aloisa Wolfmüllera opuściło 1500 takich pojazdów, a dalszych tysięcy wyprodukowano we Francji na podstawie licencji. Motocykl zaopatrzony był w pneumatyczne ogumienie wynalezione przez Johna Boyda Dunlopa w 1888 roku. Pomysł podchwyciło wielu innych przedsiębiorców i w kolejnych latach rozpoczęła się produkcja różnych modeli motocykli, a popyt na tego typu pojazdy wzrastał. Wymienić należy tu takie firmy, jak Metz Company, Triumph (od 1898 r.), Royal Enfield (od 1899 r.), Indian (od 1901 r.), NSU (od 1901 r.), Norton (od 1902 r.), Harley-Davidson (od 1903 r.) oraz Birmingham Small Arms (od 1910 r.).

Firma Dedion-Buton, założona przez francuskiego pioniera motoryzacji Jules'a-Alberta de Diona wprowadza innowacyjny silnik czterosuwowy z zapłonem elektrycznym. Silnik osiągał wysoką jak na tamte czasy prędkość 1800 obr./min. Powstały wersje o pojemności 185, 215 lub 250 cm³. Ich moc wynosiła odpowiednio: 0,75, 1,25 i 1,75 KM. W układzie zasilania stosowano gaźnik powierzchniowy. Urządzenie to działało w niezwykle prosty sposób. Strumień zasysanego przez tłok powietrza prowadzony był specjalnym przewodem przez zbiornik paliwa, gdzie rozdzielone powietrze, przelatując nad lustrem cieczy, mieszało się z oparami benzyny. Utworzona takim sposobem mieszanka wędrowała dalej do cylindra, gdzie następował jej zapłon. Napęd obu tylnych kół realizowano za pomocą przekładni zębatej, mechanizmu różnicowego i odciążonych półosi. Zastosowanie tego rozwiązania tej firmy winduje produkcję motocykli na świecie.

William Harley i jego wspólnicy, Arthur oraz Walter Davidsonowie, zakładają Harley-Davidson Motorcycle Company (5). Chociaż firma początkowo zamierzała sprzedawać swoje motocykle jako pojazdy transportowe, ich silnik okazał się szybki, dzięki czemu stale wygrywał wyścigi.



4. Motocykl Wolfmüllera i Hildebranda, 5. Motocykle Harley-Davidson, 6. Pierwszy model Vespy z przyczepą, 7. NSU SportMax z lat 50-tych XX wieku, 8. Honda CB750, 9. Hybryda ET-120 firmy Eko Vehicle



1932

Pierwszy polski motocykl został skonstruowany przez inż. Tadeusza Rudawskiego, a jego produkcję podjęły w tym samym roku Centralne Warsztaty Samochodowe w Warszawie. Był to motocykl Sokół, wyposażony w silnik czterosuwowy, dwucylindrowy o pojemności 995 cm sześciennych, o mocy 21 KM i prędkości jazdy 105 km/h. W 1935 roku powstaje jego odmiana Sokół 600 RT, o pojemności 575 cm sześciennych, a w 1936 roku rozpoczęła się produkcja motocykla SHL (Suchedniowska Fabryka Odlewów i Huta Ludwików). Konstrukтором SHL-ki był inżynier Rafał Ekielski.

1946

Włoski projektant Piaggio wprowadza na rynek model skutera Vespa (6), który od razu zyskuje dużą popularność.

lata 50. XX wieku

W rozwoju motocykli wyścigowych coraz większą rolę odgrywała optymalizacja aerodynamiki. Nowego typu owiewki dawały możliwość radykalnych zmian w konstrukcji motocykli. NSU i Moto Guzzi były w awangardzie tych rozwiązań i oba stworzyły bardzo radykalne projekty. NSU wyprodukowało najbardziej zaawansowany projekt (7), ale po śmierci czterech zawodników NSU w sezonach 1954-1956, zrezygnowało z dalszego rozwoju i wycofało się z wyścigów motocyklowych Grand Prix. Moto Guzzi produkowało konkurencyjne maszyny wyścigowe i do końca 1957 roku odnosilo kolejne zwycięstwa. W następnym roku, 1958, pełne owiewki zostały zakazane w wyścigach w związku z obawami o bezpieczeństwo.

1952–60

Na rynek motocyklowy wlewa się fala japońskich konstrukcji. Jako pierwsze w 1952 pojawiają się motocykle Suzuki. Trzy lata później na rynku pojawiają się również Honda i Yamaha. W 1960 dołącza do nich Kawasaki, które wkrótce zaczyna robić furorę w branży wyścigowej.

1969

Honda tworzy czterocylindrowy motocykl, określany jak pierwszy „supermotocykl”. Był to model CB750 z silnikiem rzędowym czterocylindrowym SOHC (8). W dodatku był to produkt relatywnie niedrogi i od razu odniósł sukces. Ustanowił on konfigurację silnika typu „across-the-frame-four”, konstrukcję o ogromnym potencjale mocy i osiągów.

lata 80. XX wieku

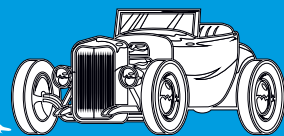
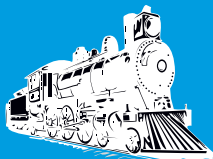
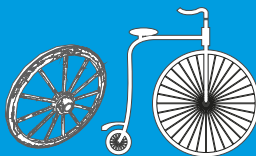
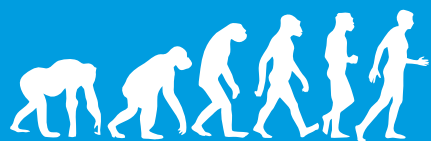
Honda i Kawasaki jako pierwsi producenci motocykli wprowadzają maszyny wyposażone w elektroniczne systemy wtrysku paliwa. System ten stanie się później normą dla wielu innych konstrukcji.

2009

Pojawia się pierwsza rynkowo dostępna konstrukcja motocykla hybrydowego, ET-120 firmy Eko Vehicle (9), który trafił do sprzedaży w Indiach.

2018–21

Harley-Davidson ogłasza plany produkcji motocykla elektrycznego, LiveWire, który ma zadebiutować w 2019 roku. W 2021 roku Harley-Davidson przekształca swój motocykl elektryczny LiveWire w samodzielną markę, odrębną od HD.



Rodzaje motocykli

Motocykle sportowe (ścigacze)

Mają zazwyczaj dużą moc i sportową sylwetkę. Wyposażone w owiewki zwiększające właściwości aerodynamiczne. Typowe dla tego typu są też twarde zawieszenie i ostre hamulce.

Motocykle turystyczne i adventure

Przeznaczone do długich wypraw. Masywne, z silnikiem o dużej pojemności oraz dużym zbiornikiem paliwa, umożliwiającym dłuższą jazdę bez tankowania. Nacisk kładzie się tu na wygodną pozycję za kierownicą. Klasa Adventure to odmiana motocykli turystycznych, które radzą sobie również dobrze na drogach nieutwardzonych. Przykładem motocykla turystycznego jest Honda Goldwing GL 1800, a turystycznego enduro – BMW R 1200 GS.

Motocykle sportowo-turystyczne

Chodzi w nich o próbę połączenia osiągnięć z większą wygodą jazdy. Przykładami mogą być Honda CBR 600 F4 lub Suzuki Hayabusa.

Nakedy

Jak sugeruje nazwa, są „nagie”, czyli pozbawione owiewek, z wygodną, bardziej wyprostowaną pozycją za kierownicą. Świetnie sprawdzają się w mieście. Przykłady to: Kawasaki Z 900, Yamaha MT-07, MT-09 lub MT-10.

Cruisery

Ich cechą charakterystyczną jest nisko położone siedzisko. Mają też zazwyczaj dużo elementów chromowanych i są stosunkowo masywne. Pozycja za kierownicą jest bardzo wygodna. Kierowca ma nogi wyciągnięte lekko do przodu. Ikonomi tego typu są maszyny Harley-Davidson.

Choppery

Chodzi w nich o tworzenie odchudzonych wersji cruiserów. Ich cechą charakterystyczną jest też długi przedni widelec oraz wysoko umieszczona kierownica.

Motocykle klasyczne

Motocykle z dawnych lat lub tzw. neoklasyki, czyli nowe maszyny stylizowane na konstrukcję z dawnych lat. Przykładem może być Triumph Bonneville T100.

Customy

D tej kategorii należą różnego rodzaju motocykle przerabiane. Panuje w tej kategorii duża różnorodność. Mogą to być np. bobbery, maszyny stylizowane na okres II wojny światowej, miejskie, wymyślne streetfightery lub café racer, stylizowane na wyścigowe motocykle z lat 60-tych.



Enduro

To cała grupa motocykli o zdolnościach terenowych. Charakteryzują się dużym kołem z przodu – najczęściej średnicy 21 cali, choć także 19 cali, koła są szprychowe, a silniki najczęściej jednocylindrowe. Do tego wąska kanapa, wysoki skok zawieszenia i szeroka kierownica.

Trajki (trike)

Motocykle trzykołowe. Jak wiadomo niektórzy najstarsi pionierzy budowy motocykli stosowali trzy koła.

Crossy

Typowo wyczynowe motocykle, stworzone do poruszania się w terenie i przeznaczone do sportów off-roadowych. Są pozbawione niepotrzebnych elementów, zwiększających masę maszyny, takich jak np. światła, kierunkowskazy, lusterka – przez co nie można się nimi poruszać w ruchu ulicznym (nie mają zazwyczaj homologacji, czym się różnią od enduro, które są dopuszczone także do ruchu ulicznego, gdyż mają światła i inne elementy wyposażenia wymaganego na drodze).

Skutery

Tak nazywa się rodzaj motoroweru ale również motocykla, który prowadzi się bez obejmowania go nogami, z silnikiem i zbiornikiem paliwa najczęściej umieszczonymi pod kanapą. Ma stałe osłony nóg, obniżoną ramę ułatwiającą wsiadanie, a zamiast podnóżków podesty dające oparcie całej stopie. ■

M.U.

*** Pisownia oryginalna ***



PRZEGLĄD TECHNICZNY

Beton zbrojony drzewem

Dotkliwy brak żelaza podczas wojny wywołał duże zastosowanie betonu a także dążność do zbrojenia go innymi tworzywami niż żelazem. Inżynier włoski, Mario Viscardini na początku r. 1918 wykonywał belki betonowe wzmocnione szkieletem drewnianym. Próby tych belek wykazały ich praktyczność podobną, jak belek żelbetonowych, oczywiście z innymi wartościami dla granicy sprężystości i dla wytrzymałości. W razie przeciążenia belki drewniano-betonowej wykazują pęknięcia szersze i wyraźniejsze niż w belkach żelbetonowych; naprawa pęknięć jest również łatwa o ile natężenie nie przekroczyło granicy sprężystości. Podobne wyniki osiągnięte przy próbach, czynionych w Austrii, v. Emperger, który stwierdził między innymi, że stosunek współczynnika sprężystości betonu i drzewa można w praktyce przyjąć równym jedności. Trudność zbrojenia betonu drzewem polega na tem, że części uzbrojenia muszą być proste i że strzemiona podtrzymujące zbrojenie nie przyczyniają się do powiększenia przyczepności między niem a betonem. Smołowanie drzewa zmniejsza przyczepność, pomalowanie zaś drzewa wapnem lub cementem znacznie ją powiększa.

7 grudnia 1921

Nowe dworce osobowe Dyrekcji Kolejowej Warszawskiej

Wkrótce po odebraniu kolei żelaznej od Niemców, Warszawska Dyrekcja Kolejowa stanęła

wobec palącego zagadnienia – odbudowy dworców osobowych. Większość z tych dworców, zrujnowana podczas odwrotu wojsk rosyjskich w r. 1915, leżała jeszcze w gruzach, gdyż okupanci, ignorując wyгоды publiczności, pobudowali na stacjach jedynie czasowe baraki dla służby kolejowej. Podróżni, w pogodę i niepogodę, latem i zimą, byli zmuszeni do oczekiwania na stacjach pociągów pod gołym niebem. W celu doraźnego zaradzenia złemu, Dyrekcja Kolejowa wybudowała przedewszystkiem szereg czasowych dworców kolejowych typu barakowego na tych stacjach, gdzie był zupełny brak schroniska dla podróżnych (Kołuszki, Warszawa-Gdańska, Ruda-Tałubaska) i następnie przystąpiła do odbudowy dworców osobowych tego typu. W owym czasie, na wiosnę r. 1919, ruch budowlany zamartwiał zupełnie: cegielnie stały beczynne, brak było na rynku drzewa budowlanego. Inicjatywa budowlana paraliżowaną była w zarodku przez tendencję do ciągłej wyższości cen robocizny. W tych niepewnych warunkach trzeba było dużo zapłacić ze strony Ministerstwa Kolei Żelaznych, a wreszcie obywatelskiego stanowiska kierowników Stowarzyszenia Przemysłowców Budowlanych w Warszawie, ażeby doprowadzić do układu, na mocy którego stowarzyszone firmy budowlane podjęły się wykonania naznaczonych do odbudowy dworców. Były to pierwsze, pod względem czasu, budowle rozpoczęte w szerszym zakresie w wyzwolonej Polsce. I choć prowadzenie ich spotykało niezwykle przeszkody wskutek braku materiałów budowlanych i trudności dowozu, spowodowanej panującym wtedy stanem wojny, a także wskutek nieporozumień z robotnikami na tle bezwzględnej walki o coraz wyższe płace, jednakże robót nie zaniedbano i dziś dobiega końca budowa kilku ostatnich z liczby dwunastu dworców stanowiących pierwszą serię robót przedsięwziętych. Projekty nowych

dworców wykonało Biuro Architektoniczne Wydziału Drogowego pod kierunkiem architekta ś. p. Bronisława Rogoyskiego, przy czynnym współudziale w zakresie projektowania architekta Romualda Millera. Wszystkie projekty traktowane są w charakterze renesansu polskiego w kształtach monumentalnych, ażeby tem uwydatnić powagę chwili odrodzenia kolejnictwa polskiego. Kształty zewnętrzne budowli zastosowano dokładnie do wymagań technicznych rozplanowania ich wnętrza, przyczem główną uwagę zwrócono na wygodę publiczności, tak wydatnie korzystającej z dworców w naszych stosunkach. Serię dwunastu dworców stanowią dworce na stacjach: Pruszków, Żyrardów, Grodzisk, Radziwiłłów, Skierniewice, Teresin, Modlin, Zieloniec, Urle, Biała, Chotyłów i Terespol. (...) widać, że przy projektowaniu strony zewnętrznej kierowano się głównie względami estetyki. Stało się to ze względu na wyżej wymienioną chęć uwydatnienia epoki, w której te budowle wykonano. W dalszym ciągu Dyrekcja przy odbudowie zamierza powodować się przedewszystkiem względami oszczędności. Zresztą, wobec zawrotnego wzrostu cen robocizny i materiałów stało się to prostą koniecznością.

13 grudnia 1921

PRZEGLĄD

ELEKTROTECHNICZNY

Wybór systemu prądu dla traktacji elektrycznej w Polsce

Wobec znaczenia, jakie ma dla elektryfikacji kraju elektryfikacja kolei, obecnie przez Rząd studjowana, Rada elektrotechniczna wyłoniła w lecie r. b. komisję, która ma rozpatrzyć jaki system prądu byłby ze względu na ogólną elektryfikację kraju najbardziej pożądany, oraz czy koleje mają czerpać prąd z sieci ogólnej, czy też ze specjalnych elektrowni. Komisja ta ma zakończyć swe prace w połowie grudnia. Z uwagi na to, że jest to sprawa pierwszorzędnej wagi, Zarząd Warsz. Koła

Stow. Elekt. zdecydował poświęcić jej jedno z posiedzeń technicznych i tym celu porozumiał się z członkiem Rady elektrotechnicznej inż. r. Podolskim, który dn. 13 grudnia wygłosił w tej sprawie referat.

1 grudnia 1921

Radjotelefony w pociągach pośpiesznych w Niemczech

Pismo amerykańskie „Telegraph and Telephone Age” w numerze z dn. 1 października ogłasza, że w wielu pociągach pośpiesznych w Niemczech zainstalowano aparaty radjotelefoniczne, które umożliwiają komunikację telefoniczną z hotelami i ambasadami w Berlinie. Próby skuteczne w pociągu, będącym w ruchu w obecności inżynierów, attachés wojskowych i przedstawicieli dyplomatycznych Stanów Zjednoczonych i Szwecji dały pomyślne wyniki. W krótkim czasie urządzenia te miały być oddane do użytku podróżujących.

1 grudnia 1921

Telefonia bez drutu

Chcąc się przekonać, na jaką odległość można obecnie przysyłać rozmowy telefoniczne bez drutu, wykonano cały szereg prób systematycznych na statku amerykańskim Bahía Blanca w drodze z Europy do Ameryki. Próby te wykazały, że słowa wysłane ze stacji w Königswusterhausen za pośrednictwem stacji lampowej o mocy 10 kW, można jeszcze było dobrze słyszeć w odległości 3500 km., zaś – ze stacji w Nauen, wysłaną za pośrednictwem maszyny o wysokiej częstotliwości o mocy 130 kW, słychać było w odległości 4340 km. Odbiór byłby możliwy na odległości większej jeszcze, gdyby nie zakłócenia atmosferyczne. Stacja w Nauen korzystała podczas tych prób tylko z części energii, jaką ma do rozporządzenia, jest przeto pewne, że komunikacja radiotelefoniczna możliwa jest na odległość dalszą od wskazanych przy użyciu generatorów o większej mocy.

1 grudnia 1921



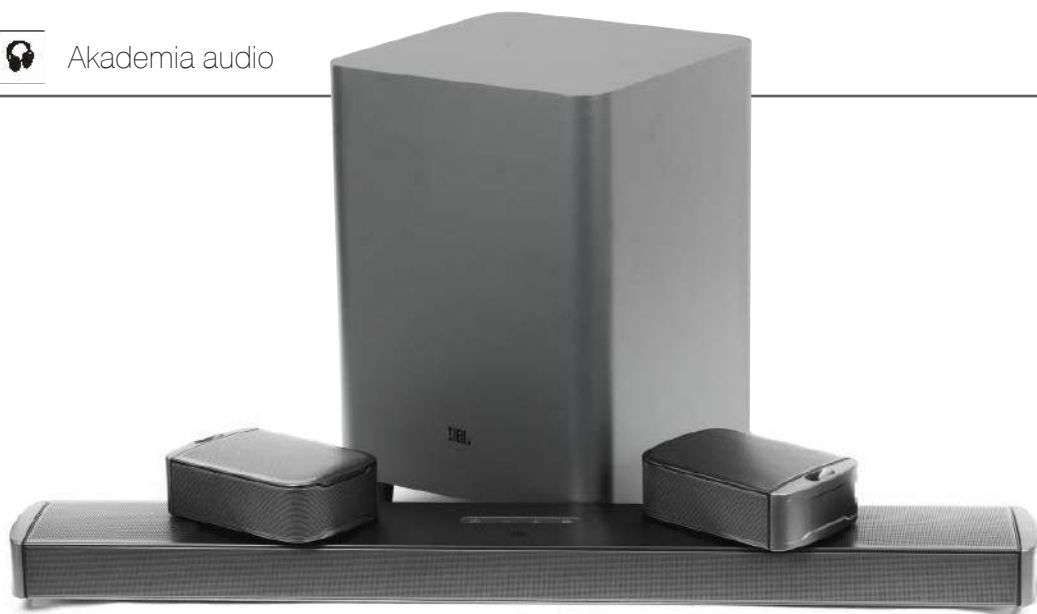
KREATYWNE
fotografowanie



Jasper Doest

Ten wielokrotnie nagradzany fotograf opowiada nam o swojej przygodzie z fotografią, inspirowaniu do działania innych, oraz... o flamingu imieniem Bob!





Nowoczesne soundbary (2)

W poprzednim numerze zrobiliśmy obszerny wstęp, przedstawiając historię, ogólną koncepcję i najnowsze trendy. Teraz trzy przykłady, których pełny test ukazał się w „Audio” 9/2021.

JBL Bar 9.1

Bar 9.1 to najlepszy soundbar JBL-a. Oznaczenie 9.1 może być mylące, bo kojarzy się raczej z klasycznym systemem „poziomym”, bez kanałów sufitowych systemu Dolby Atmos. To tylko zmyłka, która prawdopodobnie ma zwrócić naszą uwagę na obecność aż dziewięciu kanałów, system jest jednak dostosowany do wymogów atmosfowych i formalnie powinien być oznaczony 5.1.4. Punktem wyjścia jest podstawowy układ 5.1, typowy dla zaawansowanych soundbarów, do czego dodano cztery atmosfowe kanały „sufitowe” (parę przednich i parę tylnych).

W głównej części listwy znajdują się głośniki kanałów przednich – lewego, prawego, centralnego i pary przednich „sufitowych” Atmosów. Z listwą połączone są bezprzewodowe satelity, które obsługują kanały efektowe „zwykłe” (np. systemu Dolby Digital) i tylne „sufitowe” Atmosy.

Pod względem formy Bar 9.1 jest podobny do Philipsa B97. Cechą wyróżniającą te systemy są głośniki satelitarne, połączone z listwą lub nie, zasilane z ogniw akumulatorowych (gdy są odłączone). Do naładowania trzeba je od czasu do czasu podłączyć do listwy, akumulator wystarczy na ok. 10 godzin pracy.

W samej listwie mamy więc konfigurację 3.0.2. Kanały lewy i prawy obsługują systemy

dwudrożne, w każdym jeden przetwornik owalny typu RaceTrack oraz 20-mm kopułka wysokotonowa; w kanale centralnym razem z tweeterem pracują dwa przetworniki RaceTrack – bardzo słusznie, że kanał dialogowy w tak rozbudowanym systemie został wzmocniony. Natomiast kanały sufitowe są obsługiwane przez pojedyncze, owalne przetworniki szerokopasmowe.

W takiej samej roli występują one w głośnikach satelitarnych, ale „zwykle” kanały surround mają do dyspozycji tylko pojedyncze kopułki wysokotonowe.

Bezprzewodowy subwoofer jest bardzo duży (jak na element systemu soundbarowego), stoi na wysokich nóżkach, bowiem aż 25-cm głośnik umieszczono na dolnej ściance, a z tyłu wyprowadzono imponujący tunel bas-refleks.

Układ 5.1.4 oczywiście wymaga dekodów Dolby Atmos, jest też DTS:X.

Bar 9.1 wykorzystuje cały swój potencjał do kreowania przestrzeni w każdej sytuacji. U konsumentów ustawieniem wyjściowym jest tryb standardowy, w którym każdy sygnał jest traktowany zgodnie z jego charakterem, więc np. muzyka stereo odtwarzana jest dwukanałowo (co możemy zmienić). Bar 9.1 od razu uruchamia całą dostępną artylerię przestrzenną. Jednak po wnikliwym

przestudiowaniu dostępnych ustawień okazuje się, że tak być nie musi; dostępny jest „oszczędny” tryb standardowy, trzeba po niego sięgnąć do ukrytego menu (dostęp przez naciśnięcie i dłuższe przytrzymanie przycisku), a po wyłączeniu i ponownym włączeniu urządzenia automatycznie wróci rozwinęta przestrzenność. Projektanci założyli, że dysponując takim potencjałem akustycznym, należy z niego intensywnie korzystać.

Bar 9.1 ma jedno wejście i jedno wyjście HDMI (z eARC), wejście optyczne i dość nietypowe w soundbarach złącze sieciowe LAN. Wejście USB teoretycznie pozwala odtwarzać pliki audio z nośników pamięci, ale tylko w przypadku egzemplarzy przeznaczonych na rynek amerykański, a my mamy do dyspozycji lepsze (niż USB) sposoby odtwarzania muzyki. Przygotowano sieciowe systemy Chromecast i AirPlay 2, co właściwie gwarantuje już pełną swobodę, w rezerwie jest jeszcze Bluetooth (choć tylko z podstawowym kodowaniem SBC).

JBL nie opracował natomiast własnej aplikacji mobilnej. System Google Chromecast „podpada” wprawdzie pod mobilny sterownik samego Google, jest to jednak ogólne narzędzie, które nie wgryzie się w bardziej zaawansowane funkcje – trzeba będzie sięgnąć po klasyczny pilot, który oczywiście jest w zestawie.

Bar 9.1 to jedyny soundbar w tym teście, który samodzielnie (bez użycia dodatkowych urządzeń) przeprowadza pełną kalibrację akustyczną, ustawiając na podstawie wykonywanych pomiarów niezbędne parametry. W przeciwieństwie do wywoływania trybów przestrzennych, informacja o układach kalibracyjnych nie jest schowana w zakamarki menu. O takiej możliwości dowiemy się biorąc do ręki pilot

zdalnego sterowania – ułotkę przyklejono do pokrywy baterii.

Przyglądając się listwie, dojrzymy charakterystyczne otwory mikrofonów; ponieważ BAR 9.1 nie obsługuje samodzielnie asystentów głosowych, więc służą one właśnie systemowi automatycznej kalibracji. Proces kalibracji jest przeprowadzany w dwóch etapach: w pierwszym głośniki satelitarne powinniśmy umieścić w miejscu, z którego słuchamy, co sugeruje, że również w nich znajdują się mikrofony kalibracyjne; na kolejnych etapach satelity powinniśmy ustawić już za słuchaczem (czyli tam, gdzie docelowo będą się znajdowały).

Philips B97

Niedawno Philips reaktywował prestiżową serię Fidelio, do której należy również B97. Podobnie jak w przypadku Bar 9.1 JBL-a, skrajne elementy listwy można oddzielić, co jest rekomendowane do osiągnięcia najlepszych rezultatów brzmieniowych – będą pracowały w roli surroundów. Musimy je podłączyć tylko na czas ładowania akumulatorów, bowiem działają całkowicie bezprzewodowo – bez kabli zasilających i sygnałowych – na jednym ładowaniu do 10 godzin. W formie zintegrowanej system też może pracować, wciąż z udziałem doczepionych modułów, przestrzeń nie będzie wtedy tak „dokólna”, ale pomogą układy dźwięku wirtualnego.

Formalnie rzecz ujmując, mamy do czynienia z układem 7.1.2. Para kanałów sufitowych (przeznaczonych dla nich głośników szerokopasmowych) została ułożona w listwie i oczywiście skierowana do góry (tak by dźwięk odbijał się od sufitu). Subwoofer jest oczywiście aktywny i bezprzewodowy.

Zewnętrzne moduły zajmują się kanałami tylnymi, mając do dyspozycji nawet układy





dwudrożne (owalny nisko-średnionowy i wysokotonowa kopułka), podobnie jak przednie kanały lewy i prawy, natomiast centralny – parę szerokopasmowych. Z kolei kanały boczne promieniują tylko kopułkami wysokotonowymi, znajdującymi się na zakończeniach (głównej części) listwy, aby kierować promieniowanie głównie na bok, w stronę ścian, i tam wywołać odbicie. Zostaną więc one zakryte i tym samym wyłączone, gdy dopniemy zewnętrzne moduły efektowe, wtedy zintegrowana całość będzie pracować w konfiguracji 5.1.2.

20-cm przetwornik niskotonowy zainstalowano w dolnej ścianie subwoofera, co tłumaczy zastosowanie wysokich nóżek. Tunel bas-refleks wyprowadzono do tyłu.

Listwa ma dwa wejścia HDMI i jedno wyjście (z eARC), złącze optyczne, a nawet analogowe (mini-jack). Wi-Fi jest dwuzakresowe, połączenie najlepiej uruchomić w asyście aplikacji mobilnej, którą Philips przygotował we współpracy z firmą DTS, więc pachnie to strumieniową platformą DTS Play-Fi, dzięki której z łatwością zaprzęgniemy do pracy popularne usługi internetowe i utworzymy system multiroom (wraz z kompatybilnymi urządzeniami).

Wstępna konfiguracja Wi-Fi pójdzie chyba najłatwiej posiadaczom sprzętu Apple, bowiem B97 wspiera system AirPlay 2, a ten rozprawia się z takimi zadaniami bez stresu, przenikania się menu, ustawień i aplikacji. AirPlay 2 to nie tylko początkowa konfiguracja, ale także strumieniowanie. B97 proponuje także system Chromecast. Ta dwoistość (DTS Play-Fi i Google Chromecast) może wydawać się niepotrzebnym dublowaniem funkcji i komplikacją, ale dzięki temu każdy znajdzie coś dla siebie. Wiąże się z tym także wsparcie dla asystentów głosowych, chociaż soundbar nie ma wbudowanego mikrofonu, więc niezbędna jest współpraca sprzętu mobilnego.

Jest też usługa dla serwisu Spotify (Spotify Connect), oraz Bluetooth (z podstawowym kodowaniem SBC).

Nie przygotowano automatycznej kalibracji, zamiast niej producent przygotował dokładne schematy (ze zdefiniowanymi odległościami!) rekomendowanego ustawienia listwy, subwoofera, głośników tylnych, kanapy, a nawet sufitu. Można eksperymentować z korekcją (gotowe tryby EQ), poziomem basu i wysokich tonów. Wirtualne tryby przestrzenne obejmują trzy praktyczne ustawienia – standardowe (podąża za oryginalną

specyfikacją sygnału), pełną wirtualizację Upmix (wszystkie sygnały przetwarzane są do przestrzeni 7.1.2 – nawet sygnały stereo) oraz modną, sztuczną inteligencję A-I, która na bieżąco (w zależności od odtwarzanych treści) zmienia parametry przestrzenne, a także koryguje inne parametry, jak np. natężenie dialogów w kanale centralnym.

W pakiecie dekodów surround znalazło się oczywiście Dolby Atmos, a także DTS:X i wszystkie bardziej “podstawowe” systemy.

Do obsługi przewidziano tradycyjny pilot, sterowniki mobilne są również dostępne, nawet dwa – jeden do funkcji podstawowych, głównie strumieniowania, drugi do ustawień zaawansowanych.

W konstrukcji soundbarów zgromadzonych w tym teście zwraca uwagę niekonwencjonalne podejście do budowy samej listwy, której towarzyszą dodatkowe moduły kanałów efektowych. W przypadku Philipsa B97 i JBL-a Bar 9.1 można, a nawet trzeba (od czasu do czasu) je “podczepić”, aby naładować akumulatory. W Samsungu Q950A nie musimy, bo mają niezależne zasilanie, co jednak wiąże się ze stałym podłączeniem do sieci 230 V.

Samsung Q950A

Za pomocą Q950A Samsung postanowił pobić wszystkie, także własne wcześniejsze soundbarowe rekordy. Q950A ma najwięcej kanałów w najbardziej rozbudowanej konfiguracji, jaką kiedykolwiek spotkaliśmy. Sama listwa mierzy 123 cm – i to bez żadnych odłączanych elementów.

Soundbary rosną wraz z przekątnymi telewizorów: im dłuższe, tym lepiej spełniają swoją rolę, zapewniając szerszą bazę dla przetworników mających przecież tworzyć dźwięk przestrzenny.

Oprócz listwy w zestawie jest bezprzewodowy subwoofer oraz bezprzewodowe głośniki efektowe. Ale o jakie efekty tutaj chodzi? Q950A imponuje konfiguracją 11.1.4. Czwórka oznacza dwie pary – przednią i tylną – kanałów „sufitowych”. Pierwsza z nich znajduje się w samej listwie, a druga – w dodatkowych głośnikach. Promieniują one także dźwięk aż czterech kanałów tylnych „poziomych” (należących już do jedynastu kanałów zapisanych na początku formatu). W samej listwie znajduje się więc (nie licząc „sufitowych”) jeszcze siedem kanałów „poziomych”. Do trzech oczywistych kanałów przednich (lewy, prawy, centralny) dodano aż dwie pary bocznych, odchylonych pod różnymi kątami – tak aby jak najlepiej wypełnić przestrzeń (za pomocą odbić). Sama listwa jest więc formatu 7.0.2.



W dodatkowych (tylnych) głośnikach efektywnych, przetworniki rozmieszczono na dwóch sąsiednich ściankach, tak aby ich promieniowanie skierować do sufitu i ścian bocznych, podobnie jak z przetworników analogicznych kanałów w listwie. Nie licząc subwoofera, mamy więc układ aż 15 kanałów, w którym pracują aż 21 przetworniki. We wszystkich kanałach efektywnych (górnych i dolnych) są to przetworniki szerokopasmowe, a w kanałach głównych układy dwudrożne.

Przetwornik subwoofera znajduje się na bocznej ścianie, a bas-refleks wyprowadzono z tyłu.

Głośniki tylne są bezprzewodowe dla sygnału audio, ale musimy je podłączyć do zasilania 230 V.

Konfiguracja 11.1.4 wynika się standardom, nawet atmosfowym, dodatkowe tylne kanały efektowe (te promieniujące na boki, w stronę ścian) są aktywne tylko w specjalnych, „samsungowych” trybach przestrzennych.

Ciekawostką jest mikrofon ulokowany w subwooferze, pracujący w systemie automatycznej kalibracji. Koryguje więc dźwięk w bezpośrednim otoczeniu subwoofera, a nie w miejscu odsłuchowym. Jest także szerokopasmowy system automatycznej kalibracji, ale dostęp do niej mają tylko posiadacze (niektórych modeli) telewizorów Samsunga, bowiem system zaprzęga do pracy mikrofony w telewizorze (które służą tam również do innych celów). Również te mikrofony nie znajdują się w miejscu odsłuchowym i tak prowadzona kalibracja albo nie może być dokładna albo... posługuje się bardzo zaawansowanymi, nieznanyymi nam algorytmami.

Można też próbować kalibracji ręcznej, obejmującej poziomy dla poszczególnych kanałów (ale już nie opóźnienia) i korekcję częstotliwościową – pięciopasmową lub tradycyjną, niskich i wysokich.

Podstawowym ustawieniem przestrzennym jest „Standard”, w którym wszystkie sygnały zostaną odtworzone w oryginalnej konfiguracji, wszystko można rozwinąć w ustawieniu „Surround”. Jest też tryb dla graczy („Game”) oraz inteligentne ustawienie „Adaptive Sound”, które na bieżąco analizuje sygnał i modyfikuje efekty przestrzenne.

Są oczywiście wejścia HDMI oraz jedno wyjście z eARC, a także wejście optyczne. Jest Wi-Fi i Bluetooth, a jeśli posiadamy odpowiedni model telewizora Samsunga, możemy połączyć go z soundbarem bezprzewodowo. .

Do strumieniowania muzyki Samsung ma własne rozwiązania, udostępnia także Spotify Connect oraz (co jest nowością w tym modelu) Apple AirPlay 2. Nie ma natomiast Google Chromecast.

Wi-Fi zwykle zestawiamy w asyście sprzętu mobilnego i aplikacji Samsung SmartThings, która natychmiast znajduje soundbar, ale nawet w celu przeprowadzenia wstępnej konfiguracji (włączenia sprzętu do domowej sieci Wi-Fi) trzeba utworzyć konto na platformie Samsunga. SmartThings to szeroki zakres możliwości, chociaż nie daje dostępu do zaawansowanych funkcji soundbara (np. regulacja poziomu poszczególnych kanałów możliwa jest tylko klasycznym pilotem), to obsługuje inteligentny dom Samsunga (ze sprzętem AGD).

Q950A możemy też obsługiwać za pomocą komend głosowych. Wtedy okazuje się, że soundbar... jednak ma mikrofony (ale nie do kalibracji); poza asystentem głosowym jest też inteligentny układ monitorujący otoczenie – w przypadku hałasu (np. odkurzacza czy sprzętu kuchennego) odpowiednio zmodyfikuje brzmienie. ■

Andrzej Kisiel

Nie przegap grudniowego wydania „Elektroniki dla Wszystkich”



ELPORTAL.pl

EdW możesz zamówić na
www.ulubionykiosk.pl
lub w Empikach i wszystkich
większych kioskach z prasą.

W numerze między innymi:

Własna karta USB audio

Budowa własnej karty audio od podstaw to wyzwanie i okazja do zdobycia cennego doświadczenia. Nawet jeżeli nie planujesz takich działań, zapoznaj się z artykułem oraz z wnioskami, do jakich doszedł Autor.

Miernik wzmacniaczy operacyjnych. Realizacja i konfiguracja

Bardzo interesujący miernik, który pozwala zmierzyć wszystkie kluczowe parametry wzmacniaczy operacyjnych. Trzeba go jednak odpowiednio skonfigurować, co ułatwią opisane w artykule szablon.

Inteligentny dom także dla Ciebie. Smart Home – punkt centralny?

Inteligentny dom można zrealizować na wiele sposobów. W ostatnich kilku latach zmiany w tej dziedzinie wręcz wyrzuciły wcześniejsze koncepcje. Jaki będzie punkt centralny Twojego?

NanoVNA – wykonaj dokładne pomiary

NanoVNA straszy przy pierwszym kontakcie. Jeżeli się jednak nie zniechęcisz, możesz uzyskać bardzo dokładne wyniki pomiarów. Ale zależy to od kilku istotnych czynników, które omówimy.

Efekt świetlny z diodami RGB

Układ, który powstał z potrzeby chwili. Moduł ESP12 połączony został z NeoPixel Ring60 – matrycą 60 programowalnych diod WS2812B. Dzięki gotowej bibliotece NodeMCU powstała ozdoba oferująca kilkanaście atrakcyjnych efektów świetlnych.

Ponadto w numerze:

- Filozofia sieci. Epilog
- Powolny ściemniacz-rozjaśniacz oświetlenia 1 V
- Akumulatory kwasowo-ołowiowe
- Szkoła Konstruktorów:
 - Wykorzystanie gotowego modułu pomiarowego AC
 - Mając do dyspozycji napięcie stałe 12...14 V zaproponuj sposób zasilania dwóch tańcuchów diod LED o napięciu około 30...40 V, najlepiej niezależnie za pomocą przewodu trzyżyłowego

Masz może pomysł na ciekawy artykuł lub projekt? Skonstruowałeś urządzenie, które jest godne zaprezentowania szerszej publiczności?

Możesz napisać artykuł edukacyjny? Chcesz podzielić się doświadczeniem?

W takim razie zapraszamy do współpracy na łamach Elektroniki dla Wszystkich.

Kontakt: edw@elportal.pl