

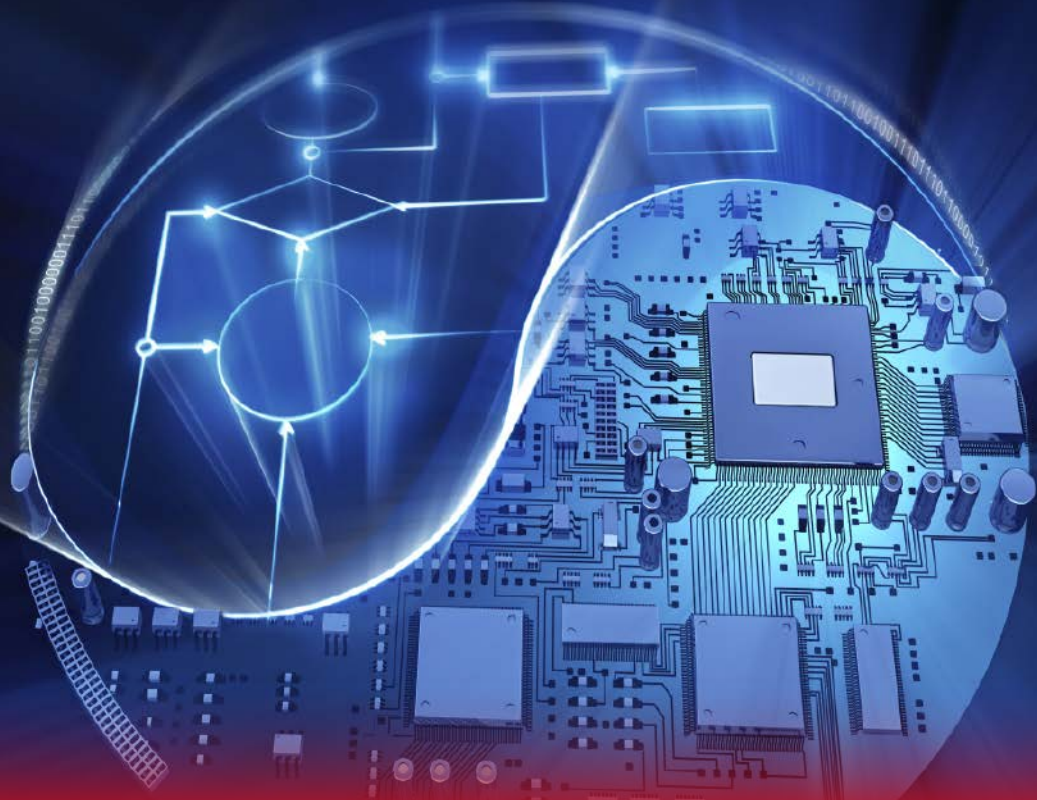


Tu przejrzyj
i kupisz ten numer

NEWS 24/7
przełóżaj codziennie
na swoim smartfonie

młody **m.technik**

Ciekawi świata są zawsze młodzi



KOMPUTERY XXI

I znów rewolucja

RAPORT: Co w fizyce wisi w powietrzu?

Piąta siła, niby-cząstki i kwantowe dziwy

ISSN 0462-9760 Indeks 365408



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 prenumeraty.

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 prenumeraty.

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 prenumeraty...”

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×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 59 59
e-mail: prenumerata@ruch.com.pl

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



Temat okładkowy

Według prognoz, do 2040 roku komputery na świecie będą zużywać więcej energii elektrycznej niż wszystkie kraje łącznie będą w stanie wyprodukować, zatem zmiany zmierzające do ograniczenia energochłonności układów obliczeniowych mają ogromne znaczenie.

Nie pytaj – co z tym krzemem? Pytaj – skąd go wziąć?

Pisaliśmy w „Młodym Techniku” nie raz i nie dwa o zmianach na rynku komputerów osobistych, rewolucji urządzeń mobilnych, a także o fundamentalnych ograniczeniach fizycznych w technice budowy procesorów. Wydarzenia ostatnich dwóch lat nie tyle zmieniły ogólne, długofalowe tendencje i przewidywania, ile przeniosły je na inny plan, bowiem dziś te sprawy nie wydają się w świecie komputerów najważniejsze.

O co konkretnie chodzi? Na przykład o to, że chociaż kwestie miniaturyzacji i zarządzania wydajnością i energią działania chipów komputerowych wciąż są wyzwaniem, jednak schodzi to na razie na dalszy plan w sytuacji, gdy owych chipów po prostu na rynku brakuje. Przyczyn braków, który, co wielu może zaskakiwać, uderza najmocniej w branżę motoryzacyjną, jest wiele i staramy się je w tym numerze w sposób pełny przedstawić,

Big Tech chce budować „własne” chipy

z pandemią na czele. Na razie więc przedstawiciele branży komputerowej mniej skłonni są do debatowania, „co z tym krzemem”, szukają raczej odpowiedzi na pytanie – „skąd wziąć krzem” w zaspokajających potrzeby ilościach.

Ostatnich kilkanaście miesięcy to również ciekawa potyczka o charakterze technologicznym i, zdaniem wielu, początek wielkiego przełomu, jeśli chodzi o sposób budowania jednostek obliczeniowych komputerów. Dzieje się tak głównie za sprawą Apple’a, który odważnie zaczął przechodzić w swoich urządzeniach na własne chipy M1 oparte na architekturze ARM, zamiast x86, w której przez dekady rządził Intel, stając się symbolem epoki krzemu.

Budowanie „własnych” procesorów staje się zresztą czymś w rodzaju „trendu” wśród potentatów Big Tech. Kwestia nowej architektury nie jest jeszcze rozstrzygnięta, bo technika ARM ma swoje wady i pojawiają się jeszcze inne pomysły na architekturę procesorów.

A to tylko hardware, połowa świata komputerów. Jeśli chodzi o oprogramowanie, software, dzieją się również ciekawe rzeczy, wśród których ekspansja sztucznej inteligencji jest najbardziej widocznym zjawiskiem. Staramy się możliwie wiernie opisać obecny stan zamieszania w świecie komputerowym. Prosimy jednak o zrozumienie, że niełatwo w tej sytuacji o prorocтва na przyszłość.

Miroslaw 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. 60

STAŁY KONKURS

Active Reader

Supernagrody!

Szczegóły na stronie 2

KSIĄŻKI

GRY

PLYTY

MODELE

NARZĘDZIA

SPRZET

AKCESORIA



Jakie są główne kierunki zmian w świecie komputerów? Obok dążenia do uzyskania coraz większej mocy obliczeniowej upowszechniają się algorytmy uczenia maszynowego a także Internet Rzeczy ze swoimi specyficznymi wymaganiami sprzętowymi. Konstruktorzy, poza szybkością i wydajnością, myślą dziś równie intensywnie o oszczędzaniu energii.



Spis treści

Temat numeru: Komputery XXI. I znów rewolucja

- 28 • Czy pisanie kodu to zawodowa przyszłość wszystkich inżynierów i naukowców? Biceps programistycznej logiki
- 33 • Grozi nam potop danych i energetyczna apokalipsa, ale pecety trwają. Komputery – co dalej?
- 37 • Programowanie i web development po nowemu. Czy da się wyrazić nowe rzeczy za pomocą starego języka?
- 42 • Wojna o krzem, którego nagle zabrakło. Procesory pożądania

Technika

- 8 Info Zoom
- 18 Dodaj do obserwowanych Horyzonty mgłą spowite
- 19 • Samolot bez ruchomych części napędzających? Jonowy wiatr w uszach
- 22 • Wyścig o udział w wyścigu na Księżyc. Nie wystarczy miliardy w kieszeni, by wyprzedzić SpaceX
- 25 • Scenariusz Carringtona. Zanim Słońce ześle nam wielką burzę
- 50 Raport MT: Co w fizyce wisi w powietrzu? Nauka, która poszła wieloma ścieżkami... oby nie w las
- 61 Nasi idole – liderzy innowacji: Konsekwentny chłopak z Warszawy – Piotr Szulczewski

m.technik

- 64 e-Technologie: Web 3.0 ponownie, ale znów inaczej. Łańcuchy, które mają nas wyzwolić

Szkola

- 67 Matematyka z ludzką twarzą: Matematyka i rower
- 72 Chemia inna niż w szkole: Obrońca stali
- 76 Edukacja przez szachy: 57. Międzynarodowy Festiwal Szachowy im. Akiby Rubinsteina
- 82 Koniec i co dalej: Era tanich dóbr. Coraz droższa droga w przyszłość
- 85 Fizyka bez tajemnic: Te niezwykle diamagnetyki
- 90 MT studiuje: Gospodarka przestrzenna Klub i Szkoła Wynalazców
- 92 • Szkoła Wynalazców, dozwolone do lat 15
- 93 • Klub Wynalazców, bez ograniczeń wieku
- 94 • Vademecum Młodego Wynalazcy
- 97 Pomysły genialne, zwariowane i takie sobie Na warsztacie
- 98 • Elektronika dla Ciebie: Regulator do prostownika
- 99 • Mechaniczne wahadło Odkryj historię wynalazków
- 106 • Ogniwa paliwowe
- 110 • Klasyfikacja ogniów paliwowych

Hobby

- 112 Akademia audio: Głośniki bezprzewodowe, czyli jak ukreślić stereo z niczego

Prezentacje

- 114 MT testuje: AOC AGON AG493UCX – gdy oczy wychodzą za uszy

- 2 Konkurs: Active Reader
- 3 Od wydawcy
- 6 Listy, Facebook
- 60 Prenumerata
- 111 Sędziwy Technik – 100 lat temu prasa pisała

Komputery XXI wieku – i znów rewolucja 27

Co w fizyce wisi w powietrzu?

50

List miesiąca

nagroda: 30 punktów AR

Szczegóły na stronie 2

Nie, Mars nigdy nie będzie naszym drugim domem

„Młody Technik” zdaje się, pisząc w sierpniowym wydaniu o eksploracji Marsa, podsycać jakieś złudzenia co do ewentualności osiedlenia się ludzi na Marsie. Niesłusznie. Powinniśmy się zmierzyć z twardą i przykrą rzeczywistością.

A brzmi ona tak, że Mars to martwa planeta, bez wody (przynajmniej w formie jakkolwiek przydatnej dla nas) i powietrza. Gleba jest toksyczna dla wszelkich form życia. Nie ma paliwa, a nawet gdyby było, potrzeba tlenu w powietrzu, żeby je spalić, a tlenu tam nie ma.

Być może kiedyś nawet setki ludzi mogłyby żyć na Marsie, ale wymagałoby to stałej linii zaopatrzenia z Ziemi w prawie wszystko, czego potrzebują. Byłoby to najprawdopodobniej bardzo podobne do utrzymywania garstki ludzi, którzy żyją dziś i pracują na Antarktydzie, nie przebywając tam zresztą na stałe, na misjach, które mają ograniczone okresy trwania. Są oni całkowicie zależni od dostaw z zewnątrz, za pomocą kosztownego transportu (a transporty na Marsa będą tysiące razy bardziej kosztowne).

Dlaczego mielibyśmy jako ludzkość opuszczać Ziemię? Załóżmy, że wybuchłaby wojna nuklearna i wszystkie obszary lądowe Ziemi zostałyby zamienione w radioaktywne pustkowia. Nawet wtedy warunki na Ziemi byłyby lepsze niż na Marsie. Mars jest jeszcze gorszym i mniej przyjaznym, zamrożonym, radioaktywnym pustkowiem, gdzie żaden człowiek nie mógłby przeżyć nawet minuty bez wyrafinowanego i bardzo drogiego sprzętu podtrzymującego życie.

Prawdopodobnie wielu ludzi mogłoby mieć ochotę tam polecieć, zobaczyć i doświadczyć tego zamrożonego radioaktywnego pustkowia. Mowa o ekstremalnej i ekstremalnie drogiej turystyce. Ale nie żadnych szans na to, aby miliardy ludzi mogły tam żyć bez żadnego wsparcia z Ziemi. Kto mówi, że jest inaczej, to albo kłamie, albo nie wie o czym mówi.

Mars zawsze będzie miejscem, gdzie potrzebny będzie pełny kombinezon kosmiczny i pełne przeszkolenie astronautyczne, które pozwoli z niego umiejętnie korzystać. Zatem tylko specjaliści mogliby tam wyjść na zewnątrz i chodzić po powierzchni. Życie innych ludzi na Marsie oznaczałoby bytowanie w tunelach pod ziemią, w których WSZYSTKO będzie rzadkim i drogim zasobem. Nawet powietrze, którym będą oddychać, kosztować będzie w ciągu roku więcej, niż większość ludzi zarabia przez ten sam okres.

Możemy mieć bazę na Marsie na początku 2200 roku, ale nie miasto i na pewno ludzie nie będą tam migrować masowo.

Szczepan Moneta z Otwocka

Kolonizacja Marsa?

Był taki moment, po lądowaniach na Księżycu, gdy wydawało nam się, że Układu Słoneczny mamy już w zasięgu ręki, a na Marsa już tylko jeden mały krok... ludzkości. Nawet niektórzy z astronautów Apollo oddali się temu entuzjazmowi i głośno zapowiedzieli rychłą misję na Czerwoną Planetę.

Niestety, zamiast polecieć na Marsa, astronauta ci stali się ostatnimi ludźmi, którzy dotarli dalej niż niska orbita okołozemska. W kolejnych dekadach, choć wspaniała armada sond kosmicznych zapuszczała

się coraz dalej w głąb Układu Słonecznego, w okolice Marsa, Saturna i Plutona, ludzie nie ruszyli się nigdzie dalej, latając co najwyżej wahadłowcami i budując stacje kosmiczne na orbicie własnej planety.

Po przeczytaniu „marsjańskiego” numeru „Młodego Technika” zacząłem zastanawiać się, czy w ogóle można myśleć o przeniesieniu ludzkości albo chociaż dużej liczby ludzi na Marsa? Czy byłoby to wykonalne, gdyby np. Ziemi i naszemu pobytowi na niej zagroziłyby jakieś straszliwe kosmiczne niebezpieczeństwo?

W dającej się przewidzieć przyszłości niestety nie ma na to szans. Przynajmniej ja nie widzę żadnego wykonalnego sposobu na kolonizację Marsa.

Chociaż rasa ludzka może z czasem przenieść tam niektórych swoich przedstawicieli, środowisko marsjańskie wymaga o wiele więcej zachodu i energii do ludzkiego przetrwania niż to, z czym mamy do czynienia na Ziemi. Utrzymanie choćby niewielkiej grupy przedstawicieli gatunku Homo sapiens na Czerwonej Planecie byłoby nad wyraz kosztowne.

Prawdę mówiąc, jedynym racjonalnym powodem, dla którego ludzkość miałaby masowo próbować przenosić się na Marsa, byłaby jakaś grożąca nam wielka katastrofa. Z drugiej strony – skoro dysponowalibyśmy możliwościami technicznymi i naukowymi, by z jej powodu migrować na Marsa, to dlaczego nie moglibyśmy jej zapobiec, wykorzystując te same możliwości techniczne, siły i środki, które mielibyśmy wykorzystać do kolonizacji Marsa.

Oczywiście, jeśli pomyślimy o tym, że mamy jeszcze cztery miliardy lat, by zdobyć owe możliwości techniczne, to taka perspektywa jest tak odległa, że wszystko wydaje się możliwe. Sensowniej jednak snuć rozważania w perspektywie kolejnego wieku lub dwóch.

Tzw. terraforming Marsa przekracza w tej chwili nasze zdolności technologiczne. I raczej nikt nie spodziewa się, abyśmy uzyskali te moce w ciągu zaledwie stulecia. Pamiętajmy, że przy budowie siedlisk, które miałyby pomieścić miliardy ludzi na Marsie, wszystkie dotychczasowe największe projekty inżynierskie na Ziemi w całej jej historii wyglądałyby jak dziecinna igraszka.

Z drugiej strony, gdyby zastanowić się nad tym, ile sił i środków jest dziś przeznaczanych na rzeczy, które służą wyłącznie do wojny i rywalizacji pomiędzy państwami, to zaczyna wyglądać to inaczej. Przecież projekt kolonizacji Marsa, zwłaszcza jeśli ratowalibyśmy się w ten sposób przed katastrofą i tak wymaga współpracy wszystkich państw, nieprawdą?

Wyobraźmy sobie, że wszystkie pieniądze przeznaczone dziś na cele obronne (czyli na wojsko, zbieranie danych wywiadowczych, bezpieczeństwo wewnętrzne, cyberwojnę itp.) mogłyby zostać przeznaczone na badania i rozwój nauki służącej do realizacji naszego



kosmicznego celu. Wystarczy, by te pieniądze pomnożyć przez lata i wieki. Byłoby tego prawdopodobnie tyle, że szybko stalibyśmy się zdolni do podbijania Marsa i dalszych połaci kosmosu.

Na ewentualną kolonizację na masową skalę patrzmy głównie w kontekście jakiegoś wielkiego niebezpieczeństwa, np. wielkiej asteroidy czy komety, które zagroziłyby planecie Ziemia i zamieszkującej ją ludzkości.

Pamiętajmy jednak, że groźba takiego impaktu nie jest jedyną rzeczą, która może nas wypędzić z planety. Za kilka miliardów lat Słońce rozszerzy się do tego stopnia, że warunki na Ziemi będą bardziej przypominać te na Merkury. Kiedy to się stanie, Mars znajdzie się prawdopodobnie w „Strefie Złotowłosej”, a przeprowadzka i tak będzie konieczna.

Zresztą zapewne już wcześniej, bo Ziemia może być niezdadna do życia już za około miliarda lat z powodu wzrostu temperatury i zniknięcia dwutlenku węgla.

W tak dalekich perspektywach trudno cokolwiek przewidywać. Ogólnie możemy założyć, że będziemy dążyć do tego, by stać się cywilizacją przemierzającą przestrzeń kosmiczną. Mars może być jakimś etapem, a może jednak wybierzemy bardziej przyjazne miejsce, jakąś bardziej podobną do Ziemi egzoplanetę, mając oczywiście już w zanadru technikę podróży kosmicznych przeprowadzanych w rozsądnym przedziale czasowym.

Kto wie?

Przemysław Krabień, Wschowa

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.





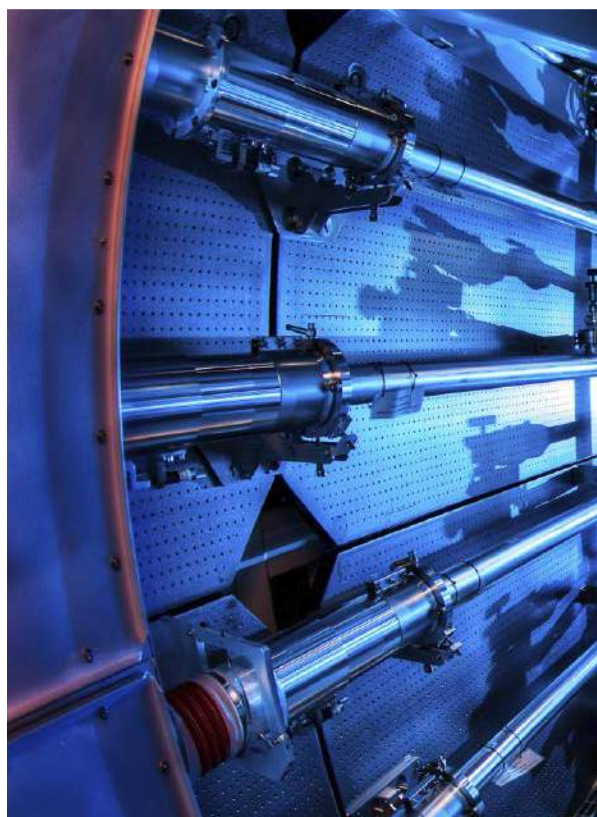
GRAFEN

Beza grafenowa zapowiada rewolucję w wyciszaniu hałasu

Oparty na grafenie supermateriał, opracowany w Centrum Materiałów i Struktur na brytyjskim uniwersytecie w Bath, waży jedynie 2,1 kilograma na metr sześcienny objętości i zdaniem specjalistów może w niektórych dziedzinach, np. w izolacji akustycznej wyciszającej silniki samolotów, okazać się prawdziwie rewolucyjnym wynalazkiem.

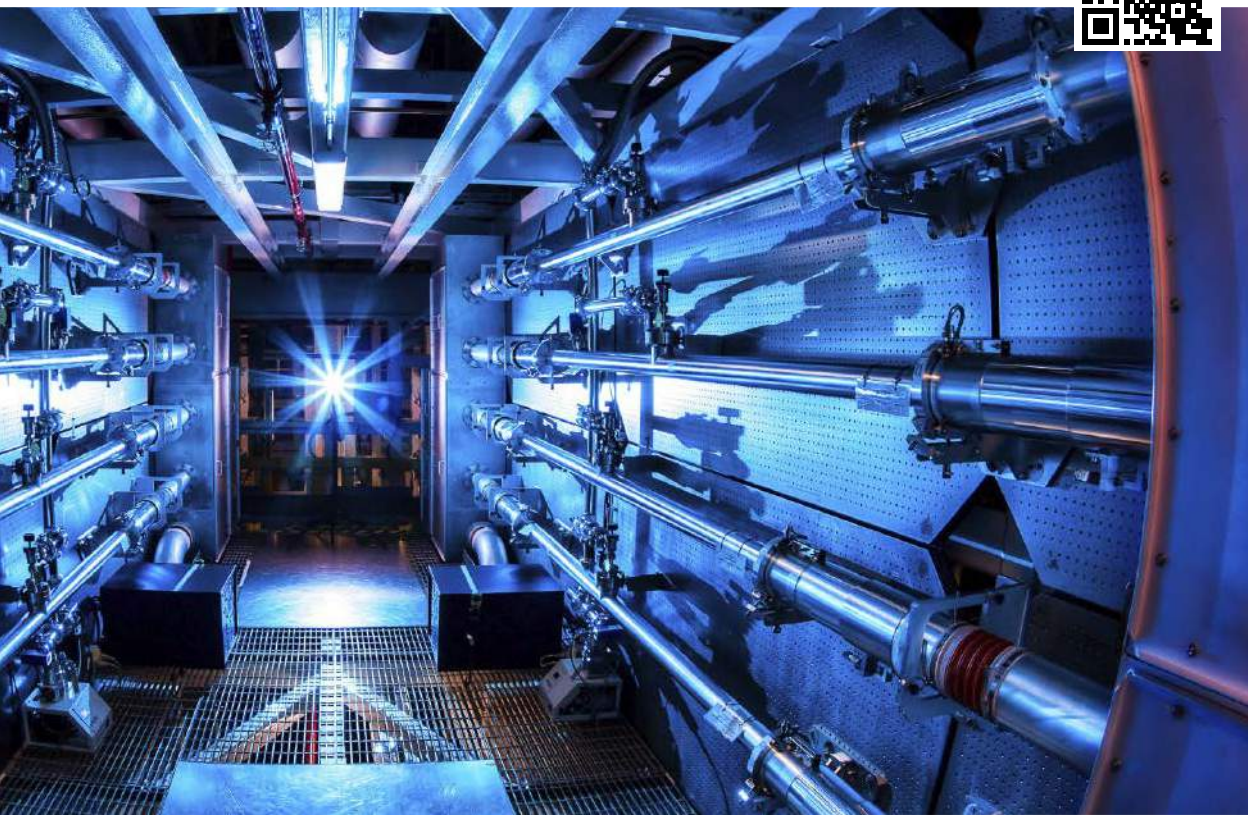
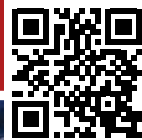
Aerożel z tlenku grafenu i alkoholu poliwinylowego konsystencją przypomina bezę i jest najlżejszą substancją dźwiękochłonną, jaką kiedykolwiek wyprodukowano. Zdaniem ekspertów, z powodzeniem może być stosowany jako izolacja w silnikach lotniczych, zmniejszając hałas nawet o 16 decybeli, co jest znaczącą redukcją, jeśli weźmie się pod uwagę, że sięgający ponad stu decybeli huk startującego silnika odrzutowego wycisza to do dźwięku o intensywności zbliżonej do suszarki do włosów. Oczywiście są materiały, które wyciszają jeszcze lepiej, ale nowy materiał jest przy okazji bardzo lekki, co czyni go obiecującym nie tylko w lotnictwie.

Naukowcy z uczelni w Bath opublikowali opis metody wytwarzania tego rodzaju materiału w czasopiśmie „Nature Scientific Reports”. Z ich pracy wynika m.in., że technika wytwarzania nowatorskiego aerożelu może być porównana do ubijania białek jaj, aby uzyskać bezę. Jak oceniają, materiał może pojawić się jako normalny produkt rynkowy za ok. półtora roku.



Laboratorium Lawrence Livermore National Laboratory w Kalifornii ogłosiło, że 8 sierpnia udało mu się wyprodukować ok. 1,3 megadżuła energii syntezy termojądrowej w National Ignition Facility, choć proces fuzji trwał bardzo krótko, zaledwie ok. sto trylionowych części sekundy. „Ten wynik jest historycznym przełomem w badaniach nad fuzją ocenił Kim Budil, dyrektor Laboratorium w komunikacie.

W ramach eksperymentu badacze skoncentrowali wielki kompleks stu dziewięćdziesięciu dwu wiązek laserowych w jednym bardzo małym punkcie, przez co udało się stworzyć potężny impuls energetyczny, o mocy osiem razy wyższej niż kiedykolwiek wcześniej osiągnięto. Jednak, jak podkreślają przedstawiciele kalifornijskiego ośrodka w komentarzach do ogłoszonego wyniku eksperymentu, należy unikać przesady, gdyż produkcja energii trwała bardzo krótko, a syntezę termojądrową należy traktować wciąż jako źródło przyszłości. Wyniki nie zostały jeszcze opublikowane w recenzowanym czasopiśmie naukowym, ale Laboratorium Narodowe Lawrence Livermore zapowiada rychłą publikację wyników eksperymentu. Choć tego rodzaju informacje zazwyczaj związane są z naukowymi publikacjami, to jednak jak powiedział serwisowi CNBC przedstawiciel

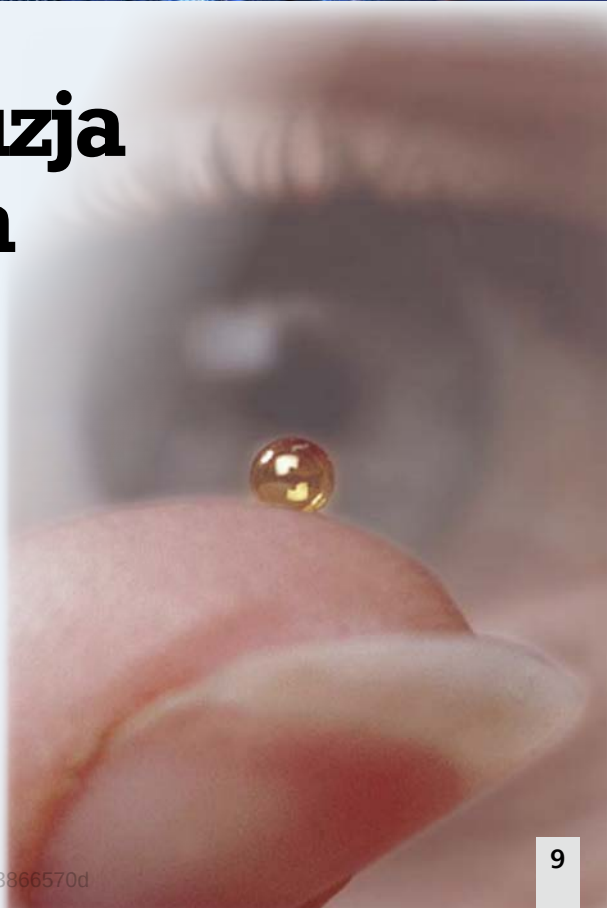


ENERGIA

Rekordowa fuzja termojądrowa w Kalifornii

ośrodka, „wiadomości już zaczęły rozprzestrzeniać się z powodu skali tego osiągnięcia, więc czuliśmy, że ważne jest, aby przedstawić fakty”.

Naukowcy z National Ignition Facility w swoim podejściu wykorzystują system laserów do podgrzewania paliwa zawierającego deuter i tryt. Kapsułki paliwa muszą zostać podgrzane do wysokich temperatur pod ogromnym ciśnieniem. Na początku badań strzały wiązką laserową dawały ok. 1 kJ energii. Po późniejszych ulepszeniach laser wytwarzał już 100 kJ energii. Z początkiem 2021 roku udało się osiągnąć 170 kJ. Zaś rekordowy strzał wiązką laserową wygenerował już 1350 kJ energii.





TECHNIKA WOJSKOWA

Izraelski laser samolotowy skutecznie niszczy drony

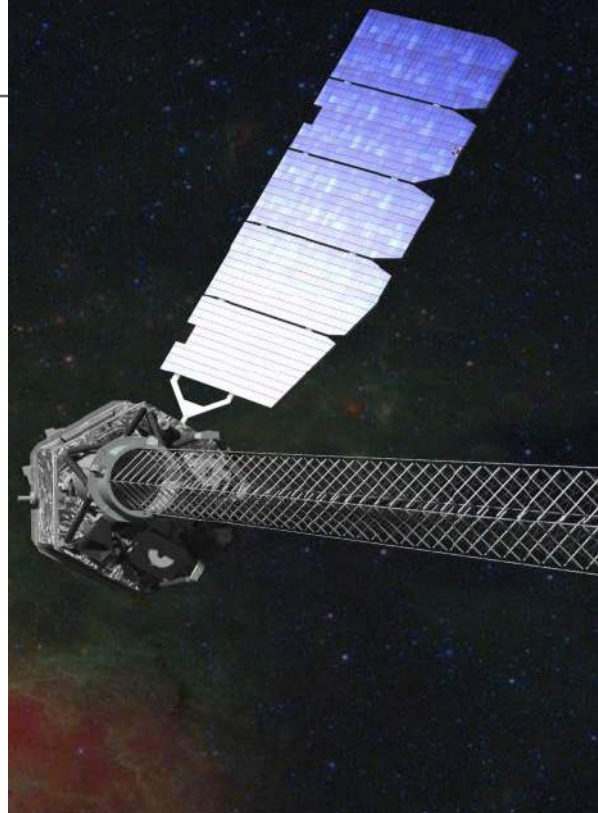
Broń laserową zdolną do zestrzeliwania dronów w locie zademonstrowało izraelskie wojsko. High-Power Laser Weapon System został zamontowany w samolocie. System działa namierza i unieszkodliwia bezałocowce w powietrzu. W materiale wideo, udostępnionym na Twitterze przez Ministerstwo Obrony Izraela, można obejrzeć, jak potężny laser wypala dziury w dronach, powodując ich zniszczenie i spalanie.

Izraelski generał brygady i szef działu badań izraelskich sił zbrojnych, Yaniv Rotem, powiedział dziennikowi „The Jerusalem Post”, że oczekuje się, iż laser będzie działał w każdych warunkach pogodowych, choć nie wiadomo, jaka będzie wówczas skuteczność namierzania celów przez pilotów samolotu wyposażonego w laser. Na razie wiązka może zestrzelić drona z odległości jednego kilometra. Rotem ma nadzieję, że w najbliższych latach uda się zwiększyć jej zasięg dwudziestokrotnie.

Podczas testów, przeprowadzonych przez izraelskie Ministerstwo Obrony i firmę Elbit Systems, samolot, jak głosi komunikat, „z powodzeniem przechwycił i zniszczył 100% jednostek UAV wystrzelonych w ramach ćwiczeń”. Tamtejsze siły zbrojne postrzegają broń laserową jako znacznie tańszą alternatywę dla pocisków przechwytyjących, operujących w systemie obrony powietrznej Izraela, nazywanym Żelazną Kopułą.



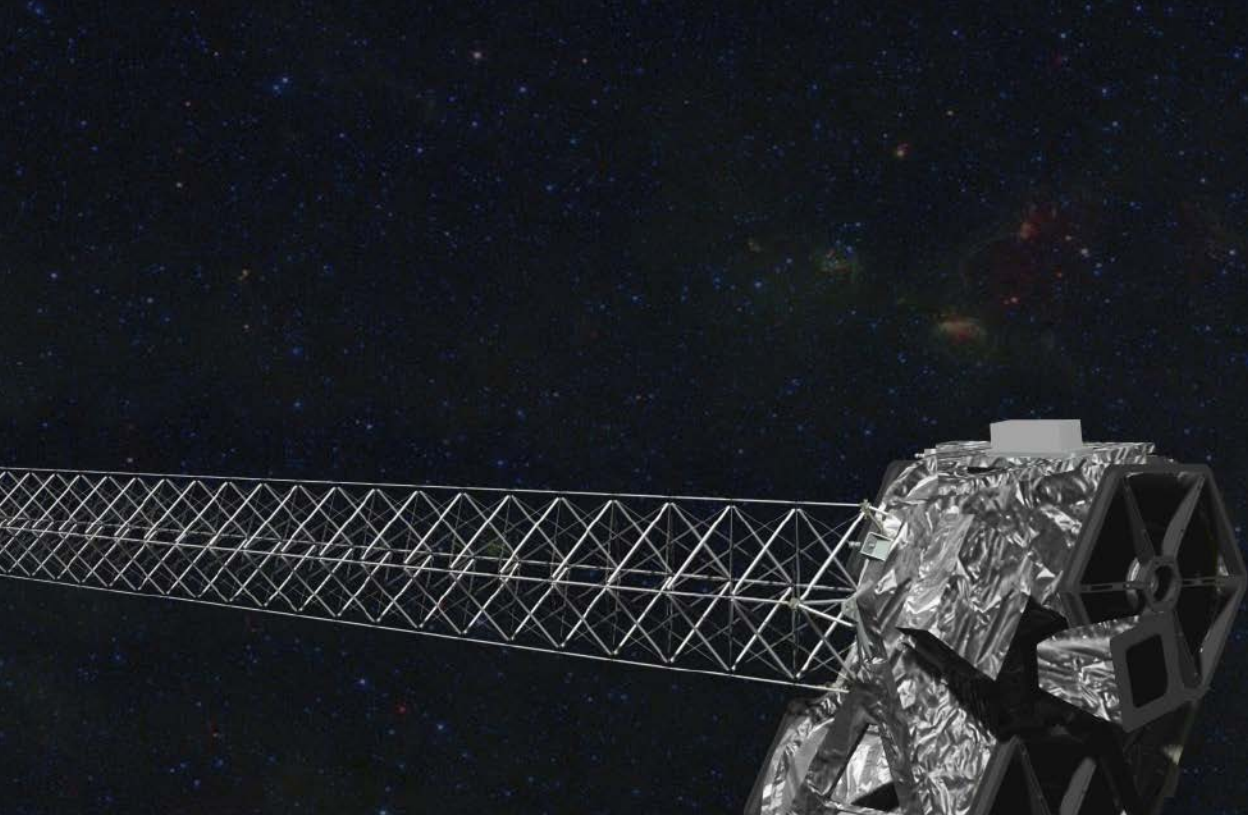
Film z Twittera demonstrujący High-Power Laser Weapon System w działaniu:
<https://bit.ly/2XtlcGk>



Po raz pierwszy w historii badań astronomowie zaobserwowali odbicie światła z za supermasywnej czarnej dziury, oddalonej od Ziemi o 800 milionów lat świetlnych. Według badań opublikowanych w „Nature”, owe „echa” świetlne miały formę błysków promieniowania rentgenowskiego. Uczni podkreślają, że odkrycie potwierdza teorie Alberta Einsteina o zakrzywianiu światła przez obiekty o bardzo silnej grawitacji a czarne dziury są właśnie takimi obiektami.

Astronomowie już wcześniej widzieli światło zaginające się wokół czarnej dziury, jednak po raz pierwszy udało im się zaobserwować to zjawisko dokładnie po drugiej stronie takiego obiektu. Dokonał tego Dan Wilkins, astrofizyk z Uniwersytetu Stanforda z zespołem badającym rozbłyski rentgenowskie wyrzucane z supermasywnej czarnej dziury w centrum galaktyki zwanej I Zwicky 1. Badacze użyli w tym celu należącego do NASA satelity Nuclear Spectroscopic Telescope Array oraz teleskopu XMM-Newton Europejskiej Agencji Kosmicznej.

Zespół badawczy wykrył słabsze rozbłyski promieniowania X, które miały inną długość fali niż te większe, co wskazywało na to, że odbiły się one od dysku akrecyjnego z za czarnej dziury. Dzieje się tak, gdy niektórym promieniom X udaje się ominąć potężne przyciąganie grawitacyjne czarnej dziury, po czym zostają one „wessane” z powrotem. Niektóre z tych uciekających promieni rentgenowskich



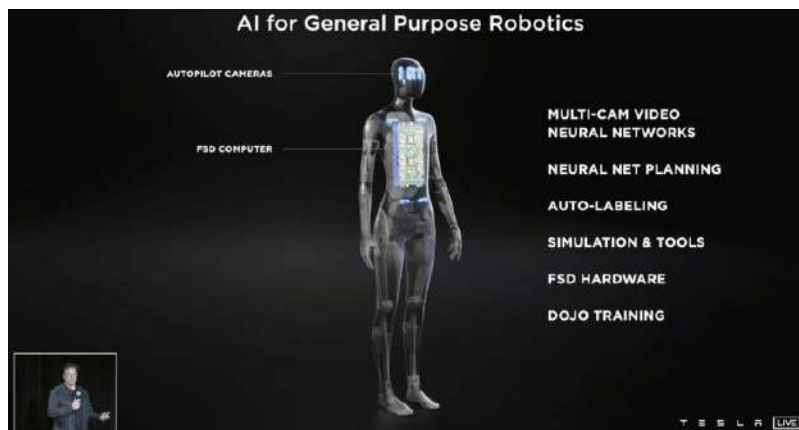
CZARNE DZIURY

Udało się zobaczyć światło zza czarnej dziury – to pierwszy taki przypadek

odbijają się od tylnej części dysku akrecyjnego i są zagnane wokół czarnej dziury przez jej grawitację.

Właśnie to zjawisko pozwoliło na wykrycie „echa” promieniowania rentgenowskiego z drugiej strony.

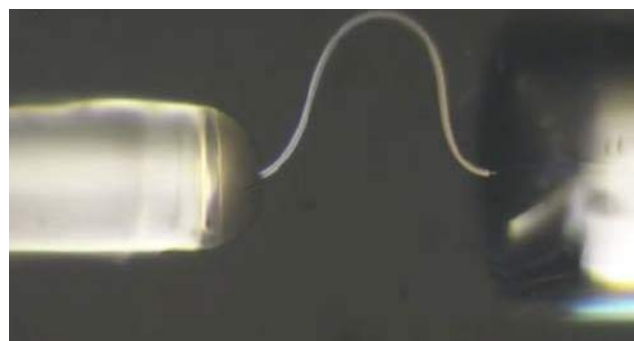




ROBOTY

Humanoidalny Tesla Bot

Szef Tesli, Elon Musk, zaprezentował podczas wydarzenia o nazwie „Tesla’s AI Day” humanoidalnego robota o nazwie Tesla Bot, który działa, wykorzystując tę samą sztuczną inteligencję, która jest bazą autonomicznych pojazdów jego firmy. Podczas prezentacji



TECHNIKI MATERIAŁOWE

Elastyczny lód

Limin Tong, fizyk z Uniwersytetu Zhejiang w Chinach, i jego koledzy wytworzyli mikroskopijne kryształy lodu, które w odróżnieniu od twardego i kruchego lodu charakteryzują się giętkością i elastycznością. Ponadto nadzwyczaj dobrze przepuszczają światło wzdłuż swojej długości. Uчени, którzy opisali swoje badania w „Science”, twierdzą, że w dążeniu do stworzenia „elastycznego lodu”

Muska nie pokazano samego robota. Zamiast tego na scenie pokazano postać ubraną w specyficzny kostium.

Znane są jednak pewne parametry nowej konstrukcji Tesli. Humanoidalny robot ma mieć od 1,5 do niemal 2,5 metra wzrostu i być w stanie dźwigać ciężary powyżej 60 kilogramów. Jego głowa będzie wyposażona w kamery autopilota używane

przez pojazdy Tesli do rozpoznawania otoczenia i będzie mieć ekran do wyświetlania informacji. W środku będzie działał komputer Full Self-Driving (FSD) Tesli. Jak podkreślił Musk, robot ma być „przyjazny”.

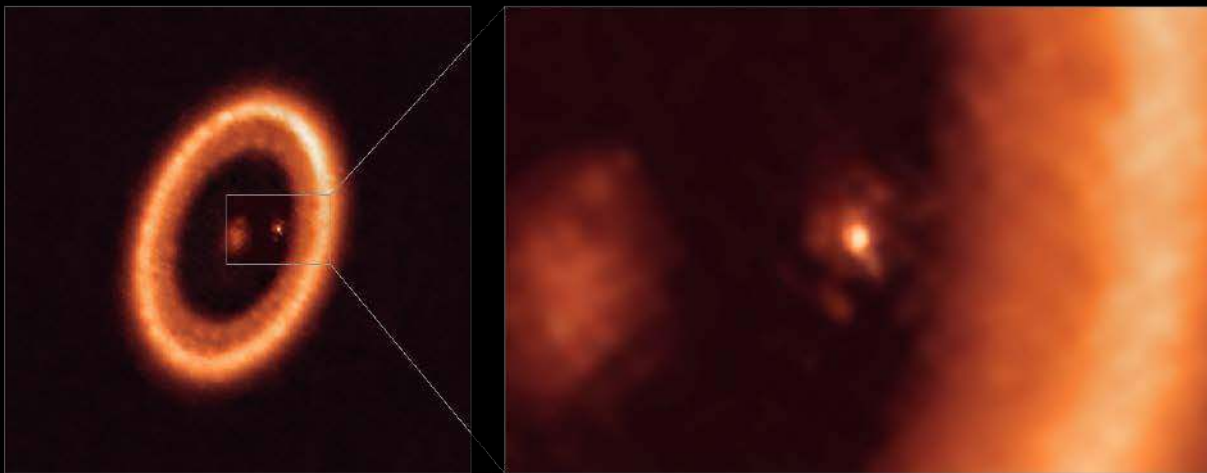
Podczas prezentacji Musk mówił również, że Tesla Bot miałby uwolnić ludzi od „niebezpiecznych, powtarzalnych i nudnych zadań”, podając przykład możliwego zlecenia maszynie robienia zakupów spożywczych. Jak uważa prezes Tesli, tworzenie robota w humanoidalnej postaci ma sens, gdyż uczłowiecza technologię.

inspirowali się szkłem, które w światłowodach również jest giętkie.

Aby otrzymać lód w tej niezwyklej postaci, naukowcy wydrukowali w drukarce 3D komorę o średnicy nieco blisko trzech centymetrów. Używając ciekłego azotu, schłodzili przestrzeń wewnątrz komory do temperatury -50°C . Następnie umieścili tam metalową igłę, do której przyłożono prąd o napięciu 2000 woltów. Napięcie wytworzyło pole elektryczne, a cząsteczki wody obecne w powietrzu zareagowały na to pole, osiadając

na igłę. Bardzo powoli, w tempie około jednej setnej cała na sekundę, z końcówki igły wyrastały mikrowłókna lodu z atomami cząsteczek wody układającymi się w struktury podobne do plastrów miodu.

Uzyskane w ten sposób mikrowłókna tworzyły pręciki, które, jak się okazało, nie tylko charakteryzują się elastycznością, ale również doskonale przewodzą światło, niczym światłowody. Okazało się, że na końcu pręcika lodowego rejestrowane jest 99% światła wpadającego z drugiej strony. Uчени twierdzą, że struktury te można wykorzystać do badania czystości powietrza.



ASTRONOMIA

Odkrycie egzoplanety z dyskiem, w którym powstają księżyce

Po raz pierwszy w historii naukowcom udało się w sposób wyraźny zidentyfikować pierścień gazu i pyłu okrążający planetę spoza naszego Układu Słonecznego. Odkrycie może pomóc wyjaśnić procesy formowania planet i księżyców wokół większych obiektów w kosmosie. Dysk, o którym mowa, otacza egzoplanetę o nazwie PDS 70c, jednego z dwóch gazowych olbrzymów o masie i rozmiarze podobnym do Jowisza, krążących wokół gwiazdy PDS 70, prawie 400 lat świetlnych od naszego Układu Słonecznego.

Astronomowie z Europejskiego Obserwatorium Południowego odkryli tę egzoplanetę w 2019 roku za pomocą Bardzo Dużego Teleskopu (VLT). Obserwacje te w połączeniu z wysokiej rozdzielczości obrazami z teleskopu ALMA, również znajdującego

się w Chile, pozwoliły na wyciągnięcie konkluzji, że dysk PDS 70c zawiera materiał, który pozwala na formowanie się księżyców wokół planety. Wyniki te zostały opublikowane w pracy zamieszczonej w „The Astrophysical Journal Letters”.

Obie planety odkryte w tym układzie są bardzo interesujące dla naukowców, ponieważ należą do młodego systemu gwiazdnego. Gwiazda PDS 70 ma 5,4 miliona lat, czyli jest, wieku niemowlęcym w porównaniu z naszym Słońcem, które istnieje od 4,6 miliarda lat. „Egzoplanety PDS 70b i PDS 70c tworzą układ przypominający parę Jowisz-Saturn”, pisze Miriam Keppler, badaczka z instytutu Maxa Plancka, która odkryła planetę PDS 70b w 2018 roku, w opublikowanym komunikacie z badań.

319 terabitów na sekundę wynosi rekord prędkości transmisji internetowej pobity latem 2021 roku przez japońskich naukowców na odcinku światłowodu o długości trzech tysięcy kilometrów.



BEZPIECZEŃSTWO

Sposób na pożary baterii

Znaleziono nowy sposób na zapłony w ogniach elektrochemicznych. Zespołowi naukowców z Koreańskiego Instytutu Nauki i Technologii (KIST) pod kierownictwem Joonga Kee Lee z Centrum Badań nad Magazynowaniem Energii przez tworzenie ochronnych półprzewodnikowych warstw pasywacyjnych na powierzchni elektrod litowych udało się zahamować wzrost dendrytów, kryształów o wielu rozgałęzieniach, które powodują pożary akumulatorów pojazdów elektrycznych.

Aby zapobiec tworzeniu się dendrytów, zespół badawczy poddał fuleren (C60), materiał o wysokim przewodnictwie elektrycznym, działaniu plazmy, co spowodowało utworzenie półprzewodzących pasywacyjnych warstw węglowych pomiędzy elektrodą litową a elektrolitem. Półprzewodzące warstwy pozwalają na przejście jonów litowych, blokując jednocześnie elektrony dzięki wytworzeniu tzw. bariery Schottky'ego, co zapobiega interakcji elektronów i jonów na powierzchni elektrody, powstrzymując zarazem tworzenie się kryształów litu i w konsekwencji wzrost dendrytów. W znanych konstrukcjach akumulatorów litowe dendryty są odpowiedzialne za niekontrolowane fluktuacje objętościowe i prowadzą do reakcji pomiędzy elektrodą stałą a ciekłym elektrolitem, co powoduje pożary.

Nowo opracowane elektrody wykazują znacznie zwiększoną stabilność, dzięki czemu wzrost dendrytów Li został zahamowany do 1200 cykli. Ponadto, przy zastosowaniu katody z tlenku litowo-kobaltowego (LiCoO₂) jako dodatku do opracowanej elektrody, po 500 cyklach utrzymano około 81 proc. początkowej pojemności baterii, co stanowi poprawę o około 60% w stosunku do konwencjonalnych elektrod litowych.



ROBOTYKA

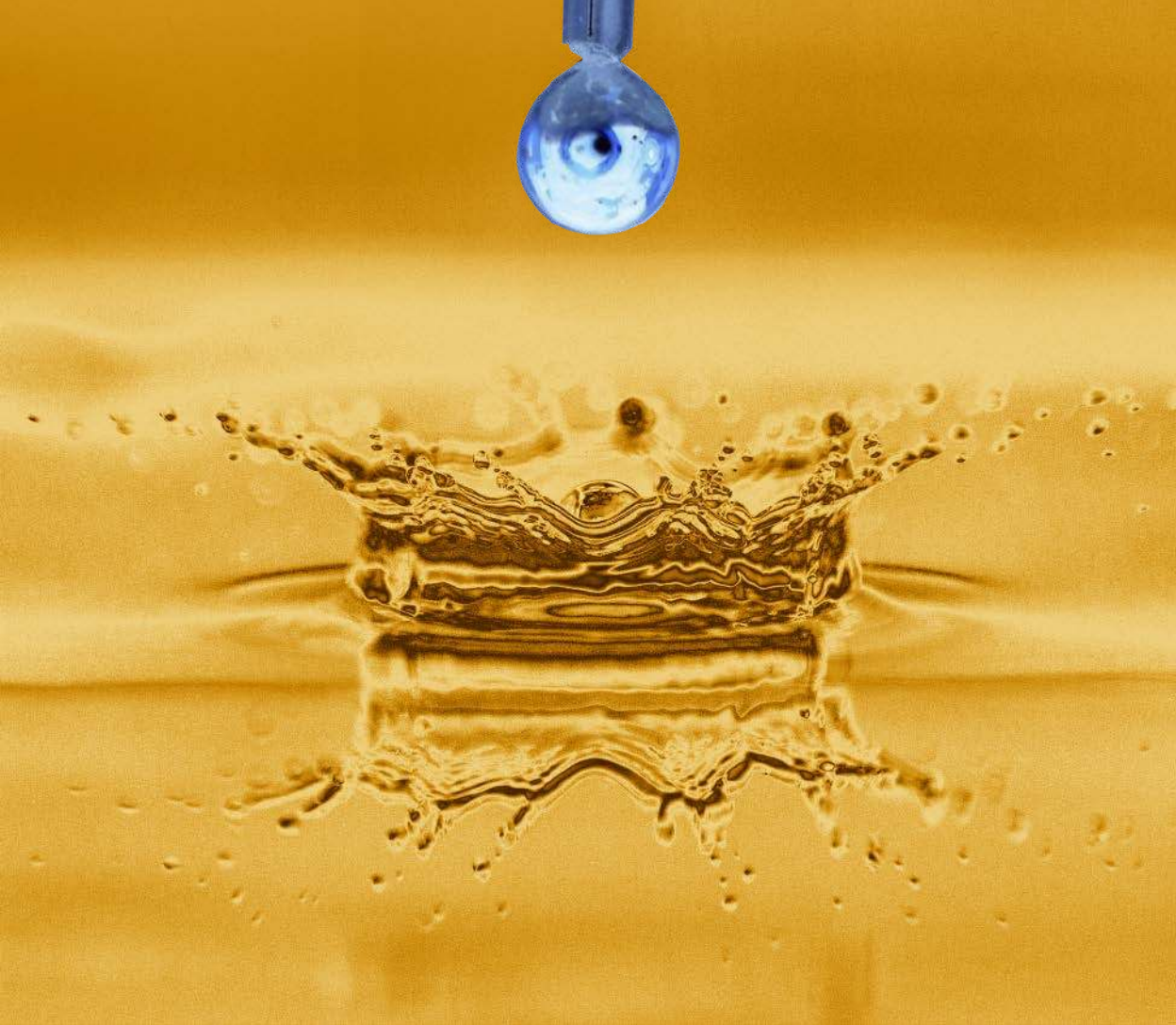
Autonomiczne barki-roboty do tankowania wojskowych maszyn

Amerykańskie siły zbrojne w ramach projektu o nazwie Mayflower testują zrobotyzowane autonomiczne platformy, na których lądować mogą jednostki latające pionowego startu i lądowania, czyli śmigłowce i samoloty

typu VTOL. Jako główną funkcję tych autonomicznych platform wymienia się tankowanie latających maszyn.

Głównym partnerem technicznym w tym projekcie jest Sea Machines, firma z Bostonu, specjalizująca się w oprogramowaniu i systemach autonomicznych statków. Ogłosiła niedawno rozpoczęcie drugiej fazy swojego wieloletniego kontraktu z Departamentem Obrony Stanów Zjednoczonych. W ubiegłym roku firma z powodzeniem zademonstrowała zestaw oparty na autonomicznym systemie sterowania SM300 oraz technologii detekcji opartej na zaawansowanym systemie widzenia komputerowego. Podczas drugiej fazy zostanie opracowany i ostatecznie wdrożony zestaw autonomicznego dowodzenia i kontroli na pełnowymiarowej platformie robotycznej.

Jedną z głównych zalet tego systemu jest to, że nie wymaga on budowania nowych jednostek pływających, bo korzysta z istniejących już i dostępnych komercyjnych barek, na których mogą lądować i być tankowane wojskowe samoloty VTOL i helikoptery. Technika jest opracowywana przy wsparciu wielkich znanych firm z branży zbrojeniowej, takich jak FOSS Maritime, Huntington Ingalls i Bell Flight.



NOWE MATERIAŁY

Przemiana wody w złoty metal

Czeskim naukowcom udało się przeobrazić czystą wodę w stałą substancję o właściwościach metalicznych, czyli np. przewodzącą prąd elektryczny, co dla wody destylowanej jest niemożliwe. Choć możliwe jest osiągnięcie przez wodę takiego stanu w naturze, to jednak wymagane do tego ciśnienie rzędu piętnastu milionów atmosfer jest osiągalne naturalnie jedynie w jądrach wielkich planet, np. Jowisza czy Neptuna.

Uczni zdecydowali się na doprowadzenie do tego inną metodą. Użyli metali alkalicznych, do których należą sód i potas, mające tylko jeden elektron w swojej powłoce walencyjnej. Metale alkaliczne mają tendencję do oddawania tego elektronu innym atomom podczas tworzenia wiązań chemicznych, ponieważ

utrata tego samotnego elektronu sprawia, że metal alkaliczny jest bardziej stabilny.

Jednak kontakt metali alkalicznych z wodą prowadzi do eksplozji. By jej uniknąć w eksperymencie, opisanym w czasopiśmie „Nature”, zespół użył strzykawki wypełnionej sodem i potasem w komorze próżniowej, z której wyciśnięto drobne krople metaliczne a następnie wystawiono je na działanie niewielkiej ilości pary wodnej. Woda utworzyła warstwę o grubości 0,1 mikrometra na powierzchni kropli metalu, zaś elektrony zaczęły przechodzić do wody bez efektu eksplozji. W widocznej gołym okiem przemianie fazowej powstała zadziwiająca, metaliczna woda o złotym kolorze.



Konkurs Explory jest przeznaczony dla młodych ludzi interesujących się nauką, którzy chcą rozwijać swoje zainteresowania, spotkać inspirujących rówieśników i poznać tajniki pracy badacza. W konkursie można prezentować pomysły z różnych dziedzin nauki. Skierowany jest do osób w wieku 13–20 lat i składa się z trzech etapów: zgłoszenie projektu polegające na wysłaniu formularza online, udział w regionalnych eliminacjach

KONKURS

Poszukiwani młodzi naukowcy – ruszyła rekrutacja do Konkursu Explory

W Konkursie co roku startują młodzi i ambitni uczniowie, którzy mają ciekawe pomysły na projekty i szukają szansy na ich rozwój. Udział w konkursie daje możliwość spotkania wybitnych autorytetów ze świata nauki i biznesu oraz zdobycia cennych nagród. Autorzy najlepszych projektów otrzymują

stypendia naukowe i udział w prestiżowych konkursach międzynarodowych. Nabór do Konkursu trwa do 7 stycznia 2022. „Młody Technik” jest patronem medialnym programu Explory i konkursu dla młodych naukowców.



Więcej o Programie Explory:
https://bit.ly/explory_pl

odbywających się w ramach Festiwalu Explory oraz finał podczas wydarzenia Gdynia Explory Week. Autorzy najlepszych projektów otrzymają wsparcie ze strony ekspertów w zaplanowaniu najbliższych etapów kariery zawodowej, stypendia naukowe na rozwój swoich projektów ufundowane przez mecenasa programu Explory – Grupę LOTOS i innych partnerów oraz udział w konkursach międzynarodowych. Jednym z nich jest najbardziej prestiżowy konkurs dla młodych naukowców na świecie – Regeneron International Science and Engineering Fair w USA, w którym swoją karierę naukową rozpoczynało do tej pory aż 27 noblistów.

Explory od 10 lat odkrywa, wspiera i rozwija młode talenty. W tym czasie udało się stworzyć społeczność ponad 2400 młodych innowatorów, którzy zgłosili do Konkursu Explory ponad 1500 projektów naukowych i wynalazków. Pierwsze miejsce w zeszłorocznym finale Konkursu Explory zdobył algorytm szyfrowania stworzony przez 14-letniego Szymona Perlickiego. Inne godne uwagi projekty uczestniczące w konkursie Explory w ostatnich latach to m.in. nanokrystaliczne ogniwo słoneczne Anny Skierskiej lub prototyp automatycznego pojazdu do prewencyjnej analizy stanu gleby na polach uprawnych Piotra Lazarka.

62 831 853 071 796 cyfr po przecinku ma rekordowe obliczenie przybliżenia liczby π dokonane przez naukowców z Uniwersytetu Nauk Stosowanych w Graubünden w Szwajcarii.



KULTURA I NAUKA

Kongres Futurologiczny w Krakowie



Więcej informacji
i szczegółowy program
Kongresu:
<https://bit.ly/39fcQ70>

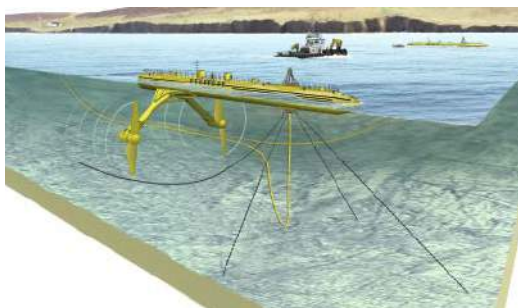
W ramach oficjalnych obchodów jubileuszu stulecia urodzin Stanisława Lema w Krakowie zebrał się Kongres Futurologiczny. Trzydniowa konferencja miała bogaty program wystąpień i dyskusji w pasmach tematycznych bloku kultury i bloku nauki. Celem wydarzenia jest m.in. promowanie rodzimej twórczości fantastycznonaukowej, inspirowanie młodych autorów, a także integracja naukowców, artystów oraz przedstawicieli branży technologicznej. „Młody Technik” objął imprezę patronatem medialnym.

W ciągu trzech dni, wypełnionych spotkaniami, uczestnicy kongresu mieli okazję wysłuchać wystąpień i dyskusji z udziałem reprezentantów polskiego sektora kosmicznego, specjalistów zajmujących się sztuczną inteligencją i rzeczywistością wirtualną, jak również pisarzy fantastyki, krytyków, a także twórców oraz producentów gier. Wśród zaproszonych gości znajdowali się między innymi Mirosław Hermaszewski, pierwszy i jedyny polski kosmonauta, Wiktor Niedzicki, twórca i prowadzący legendarnego programu telewizyjnego „Laboratorium” oraz Jacek Rodek, autor scenariusza kultowego polskiego komiksu „Funky Koval”. Wydarzenie odbywa się zarówno w formie stacjonarnej z wolnym wstępem dla wszystkich zainteresowanych, jak też jest transmitowane w mediach społecznościowych.

Podczas kongresu uczestnicy mogli również obejrzeć wystawę zdjęć poświęconych 60. rocznicy pierwszego lotu człowieka w kosmos, adaptację filmową Hanki Brulińskiej na podstawie opowiadania „Maska” Stanisława Lema oraz ekspozycję fantastycznonaukowych

prac konkursu organizowanego przez krakowską Fundację IDEANOVA w ramach projektu Kapsuła Czasu – mającego gromadzić wizję świata za sto lat. W trakcie konferencji premierę ma też wyjątkowy zbiór „Ku gwiazdom: Antologia polskiej fantastyki naukowej 2021”, stanowiący pokłosie pierwszego konkursu literackiego i współpracy Polskiej Fundacji Fantastyki Naukowej z krakowskim Wydawnictwem IX.





ENERGIA

♦ Firma Orbital Marine Power ogłosiła, że zainstalowała w pobliżu brzegu jednej z wysp na szkockim archipelagu Orkady ważącą blisko 700 ton i umieszczoną na 73-metrowej platformie gigantyczną turbinę pływającą o nazwie O2, która ma wytwarzać do dwóch megawatów mocy, co pozwoli zasilić nawet dwa tysiące gospodarstw domowych. ♦ Naukowcy z amerykańskiego Uniwersytetu Stanowego Waszyngton opracowali proces utylizowania plastikowych odpadów z wykorzystaniem katalizatora rutenowego na węglu, który w 90 proc. przetwarza polietylen na składniki paliwa do odrzutowców i inne wartościowe węglowodory. ♦

PODOBÓJ KOSMOSU

♦ Z propozycją wynoszenia satelitów na orbitę za pomocą wiązek mikrofalowych zamiast rakiet wystąpiła grupa inżynierów z japońskiego Uniwersytetu Tsukuba w Nagoi, która w ramach demonstracji przeprowadziła próbną lot drona zasilanego z powierzchni ziemi za pomocą silnego emitera fal. ♦ Rosyjska agencja kosmiczna Roskosmos zapowiada uruchomienie do 2030 r. międzyplanetarnej misji, najpierw na orbitę Księżyca, potem Wenus, by w końcu dolecieć do Jowisza z wykorzystaniem statku-holownika o nazwie „Zeus” napędzanego reaktorem jądrowym o mocy 500 kilowatów. ♦



SZTUCZNA INTELIGENCJA

♦ James Noeckel z Uniwersytetu stanu Waszyngton w Seattle wraz ze współpracownikami stworzył algorytm, który potrafi przekształcać zdjęcia dowolnych drewnianych przedmiotów w trójwymiarowe modele, które są wystarczająco szczegółowe, aby wykwalifikowany stolarz mógł je odtworzyć w rzeczywistości, z drewna. ♦ Amerykańskie wojsko testuje i rozwija system Global Information Dominance Experiments (GIDE), który polega na połączeniu sztucznej inteligencji, przetwarzania danych w chmurze i odczytów z czujników, by przez obserwację zmian w nieprzetworzonych danych w czasie rzeczywistym wyszukiwać i identyfikować te istotne i mogące wpłynąć na przyszłe wydarzenia, np. inwazje lub ataki, co można określić jako „przewidywanie przyszłości”. ♦

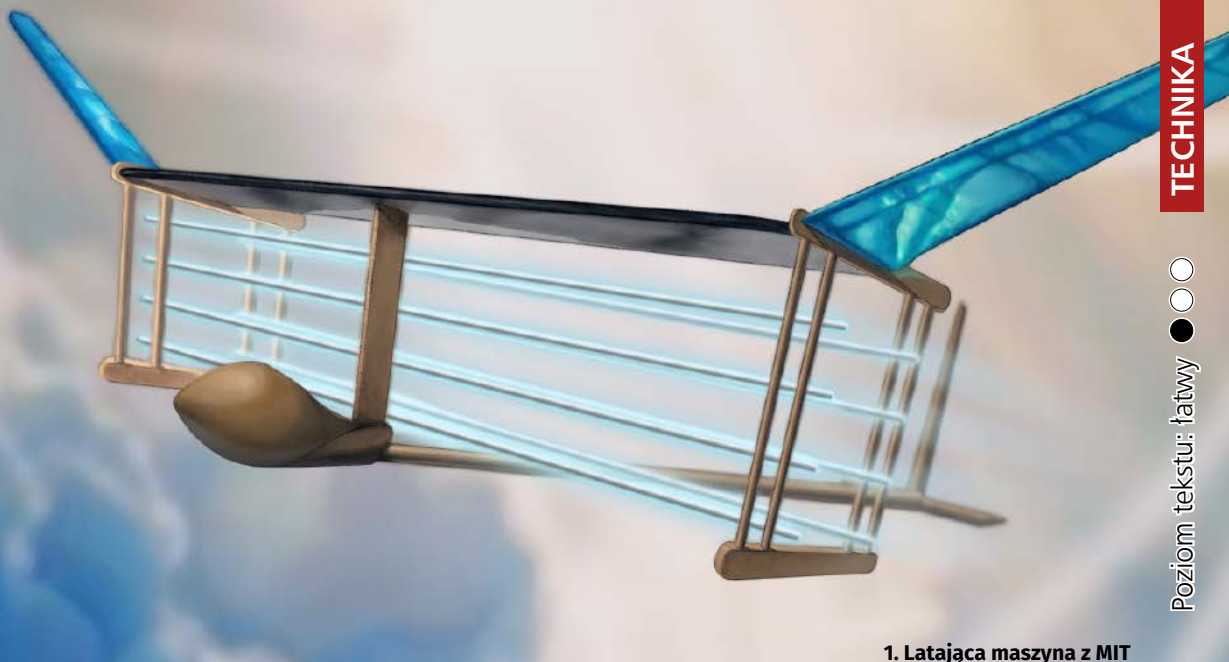
CYBERBEZPIECZEŃSTWO

♦ Zespół badaczy ze stanowego uniwersytetu w Cleveland opracował protokół zabezpieczający dane, w którym „hasłem” jest cykl oddechu człowieka – co kilka sekund w systemie tworzony jest 256-bitowy klucz szyfrujący. ♦ Po tym jak możliwość taką uzyskali dzięki Apple użytkownicy iPhone’ów, Google zamierza udostępnić korzystającym z systemu Android w ciągu kilku miesięcy opcję wyłączenia możliwości śledzenia przez reklamodawców w aplikacjach smartfonowych. ♦

FIZYKA

♦ Według wstępnej publikacji (tzw. pre-printu) badaczom z Google’a udało się za pomocą komputera kwantowego Sycamore stworzyć „nowy stan materii” zwany „kryształ czasu”, który, według opisów kolportowanych w mediach, ma „przeczyć drugiej zasadzie termodynamiki”, pozostając w dwóch stanach termicznych jednocześnie. ♦ Jak podało czasopismo „Science”, grupie izraelskich naukowców udało się z monowarstwowych struktur boru i azotu stworzyć najcieńsze znane fizyce warstwy krystaliczne – o grubości zaledwie dwóch atomów, które mogą służyć do rekordowo gęstego i wydajnego zapisu informacji. ■

M. U.



1. Latająca maszyna z MIT

Samolot bez ruchomych części napędzających?

Jonowy wiatr w uszach

Pod koniec 2018 roku inżynierowie z renomowanego Massachusetts Institute of Technology zaprezentowali samolot z napędem jonowym (1). Od tego czasu sprawa nieco przycichła, jak wiele różnych spraw podczas pandemii. Jednak warto przyrzeć się po opadnięciu kurzu technice lotniczej, która, zdaniem wielu ekspertów, może sporo namieszać w przyszłości.

Zamiast ruchomych śmigieł czy turbin lekki samolot napędzany jest ciągiem strumienia jonów, który jest wytwarzany na pokładzie samolotu. Zdaniem konstruktorów moc takiego rozwiązania może pozwolić na długotrwały lot ze stałą prędkością dronów a nawet lekkich samolotów pasażerskich. W przeciwieństwie do samolotów z napędem turbinowym samolot nie jest uzależniony od paliw kopalnych. A także, w przeciwieństwie do dronów napędzanych śmigłem, nowa konstrukcja jest całkowicie bezgłówna.

Steven Barrett (2), profesor aeronautyki i astronautyki na MIT, który kierował zespołem twórców

latającej maszyny nowego typu, snuje w wypowiedziach dla mediów bardziej ambitne plany. Jego zdaniem, napęd jonowy w połączeniu z konwencjonalnymi systemami spalinowymi pozwoli na stworzenie paliwooszczędnych, hybrydowych samolotów pasażerskich i innych dużych statków powietrznych. Opis konstrukcji i badań zespołu Barretta ukazał się w czasopiśmie „Nature”.

Ciąg pomiędzy elektrodami

Około dziewięciu lat temu Barrett zaczął szukać sposobów na zaprojektowanie systemu napędowego



dla samolotów pozbawionych ruchomych części. Inspirowały go wahadłowce, które oglądał w serialu „Star Trek”. W końcu natrafił na zjawisko nazywane „wiatrem jonowym”, znane również jako „ciąg elektro-aerodynamiczny” (EHD). To fizyczne zjawisko, która zostało po raz pierwszy zidentyfikowane w latach 20. XX wieku, opisuje ciąg powietrzny, który może być wytwarzany, gdy prąd przepływa pomiędzy cienką i grubą elektrodą. Przez przykładanie napięcia do pary elektrod elektrony są usuwane z pobliskich cząsteczek powietrza, a zjonizowane powietrze zderza się z cząsteczkami neutralnego powietrza, gdy przemieszcza się z jednej elektrody do drugiej. Jeśli zastosuje się wystarczające napięcie, powietrze pomiędzy elektrodami może wytworzyć ciąg.

Amerykański eksperymentator Thomas Townsend Brown spędził dużą część swojego życia, pracując nad jonowym napędem, będąc zresztą mylnie przekonany, że jest to efekt antygrawitacyjny (3), który nazwał efektem Biefelda–Brown. Ponieważ jego urządzenia wytwarzały ciąg w kierunku gradientu pola, niezależnie od kierunku grawitacji, i nie działały w próżni, inni pracownicy zdali sobie sprawę, że efekt ten był spowodowany przez „wiatr jonowy”.

W latach 50. i 60. XX wieku amerykański konstruktor samolotów major Alexander Prokofieff de Seversky badał zastosowanie napędu EHD do unoszenia statków powietrznych. Zgłosił nawet patent na „jonolot” w 1959 r. Zbudował i oblatywał model jonolotu VTOL zdolnego do manewrowania na boki poprzez zmianę napięcia przyłożonego w różnych obszarach, choć zasilanie pozostało zewnętrzne. W późniejszym okresie na uwagę zasługuje konstrukcja Wingless Electromagnetic Air Vehicle (WEAV) z 2008 roku, badana przez zespół naukowców kierowany przez Subratę Roya na Uniwersytecie Florydy. W systemie napędowym zastosowano wiele innowacji, w tym pole magnetyczne do zwiększenia wydajności jonizacji.

Pierwszym samolotem z napędem jonowym, który wystartował i poleciał z wykorzystaniem własnego zasilania pokładowego, był statek VTOL opracowany przez Ethana Kraussa z firmy Electron Air w 2006 roku. Jednostka rozwijała ciąg



Reportaż o napędzie jonowym samolotu:
<https://bit.ly/3ADr2mn>



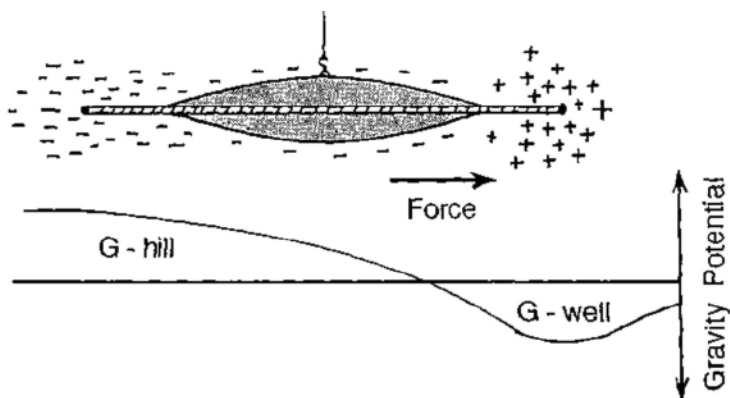
2. Steven Barrett

wystarczający do szybkiego wznoszenia się lub do lotu poziomego przez kilka minut.

Rekord w sali gimnastycznej

Dawniej często uważano, że niemożliwe będzie wytworzenie wystarczającej ilości wiatru jonowego, aby napędzić większy samolot w trakcie długotrwałego lotu. Zjawisko to było raczej przedmiotem zabawy hobbystów. W końcu Barrett dokonał obliczeń, z których wynikało, że jednak możliwe jest wykorzystanie tego zjawiska do napędu cięższej jednostki.

Ostateczny projekt jego zespołu przypomina szybowiec. Testowana jednostka latająca ważyła około 2,5 kg i miała pięciometrowej rozpiętości skrzydła. Konstrukcję przenika sieć cienkich przewodów, które rozciągają się wzdłuż i pod przednią częścią skrzydła samolotu. Działają jak dodatnio naładowane elektrody, zaś podobnie ułożone grubsze przewody, biegnące wzdłuż tylnej części skrzydła samolotu, służą jako



3. Rysunek jednego z układów Thomasa Townsenda Browna

elektrody ujemne. W kadłubie samolotu znajduje się stos akumulatorów litowo-polimerowych.

W skład zespołu Barretta pracującego nad samolotem jonowym wchodził członek grupy badawczej Power Electronics profesora Davida Perreault z Laboratorium Badawczego Elektroniki, którzy zaprojektowali zasilacz przekształcający moc akumulatorów na wystarczająco wysokie napięcie. W ten sposób akumulatory dostarczały prąd o napięciu 40 tysięcy woltów, który ładował przewody za pomocą lekkiego konwertera. Kiedy przewody zostają już naładowane, przyciągają i usuwają ujemnie naładowane elektrony z otaczających je cząsteczek powietrza, tak jak magnes przyciąga opilkę żelaza. Pozostałe cząsteczki powietrza są zjonizowane i przyciągane do ujemnie naładowanych elektrod z tyłu samolotu. Czyli inaczej mówiąc, chmura jonów płynie w kierunku ujemnie naładowanych przewodów, a każdy jon zderza się miliony razy z innymi cząsteczkami powietrza, tworząc ciąg, który napędza samolot do przodu.

Zespół wykonał wiele lotów próbnych w sali gimnastycznej w centrum sportowym duPont MIT, czyli w największej przestrzeni wewnętrznej, jaką udało się znaleźć do przeprowadzenia eksperymentów. Mały samolot pokonał 60 metrów (więcej w sali gimnastycznej nie mógł), osiągając ok. 17 km/h prędkości. Loty te powtarzane były wielokrotnie. Na podstawie tych eksperymentów wysunięto wniosek, że ciąg jonowy jest wystarczający do napędu małych jednostek na długich dystansach. Osiągnięcie było znaczące, zwłaszcza że wcześniejsi eksperymenci nie osiągnęli nic więcej niż krótkotrwałe loty kilkugramowych modeli napędzanych jonowo.

Po przeprowadzeniu tych udanych prób zespół Barretta pracował nad zwiększeniem wydajności swojej konstrukcji, czyli produkowaniem silniejszego wiatru jonowego przy mniejszym napięciu. Badacze mają również nadzieję na zwiększenie gęstości ciągu

– ilości ciągu generowanego na jednostkę powierzchni. Obecnie latanie lekkim samolotem zespołu wymaga dużej powierzchni elektrod. W idealnej sytuacji Barrett chciałby zaprojektować samolot bez widocznego układu napędowego a ponadto, zdaniem Barretta, napęd jonowy mógłby nawet zostać wykorzystany do sterowania samolotem, dzięki czemu nie potrzebowałby on tradycyjnych powierzchni sterujących. W ten sposób silniki półprzewodnikowe nie tylko napędzałyby samolot, ale także kontrolowałyby jego kierunek.

Warto zwrócić uwagę, że „wiatr jonowy” wykorzystywany w konstrukcji MIT, choć opiera się na ciągu naładowanych atomów, jednak nie jest tym samym typem napędu, co raketowe silniki jonowe wykorzystywane w kosmonautyce. Choć w raketowym silniku jonowym czynnikiem nośnym również są jony, jednak są one niejako „sztucznie” rozpędzane. Energia wyrzucająca gaz z silnika pochodzi w stosowanych w kosmosie konstrukcjach z zewnętrznego źródła (najczęściej z baterii słonecznych). Najpierw atomy gazu (ksenonu) pozbawiane są ładunku ujemnego, czyli zostają przekształcone w jony dodatnie. Następnie są przyspieszane pod wpływem pola elektrycznego lub magnetycznego, osiągając prędkość nawet kilkudziesięciu km/s. Duża prędkość wyrzucanego czynnika daje dużą siłę ciągu przypadającą na jednostkę masy wyrzucanej substancji. Jednak ze względu na małą moc układu zasilającego masę wyrzucanego czynnika zwykle nie jest duża, zmniejszając przez to sumaryczną siłę ciągu rakiety.

Zarówno powietrzny, jak i kosmiczny napęd jonowy jest uważany za coś niezwykle obiecującego. Oba jednak typy „jonolotów” mają wciąż niezadowolające parametry, choć rakiety jonowe są jednak na innym etapie niż samoloty, gdyż korzysta się z nich już w normalnych misjach kosmicznych, np. w sondzie Dawn, która badała planetoidę Ceres. ■

Miroslaw Usidus

Bądź sexy i nie daj się zamordować. Rady na różne okazje

Karen Kilgarriff, Georgia Hardstark

Wydawnictwo Kompania Mediowa, cena: 44,90 zł

Kilgarriff i Hardstark przytaczają całą serię najgorszych błędów, które zdążyły popełnić w życiu, opowiadając o nich w sposób pełen humoru, wzajemnego zrozumienia i z dystansem do swoich słabości i lęków. Jednocześnie pokazują uwikłanie w świat funkcjonujący na męskich zasadach i codzienne trudności z tym związane. Karen i Georgia sięgają do przeszłości, by na przykładzie własnych przeżyć oraz prawdziwych historii kryminalnych pokazać kulturowe i społeczne schematy, z którymi od lat muszą zmagać się kobiety na całym świecie. Piszą szczerze i z ogromnym dystansem do siebie, z wielką dozą zdrowego rozsądku i bardzo obrazowo. To feministyczna, słodko-gorzka, zabarwiona czarnym humorem i dużą ilością empatii opowieść o tym, jak walczyć o bycie sobą. Na własnych regułach.



Wyścig o udział
w wyścigu na Księżyc

Nie wystarczą miliardy w kieszeni, by wyprzedzić SpaceX

Biuro Government Accountability Office (GAO), instytucja kontrolna Kongresu Stanów Zjednoczonych, rozpatrując skargę Jeffa Bezosa i firmy Dynetics na decyzję agencji o wyborze systemu lądowania na Księżycu, orzekło, że NASA miała pełne prawo do wyboru jednego tylko wykonawcy, SpaceX (1).



1. Wizualizacja porównująca wizje lądowników księżycowych, od lewej SpaceX, Dynetics i lądownika konsorcjum z Blue Origin w składzie

eprasa.pl d03866570d

Nagłośniony w mediach lot Bezosa z załogą i to akurat 20 lipca, w rocznicę pierwszego lądowania człowieka na Księżycu, był interpretowany jako rywalizacja z Richardem Bransonem, ale tak naprawdę należy spoglądać na niego w kontekście rozstrzygniętego wiosną tego roku przetargu NASA na lądownik księżycowy, który ma pozwolić Amerykanom w ramach programu Artemis powrócić na naszego największego naturalnego satelitę. Po tym, jak agencja wybrała SpaceX jako jedyne partnera w księżycowym programie, Bezos być może chciał zademonstrować wszystkim, ale przede wszystkim amerykańskiej administracji, że on i jego Blue Origin mają „kompetencje kosmiczne”. Może większe wrażenie zrobiłoby, gdyby założyciel Amazona i jego załoga wylądowali na pokładzie odzyskiwanej rakiety, bo takie mniej więcej są wymagania lądowania na Księżycu. Niestety członkowie załogi Bezosa powrócili na Ziemię za pomocą tradycyjnej kapsuły ze spadochronem, metodą, która na Księżycu jest niewykonalna z powodu braku atmosfery.

Pomimo niekorzystnych decyzji NASA a potem GAO, Blue Origin Bezosa nie składa broni. O tym, jak bardzo twórcy Amazona zależy na księżycowym kontrakcie, świadczy jego deklaracja, tydzień po tym słynnym już locie na orbitę, że wykona w ramach kontraktu prace warte dwa miliardy dolarów, co będzie dużą finansową ulgą dla NASA, aby tylko Blue Origin został dopuszczony jako drugi alternatywny wykonawca infrastruktury załogowej na Księżyc.

Human Landing System to zarazem projekt oraz intratny kontrakt na budowę lądownika zdolnego do przeniesienia astronautów na powierzchnię Księżyca. Wybrana do jego realizacji została firma Elona Muska, jednak na początku maja NASA nakazała SpaceX zatrzymać wszelkie prace nad tym projektem po tym, jak firmy Blue Origin i Dynetics nie zgodziły się z wyborem SpaceX, wstrzymując kontrakt o wartości blisko trzech miliardów dolarów. Argumentacja skarżących skupiała się nie tyle na samej SpaceX, ile na fakcie, że firma Muska dostała wyłączność, co, jak argumentowali, jest sprzeczne z wcześniejszą praktyką NASA dobierania kilku podmiotów, co zwiększało bezpieczeństwo projektu. Jednak tym razem nie zdecydowała się na wybranie „wachlarza”, głównie ze względów finansowych.

Starship może być czymś więcej niż lądownikiem księżycowym?

Biuro GAO wprost stwierdziło, że choć NASA pierwotnie zapowiadała, że zamierza wybrać do kontraktu więcej wykonawców, to po prostu nie miała

wystarczająco dużo pieniędzy w budżecie, aby podjąć taką decyzję, zaś wybranie tylko jednego wykonawcy kontraktu było zgodne z prawem. Ponadto jak czytamy w oświadczeniu przedstawiciela GAO Kennetha Pattona, ocena wszystkich trzech ofert dokonana przez NASA „była rozsądna i zgodna z obowiązującymi przepisami prawa zamówień publicznych oraz warunkami ogłoszenia”.

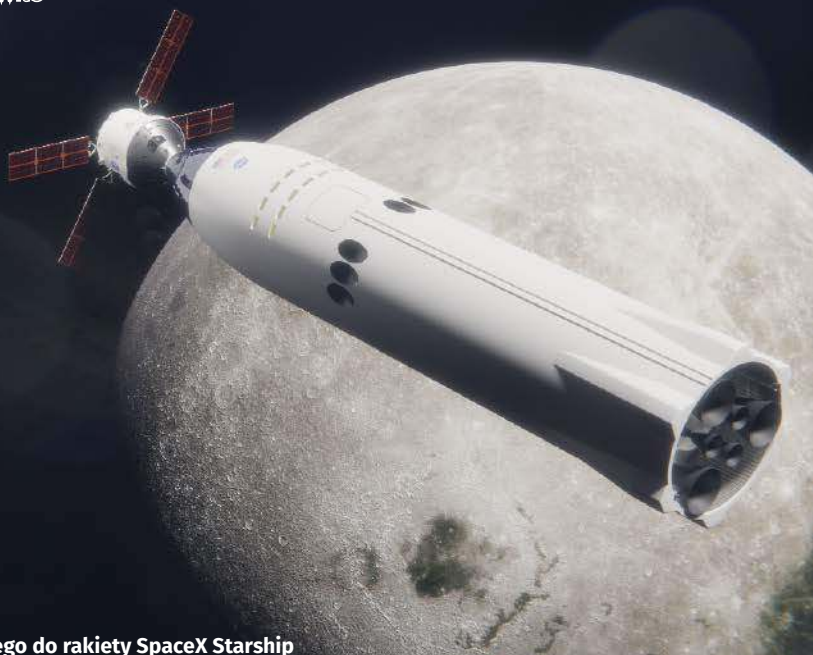
Lądownik księżycowy jest częścią planów NASA na najbliższe lata. Program Artemis przewiduje wykorzystanie dużej rakiety (teoretycznie budowany od lat system SLS, ale w obliczu problemów technicznych, które wynikły w testach, nie jest to pewne), która wyniosłaby czterech astronautów na pokładzie kapsuły kosmicznej Orion na orbitę Księżyca. Lądownik zabrałby dwóch astronautów na powierzchnię Księżyca, gdzie badaliby go przez około tydzień, wrócili na orbitę Księżyca, zaokrętowali z powrotem na Oriona i wrócili na Ziemię.

Lądownik SpaceX, nazwany Starship, „zawiera przestronną kabinę” i może zostać rozbudowany do systemu nośnego wielokrotnego użytku, umożliwiającego podróże na Księżyc, Marsa i w inne miejsca. Takie stwierdzenia możemy znaleźć w komunikacie NASA ogłaszającym wyniki przetargu. Próbnny lot kapsuły, bez astronautów na pokładzie, jest zaplanowany na ten rok, zaś próbnny lot astronautów na Księżyc, ale bez lądowania, planuje się na 2023 r.

Informacje te i decyzje NASA zdają się układać w pewną całość, choć oczywiście oficjalnie systemem wynoszenia statku załogowego Orion jest opracowywana przez NASA rakietą SLS. Jednak wieści z poligonu doświadczalnego tej rakiety, przynajmniej jeszcze na początku 2021 r., nie były najlepsze.

W styczniu przeprowadzono pierwszy test silników RS-25 zintegrowanych z pierwszym egzemplarzem głównego członu rakiety SLS. Test przerwano przedwcześnie z powodu złych parametrów. W kolejnych testach przeszkadzały problemy z zaworami kriogenicznymi. Doniesienia z lata były bardziej optymistyczne – udało się wreszcie odpalić silniki na więcej niż osiem minut, co jest symulacją lotu w kosmos. Załadowano do rakiety oprogramowanie sterujące. Czy jednak planowany wcześniej na listopad 2021 roku bezzałogowy lot SLS na orbitę odbędzie się w pierwotnym terminie? Zobaczymy.

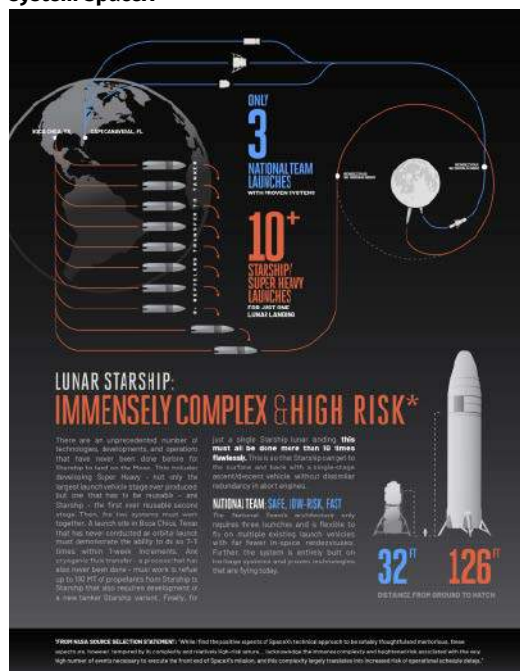
Warto pamiętać, że choć w kontekście przetargu NASA mówimy o Starshipie jako lądowniku, jest to także, a może przede wszystkim w zbiorowej wyobraźni – rakietą. W planowanych rozbudowanych konfiguracjach ma posłużyć nawet do lotów załogowych na Marsa. Przystosowanie jej do współpracy



2. Wizualizacja Oriona zadokowanego do rakiety SpaceX Starship

ze statkiem Orion (2) nie wydaje się niemożliwe, a ponadto właściwie system transportowy Starship w rozbudowanych wersjach mógłby w ogóle nie potrzebować odrębnej kapsuły. Zatem, choć wciąż głównym planem misji załogowej NASA na Księżyc

3. Infografika kolportowana przez Blue Origin przekonująca, jak ryzykownym rozwiązaniem jest system SpaceX



jest rakieta wynosząca SLS i statek Orion, produkty firmy Muska sprawiają wrażenie opcji rezerwowej.

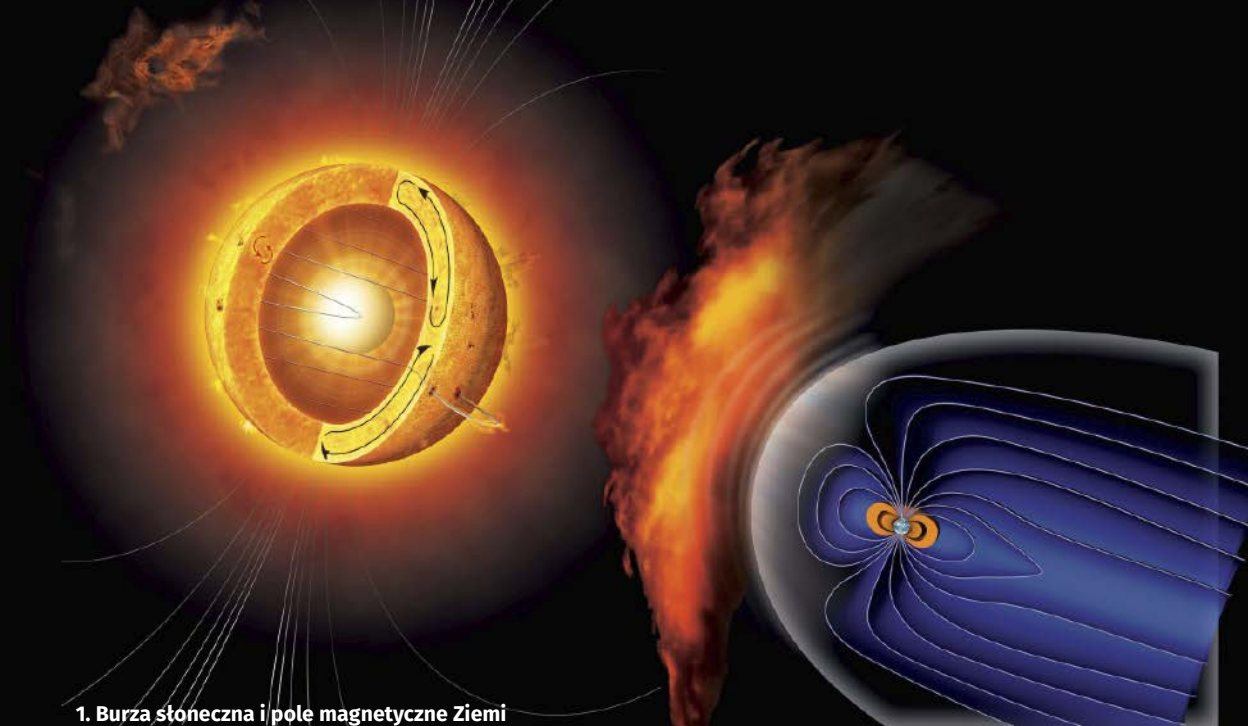
Blue Origin uderza infografiką

Już po decyzji GAO, firma Blue Origin opublikowała infografikę (3), która miała przekonywać opinię publiczną, że system opracowany przez SpaceX jest „skomplikowany i ryzykowny” w porównaniu z „bezpiecznym” rozwiązaniem konsorcjum „National Team”, do którego należy firma Bezosa.

Jednak komentatorzy traktują owe wizualizacje raczej jako czczą propagandę, pytając, skąd Blue Origin wie, ile tankowań na orbicie będzie potrzebne SpaceX do wykonania misji księżycowej. Firma Muska nie podała takich informacji, zaś same tankowania w przestrzeni kosmicznej mogą z punktu widzenia NASA być dodatkową zaletą, a nie wadą, projektu. Byłoby to bowiem testowanie zupełnie nowych, dających całkiem nowe możliwości rozwiązań transportu kosmicznego, także w kontekście misji dalszych niż księżycowe, np. marsjańskich.

Nie da się też ukryć, że Starship i rozwiązania lądowania dużą rakieta SpaceX są już zbudowane i testowane, czego nie można powiedzieć o lądowniku Blue Moon firmy Bezosa, który jest na znacznie wcześniejszym etapie rozwoju. Zatem NASA przedkłada dostawcę, który już coś zrobił, np. statek załogowy na orbitę Ziemi i naprawdę dobre rakiety, nad firmę, która wprawdzie potrafi robić szum medialny, ale jeszcze wciąż ma niewiele namacalnych rzeczy do zaoferowania. ■

Mirosław Usidus



1. Burza słoneczna i pole magnetyczne Ziemi

Scenariusz Carringtona

Zanim Słońce ześle nam wielką burzę

Na połowie powierzchni Ziemi dochodzi do gwałtownych zaburzeń w sieciach energetycznych. Prądy i napięcia ulegają dzikim wahaniom, grożąc przeciążeniem urządzeń kluczowych dla współczesnej gospodarki, a nawet w pewnym sensie cywilizacji. Wkrótce po tym sieci elektryczne zaczynają padać, a świat pogrąża się w chaosie.

Co dalej? Naprawa uszkodzeń systemu energetycznego trwa długie miesiące. Bez stałego dostępu do energii elektrycznej, która jest siłą napędową współczesnego życia, społeczeństwa i gospodarki, wszystko zatrzymuje się. Rynki załamują się, gdy brak prądu przerywa dostawy żywności, paliwa i usługi transportowe. W ciągu kilku tygodni nieuniknione stają się: zapaść gospodarcza, globalne wstrząsy społeczne i masowy głód.

To scenariusz katastrofy nazywanej efektem Carringtona, lub podobnie, na cześć Richarda Carringtona, badacza nieba, który w 1859 roku zauważył dwie plamy jasnego światła odrywające się od powierzchni

Słońca i zmierzające z ogromną prędkością ku Ziemi. Rozbłysk z 1859 roku, choć nastąpił jeszcze przed wielkim rozkwitem cywilizacji technicznej opartej na elektryczności, był pierwszym, który został zauważony nie tylko jako zdarzenie astronomiczne, ale także burza elektromagnetyczna mająca wpływ na pierwsze urządzenia oparte na przepływach prądu elektrycznego.

Wówczas pojawiły się też pierwsze sygnały ostrzegawcze przed zakłóceniami. Latem 1859 roku astronomowie odnotowali niezwykle aktywność na Słońcu. Pojawiła się duża liczba plam słonecznych. Pod koniec sierpnia obserwatorzy na całym świecie



obserwowali zdumiewająco jasne zorze polarne. Gazety donosiły o wspaniałych pokazach na niebie w północnej Australii i Stanach Zjednoczonych. Żeglarze na Oceanie Spokojnym pisali o nocnym niebie w tonach pełnych zachwytu.

Dzień po tym, jak Carrington zauważył odrywające się od Słońca wyrzuty materii, te wielkie polacie silnie namagnetyzowanej plazmy słonecznej uderzyły w Ziemię. Zorze znów były widoczne na całej planecie, nawet w krajach położonych blisko równika. Nocne niebo było tak jasne, że niektórzy mylili je ze światem. W Ameryce Północnej iluminacja tworzona przez zorze była wystarczająca do czytania w nocy.

Operatorzy telegrafów, pracujący nad wczesnymi sieciami komunikacyjnymi obejmującymi Europę i Amerykę Północną, donosili o zadziwiających zjawiskach. W Europie zakłócenia magnetyczne wywoływały niekontrolowane prądy elektryczne w całej sieci. Na wiele godzin komunikacja między europejskimi stolicami całkowicie się zatrzymała. W Ameryce operatorzy odnotowali przeciążenia oraz iskrzenie sprzętu, a w niektórych miejscach nawet samozapłony. Były przypadki, gdy operatorzy, którzy odcięli dopływ prądu do linii, ze zdumieniem zauważali, że nadal mogą się komunikować dzięki elektryczności wzbudzonej przez... wtedy nie rozumiano jeszcze dokładnie tego mechanizmu. Jednak ogólnie szkody i zniszczenia miały dość ograniczony zasięg. Były to raczej ciekawe, zadziwiające ówczesnych ludzi zjawiska niż wielka katastrofa.

Potężne rozbłyski słoneczne są szkodliwe z powodu silnych pól magnetycznych i elektrycznych, które ze sobą niosą. Na szczęście dla nas Ziemia ma wbudowaną osłonę. Wirujące, stopione żelazne jądro naszej planety tworzy potężne pole magnetyczne wokół planety (1). Zwykle jedynym znakiem, że jesteśmy smagani plazmą słoneczną, są zorze polarne wokół

biegunów. Jednak gdy scenariusz Carringtona zacznie się ziszczać, ucierpią nie tylko transformatory i sieci energetyczne (co już się zdarzało np. w 1989 roku), ale również np. GPS, co może oznaczać paraliż transportowy i komunikacyjny na Ziemi.

Dokładna geneza rozbłysków słonecznych wciąż nie jest dobrze poznana. Wiadomo, że Słońce porusza się w cyklu jedenastoletnim, na przemian w okresach niskiej i wysokiej aktywności. Choć rozbłyski słoneczne mogą się zdarzyć w każdym momencie cyklu, prawdopodobieństwo ich wystąpienia jest znacznie większe w latach aktywnych. Obecny cykl osiągnął minimum około końca 2019 roku, a maksimum osiągnie prawdopodobnie między 2023 a 2026 rokiem.

Duże zakłócenia magnetyczne zdarzały się już wcześniej. Zdarzą się ponownie. Świat i operatorzy sieci energetycznych muszą się przygotować albo staną w obliczu katastrofy. Pewne kroki w tym kierunku zostały już podjęte. Od 1996 roku sonda kosmiczna SOHO stale monitoruje Słońce, dostarczając kilkugodzinnych informacji o nadchodzących rozbłyskach. Od tego czasu do SOHO dołączyły inne sondy kosmiczne, dając nam coraz bardziej szczegółowe obrazy naszej gwiazdy. Dzięki ostrzeżeniom z tych systemów operatorzy sieci energetycznych mają teraz trochę czasu, aby zabezpieczyć wrażliwe części sieci.

Poza samym monitorowaniem Słońca, operatorzy sieci energetycznych mogą również pracować nad uodparnianiem swoich systemów na uderzenie z kosmosu. Można np. zainstalować nowe urządzenia, które będą absorbować lub przekierowywać wahania mocy. To jednak kosztuje. Szacunki są różne, ale zwykle mówi się o kosztach rzędu 10–20 miliardów dolarów dla samej sieci w USA. Jednak w porównaniu z ogromnymi szkodami, jakie może spowodować taka flara, koszt przygotowań wydaje się niewielki. ■

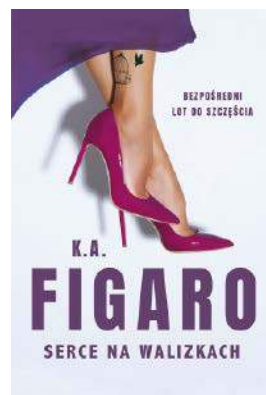
Mirosław Usidus

Serce na walizkach

K.A. Figaro

Wydawnictwo Lipstick Books, liczba stron: 368, cena: 39,99 zł

Czy chociaż raz zastanawiałaś się, jak wyglądałoby twoje życie, gdybyś...? Czy czułaś się kiedyś zdeptana, oszukana, rozczarowana i odepchnięta przez męża? Czy żyjąc w idealnym związku, chociaż raz zatraciłaś obraz własnej siebie? Czy przez problemy i samo życie zapomniałaś o gorących i namiętnych chwilach w sypialni, czy też o motylach szalejących w brzuchu? Co musiałoby się stać, abyś pozwoliła sobie na zapomnienie i uległa obcemu, bezwstydному i inteligentnemu nieznajomemu? Małgorzata Piłat to kobieta żyjąca w związku, który dla otoczenia wydaje się idealny, lecz w rzeczywistości wyniszcza ją od środka, mimo że ma wszystko. Jej mąż już dawno zapomniał, że miłość składa się z drobnych rzeczy, a nie z góry pieniędzy, które zdobią ich konta. Co robi Gośka, gdy na jej drodze stanie mężczyzna, przy którym poczuje się znów sobą, kochana i szczęśliwa? Czy zatraci się w obcych ramionach, czy będzie trwać na straży swojego małżeństwa? Tyle pytań... A odpowiedzi są na wyciągnięcie ręki, w książce „Serce na walizkach”.





**KOMPUTERY XXI W.
I znów rewolucja**

Czy można być dobrym inżynierem lub naukowcem, jeśli nie wie się, co komputer może, a czego nie może? Wydaje się, że dziś już nie bardzo. Ponadto wiedza o algorytmach, technikach programowania i ograniczeniach maszyn przydaje się choćby po to, by umieć dokonać właściwego wyboru w świecie agresywnego marketingu i sprzedaży software'u.

Czy pisanie kodu to zawodowa przyszłość wszystkich inżynierów i naukowców?

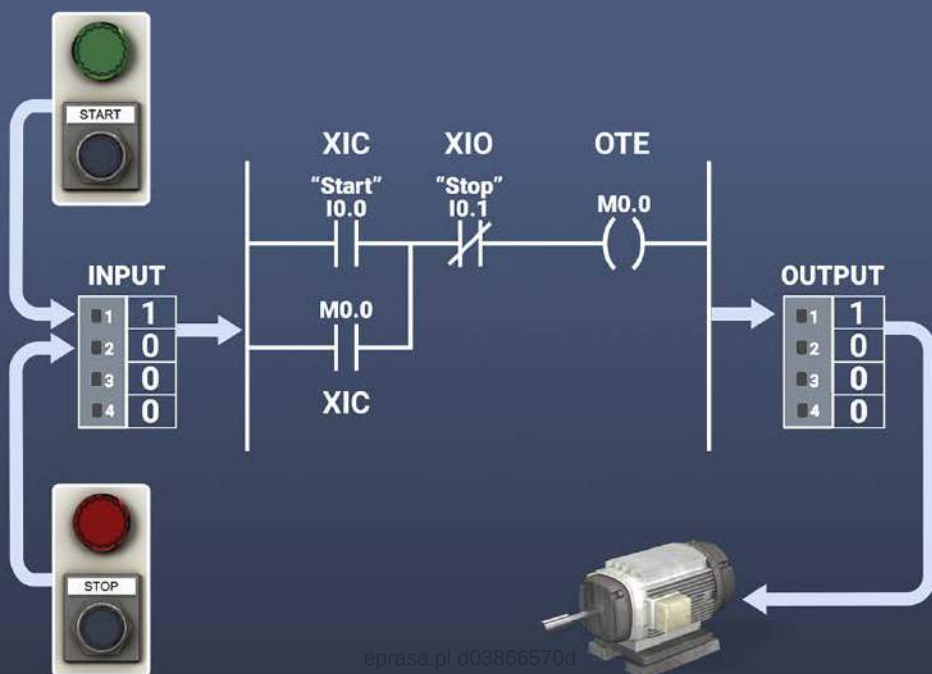
BICEPS PROGRAMI- STYCZNEJ LOGIKI

Edukacja zwykle nie nadąża za rozwojem technologii. Współcześni inżynierowie mechanicy wciąż raczej nie są programistami, a przynajmniej nie są edukowani z naciskiem na tego rodzaju umiejętności. Jednak przykłady i zachodzące zmiany sygnalizują, że znajomość programowania staje się właściwie niezbędna.

Inżynierowie z tradycyjnych dziedzin techniki np. wspomnianej mechaniki często mają do czynienia z tzw. logiką drabinkową. Język drabinkowy/logika drabinkowa (ang. Ladder logic, LD, LAD) – tak nazywany jest graficzny język programowania sterowników PLC (1). Pierwotnie był pisemną metodą dokumentacji sterowania przekąźnikowego, stosowanego w produkcji i kontroli procesów przemysłowych. Urządzenia w szafie przekąźnikowej są reprezentowane przez symbole na schemacie drabinkowym razem z połączeniami między nimi. Jako język graficzny jest prostszy i łatwiejszy do opanowania niż język programowania sterowników.

Ponieważ jednak programowanie skryptów wyższego poziomu zaczyna coraz częściej integrować się ze sterownikami i napędami, a urządzenia korzystające z języków wyższego poziomu dają konkurencyjną przewagę, coraz bardziej pożądana jest wiedza z zakresu programowania.

1. Przykład schematu w logice drabinkowej



Inżynier może pracować w fabryce, w której maszyny działają, opierając się na logice drabinkowej, ale co się stanie, gdy firma podejmie projekt integracji jednej z linii produkcyjnych z ramionami robotycznymi (2), które działają w oparciu o język programowania wyższego poziomu? Logika drabinkowa nie zniknie, owszem, a wykształceni inżynierowie zawsze będą mieli pracę, pojawiają się jednak nowe wymagania, którym warto sprostać w ramach zawodowego rozwoju. Wciąż nie ma bezwzględnego wymogu, by inżynierowie mechanicy byli programistami. Wciąż większość linii produkcyjnych to urządzenia mechaniczne. Jednak jeśli chce się rozwiązywać nowoczesne problemy, to trzeba znać nowoczesne metody rozwiązań. A rozwiązywanie problemów to niezmiennie jedno z głównym oczekiwań wobec inżynierów.

Programowanie rozwija i pozwala lepiej się zrozumieć

Wyjdźmy poza hale produkcyjne na obszary badań, projektowania, analizy danych. Współcześni naukowcy i inżynierowie coraz większą część dnia pracy spędzają przed komputerem. Reguła ta dotyczy najróżniejszych dziedzin: astronomii, biologii, fizyki, inżynierii lotniczej i kosmicznej, ekonomii, genetyki, ekologii, inżynierii środowiska, neurobiologii... lista jest długa. Współczesna nauka i inżynieria polegają na przetwarzaniu, analizowaniu i wyciąganiu wniosków z danych.

Dość banalnie brzmi dziś stwierdzenie, że pracę badawczą można przyspieszyć wielokrotnie, pisząc programy komputerowe w celu zautomatyzowania

żmudnych zadań (choćby takich jak czyszczenie i integracja danych), które w przeciwnym razie trzeba wykonywać ręcznie. Można oczywiście zwrócić się do zawodowych programistów, by przygotowali software, ale w dzisiejszym wysoko wyspecjalizowanym świecie nauki i inżynierii najefektywniejszy wydaje się pomysł, aby specjalistyczne oprogramowanie przygotowywał nie programista z zewnątrz, lecz specjalista w danej dziedzinie. To on wie najlepiej, jakie są potrzeby, zadania i cele.

Ciekawym aspektem programowania jest to, że pozwala ono specjalistom odkrywać bardziej kreatywne rozwiązania niż znane dotychczas. Pozwala bowiem wyjść poza rutynę używanych zwykle narzędzi i zestawów danych, przekroczyć ograniczenia. Każdy, kto miał do czynienia z pisaniem programów, wie, że polega ono w dużym stopniu na odkrywaniu nowych możliwości i zasobów. Istnieje mnóstwo gotowych pakietów, bibliotek, algorytmów, funkcji, programów, języków i oprogramowania do wielu różnych celów. Nie trzeba uczyć się wszystkiego, aby zostać programistą, ale trzeba określić, które z nich najlepiej pasują do celów i wyników, które chcemy osiągnąć.

Niebagatelne znaczenie ma także to, że znajomość programowania pozwala naukowcom i inżynierom skutecznie komunikować się z programistami. Nie trzeba być samemu wybitnym programistą, ale samo rozumienie, na czym polega programowanie, wybitnie zwiększa efektywność współpracy z członkami zespołu zajmującymi się software'em.

2. Robotyczne ramię w fabryce mebli



Współcześnie zarówno naukowcy, jak też inżynierowieprojektanci w swojej pracy mają do czynienia w wielkimi zasobami danych. Pojawia się więc kolejna pożądana kompetencja, wybiegająca nawet poza

zestaw typowych umiejętności programisty, polegająca na radzeniu sobie z big data, a najlepiej biegłym opanowaniu technik ich przetwarzania w użyteczne informacje. Trudno oczekiwać, że każdy od razu

Języki programowania, które uchodzą za najbardziej przydatne dla specjalistów w dziedzinie nauki i techniki

1. Java

Java jest własnością firmy Oracle. Działa na urządzeniach mobilnych, a konkretnie w aplikacjach na Androida, niektórych lub wszystkich aplikacjach desktopowych, aplikacjach internetowych, serwerach, grach, bazach danych i wielu innych. Java może być używana także w komputerach z systemem Linux, Raspberry Pi, Mac i oczywiście Windows. Uchodzi za język łatwy do nauczania.

2. C

Język programowania ogólnego przeznaczenia, który obsługuje programowanie proceduralne i strukturalne, a także rekurencję i leksykalny zakres zmiennych. C jest powszechnie używany do tworzenia aplikacji desktopowych. Przykład języka programowania niskiego poziomu, nadającego się jednak do realizacji różnorodnych celów.

3. Python

Python jest kolejnym przykładem języka programowania ogólnego przeznaczenia. Może być stosowany do tworzenia systemów back-end lub pisania oprogramowania użytkowego, do data science lub pisania skryptów systemowych. Python kładzie nacisk na czytelność kodu. Jest to obecnie jeden z najpopularniejszych języków programowania.

4. C++

Jest to język programowania ogólnego przeznaczenia. Może działać na komputerach z takimi systemami operacyjnymi jak Windows, kilku wersjach UNIX-a czy Mac OS. Obsługuje programowanie obiektowe, proceduralne i generyczne. Używany głównie do tworzenia systemów operacyjnych, przeglądarek, gier i innych aplikacji.

5. Visual Basic .NET

Kolejny przykład obiektowego języka programowania. Najlepiej sprawdza się na platformie .NET Framework firmy Microsoft. VB.NET jest strukturalnym językiem programowania. Używa on instrukcji do wskazywania działań, które mają być podjęte przez komputer. Najlepiej nadaje się do tworzenia aplikacji, które będą działać głównie w systemie Windows.

6. JavaScript

Jest jednym z najbardziej znanych języków programowania. Powszechnie używany na stronach internetowych. JavaScript nie jest trudnym do nauki językiem. Twórcy stron internetowych uczą się tego języka w połączeniu z HTML i CSS, co dziś jest niezbędnikiem web developera.

7. C#

Kolejny przykład obiektowego języka programowania, C#, czyli C Sharp, to popularny język programowania służący do tworzenia aplikacji desktopowych, aplikacji webowych i innych serwisów internetowych. Jest to również jeden z głównych języków używanych do tworzenia aplikacji w ekosystemie Microsoftu, wykorzystywany także do programowania gier na platformie Unity. Jest pochodną języków programowania C i C++. Stosunkowo łatwy do nauczania.

8. PHP

PHP lub Hypertext Preprocessor jest jednym z najpopularniejszych na świecie języków programowania ogólnego przeznaczenia typu open source. Używany głównie do tworzenia stron internetowych, często w połączeniu z HTML. PHP działa po stronie serwera, co oznacza, że przetwarza rzeczy przed wysłaniem ich do klienta (użytkownika), bez możliwości zobaczenia kodu przez klienta. PHP generuje dynamiczne strony, zbiera dane z formularzy, wysyła i odbiera pliki cookie, a nawet szyfruje dane. Łatwy do opanowania.

9. SQL

Structured Query Language jest szczególnym językiem programowania używanym do budowania, utrzymywania i manipulowania danymi w bazie danych. Za pomocą SQL operator może tworzyć zapytania, pobierać dane, wstawiać rekordy, aktualizować rekordy, usuwać rekordy, tworzyć nowe bazy danych, nowe tabele, tworzyć procedury składowane, widoki i ustawiać uprawnienia.

10. Objective-C

Objective-C jest podstawowym językiem dla osób tworzących oprogramowanie, które będzie działać w ekosystemach OS X i iOS. Objective-C rozwinął się jako obiektowe rozszerzenie języka programowania C, dziedzicząc wiele cech tego języka.



stanie się data scientist, jednak umiejętności związane z kodowaniem algorytmów tworzących modele klasyfikujące i analizujące dane są przydatne w takich samych powodów, jak wyżej wymienione w kontekście „zwykłego” programowania.

Wysoki i niski poziom

Najprościej rzecz ujmując, język programowania to zestaw instrukcji używanych do nakazania komputerowi, aby coś zrobił. Komputery „myślą” i „rozmawiają” w systemie binarnym (zbiory 1 i 0). Języki programowania są rodzajem translatora, który przekształca ludzkie polecenia na te jedyne i zera, tak aby komputer mógł je „zrozumieć” (3). Każdy język programowania jest nieco inny pod względem sposobu, w jaki rozmawia z komputerem, używając różnych poleceń, symboli, itp. Niektóre z języków programowania są bardziej uniwersalne, inne zaprojektowane specjalnie dla konkretnych systemów operacyjnych (jak Swift dla iOS lub C# dla Windows). Można je ogólnie podzielić na dwa główne typy – pierwszy nazywany jest „językami wysokiego poziomu”, a drugi „językami niskiego poziomu”.

Języki wysokiego poziomu to te, które są, ogólnie rzecz biorąc, bliższe sposobom komunikacji międzyludzkiej. Używają czytelnych dla człowieka poleceń, takich jak na przykład „object”, „order”, „run”, „class”, „print”, itp. Z tego powodu języki wysokiego poziomu są zwykle łatwiejsze do opanowania niż języki niskiego poziomu. Aby komputer mógł zrozumieć polecenia napisane w językach wysokiego poziomu, muszą one jednak zostać przetłumaczone lub skompilowane na język niskiego poziomu, czyli język maszynowy. Często zdarza

3. Binarny kod maszynowy

się, że programiści nigdy nie widzą wyniku w języku maszynowym. Języki wysokiego poziomu działają wolniej niż języki niskiego poziomu, ponieważ potrzeba czasu na przetłumaczenie poleceń wysokiego poziomu na kod maszynowy, zanim maszyna będzie mogła je wykonać. Mowa tu o milisekundach lub nawet sekundach, jeśli są to bardzo duże programy. Szeroko znanymi przykładami języków wysokiego poziomu są PHP, Ruby i Java (4).

Języki niskiego poziomu są bliższe kodowi maszynowemu (binarnemu). Są one o wiele trudniejsze do zrozumienia dla człowieka, ale wciąż o wiele łatwiejsze niż czysty kod binarny. Można je generalnie podzielić na dwa typy: „język montażu” i „język maszynowy”. Główną zaletą tego rodzaju języków jest względna szybkość ich „tłumaczenia” na kod działania maszyny. Oferują też znacznie bardziej precyzyjną kontrolę nad wykonywaniem poleceń przez maszynę. Przykłady języków niskiego poziomu to: Fortran, COBOL, x86.

4. Kod języka programowania wysokiego poziomu Java



Kodowanie rozwija

Panuje opinia, że znajomość języków i technik programowania staje się coraz bardziej niezbędna w pracy inżynierów i innych specjalistów w dziedzinie techniki (5). Obecnie kształcą się mechanicy, budowniczcy, architekci, elektronicy i przedstawiciele wielu innych zawodów raczej na pewno i to wcześniej niż później zetkną się z potrzebą napisania jakiegoś programu, który umożliwi lub usprawni coś w ich pracy.

Dlaczego uczyć się programowania? Są argumenty odwołujące się do samorozwoju i samodoskonalenia. Nie muszą one koniecznie dotyczyć inżynierów i naukowców, bo każdemu może przydać się wzmocnienie „mięśnia logiki” przez używanie zdefiniowanych operatorów logicznych, pętli i instrukcji warunkowych. Nauka logiki programistycznej pozwala spojrzeć na procesy i problemy do rozwiązania w sposób analityczny, układający je w uporządkowane sekwencje. Rozwija się również myślenie systemowe, w którym uświadamiamy sobie, że wszystko jest połączone wzajemnie ze sobą. Analizując problem i proponując jego rozwiązanie, trzeba w myśleniu programistycznym przeanalizować, ocenić wpływy i implikacje w większej skali niż tylko segment, w którym pracujemy. Lepiej też widać, iż całość jest czymś więcej niż sumą części.

Podczas pisania kodu lub tworzenia programu do wykonania konkretnego zadania lub rozwiązania problemu, zwykle jest potrzeba rozłożenia go na mniejsze elementy. Analizowanie problemu przez mniejsze komponenty pomaga zidentyfikować przyczyny problemów, jeśli występują. Jednocześnie



5. Współczesny inżynier

zdajemy sobie sprawę, że w większości przypadków istnieje więcej niż jedno rozwiązanie danego problemu. Umiejętność programowania rozwija więc nas w sposób różnorodny i wszechstronny. ■

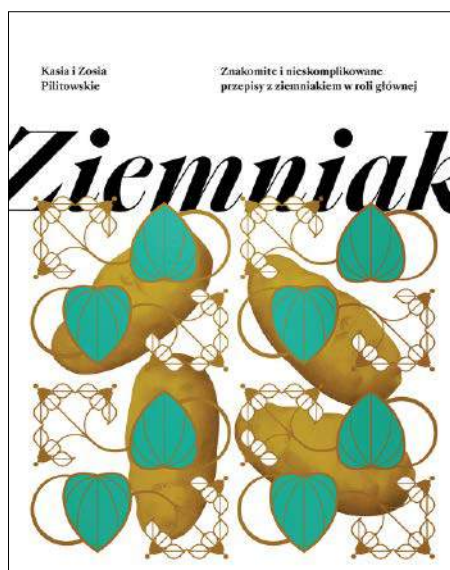
Mirosław Usidus

Ziemniak

Zofia i Katarzyna Piłtowskie

Wydawnictwo Buchmann, liczba stron: 256, cena: 49,99 zł

Znakomite i nieskomplikowane przepisy z ziemniakiem w roli głównej. Po sukcesie książki „Jajko” pozytywnie zakręcony duet Kasi i Zosi Piłtowskich, mamy i córki z krakowskiej kawiarni Ranny Ptasek, zabiera nas w ziemniaczaną podróż! I to dosłownie, bo w książce znajdziemy potrawy z ziemniakiem w roli głównej ze wszystkich stron świata. „Ziemniak” podzielony jest na pięć rozdziałów. Są tutaj przepisy podstawowe, pomysły na dania zagrabione, smażone, ulepione, z każdego zakątka świata. A na koniec ziemniaczki na słodko. Oprócz rodzimych pyszności (takich jak pyzy, cepeliny czy kluski śląskie) w „Ziemniaku” znajdziemy także przepisy na bombajskie ziemniaczki Vada Paw, afrykańską zupę z masłem orzechowym czy znane wszystkim miłośnikom nart austriackie rosti. W skrócie – ogrom inspiracji do przyrządzenia prawdziwej skrobiowej uczty!





1. Komputer biurowy stacjonarny

O roli i wszechobecności komputerów (1) w naszym życiu można dziś mówić już tylko banały. Wydaje się, że wszystko na ten temat już zostało napisane. Zaś nasze uzależnienie od komputerów, niezależnie od tego, jak się zmienia, a podobno akurat się zmieniają, będzie tylko rosnąć.

Grozi nam potop danych i energetyczna apokalipsa, ale pecety trwają

KOMPUTERY – CO DALEJ?

Jakie są główne kierunki owych właśnie zachodzących zmian w świecie komputerów? Poza dążeniem do uzyskania większej mocy obliczeniowej coraz

powszechniejsza jest adaptacja technik uczenia maszynowego a także rozwiązań IoT. Każdy z tych nurtów przenikają dążenia do redukcji zużycia energii związanej z obliczeniami, szczególnie w centrach danych. W świetle raportu Stowarzyszenia Przemysłu Półprzewodnikowego (SIA) z 2016 roku prognozującego, że do 2040 roku komputery na świecie będą zużywać więcej energii elektrycznej niż wszystkie kraje łącznie będą w stanie wyprodukować, zmiany zmierzające do ograniczenia energochłonności mają kluczowe znaczenie.

Panuje przekonanie, że nowe procesory obliczeniowe, oparte na dotychczasowej konstrukcji, nie zwiększają już ogólnej wydajności, głównie ze względu na fundamentalne ograniczenia stale kurczących się tranzystorów krzemowych. Prawo Moore'a w klasycznej postaci rzeczywiście przestało obowiązywać.

Szerzej piszemy o tym w innym artykule w tym numerze MT. Trwają poszukiwania alternatyw dla krzemu lub innego podejścia, jeśli to wciąż ma być krzem. Wszystko, ma się rozumieć, z myślą nie tylko o szybkości i wydajności, ale też o zużyciu energii.

Nanorurki, magnetyzm i światło

Są pewne komputerowo-elektroniczne innowacje a raczej kierunki innowacji, o których mówimy i piszemy od lat, ale jakoś wciąż nie możemy doczekać się praktycznych wdrożeń. Należy do nich np. pomysł integracji nanorurek węglowych z układami scalonymi, który wciąż ma status rozwiązania obiecującego, ale niezrealizowanego w praktyce, choć niektóre ośrodki i firmy, choćby Aligned Carbon, zapewniają, że już wkrótce da nam to niezrównaną wydajność w chipach nowego typu. Zobaczmy.

Skoku wydajności oczekuje się również od zastosowania w układach pamięci typu MRAM (Magnetic Random Access Memory), ultraszybkiego nośnika o dużej gęstości. Głównym ograniczeniem w skalowaniu tego rozwiązania jest obecność zanieczyszczeń strukturalnych w materiałach magnetycznych, co zmniejsza pojemność. Rozwiązać mogłaby to technika nazywana spin-Ion, w której wykorzystuje się napromieniowanie materiałów magnetycznych wiązką lekkich jonów. Rozwiązania te mają być stosowane na początek tylko w superkomputerach, ale z czasem staną się prawdopodobnie integralną częścią wielu innych platform obliczeniowych (2).

Za jedną z barier współczesnych układów przetwarzania danych w aplikacjach uważa się zbyt małą szybkość komunikacji pomiędzy pamięcią a układami przetwarzającymi. Aby wesprzeć takie zastosowania, jak uczenie maszynowe i nowe typy

generowanej komputerowo rzeczywistości (AR/VR), firmy MemComputing i UpMEM proponują zintegrowanie procesorów obliczeniowych w pamięci. Skracając czas obliczeń, jednocześnie zmniejszając ilość potrzebnej pamięci i zużycie energii.

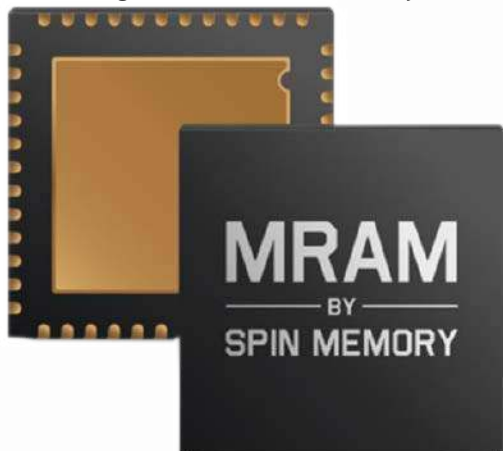
Odpowiedź na problemy energetyczne współczesnych komputerów, a przede wszystkim wielkich ich kompleksów zgromadzonych w centrach obliczeniowych, być może znajdziemy w świetle. Optoelektronika pozwala na wykonywanie obliczeń z większą prędkością przy znacznie mniejszym zużyciu energii, zapewniając tym samym wyższą wydajność na jednostkę mocy. Ponieważ tradycyjnie procesory optoelektroniczne wymagają znacznie więcej przestrzeni fizycznej, nie są one powszechnie stosowane. Jednak dla centrów danych oszczędność energii jest ważniejsza niż zajmowana przestrzeń, dzięki czemu procesory optoelektroniczne mogą stać się tam wkrótce dominującą platformą sprzętową. Czynnikiem ograniczającym chipy optoelektroniczne jest techniczna trudność w umieszczeniu wystarczającej liczby laserów na chipie w celu osiągnięcia wysokiej wydajności przy niskiej mocy. Rozwiązaniem może być technika Iris Light Technologies, pozwalająca na drukowanie setek laserów bezpośrednio na chipach przy użyciu nanotechnologicznego „tuszu”.

Na krawędzi

Specyficzne wymagania techniczne ma upowszechniający się model obliczeniowy Edge Computing, którego podstawowym założeniem jest fizyczne przybliżenie obliczeń i przechowywania danych do ich źródła, czyli do urządzeń, w których dane są generowane. Zmniejsza to opóźnienia i zwiększa autonomię urządzeń brzegowych, którymi mogą być np. współpracujące w rojach roboty lub autonomiczne drony, natychmiastowo przetwarzające informacje, aby móc sprawnie i bezpiecznie działać. Przesyłanie informacji do chmury do przeanalizowania, a następnie zwracanie ich do urządzenia brzegowego zajmuje zbyt dużo czasu, aby zagwarantować odpowiednie i wystarczająco szybkie reakcje (3). Poza tym wzrasta przez to zapotrzebowanie na zasoby obliczeniowe i pamięci masowe.

Szczególnie w przypadku aplikacji, od których wymaga się szybkiego działania i reagowania, rośnie potrzeba podejmowania decyzji bez wysyłania danych do chmury. Oznacza to jednak, że urządzenia brzegowe muszą sprostać w sensie energetycznym. Dlatego firmy takie jak Synthara pracują nad nowego typu układami przetwarzającymi, które umożliwią aplikacjom opartym na sztucznej inteligencji

2. Układ Magnetic Random Access Memory





„W systemach komputerowych, z których obecnie korzystamy, mamy CPU (centralną jednostkę obliczeniową) i pamięć do przechowywania danych”, wyjaśnia Wang w opublikowanym komunikacie. „Gdy chcemy

przetwarzać informacje, dane te muszą być przeniesione z pamięci do procesora, co nie jest problemem, gdy nie mamy dużej ilości danych. Ale jeśli chodzi o zadania z zakresu uczenia maszynowego, eksploracji danych lub specyficzne wyzwania, jak np. sekwencjonowanie genomu, to ruch danych jest bardzo duży. Pojawiają się wąskie gardła z powodu ograniczonej przestrzeni i oczywiście skokowo rośnie zużycie energii”.

Propozycja Wanga i zespołu z Illinois Tech polega na rozpoczęciu wstępnego przetwarzania danych już w pamięci przed przeniesieniem ich do procesora, co powinno zmniejszyć ilość danych, które muszą migrować do CPU i sprawi, że przetwarzanie w CPU będzie bardziej wydajne. Brzmi jak proste rozwiązanie, ale w rzeczywistości wymaga to dość złożonego połączenia wielu różnych nowych technik zarówno po stronie sprzętu, jak i oprogramowania.

Pogłoski o śmierci peceta mocno przesadzone

W ciągu ostatnich kilkunastu lat niejednokrotnie mogliśmy usłyszeć i przeczytać o tzw. erze postpeceowej. Pojęcie to powstało pod wpływem zmian obserwowanych w późnych latach 2000 i później, w ubiegłej dekadzie. Chodzi przede wszystkim o obserwowany przez lata spadek sprzedaży komputerów osobistych (PC) na rzecz urządzeń mobilnych, smartfonów i tabletów, a także innych komputerów przenośnych, sprzętu noszonego itp.

Pierwszy prognozę „końca pecetów” sformułował badacz z MIT, David D. Clark w 1999 roku, który widział przyszłość komputerów jako „nieuchronnie heterogeniczną” w „sieci wypełnionej usługami”. Clark, jako jeden z pierwszych, opisywał świat, w którym „wszystko” będzie łączyć się z Internetem, obliczenia będą wykonywane głównie przez urządzenia informatyczne, a dane będą przechowywane przez scentralizowane usługi hostingowe zamiast

na fizycznych dyskach. Bill Gates, wówczas szef Microsoftu i Steve Jobs, szef Apple, również w tym czasie przewidywali przesunięcie w kierunku urządzeń mobilnych. Jobs spopularyzował termin „post-PC” w 2007 roku podczas premiery pierwszego iPhone’a, a w 2011 roku uruchomił iCloud, usługę umożliwiającą linii produktów Apple synchronizację danych z komputerami PC za pośrednictwem usług w chmurze, uwalniając urządzenia iOS od zależności od komputera PC.

Podczas prezentacji w styczniu 2012 roku, wiceprezes Intela Tom Kilroy zakwestionował nadejście ery post-PC, powołując się na badanie studentów college’u, w którym 66% respondentów nadal uważało komputer osobisty za „najważniejsze urządzenie w ich codziennym życiu”. Wkrótce także media zaczęły kwestionować nadejście ery post-PC. Najczęściej powtarzana opinia głosiła, że te podziały i klasyfikacje mają już niewielki sens w sytuacji, gdy urządzenia przenośne, choćby smartfony, są w zasadzie komputerami, w dodatku o mocy i możliwościach większych niż „definicyjne” pecety sprzed dwu dekad. Poza tym nie słabła popularność nie tylko laptopów, ale nawet ciężkich komputerów stacjonarnych w niektórych zastosowaniach, np. w grach.

W dodatku właśnie ostatnio widzimy ponowny wzrost sprzedaży tych pecetów, co to „miały odejść”. Różne raporty rynkowe mówią o kilkunastoprocentowym wzroście w 2020 roku w porównaniu z rokiem poprzednim. Oczywiście czynniki napędzające ten wzrost mają związek z pandemicznymi warunkami, pracą w domu i potrzebami zdalnego nauczania. Jednak eksperci uważają, że nie tylko o to chodzi, i pecety, laptopy, notebooki, czy jakieś inne produkty w jeszcze niesprecyzowanej postaci, ale o podobnym kształcie i charakterze, pozostaną z nami, bo są po prostu wygodnym narzędziem. ■

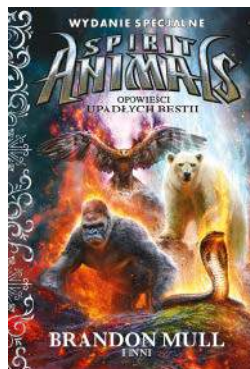
Mirosław Usidus

Spirit Animals. Opowieści upadłych bestii

Brandon Mull

Wydawnictwo Wilga, liczba stron: 208, cena: 36,99 zł

Tom specjalny popularnej serii, do której autorów zalicza się bestsellerowy pisarz fantasy – Brandon Mull! Upadłe wielkie bestie – niegdyś najbardziej potężne stworzenia Erdasu, które zdradziły inne Wielkie Bestie i których knowania kosztowały życie tysięcy istot – odradzają się jako zwierzoduchy wyjątkowych dzieci. Grozi im jednak wielkie niebezpieczeństwo. Tajemniczy myśliwy próbuje odebrać im życie i bezlitośnie walczy z każdym, kto staje w ich obronie. Poznajcie pięć opowieści o legendarnych wielkich bestiach – oraz o młodych bohaterach, którzy nie cofną się przed niczym, aby je ocalić.



Czy programiści podobnie jak wiele innych profesji powinni czuć się zagrożeni z powodu ekspansji sztucznej inteligencji? Jeśli nawet, to na pewno nie w dającej się przewidzieć przyszłości.

Programowanie i web development po nowemu

CZY DA SIĘ WYRAZIĆ NOWE RZECZY ZA POMOCĄ STAREGO JĘZYKA?

Założona przez Elona Muska firma OpenAI zaprezentowała niedawno algorytm o nazwie Codex (1), który potrafi interpretować polecenia pisane po angielsku i zamieniać je we fragmenty użytecznego kodu, przekształcając je w oprogramowanie dla prostych gier lub stron internetowych. Jak donosił „The Verge”, użytkownik może opisać „swoimi słowami” podstawowy wygląd lub funkcjonalność strony internetowej, którą chce zbudować, w zwykłym języku angielskim, a Codex generuje prosty projekt oparty na swojej interpretacji opisu. Codex wykorzystuje GPT-3, osławiony algorytm do generowania tekstu również autorstwa OpenAI. Działa jak asystent lub zastępca programisty. Przyjmuje koncepcje i pomysły, tworząc następnie kod.

„Widzimy to jako narzędzie do wspomagania programistów”, wyjaśniał Greg Brockman, współzałożyciel

OpenAI, w rozmowie z „The Verge”. „Na programowanie składają się dwa elementy – analiza problemu z dążeniem do zrozumienia go i mapowanie kodu. To „właśnie ta druga część”, dodał Brockman, „którą ludzie uważają za monotonną, bierze na siebie Codex”.

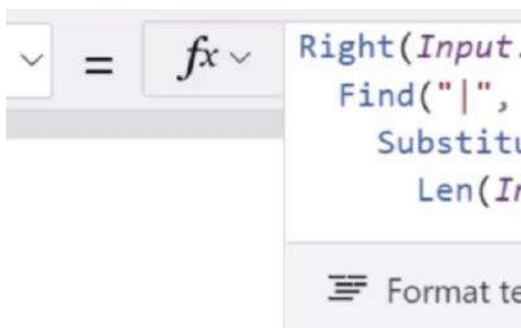
Pomysł wykorzystania GPT-3 do generowania kodu mieli wcześniej inni. Sharif Shameem poinformował w lipcu 2020 r., że użył GPT-3 do ułożenia strony internetowej poprzez wprowadzenie opisów. GPT-3 przekształcił je w kod JSX, będący rozszerzeniem składni JavaScript, która produkuje strony internetowe z użyciem Reacta, open source’owej biblioteki javascriptowej służącej do budowania interfejsów użytkownika lub komponentów stron.

W 2020 r. na konferencji Build dyrektor technologiczny Microsoftu Kevin Scott opowiedział o eksperymentalnym projekcie, w którym sztuczna inteligencja, wytrenowana na kodzie zmagazynowanym w GitHubie, tworzy programy, generując funkcje na podstawie opisowego komentarza. Scott akcentował głównie oszczędności czasu traconego przez programistów na nudne, powtarzalne zadania. Choć demo Microsoftu może wskazywać, że programiści mogą w końcu zostać uwolnieni od pracy przy kodowaniu prostych funkcji, silnik generujący kod został z pewnością zbudowany przez zespół programistów, i to prawdopodobnie dużo.

Przykładem dążenia do odciążenia programistów, choć nieco innego rodzaju niż opisywane

1. Prezentacja rozwiązania OpenAI Codex





2. Power Fx Microsoftu

wyżej rozwiązania AI, jest udostępniony znów przez Microsoft w marcu 2021 r. Power Fx, nowy język typu tzw. niskokodowego, który czerpie inspirację z popularnych formuł Excela (2). Jest dostępny na zasadzie open source. Firma ma nadzieję, że pomoże w rozwoju oferowanych przez nią platform Power, takich jak Power Automate czy Power Virtual Agents i stanie się standardem dla tego rodzaju zastosowań. Nowy język kierowany ma być przede wszystkim do użytkowników programu Excel i Power Platform. Oparty przede wszystkim na wykorzystaniu formuł, ma być użytkowany głównie przez ludzi o niewielkim doświadczeniu w kodowaniu. Umiejętności związane z obsługą Excela są w świecie biznesu rozpowszechnione w dużo większym stopniu niż biegła znajomość języków programowania. Platforma niskokodowa (ang. low-code development platform, LCDP) jest oprogramowaniem umożliwiającym budowę aplikacji w sposób wizualny, za pomocą diagramów, grafów, formuł czy formularzy bez znajomości języków programowania.

Automatyzacja, choć nieuchronna, to jednak niestraszna

Według badań firmy Evans Data Corp, na świecie są 23 miliony programistów. Do 2023 roku

3. Robot przy laptopie

liczba ta ma wzrosnąć do 27,7 miliona, choć według wielu opinii zawód ten jest zagrożony w dużym stopniu przez AI (3) i automatyzację. Obawy te są do pewnego stopnia zrozumiałe, jeśli uświadomimy sobie konsekwencje rozwoju technik opisanych wyżej lub inne, dostępne już od lat narzędzia dla twórców stron internetowych, np. Wix i Squarespace, które pozwalają każdemu zaprojektować i „zbudować” stronę internetową bez konieczności uczenia się HTML i CSS.

Jednak dane z ostatnich lat oraz cytowane badania wskazują na co innego. Choć już od dawna automatyzujemy mnóstwo prac i zadań, zapotrzebowanie na inżynierów oprogramowania nieustannie wzrasta. Ponadto AI, mimo różnych fantazji, nie programuje się sama. AI byłaby niczym bez ludzi pracujących nad algorytmami, wymyślających je, ulepszających, dopracowujących. Samo Google zatrudnia obecnie około 30 tysięcy osób, które pracują nad różnymi platformami AI.

Wreszcie, pojawienie się nowych technologii, takich jak IoT, autonomiczne pojazdy, wirtualna rzeczywistość, itp. stawia nowy zestaw wyzwań, które wymagają nowych inżynierów oprogramowania, specjalistów IT, inżynierów systemowych i wielu innych specjalistów z nowymi umiejętnościami do pracy nad nimi. Chociaż narzędzia AI, które potrafią pisać prosty kod, już istnieją, wciąż nie umiemy określić, którym funkcjom nadać priorytet

lub jaki problem miałby rozwiązać tworzony

fragment oprogramowania. Wcześniej uważano, że rola programistów może się zmienić dopiero w miarę dalszego doskonalenia systemów AI i to pozostaje prawdą. W przyszłości być może zamiast pisania kodu, programiści będą odpowiedzialni za analizowanie i gromadzenie danych wejściowych do algorytmów AI, które następnie będą samodzielnie tworzyć

oprogramowanie. Ale to jeszcze nie teraz.

Sztuczna inteligencja niewątpliwie jednak już potrafi tworzyć kod i z pewnością będzie stopniowo zastępować prozaiczne i uciążliwe zadania związane z programowaniem. Jednak po to, aby



„programiści przestali być potrzebni”, potrzebna byłaby o wiele bardziej zaawansowana sztuczna inteligencja. Ale wtedy, na tym hipotetycznym etapie w przyszłości będziemy zadawać to samo pytanie na temat każdego zawodu, jaki znamy.

Koniec języków obiektowych?

Większość języków programowania (i technologii internetowych) używanych dzisiaj została stworzona kilkadziesiąt lat temu. Dawni projektanci języków nie posiadali takiej wiedzy, jaką posiadamy dzisiaj. Stare języki programowania były tworzone z myślą o starych komputerach, w szczególności o jednordzeniowych procesorach. Stare języki programowania nie zostały zaprojektowane do korzystania z możliwości wielordzeniowych procesorów. Stare języki programowania nie zostały zaprojektowane z myślą o niezmienności, która okazała się kluczowa dla niezawodnego oprogramowania wielordzeniowego. Niezmiennność skutkuje mniejszą liczbą błędów, mniejszą złożonością i ogólnie lepszym doświadczeniem programistycznym.

A największym problemem jest to, że języki te muszą być wstecznie kompatybilne. Oznacza to, że nowa wersja Javy zawsze będzie musiała być wstecznie kompatybilna z Javą 1.0. Porównajmy te software'owe problemy z hardware'em. Sprzęt nie musi być wstecznie kompatybilny i nie musi obsługiwać 20-letnich modułów pamięci RAM. Języki programowania muszą zachować cały swój bagaż. Dzisiejsza Java to tak naprawdę ta sama Java sprzed 25 lat, choć z pewnymi nowymi funkcjami. Duże programy zorientowane obiektowo zmagają się z rosnącą złożonością. Coraz bardziej popularne staje się programowanie z wartościami niezmiennymi. Nie jest przypadkiem, że nawet nowoczesne biblioteki UI, takie jak React, są przeznaczone do używania w stanach niemutowalnych (niezmiennych).

Czym jest ów niezmienny stan? Mówiąc najprościej, są to dane, które się nie zmieniają. Funkcje, które nie mutują (nie zmieniają) stanu, nazywane są czystymi i są znacznie łatwiejsze w przetwarzaniu i testowaniu. Kiedy pracujemy z czystymi funkcjami, nigdy nie musimy się martwić o nic poza funkcją. W przeciwieństwie do obiektów w starszych językach, o których stan zawsze trzeba się martwić.

Według radykalnych ocen niektórych, szczególnie silnie dążących do zmian programistów, jesteśmy świadkami końca ery języków zorientowanych obiektowo. Dlaczego? Choćby dlatego, że jako stara koncepcja nie zostały one zaprojektowane do pracy z nowoczesnymi wielordzeniowymi procesorami

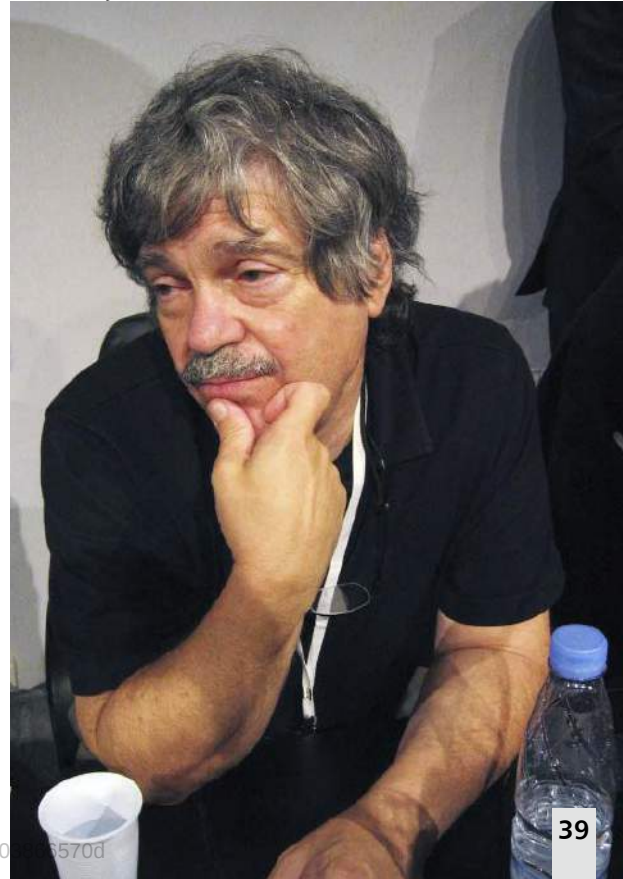
i systemami rozproszonymi. Jeden z twórców programowania obiektowego Alan Kay (4), legendarny programista zatrudniony w Xerox PARC w latach siedemdziesiątych, powiedział nie tak dawno, że niestety, współczesne programowanie obiektowe tak naprawdę skupia się na złym pomysle, zajmuje się budowaniem złożonych hierarchii obiektów, zamiast zajmować się przesyłaniem informacji. Języki obiektowe zostały stworzone, aby zarządzać złożonością kodu, jednak, jak na ironię, nie udaje im się osiągnąć nawet tego celu.

Większość innych języków programowania nie została zaprojektowana z myślą o współbieżności. Oznacza to, że pisanie kodu, który wykorzystuje wiele wątków/rdzeni procesora, jest dalekie od prostego.

Witaj świecie, w którym wszystko jest kodem

W jakim kierunku pójdzie zawód programisty, jeśli nie musi (na razie) obawiać się AI? W latach 60. i 70. ubiegłego wieku programistów nazywano „analitykami”. W książce Tima O'Reilly'ego pt. „Software Architecture Fundamentals Superstream”, Rebecca Parsons z firmy Thoughtworks przypominała, że „analitycy” zasadniczo zajmują się architekturą software'u,

4. Alan Kay



podbijają decyzje dotyczące tego, co oprogramowanie powinno robić i jak powinno być zbudowane. Analitycy analizują problem, zastanawiają się, na czym on polega i jak go skutecznie rozwiązać. Myślą o tym, jak go rozłożyć na części. Mogą nawet myśleć o tym, czy problem powinien w ogóle być rozwiązany, jak powinien być rozwiązany, jakie kwestie etyczne się z tym wiążą i jak te kwestie powinny być rozwiązane. Analitycy muszą też myśleć o tym, jak ludzie będą korzystać z oprogramowania, jaki jest interfejs użytkownika, jakie są wrażenia użytkownika, czy mogą z niego korzystać osoby niepełnosprawne itd.?

Najważniejszym nurtem w programowaniu w najbliższej dekadzie będzie wykorzystanie uczenia maszynowego i sztucznej inteligencji do zautomatyzowania większości czynności związanych z kodowaniem. O tym już była mowa. Inną wielką zmianą będzie 5G. Zwiększona przepustowość sieci i mocy obliczeniowej doprowadzi do zmian w językach programowania i powstania nowych języków programowania, które mogą wykorzystać moc obliczeniową sieci 5G i budować aplikacje programistyczne z wykorzystaniem sieci, w tym dla projektów transformacyjnych. Sieć 5G będzie wystarczająco szybka i wydajna, aby wprowadzić na rynek masowy technologie rzeczywistości rozszerzonej, wirtualnej i mieszanej.

Obecnie znanych jest na świecie wiele języków programowania, ponad siedemset, jak podaje Wikipedia. Lata temu języki programowania podzieliły się na deklaratywne (funkcyjne), takie jak Lisp, i imperatywne, czyli np. C. Podczas gdy te ostatnie dominowały przez dziesięciolecia, języki deklaratywne powracają, powiedział w mediach Jared Rosoff, znany programista, który współpracował m.in. z VMware, MongoDB i innymi firmami. „Języki imperatywne były lepiej przystosowane do kodowania logiki biznesowej dla aplikacji”, zauważył Rosoff. „Ale w infrastrukturze jako kodzie świat nie jest imperatywny. Jest sterowany regułami”.

Wciąż używamy starego języka COBOL i innych archaicznych języków programowania, wciąż wymyślamy nowe języki, z których każdy ma swoje wady i zalety. Na przykład korzystamy z Rust i C++ do programowania niskopoziomowego, wrażliwego na wydajność systemów (przy czym Rust dodaje korzyść w postaci bezpieczeństwa). Mamy Python i R do uczenia maszynowego, manipulacji danymi i wielu pokrewnych zadań. Różne narzędzia dla różnych potrzeb. Ale w miarę jak przenosimy się do świata Everything-as-Code (ang. „Wszystko-jako-kod”), coraz częściej pytamy, czy używanie tych samych języków ma sens?

Odpowiedź brzmi „nie” – mówią eksperci. Dlaczego? Ponieważ zbyt często występuje niedopasowanie pomiędzy językiem a celem, który chcemy osiągnąć. Imperatywne języki ogólnego przeznaczenia zostały zaprojektowane do budowania aplikacji i skryptów od podstaw, a nie do definiowania konfiguracji, polityk, itp. Języki deklaratywne, których odmianą jest programowanie funkcyjne, takie jak Polar i HCL, są świetne dla zastosowań takich jak konfiguracja, ponieważ pozwalają po prostu zadeklarować, jak ma wyglądać świat i nie martwić o to, co należy zrobić, aby tak się stało. Minusem jest to, że są to rozwiązania nowe, wymagają pozyskania wiedzy i ludzi, którzy pozyskali tę wiedzę, ekosystemu narzędzi itp. Skoro jest nieco za wcześnie na upowszechnienie języków deklaratywnych, to jest za wcześnie na świat Everything-as-Code, ale, prawdopodobnie nowe nadejdzie szybciej, niż myślimy.

Internet w pierwszym rzędzie – mobilny

Na koniec warto pokusić się jeszcze o mały przegląd nowych technik i tendencji (a nawet, właściwie, mód) w programowaniu i projektowaniu stron i aplikacji internetowych, gdyż obecnie jest to jeden z najbardziej widocznych i dynamicznych obszarów w szeroko pojętej informatyce.

Zacznijmy od czegoś, co jest nazywane progressive web apps (lub PWA), techniki skupionej na doskonałości doświadczeń użytkownika, do czego wykorzystuje się zarówno starsze techniki HTML, CSS i JavaScript jak też nowsze narzędzia, jak React lub Angular.

Rozwój AI i technologii uczenia maszynowego spowodował zapotrzebowanie na projekty wirtualnych asystentów nazywanych także chatbotami, zarówno do komunikacji, jak też obsługi klienta, konsultacji itp.

Głównym celem wprowadzenia Accelerated Mobile Pages (lub AMP) jest przyspieszenie działania strony i zmniejszenie ryzyka jej opuszczenia przez użytkownika. Technologia AMP jest nieco podobna do PWA. Różnica polega na tym, że strony AMP ulegają przyspieszeniu dzięki open source'owemu pluginowi opracowanemu przez Twittera i Google. Z techniką tą wiążą się kontrowersje opisywane już w MT, obawy przed nadmiernym uzależnieniem wydawców stron od Google'a.

Aplikacja jednostronicowa (SPA) to jeden z tych prądów w tworzeniu stron internetowych, który zmierza do zwiększenia wydajności i poziomu ochrony danych. SPA zyskują swoją popularność dzięki rozwojowi frameworków JavaScript. Do tego typu aplikacji należą strony Google, takie jak Gmail, Google Drive czy Google Maps, a także platformy



5. Projektowanie responsywnych stron internetowych

społecznościowe, takie jak Facebook. Są opinie, że w przyszłości większość funkcjonalnych stron internetowych będzie budowana jako SPA. Zapewniają użytkownikom natychmiastową informację zwrotną i zużywają również mniej energii i mogą działać bez kodu po stronie serwera (technologia API).

Przyszłość rozwoju stron internetowych wydaje się także w malejącym stopniu tekstowa, a w rosnącym – obrazkowa i... głosowa. Wiele firm zastanawia się obecnie nad tym, jak zoptymalizować swoje fizyczne i cyfrowe produkty pod kątem wyszukiwania głosowego i komend głosowych. Przewiduje się, że do końca 2022 roku 55% wszystkich gospodarstw domowych na świecie będzie mieć asystenta głosowego. Wzmocniona przez AI, optymalizacja wyszukiwania głosowego oszczędza czas i jest wielozadaniowa.

Choć motion user interface design jest modny od 2018 roku, to dopiero teraz dzięki technologii bibliotek SASS staje się szeroko dostępny dla użytkowników dowolnych urządzeń. Podejście to polega na niestandardowej integracji animacji i stylów zasilanych przez samodzielne biblioteki z licznymi klasami animowanych elementów. Z ich pomocą programiści spędzają mniej czasu na budowaniu produktu cyfrowego i oszczędzają na kosztach. Motion UI to także sprawdzony sposób na przyciągnięcie uwagi użytkowników.

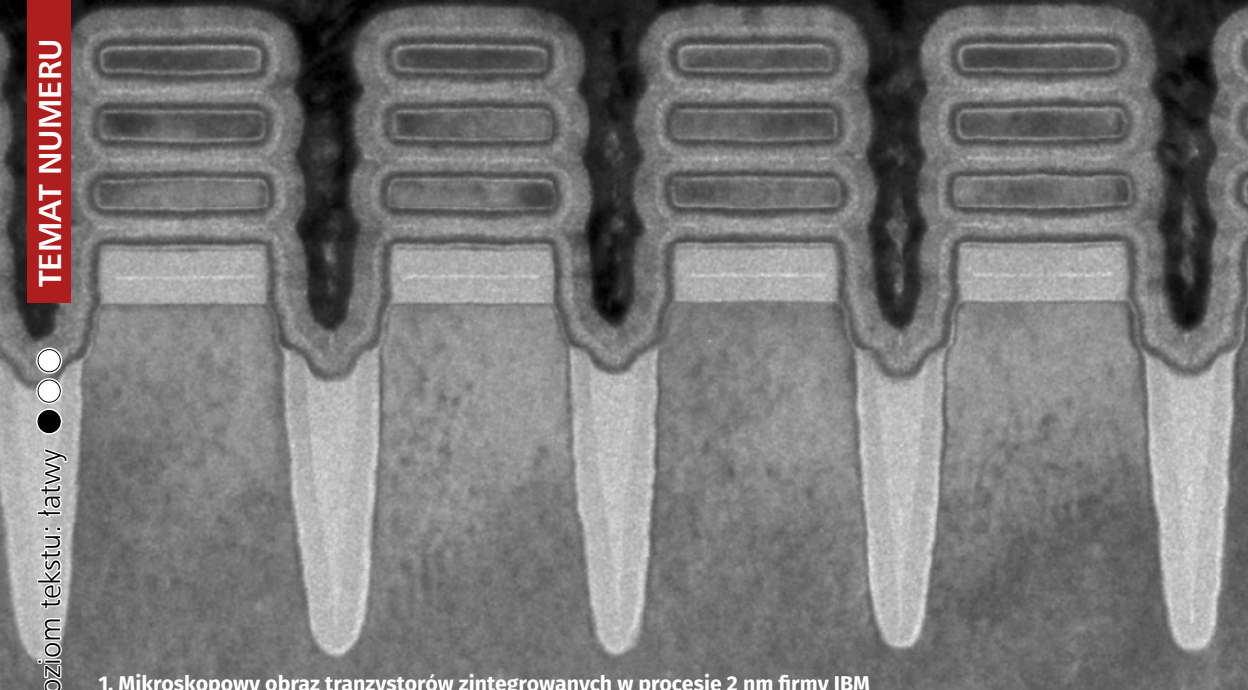
Serverless, inaczej algorytmy bezserwerowe, to koncepcja, w której serwery mogą zostać zastąpione przez chmury, które zarządzają zużyciem zasobów.

Najczęstsze zadania, które mogą być wykonywane bezserwerowo, to pobieranie kopii zapasowych plików, dostarczanie powiadomień i eksport obiektów.

Inna, nie tylko modna, ale mająca praktyczne konsekwencje koncepcja w projektowaniu stron i aplikacji, to „mobile-first development”. Polega ona w skrócie na projektowaniu strony, która nie jest zorientowana na desktop, lecz na tworzenie lekkiego, zazwyczaj minimalistycznego produktu, z którym łatwo jest wejść w interakcję za pomocą telefonu komórkowego. I to, a nie tradycyjnie rozumiany design, jest tu punktem wyjścia. Od 2018 roku Google stosuje oddzielne algorytmy w rankingu wyszukiwania dla telefonów, w których witryny mobile-first są traktowane priorytetowo. Z nurtem tym ściśle powiązana jest koncepcja responsywnych stron internetowych (5). Narodziła się kilka lat temu, kiedy urządzenia mobilne zawładnęły rynkiem. Sprowadza się do zalecenia, by deweloperzy i projektanci dopracowywali swoje produkty tak, by były wygodne dla użytkowników w dwóch formatach, desktopowym i mobilnym. Jedną z możliwości jest stworzenie strony mobile-first i dostosowanie jej do desktopu, wykorzystanie tego samego kodu HTML z CSS, który może automatycznie zmieniać wyświetlanie na stronie.

Podsumowując – nadszedł również czas, aby zapomnieć o rozdzielaniu wersji desktopowych i mobilnych i pracować nad uniwersalnym kodem dla wszystkich typów urządzeń. ■

Mirosław Usidus



1. Mikroskopowy obraz tranzystorów zintegrowanych w procesie 2 nm firmy IBM

W maju 2021 r. świat obiegnęła informacja z cyklu, który dobrze znamy od lat. Kolejny przełom w produkcji układów scalonych, tym razem w technologii 2 nm, w której, jak twierdzi IBM, można upchnąć 50 miliardów tranzystorów w „chipie wielkości paznokcia” (1). To o 20 mld więcej niż w „przełomie 5 nm” z 2017 r. Znow większa moc przy mniejszym zużyciu energii.

Wojna o krzem, którego nagle zabrakło

PROCESORY POŻĄDANIA

IBM dodał, że jego układ testowy może poprawić wydajność o 45 proc. w stosunku do obecnie dostępnych na rynku produktów wykonanych w technologii 7 nm. Zużywa zarazem 75% mniej energii. Firma głosi, że technologia ta może „czterokrotnie” zwiększyć żywotność baterii telefonów komórkowych.

Przemysł chipów komputerowych od lat używa nanometrów, jednej miliardowej części metra, do podawania fizycznych rozmiarów tranzystorów. Jednocześnie

parametr „nm” jest postrzegany jako termin marketingowy, służący do opisanie nowych generacji technologii, prowadzących do lepszej wydajności i niższej mocy. Układy do komputerów stacjonarnych oparte na procesie 7 nm, takie jak czołowe procesory Ryzen firmy AMD, nie były rynkowo dostępne aż do 2019 roku, czyli cztery lata po tym, jak IBM ogłosił, że opanował proces 7 nm. To ogólna, choć nie wiążąca, wskazówka co do tego, kiedy możemy się spodziewać procesorów 2 nm na rynku. Jak wspomniano „proces 5 nm” został ogłoszony jako opanowany przez IBM w 2017. Gdyby przyjąć czteroletnie opóźnienie jako obowiązujący interwał, to powinniśmy spodziewać się chipów tego rodzaju już w tym roku. Na razie jest zapowiedź Intela i tajwańskich zakładów TSMC wypuszczenia urządzeń w litografii 5 nm do końca 2021 r. Być może więc czteroletni cykl zostanie dotrzymany.

Jak zabrakło chipów

Doniesienia o kolejnych postępach miniaturyzacyjnych IBM pojawiają się w sytuacji światowego niedoboru chipów komputerowych i zapowiedzi



2. Niedokończone auta Volkswagena zaparkowane w zakładach w Pampelunie

radikalnych zmian w globalnym łańcuchu produkcji i dostaw chipów, sprowadzających się w skrócie do tego, by mniej uzależniać USA i inne kraje od wytwórci w Chinach i na Tajwanie. Producenci samochodów zostali zmuszeni do zawieszenia produkcji z powodu braku komponentów komputerowych, producenci smartfonów ostrzegli, że premiery produktów mogą zostać przesunięte, a zaawansowane części komputerów, takie jak karty graficzne, są trudne do znalezienia lub mają bardzo wysokie ceny.

Tymczasowe zamknięcia fabryk z powodu pandemii również spowodowały przerwy w dostawach. Gdy fabryki zostały ponownie otwarte, producenci dóbr elektronicznych składali wielkie zamówienia „nadrabiające zaległości”, tworząc coraz większe luki w dostawach chipów. Pandemia nie była jedynym czynnikiem wpływającym na tę sytuację. Z powodu zjawisk atmosferycznych w kilku zakładach w Teksasie wstrzymano na pewien czas produkcję, a w marcu 2020 r. pożar strawił jedną z japońskich fabryk. W sierpniu 2020 roku, w związku z zarzutami o szpiegostwo, Stany Zjednoczone zakazały zagranicznym firmom, których chipy wykorzystują amerykańską technologię, sprzedaży komponentów chińskiemu gigantowi technologicznemu Huawei. Chiński potentat rozpoczął gromadzenie zapasów komponentów półprzewodnikowych jeszcze przed wejściem w życie sankcji, a inne firmy poszły w jego ślady, co jeszcze bardziej ograniczyło dostawy.

Ponieważ producenci samochodów zmniejszyli produkcję na początku pandemii, ich dostawcy

układów scalonych zwrócili się ze swoją ofertą do klientów z innych sektorów, mianowicie do producentów artykułów elektronicznych, na które wzrósł duży popyt z powodu pandemii. Po pewnym czasie wielkie marki samochodowe, od Volkswagena po Volvo, gdy sprzedaż aut zaczęła znów wzrastać, zaczęły mieć problemy z brakami produktów półprzewodnikowych. Tysiące niedokończonych samochodów zaparkowanych w fabryce Volkswagena w Pampelunie, w Hiszpanii, to symboliczna ilustracja tej sytuacji (2). Jean-Marc Chery, dyrektor generalny francusko-włoskiego producenta chipów STMicroelectronics, powiedział, że zamówienia na przyszły rok już przekroczyły możliwości produkcyjne jego firmy. SEB, francuski producent sprzętu AGD ostrzegł, że jest zmuszony do podniesienia cen. Producenci smartfonów byli w tym okresie względnie bezpieczni, ponieważ mieli zapasy chipów, ale z czasem i oni zaczęli odczuwać niedobory. Szef Apple Tim Cook powiedział, że braki te mogą uderzyć w produkcję iPhone'ów i iPadów. Mocniej ucierpieli mniejsi producenci smartfonów. Oczywiście uderzyło to także w producentów sprzętu do gier, np. konsoli, takich jak PlayStation 5 i Xbox Series X.

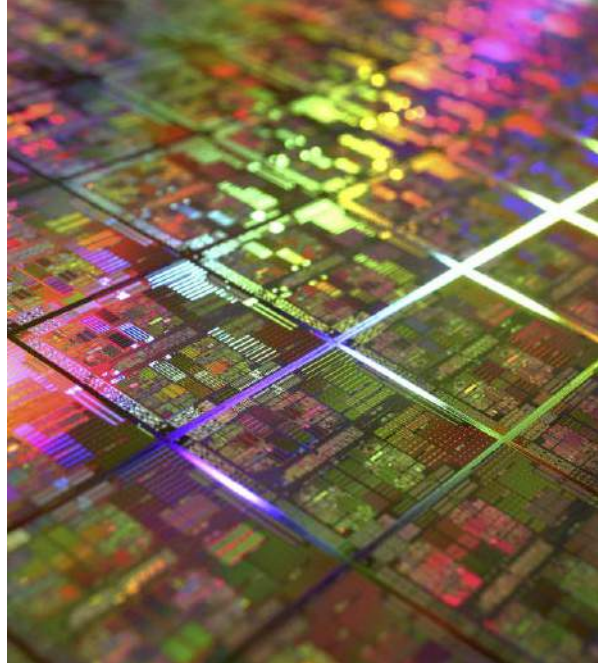
Problemy te dla wielu ludzi w krajach technologicznie rozwiniętych były jasnym sygnałem, że w trudnych sytuacjach globalizacja i przesunięcie zakładów wytwarzających komponenty gdzieś daleko do Azji może oznaczać kłopoty. Krzemowe produkty są obecnie zbyt ważne dla gospodarki, by zdawać się na niepewne dalekie źródła i długie

łańcuchy dostaw. Szef Intel'a Pat Gelsinger wprost ocenił w rozmowie z BBC, że ulokowanie 80% światowych źródeł dostaw chipów w Azji nie jest dobrym pomysłem. Kto może więc, czyli ma know-how, zdolności produkcyjne, zaplecze, no i oczywiście odpowiednie możliwości finansowe, podejmuje działania, by zapewnić produkcję krzemowych komponentów „u siebie”.

Wiosną 2021 r. Korea Południowa ogłosiła podjęcie gigantycznych inwestycji o wartości 451 miliardów dolarów, które mają docelowo przekształcić ten kraj w „giganta półprzewodników”. W tym samym czasie amerykański Senat uchwalił wydatki w wysokości 52 miliardów dolarów na dotacje dla rozwoju wytwórni układów scalonych, znanych w USA jako „fabs”. Prezydent Biden podpisał rozporządzenie wykonawcze inicjujące strategię narodową, której celem jest zapobieżenie niedoborom półprzewodników w przyszłości. Działania amerykańskich władz powiązane są z planami Intel'a, który chce wydać 20 mld dolarów na nowe fabryki i rozbudowę istniejących na terenie USA. Unia Europejska zapowiada, że dąży do podwojenia swojego udziału w globalnej produkcji układów scalonych do 20% globalnego rynku do roku 2030.

Oczywiście zakłady nie mogą zostać otwarte z dnia na dzień, szczególnie te, które produkują w procesach półprzewodnikowych opartych na trawieniu i obróbce wafla krzemowych (3), co jest znane jako jedno z najbardziej wymagających technik wytwórczych, jakie zna przemysł. „Budowa nowych mocy produkcyjnych wymaga czasu, w przypadku nowej fabryki krzemowej ponad dwóch i pół roku, więc większość przedsięwzięć, które zaczynają się teraz, nie zwiększy dostępnych mocy produkcyjnych do 2023 roku”, wyjaśniał w wypowiedzi dla prasy Ondrej Burkacky z firmy konsultingowej McKinsey.

4. Konkurujące na rynku globalnym procesory

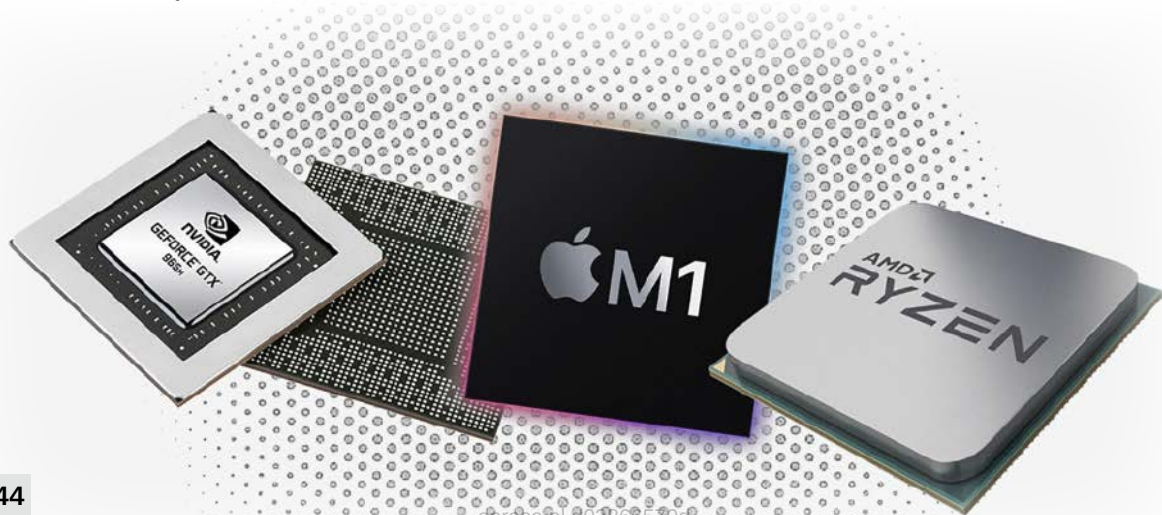


3. Krzemowe wafle

Nowi chętni do robienia procesorów

Do zwiększenia tumultu na rynku układów scalonych przyczynia się zaostrzająca się rywalizacja firm projektujących i produkujących procesory (4), takich jak AMD, Intel, Nvidia, do których chcą dołączyć, jako równorzędni gracze na rynku półprzewodnikowym, firmy znane dotychczas raczej z produkcji produktów finalnych i korzystające z gotowych komponentów, przede wszystkim Apple, za nim również Google i Amazon. Ten ostatni zresztą już wytwarza własne chipy.

Choć Intel przez dekady był niekwestionowanym liderem rynku, procesory Ryzen firmy AMD (Advanced Micro Devices), które od ok. 2018 roku zaczęły wyprzedzać pod względem parametrów technicznych produkty Intel'a, oraz ambicje firmy



Nvidia, która była wcześniej znana głównie z tworzenia GPU, a ostatnio nabyła producenta chipów komputerowych ARM Holdings od Softbanku za 40 miliardów dolarów, sprawiły, że o Intelu, mającym ogromne problemy z opanowaniem procesu produkcji procesorów 7 nm, zaczęto mówić z czymś w rodzaju politowania.

W listopadzie 2020 firma Apple zaproponowała w swoich niektórych urządzeniach M1 autorski procesor, który łączy w jednym układzie wiele funkcji wykonywanych wcześniej przez różne części komputera, takie jak procesor graficzny. „Do tej pory Mac potrzebował wielu różnych układów scalonych, aby zrealizować wszystkie swoje funkcje, w tym – procesora, układów wejścia/wyjścia, zabezpieczenia i pamięci”, czytamy w komunikacie wprowadzającym M1. „W M1 technologie te zostały połączone w jeden system na chipie (SoC), zapewniając nowy poziom integracji w celu uzyskania większej prostoty i niesamowitej wydajności”. W drugiej połowie 2021 roku Apple przygotowuje się do wprowadzenia kolejnych procesorów, nazywanych w przeciekach medialnych M1X a także M2. Być może będą to dwa odrębne produkty.

Amazon przez lata stopniowo stał się jednym z największych na świecie nabywców chipów komputerowych, które zasilają jego centra danych i usługi serwerowe AWS. Ostatnio firma coraz bardziej koncentruje się na projektowaniu własnych chipów. Amazon zaczął sygnalizować swoje zamiary już w 2015 roku, kiedy nabył Annapurna Labs, małego izraelskiego projektanta układów scalonych.

W mediach można znaleźć opinie klientów Amazona zadowolonych z korzystania z chipów własnych tej firmy. Smugmug, internetowy serwis fotograficzny, który używa AWS, twierdzi, że obniżył swoje koszty aż o 40% tylko dzięki przejściu na usługę AWS, opartą na własnym procesorze Amazona, nazwanym Graviton, opartym na architekturze ARM, podobnie jak M1 Apple'a. AWS w większości nadal opiera się na produktach Intela, ale Smugmug płaci mniej dzięki temu właśnie, że korzysta z usług opartych na własnym sprzęcie Amazona. „W rezultacie mamy znacznie niższy rachunek bez konieczności podejmowania jakichkolwiek działań”, zachwalał Don MacAskill, dyrektor generalny Smugmug.

Nad własnymi wyspecjalizowanymi chipami pracują również Microsoft Corp. i Google Alphabet Inc., skupiając się na mniejszym zużyciu energii w swoich centrach danych. Według Applied Materials, największego producenta sprzętu do produkcji chipów, do 2025 roku centra danych będą zużywać 15%

światowej energii elektrycznej, w porównaniu z około 2 proc. w 2020 roku. Utrzymanie niskiego zużycia energii staje się dla właścicieli centrów danych ważniejsze niż koszt sprzętu.

Technologia leżąca u podstaw niskoenergetycznych chipów do smartfonów jest tworzona przez ARM Ltd., brytyjską firmę półprzewodnikową, która udziela licencji, ale sama nie produkuje chipów. Amazon i Microsoft używają ARM jako podstawy do swoich wewnętrznych projektów chipów. Nvidia, która przejmie ARM, obiecuje utrzymać otwarty dostęp do technologii ARM. Jednak niektórzy z klientów ARM już zgłosili zastrzeżenia do regulatorów rozważających zatwierdzających tę transakcję. Głośno krytykuje to przejście na przykład Elon Musk. Agencja Bloomberg News donosiła, że wśród skarżących są m.in. Google, Microsoft i Qualcomm.

Koniec „prawa Moore’a” czy nieprzerwany postęp do 2050 roku?

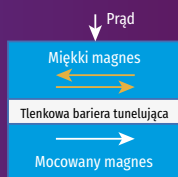
Osiemdziesiąt lat temu Alan Turing stworzył matematyczne podstawy obliczeń komputerowych. Dekadę później John von Neumann stworzył praktyczną architekturę komputera. Od tego czasu postępy w technologii informacyjnej zyskały fundamentalne znaczenie dla naszego życia i światowej gospodarki. Udało nam się przez parę dekad utrzymywać wykładniczy wzrost wydajności przy wykładniczo malejących kosztach. Opisuje to prawo, a raczej prognoza Gordona Moore'a z 1971 r., głosząca, że liczba tranzystorów, które możemy umieścić w mikroprocesorze, będzie się podwajać co 18–24 miesiące. I po latach doszliśmy do tego, że układ pamięci we współczesnym telefonie komórkowym ma milion razy większą pojemność niż komputer, który John von Neumann zbudował dla Institute for Advanced Study ok. 1951 r.

Choć od ponad dekady wykładniczy wzrost wydajności techniki procesorowej ulega spowolnieniu (przynajmniej w klasycznym rozumieniu, ale do tego wrócimy), wykładniczy wzrost popytu utrzymuje się na niezmiennym poziomie. Ilość informacji, które generujemy jako gatunek, podwaja się co dwa lata. Zdecydowana większość tych informacji jest obecnie zapisywana w centrach danych i to ich wymogi, przede wszystkim wyzwania energetyczne, dyktują dziś postęp techniczny w krzemowej branży.

Patrick Moorhead, uznawany za czołowego analityka przemysłu półprzewodnikowego, powiedział kiedyś, że prawo Moore'a, według najściślejszej definicji mówiącej o podwajaniu gęstości chipów co dwa lata, już nie obowiązuje. Zarówno on jednak, jak i wielu ekspertów zwracają uwagę na sformułowanie

NOWE TYPY PAMIĘCI DLA UKŁADÓW INTEGRUJĄCYCH PRZETWARZANIE I PAMIĘĆ

Nieulotna pamięć o dostępie swobodnym, bez wymazywania przed zapisem, integracja on-chip



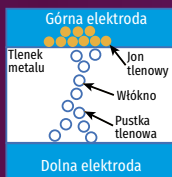
STT-MRAM

Magnetyczna pamięć o dostępie swobodnym z przeniesieniem momentu obrotowego



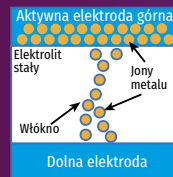
PCM

Pamięć zmiennofazowa



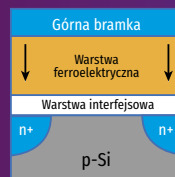
RRAM

Pamięć o dostępie swobodnym z przełączaniem rezystancyjnym



CBRAM

Pamięć losowego dostępu z mostkiem przewodzącym



FERAM

Ferroelektryczna pamięć o dostępie swobodnym

5. Integracja układów przetwarzania i pamięci

„według najściślejszej definicji”. W rzeczywistości producenci, od Intela po Apple, od wielu lat osiągają wzrosty wydajności procesorów innymi sposobami niż zagęszczanie liczby tranzystorów na chipie. Trudno zliczyć wszystkie innowacyjne techniki, swoiste „wybiegi” w architekturze i w połączeniach jednostek przetwarzających, stosowane przez producentów. Zatrzymajmy się na chwilę przy nowych rozwiązaniach proponowanych, co samo w sobie jest interesujące, przez największego producenta w branży krzemowej, tajwańską firmę TSMC, od której fabryk, jak wspominaliśmy, wiele państw chce się „uwolnić”.

Jeszcze przed pandemią, w sierpniu 2019 roku Philip Wong, szef działu badań w Taiwan Semiconductor Manufacturing Corp., na konferencji Hot Chips na Uniwersytecie Stanforda mówił o najnowszych osiągnięciach w rozwoju techniki układów scalonych, wśród nich m.in. umieszczaniu pamięci bezpośrednio na procesorze (5), które znacznie zwiększą wydajność i spowodują, że „prawo Moore’a” będzie wciąż żywe”, choć zapewne już nie „według najściślejszej definicji”. Takie połączenie radykalnie zwiększy szybkość i wydajność, ponieważ układy logiczne na chipie będą szybciej otrzymywać potrzebne dane, bez niepotrzebnych „przestojów”. Wong na slajdach swojej prezentacji, jak oceniano, „w sposób bezczelny prognozował nieprzerwany postęp w chipach aż do 2050 roku”, podkreślając, że „postęp nie zatrzyma się w miejscu, gdy elektronika osiągnie fundamentalne granice wielkości atomów”.

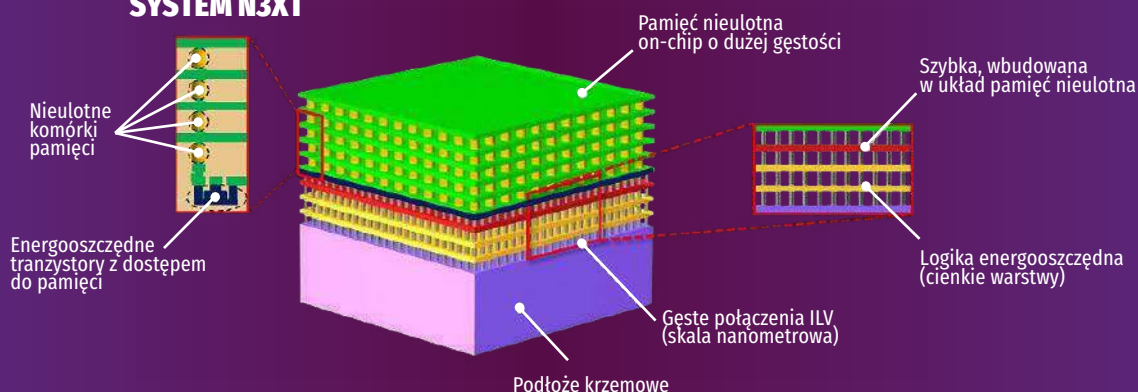
Przedstawiciel tajwańskiej firmy mówił o technikach, nad którymi prace badawcze trwają od lat, w tym o zastosowaniu w chipach nanorurek węglowych, że obecnie istnieją już praktyczne możliwości ich wdrożenia. Innowacja polegająca

na zastosowaniu dwuwymiarowych materiałów warstwowych, pozwalająca elektronom na łatwiejszy przepływ przez układy scalone, również, jak twierdził Wong, jest już produkcyjnie dostępna. Kolejna nowa technika układania w architekturze 3D (6) oznaczać ma, że funkcje procesora komputerowego, które obecnie są odizolowane, mogą być umieszczone w wielu warstwach, połączonych szybkimi ścieżkami danych, zwanymi przelotkami krzemowymi. „W tego typu systemach, w których wiele warstw logiki i pamięci jest zintegrowanych w drobnoziarnisty sposób, łączność jest kluczowa”, mówił Wong.

Przygotowania do walnej bitwy

Przez ok. dwie i pół dekady branża procesorów komputerowych była zdominowana przez jedną architekturę – x86 – utożsamianą z imperium Intela. Lata 90. to powstanie kilku architektur technicznie konkurujących z Intelem, ale ostatecznie pod koniec dekady AMD było osamotnione w walce z ówczesnym dominatorem. W połowie pierwszej dekady XXI wieku IBM zrezygnował z koncepcji G5. Wydawało się, że Intel ostatecznie wygrał. Po latach wyglądało to bardziej jak chwilowa przewaga niż trwałe zwycięstwo. Rynek procesorów rozgrzewa się teraz znów do temperatury, której nie zazналиśmy od dawna. Znowu pojawiają się chipowe startupy z buńczucznymi zapowiedziami, np. firma SiFive, która twierdzi, że zbuduje układy RISC-V do komputerów stacjonarnych. Inny startup, Nuvia, przejęty niedawno przez Qualcomm, ma zaoferować zupełnie nowe, niestandardowe układy krzemowe do 2023 roku. Wielcy też czują powiewy wiatru zmian – zarówno AMD, jak i Intel zadebiutowały ostatnio z nowymi architekturami. AMD rozszerzyło swój projekt Zen 3

SYSTEM N3XT



Źródło: W. Hwang, W. Wan, Y. Malviya, H. Li, M. Lee, M. Aly, H.-S. P. Wong, S. Mitra. Work in progress 2017 – 2019 w/ TSMC

6. Integracja przetwarzania i pamięci w procesorze

na servery i urządzenia przenośne, a Intel wprowadził nowe układy Rocket Lake i Ice Lake Server.

Premiera Ryzena w kwietniu 2017 roku była początkiem nowej ery w skalowaniu wydajności procesorów. Architektury Zen, Zen+, Zen 2 a potem Zen 3 firmy AMD wielokrotnie wybiły się ponad swoją kategorię wagową. Do gry weszło Apple z własnym M1. Choć procesor ten nie jest nokautem dla AMD i Intela, to wyznacza nową poprzeczkę w wydajności i efektywności energetycznej, z którą obie firmy bazujące w swoich głównych produktach na x86 będą musiały się zmierzyć w dłuższej perspektywie. Apple nie zamierza jednak przejąć masowego rynku procesorów. Przynajmniej tak się wydaje. Inaczej Qualcomm, który ogłosił, że zamierza wprowadzić na rynek laptopy oparte na nowej konstrukcji procesora Nuvia. Intel i AMD mają jeszcze trochę swobody we wprowadzaniu nowego sprzętu na rynek, ale długo odkładana walna bitwa pomiędzy x86 i ARM zbliża się wielkimi krokami.

Pod kierownictwem Pata Gelsingera, nowego dyrektora generalnego Intela, firma zwraca się w kierunku budowania nowych jednostek CPU i GPU. Jednak nawet w najlepszym scenariuszu dla Intela, odzyskanie dominującej pozycji zajmie firmie wiele lat. Gelsinger obiecał, że Alder Lake pojawi się na desktopach pod koniec 2021 roku, zastępując Rocket Lake. Procesory te zbudowane są w technologii 10 nm. Firma wcześniej potwierdziła wprowadzenie rodziny mikroprocesorów 7 nm, obecnie nazywanej Intel 4, o nazwie Meteor Lake, która ma się ukazać w 2023 r.

Jeśli Intel nie będzie w stanie dostatecznie szybko wprowadzić konkurencyjnych konstrukcji, to w pewnym sensie będzie potrzebował produktów AMD, by wykazać, że x86 może konkurować z architekturą ARM. Jeśli AMD nie sprostą zadaniu, obaj wielcy producenci x86 mogą zostać zepchnięci na margines. AMD może być teraz gwiazdą rynku procesorów, ale potrzebuje wolumenu sprzedaży i ogromnej bazy instalacyjnej Intela, aby

Efekt pandy

Marta Kisiel

Wydawnictwo W.A.B., liczba stron: 352, cena: 41,99 zł

SPA i smog, romanse i morderstwa, bzdziągwa w futrze oraz zaginiony trup. Za namową męża Tereska Trawna raz jeszcze daje się skusić wizją świętego spokoju z widoczkami. Tym sposobem cztery kobiety i psica wyruszają na babski wypad do SPA. Niestety, zaplanowane wellness i błogi relaks w górskich okolicznościach przyrody już pierwszego dnia idą w niepamięć, gdy na horyzoncie pojawia się przyczajony zbrodnień. Cóż. Niektórzy nigdy nie uczą się na błędach.





7. Historia chipów ARM

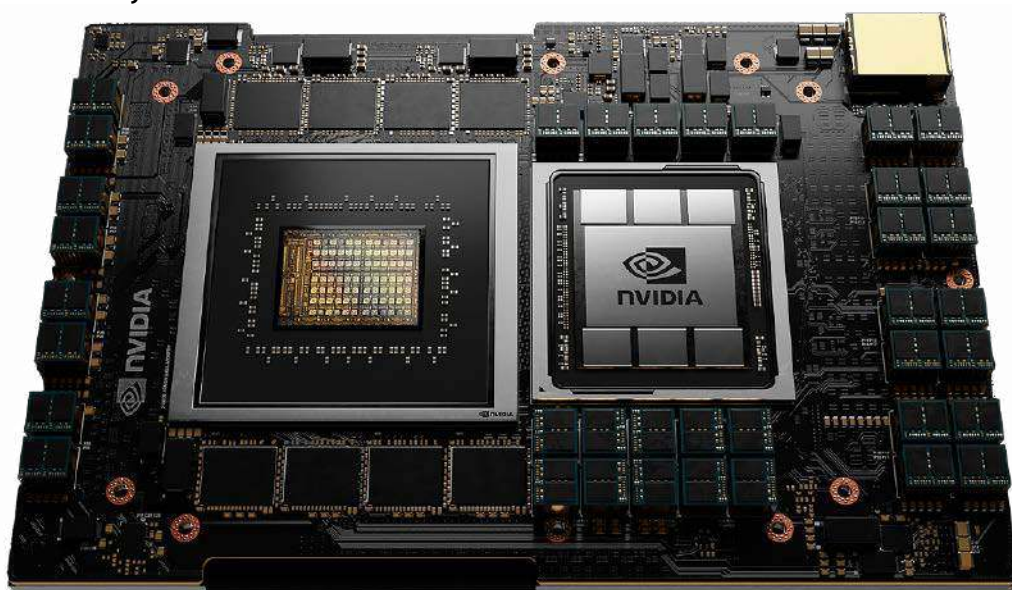
pozycjonować swoje procesory jako atrakcyjne alternatywy wobec rozwiązań ARM przedstawianych jako wciąż niepewna alternatywa. Zen 3, który wszedł w 2020 r., nie zwiększył drastycznie liczby rdzeni ani częstotliwości taktowania, jednak ma znacznie lepsze parametry energetyczne niż poprzednie wersje chipów AMD. Zen 4, który według przewidywań ma mieć przeprojektowaną płytę główną, nowy proces produkcyjny i obsługę szybkich pamięci DDR5, może pojawić się w 2022 roku.

Wstępne testy nowych macbooków potwierdzają, że oparte na architekturze ARM M1 (7) oferują dużą przewagę wydajnościową nad tradycyjnymi

laptopami z procesorami Intel'a i niemal dwukrotnie dłuższy czas pracy na baterii. Projektując swój pierwszy procesor dla pecetów, firma Apple zdobyła przewagę, ale producenci komputerów PC z systemem Windows nie chcą pozostać w tyle. Dell, HP i inni producenci komputerów coraz głośniejsze domagają się od Intel'a lepszych procesorów, które będą w stanie konkurować z M1. Jeśli dawny imperator krzemowy odpowiednio szybko tego nie zrobi, to będą szukać alternatywnych dostawców.

Kluczową technologią, na której opiera się nowy układ Apple, jest zestaw instrukcji ARM. Apple korzysta zarazem z doświadczeń dekady projektowania

8. CPU Grace firmy NVIDIA



niestandardowych procesorów do iPhone'ów. Procesory te muszą być bardzo energooszczędne, aby nie „zjadać” zbyt szybko baterii smartfona. Przeniesienie tej wydajności do laptopów umożliwia duży wzrost wydajności i wydłużenie czasu pracy na ładowanie. Procesory Intel dla komputerów tej klasy po prostu nie są tak wydajne. Gdy są umieszczone w cienkim i lekkim laptopie, muszą obniżyć szybkość pracy, aby uniknąć przegrzewania. W rezultacie, w porównaniu z MacBookiem z procesorem Intela wydanym na początku tego roku, energooszczędny procesor M1 zapewnia o 40 proc. lepsze wyniki w popularnym benchmarku CPU i o 70 proc. dłuższy czas pracy na baterii podczas przeglądania stron internetowych.

Przejście na platformę ARM staje się szerzej widoczną tendencją. Microsoft np. przeprojektował system Windows na platformę ARM, umożliwiając firmom takim jak Acer, HP i Lenovo wypuszczenie na rynek „komputerów hybrydowych” (małych komputerów typu PC-tablet) opartych na procesorach o architekturze ARM firmy Qualcomm. ARM zaczyna również pojawiać się w centrach danych. Amazon twierdzi, że Graviton oferuje o 20 proc. lepszą wydajność niż porównywalny procesor Xeon przy 20% niższych kosztach.

Zmiana z zestawu instrukcji x86 Intela na ARM stwarza jednak problemy z oprogramowaniem. Podobnie jak Microsoft, Apple musiało „przełączyć” swój system operacyjny MacOS, aby działał na ARM, wraz ze wszystkimi aplikacjami firmowymi. Natywne aplikacje w pełni wykorzystują wyższą wydajność M1 i czas pracy na baterii. Wiele aplikacji firm trzecich jest jednak zaprojektowanych wyłącznie do procesorów Intela. Aby zapewnić kompatybilność z tymi aplikacjami, Apple udostępnia emulator (o nazwie Rosetta 2), który tłumaczy z x86 na ARM, ale aplikacje działają z szybkością przeciętnie o połowę mniej niż aplikacje natywne.

Apple nie sprzedaje swoich procesorów innym producentom komputerów. Aby uzyskać podobne procesory oparte na architekturze ARM, firmy te muszą zwrócić się do dostawców układów scalonych, takich jak Nvidia i Qualcomm. Nvidia opracowała już wysokowydajny procesor kompatybilny z ARM (znany jako Carmel), który mogłaby połączyć ze swoją technologią graficzną, tworząc procesory dla komputerów PC podobne do procesorów Apple. Qualcomm, jak wspomniano, dostarcza procesory oparte na ARM dla małych pecetów, ale nie spełniają one na razie bardziej rygorystycznych wymagań laptopów.

NVIDIA wprowadza na rynek procesor Grace (8), który zwiększy wydajność procesorów ARM w ogromnych obciążeniach związanych ze sztuczną inteligencją, w superkomputerach AI wykorzystujących ogromne zbiory danych. Procesor NVIDIA Grace został zaprezentowany na konferencji NVIDIA GTC 2021 i od razu zrobił furorę w świecie obliczeń wymagających wysokiej wydajności (HPC). Grace będzie wykorzystywany w potężnych nowych superkomputerach opracowywanych przez Narodowe Laboratorium Los Alamos i Szwajcarskie Narodowe Centrum Superkomputerowe. NVIDIA twierdzi, że system oparty na Grace będzie w stanie wytrenować model przetwarzania języka naturalnego (NLP) o bilionie parametrów dziesięć razy szybciej niż współczesne systemy oparte na NVIDIA DGX, które działają na procesorach x86.

Jak widać, nowa wielka wojna krzemowa zaostrza się. Dobra wiadomość dla zwykłych użytkowników komputerów jest taka, że powinna ona doprowadzić do tego, że laptopy odnotują kolejny duży wzrost szybkości działania, wydajności i żywotności baterii. I to wszystko pomimo faktu, że „prawo Moore'a już nie obowiązuje”. ■

Miroslaw Usidus

Wina wina **Małgorzata Starosta**

Wydawnictwo W.A.B., liczba stron: 352, cena: 39,99 zł

Kiedy Agata Śródka, „panna z odzysku” – restauratorka i celebrytka, postanawia założyć winnicę w gościnieckim pałacu, nie spodziewa się, że ta przygoda zaprowadzi ją w sam środek morderczej intrygi, której ofiarami padną jej były mąż, młoda rywalka i leciwa dama. Niewiele zabraknie, żeby i ona sama straciła życie. Skrywane przez lata sekrety wylewają się wraz ze strumieniami znakomitego wina, nie każdy okazuje się tym, za kogo się podaje, a prowadzący śledztwo gubią się w gąszczu tajemnic i kłamstw. Legenda cysterskiego zakonu odżywa i niespodziewanie ujawnia prawdziwe pragnienie winiarza, pomagając w zdemaskowaniu mordercy, choć wszystko to i tak wina wina.





Co w fizyce wisi w powietrzu?

RAPORT

Nauka, która poszła wieloma ścieżkami... oby nie w las

Zwykle o tej porze roku, gdy emocjonujemy się decyzjami Komitetu Noblowskiego, myśli nasze zmierzają w kierunku stanu współczesnej fizyki. MT od wielu lat opisuje nowe odkrycia, które, jak mają nadzieję fizycy, zrewolucjonizują znane i przyjęte powszechnie modele, z drugiej strony nie ukrywamy wrażenia stagnacji, gdy ustalony stan wiedzy wydaje się ani drgnąć.

Pisaliśmy niejednemu raz, że zarówno Model Standardowy fizyki cząstek elementarnych (1), jak też teorie Einsteina raz po raz zyskują nowe potwierdzenia. Nie ukrywaliśmy, że kolejne odkrycia nie satysfakcjonują fizyków, bo wciąż nie mogą połączyć dwu fundamentalnych teorii – mechaniki kwantowej i względności, w jeden zuniifikowany model. Przecież na poziomie cząstek i we Wszechświecie wielkich obiektów kosmicznych ostatecznie powinna obowiązywać ta sama fizyka, sądzą uczeni. To sprawia, że wielu z nich zaczęło marzyć o czymś, co podważy z jednej strony Model Standardowy, a z drugiej – teorie względności opisujące grawitację.

Nic takiego, w sensie dobrze potwierdzonego odkrycia i alternatywnego opisu naukowego, jak dotąd nie znaleźliśmy. Różnego rodzaju sensacyjne doniesienia o odkryciach „otwierających drzwi do nowej fizyki” są często przesadzone, przedwczesne, a jeśli coś jest na rzeczy, to trzeba czekać na potwierdzenia, dodatkowe dane, kolejne eksperymenty. By podważyć teorie oparte na dekadach badań, potwierdzających wyniki eksperymentów, potrzeba co najmniej równie bogatego materiału dowodowego i solidnej pracy teoretycznej.

Współczesna fizyka prowadzi obecnie jednak cały szereg badań, które, jak to się czasem mówi, „mają potencjał”. Chodzi o tkwiącą w nich potencjalnie moc zrewolucjonizowania obowiązujących poglądów i nawet całych teorii fizycznych. Przyjrzyjmy się więc tym wiszącym w powietrzu sensacjom naukowym, co do których nie ma pewności, czy spadną

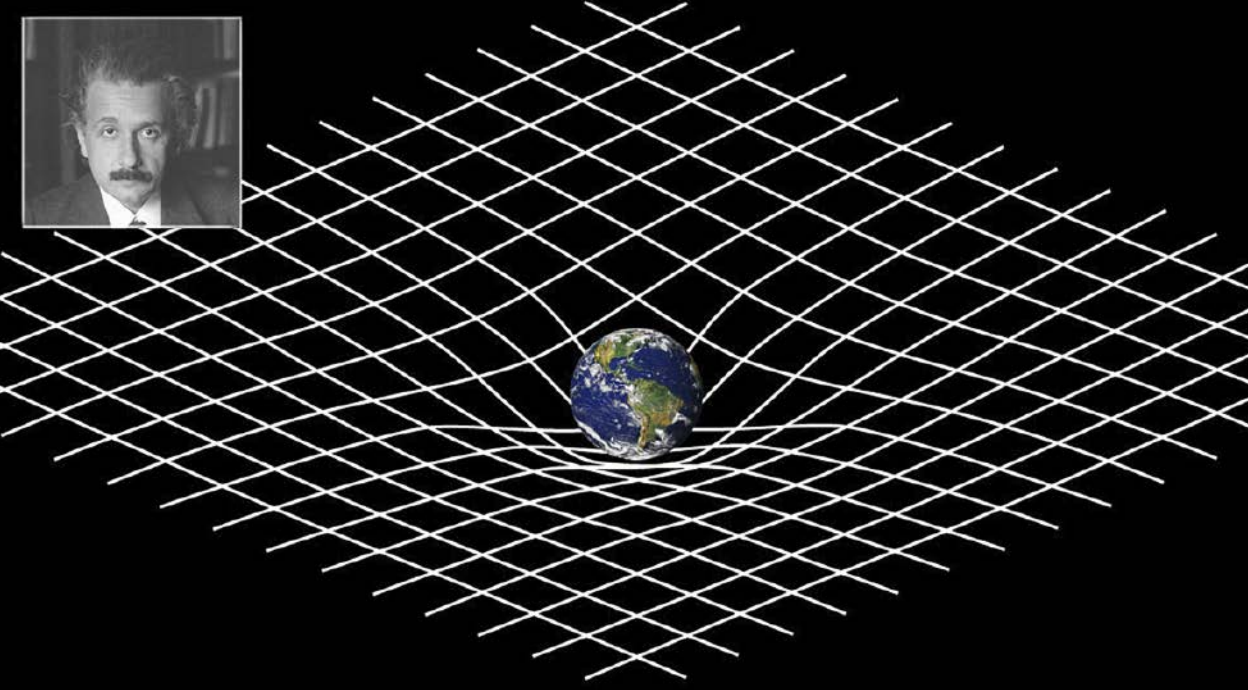
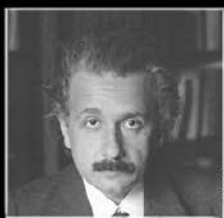
na nas z hukiem, który zburzy wszystko, co wiemy, czy może jednak rozplyną się w owym powietrzu, w ciszy zapomnienia, jak to z wieloma potencjalnie rewolucyjnymi teoriami bywa.

W poszukiwaniu „piątej siły”

Według naszej najlepszej (a może lepiej powiedzieć – standardowej) wiedzy cała materia we wszechświecie składa się z dwudziestu pięciu cząstek elementarnych (dwanaście kwarków i leptonów, ich antyodpowiedniki i bozon Higgsa). Fizycy gromadzą je w Modelu Standardowym, który można uznać za coś w rodzaju układu okresowego dla cząstek subatomowych. Cząstki modelu wchodzą w interakcje za

1. Model Standardowy fizyki cząstek elementarnych





2. Albert Einstein i zakrzywienie czasoprzestrzeni

pomocą czterech rodzajów oddziaływań (sił). Są to: grawitacja, elektromagnetyzm, silne oddziaływanie jądrowe, które odpowiada za spójność jądra atomowego, przeciwstawiając się sile elektromagnetycznej, oraz słabe oddziaływanie jądrowe, które wiąże się z rozpadem jądra atomowego. Wszystkie inne oddziaływania, które znamy, na przykład siła vander Waalsa, która utrzymuje atomy w cząsteczkach, siły tarcia, siły mięśni, to wszystko są siły emergentne, czyli wywodzą się z tych czterech fundamentalnych oddziaływań.

Einstein miałby do tego „układu sił” taką uwagę, że grawitacja nie jest siłą, lecz efektem zakrzywienia czasoprzestrzeni (2). Trzy pozostałe, po odliczeniu grawitacji, oddziaływania są do siebie podobne w tym sensie, że wiemy, iż pośredniczą w nich pewne rodzaje cząstek. Oznacza to, że jeśli istnieje oddziaływanie pomiędzy dwiema cząstkami, jak na przykład dodatnio naładowanym protonem i ujemnie naładowany elektronem, to można je rozumieć jako wymianę cząstki pomiędzy nimi. W przypadku elektromagnetyzmu tą cząstką jest foton, kwant światła. W przypadku silnych oddziaływań jądrowych są to gluony „sklejające” ze sobą kwarki, a w oddziaływaniach słabym są to bozony Z i W. Wierzymy, że również grawitacja polega na wymianie cząstek, które nazywamy „grawitonami”, ale nigdy takich nie wykryto. Jest też cząstka typu bozon, a konkretnie bozon Higgsa, która nie przenosi oddziaływań podstawowych, ale nadaje masę innym cząstkom. Są fizycy, którzy uważają, że jednak przenosi lub sam jest oddziaływaniem. Sprawa nie jest

wyjaśniona, gdyż oddziaływania podstawowe wynikają z symetrii, nawet grawitacja (choć wspomniany Einstein nie nazwałby grawitacji – „siłą”), zaś bozon Higgsa nie opiera się na symetrii, a mimo to ma cechy oddziaływania. Kto mówił, że wszystko we współczesnej fizyce jest jasne?

Wielu fizyków chętnie mówi o kolejnym, piątym typie oddziaływania, bo ich zdaniem, pozwoliliby wyjaśnić to, co nie jest w tej chwili jasne. Choć na istnienie „piątej siły natury” nie ma dowodów, często gości w nagłówkach brzmiących sensacyjnie doniesień. Z pewnością nie jest ona łatwa do wykrycia i zaobserwowania, bo gdyby tak było, już dawno przecież byśmy ją wykryli. Oznacza to, że siła ta albo staje się zauważalna dopiero przy bardzo dużych odległościach, więc można ją zaobserwować w kosmologii lub astrofizyce, albo staje się zauważalna przy bardzo małych odległościach i jest ukryta gdzieś w głębinach fizyki cząstek elementarnych.

Na przykład, anomalia dotycząca mionu $g-2$, o której rok temu pisaliśmy w MT, może być oznaką nowego nośnika oddziaływań, więc może to być owa „piąta siła”. A może nie. Istnieje też przypuszczalna anomalia w niektórych procesach jądrowych, w której mogłaby pośredniczyć nowa cząstka, zwana X_{17} , której istnienie zasygnalizował pięć lat temu węgierski uczony Attila Krasznahorkay, a MT również o tym donosił. Żadna z tych anomalii nie ma z punktu widzenia fizyków wielkiej siły przekonywania, a z pewnością nie ma żadnego silnego dowodu na „piątą siłę”.

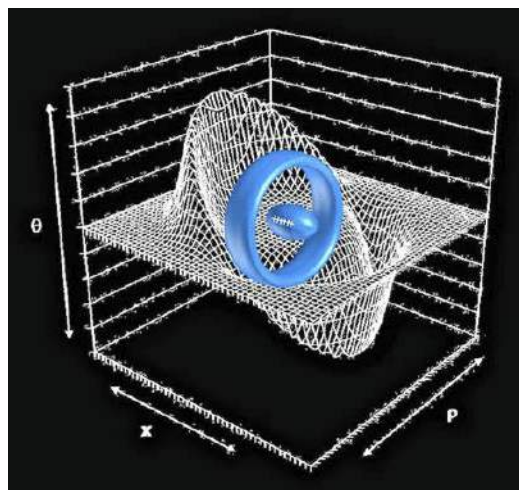
„Najbardziej wiarygodny argument przemawiający za istnieniem piątej siły pochodzi z obserwacji, które zwykle przypisujemy ciemnej materii”, pisze znana niemiecka fizyk dr Sabine Hossenfelder na swoim blogu. „Astrofizycy wprowadzają ciemną materię, ponieważ widzą siłę, która działa na normalną materię. Najszerzej akceptowaną hipotezą dla tej obserwacji jest to, że ta siła to po prostu grawitacja, a więc znany już typ oddziaływania, które działa na nieznaną jeszcze rodzaj materii. To założenie nie pasuje zbyt dobrze do wszystkich obserwacji, więc może być tak, że to nie jest tylko grawitacja, ale rzeczywiście nowa siła, i to byłaby piąta siła. Wiązana z nią bywa także ciemna energia, ale tak naprawdę nie jest to konieczne do wyjaśnienia obserwacji, przynajmniej nie w tej chwili”.

Napęd warp bez ujemnej energii i masy?

Fizyka wchodzi ostatnio na obszary, które kojarzyliśmy do tej pory jedynie z fantastyką naukową. Stamtąd znamy np. podróżowanie z prędkością większą niż światło i tajemniczo brzmiący „napęd warp”. Literatura i filmy na ogół nie wyjaśniają, jak to się dzieje. Po prostu „wchodzi się w nadświatłość” i już. Zgodnie z ogólną teorią względności Einsteina, niemożliwe jest, by cokolwiek podróżowało szybciej niż prędkość światła. A to dlatego, że gdy obiekt porusza się szybciej, jego masa wzrasta, więc zanim osiągniesz prędkość światła, masa ta zbliżałaby się do nieskończoności. Dodatkowo, aby przyspieszyć do tej prędkości, potrzebna byłaby nieskończona energia.

W 1994 r. meksykański fizyk teoretyczny Miguel Alcubierre przedstawił wizję „napędu warp”, który teoretycznie mógłby pozwolić na podróżowanie szybciej niż światło „zgodnie z fizyką”. Pomysł polega na wytworzeniu wokół obiektu bąbla ujemnej energii, w wyniku czego czasoprzestrzeń przed obiektem kurczy się, a przestrzeń za nim rozszerza (3). W środku znajduje się „płaski” obszar czasoprzestrzeni, w którym obiekt może podróżować w komfortowych warunkach, a jego mieszkańcy nie czują nawet, że się poruszają. Można sobie swobodnie mówić o generowaniu bąbla ujemnej energii, ale wymagałoby to egzotycznych form materii (ujemna masa), które nie są łatwe do zdobycia, o ile w ogóle istnieją. Wprawdzie fizyka zna coś takiego jak ujemna energia, którą możemy wygenerować z próżni w tzw. efekcie Casimira, ale te ilości energii są minimalne w stosunku do potrzeb podróży kosmicznych.

Jednak myśl ta żyje, także wśród ludzi nauki. Wiosną 2021 r. astrofizyk z uniwersytetu w Getyndze Erik Lentz opublikował nowy projekt teoretyczny, który mógłby umożliwić podróże z prędkością



3. Model napędu Alcubierre

większą niż prędkość światła, bez sięgania po egzotyczne pojęcia ujemnej energii i masy. Proponuje sposób na tworzenie „bąbli warp” z dodatnich źródeł energii. Chce do tego wykorzystać tzw. solitony, zwarte fale, które poruszają się ze stałą prędkością, nie tracąc swojego kształtu. Solitony można zaobserwować w pewnych okolicznościach w falach wodnych, ruchach atmosfery, które tworzą szczególnie formacje chmur, lub w świetle przechodzącym przez różne ośrodki. Lentz odkrył, że pewne konfiguracje solitonów mogą być formowane przy użyciu konwencjonalnych źródeł energii, nie naruszając żadnego z równań Einsteina bez korzystania z ujemnych energii. Ponieważ w centrum solitonu występują minimalne siły pływowe, czas płynąłby w tym samym tempie zarówno wewnątrz, jak i na zewnątrz bańki warp.

Lentz wykazuje, że naładowana plazma może wygenerować wymagane odkształcenia przestrzeni i czasu, aby stworzyć bąbel. Niestety, ilość potrzebnej energii jest równie ogromna, jak w przypadku rozwiązania Alcubierre. Nic, czym obecnie dysponujemy, ani Wielki Zderzacz Hadronów, ani nawet gwiazda nie są w stanie wygenerować takiej ilości energii. Lentz pozostawia swój pomysł do dalszej pracy z nadzieją, że pewnego dnia uda się wygenerować takie bańki warp, które pozwolą na podróże do gwiazd. Ponieważ rozwiązanie Alcubierre zostało jak do tej pory ulepszone o około sześćdziesiąt rzędów wielkości w stosunku do jego pierwotnej propozycji z 1994 roku, istnieje nadzieja, że dzięki modelowaniu numerycznemu i odrobinie pomysłowości wymagania Lentza również zostaną zredukowane. Reszta to inżynieria.

Niby-cząstki w prawdziwej elektronice

Skoro jesteśmy przy solitonach, które znajdują się na wydłużającej się liście „kwazicząstek”, to warto zwrócić uwagę na rosnącą popularność badań tych obiektów w fizyce. Ta ścieżka prowadzi nie tylko do odkryć czysto naukowych, ale również potencjalnie do nowych technologii, np. w świecie komputerów i danych.

Naukowcy z Laboratorium Ames Departamentu Energii USA odkryli ok. trzy lata temu subatomową kwazicząsteczkę, zwaną skyrmionem, która, według dostępnych opisów, umie „reprodukować się” zupełnie jak komórki biologiczne. Jest to także rodzaj solitonu, co kojarzy się z wcześniejszymi opisami teoretycznego napędu warp Lentza. Skyrmiony to pakiety energii i sił magnetycznych, bez masy, które samoczynnie układają się we wzory przypominające atomy ciasno ułożone wewnątrz kryształu. Naukowcy spekulują, że manipulowanie tymi subatomowymi cząstkami może doprowadzić do przełomu w technologii przechowywania i przesyłania danych, powstania nowych nadzwyczajnie wydajnych pamięci masowych i układów logicznych. W dodatku wygląda na to, że skyrmiony mają zdolność do samonaprawiania niedoskonałości we wzorze siatki poprzez replikację, stąd analogie z biologią. Teoretycznie pozwoliłoby to na zupełnie nowe typy pamięci komputerowych, odporne na uszkodzenia i ataki, które same mogłyby się odtwarzać i dowolnie powielać.

Pamięci masowe na bazie skyrmionów mogłyby mieć postać cienkich warstw magnetycznych, w których skyrmiony będą przesuwane w wyniku przepływu prądu. W takim sekwencyjnym zapisie informacji obecność (brak) skyrmionu odpowiadać będzie logicznemu zeru lub jednemu.

Elektrony wzdłuż wstęgi Möbiusa

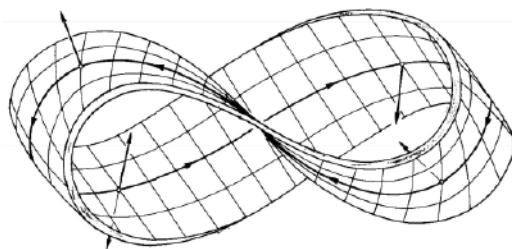
Kolejnym po kwazicząstkach „modnym”, ale zarazem obiecującym obszarem badań współczesnej fizyki są efekty topologiczne. Fizycy dawniej poświęcali niewiele uwagi topologii, która polega na matematycznej analizie kształtów i ich rozmieszczenia w przestrzeni. Ostatnia dekada przyniosła jednak wiele interesujących odkryć topologicznych, np. fakt, że niektóre izolatory mogą skrycie przewodzić prąd wzdłuż jednoatomowej warstwy na swojej powierzchni.

Topologia w matematyce to dział zajmujący się badaniem własności, które nie ulegają zmianie nawet po radykalnym zdeformowaniu obiektów (figur geometrycznych, brył i obiektów o większej liczbie wymiarów). Własności takie nazywa się niezmiennikami topologicznymi, przy czym przez deformowanie

rozumie się tutaj dowolne odkształcanie (zginanie, rozciąganie, skręcanie), ale bez rozrywania różnych części lub zlepiania różnych punktów. Odkrycie topologicznych faz materii stanowiło przełom w fizyce materii skondensowanej, prowadzący do powstania nowego opisu ciała stałego. Ponadto, niewrażliwość tych faz na zaburzenia (analogia z matematyczną topologią) czyni je potencjalnie użytecznymi w wielu zastosowaniach, w tym na przykład w informatyce kwantowej.

Coraz rzadziej można zobaczyć artykuł na temat fizyki ciała stałego, który nie zawierałby słowa topologia w tytule. Materiały te są też wykorzystywane jako wirtualne laboratoria do testowania przewidywań dotyczących egzotycznych i nieodkrytych cząstek elementarnych oraz nowych praw fizyki. Zwraca się uwagę, że topologiczny charakter mają fundamentalne zjawiska fizyki cząstek. Weźmy na przykład spin elektronu, który może być skierowany w górę lub w dół. Wydawałoby się, że obrót cząstki o 360° przywraca ją do pierwotnego stanu. Ale tak nie jest. W dziwnym świecie fizyki kwantowej elektron może być również reprezentowany jako funkcja falowa, która koduje informacje o cząstce, takie jak prawdopodobieństwo znalezienia jej w określonym stanie spinowym. Wbrew pozorom, obrót o 360° w rzeczywistości przesuwają fazę funkcji falowej, tak że szczyty fali stają się dołami i na odwrót. Potrzeba jeszcze jednego pełnego obrotu o 360° , aby ostatecznie przywrócić elektron i jego funkcję falową do ich stanów początkowych. Dokładnie to samo dzieje się w jednym z ulubionych przez matematyków obiektów topologicznych – wstędze Möbiusa (4). Jeśli mrówka przedreptałaby przez jedną pełną pętlę wstęgi, znalazłaby się po przeciwnej stronie niż na początku. Musi wykonać jeszcze jedno pełne okrążenie, zanim będzie mogła wrócić do pozycji wyjściowej.

W latach 80. teoretycy, tacy jak David Thouless z Uniwersytetu Stanu Waszyngton w Seattle, odkryli, że topologia może być odpowiedzialna za zjawisko zwane kwantowym efektem Halla. Thouless otrzymał za to odkrycie w 2016 roku Nagrodę Nobla, co jeszcze



4. Wstęga Möbiusa



silniej skoncentrowało uwagę na badaniach topologicznych. Później, w pierwszej dekadzie XXI wieku wykryto, że elektrony na powierzchni niektórych izolatorów znajdują się w stanach „topologicznie chronionych” przez magnetyzm, mogą płynąć jako prąd elektryczny prawie bez oporu. Potem okazało się, że efekty topologiczne w kryształach arsenku tantalu mogą tworzyć bezmasowe kwazicząstki, które zachowują się jak fermiony Weyla. Naukowcy mają nadzieję, że tego rodzaju materiały mogą pewnego dnia znaleźć zastosowanie w takich aplikacjach jak superszybkie tranzystory. Elektrony poruszające się w kryształach zwykle rozpraszają się, gdy trafiają na zanieczyszczenie, co spowalnia ich postęp, ale efekty topologiczne w kryształach arsenku tantalu pozwalają elektronom poruszać się bez przeszkód.

W 2005 roku Microsoft powierzył matematykowi Michaelowi Freedmanowi kierownictwo swoich projektów w dziedzinie obliczeń kwantowych. Początkowo zespół Freedmana skupiał się głównie na stronie teoretycznej. Jednak potem Microsoft zatrudnił kilku wybitnych eksperymentatorów ze środowiska akademickiego. Jednym z nich był fizyk Leo Kouwenhoven z Uniwersytetu Technicznego w Delft w Holandii, który w 2012 roku jako pierwszy potwierdził doświadczalnie, że kwazicząstki nazywane enionami pamiętają, w jaki sposób zostały zamienione. Teraz badania zmierzają do wykazania, że eniony mogą kodować kubity i wykonywać proste obliczenia kwantowe. Freedman uważa, że to topologicznych kubity, a nie obecne układy mrożone w temperaturach bliskich zera bezwzględnego, są przyszłością komputerów kwantowych.

Niezwykłe teorie fizyczne

Wszechświat holograficzny

Zgodnie z tą teorią cały trójwymiarowy wszechświat może być „zakodowany” na jego dwuwymiarowej granicy. Teoria ta ma tę zaletę, że jest to teoria testowalna naukowo. Badania z 2017 roku na uniwersytecie w Southampton w Wielkiej Brytanii wykazały, że jest ona zgodna z obserwowanym wzorcem fluktuacji CMB.

Wszechświat w stanie ustalonym raz na zawsze

Nie wszyscy naukowcy są zadowoleni z modelu Wielkiego Wybuchu i rozszerzającego się Wszechświata, więc wymyślili sposób na to, aby gęstość pozostała stała, nawet w rozszerzającym się Wszechświecie. Rozwiązanie to polega na ciągłym tworzeniu materii w tempie około trzech atomów wodoru na metr sześcienny na milion lat.

Wieloświat

W konwencjonalnym ujęciu Wielkiego Wybuchu, aby wyjaśnić jednorodność CMB, trzeba postulować wczesny zryw superszybkiej ekspansji, znany jako inflacja. Niektórzy naukowcy uważają, że kiedy nasz Wszechświat wyszedł z tej inflacyjnej fazy, był on tylko jednym małym bąbelkiem w ogromnym oceanie puchnącej przestrzeni. W tej teorii, zwanej „wieczną inflacją”, zaproponowanej przez Paula Steinhardta, inne bąbel-

kowe wszechświaty nieustannie pojawiają się w innych częściach inflacyjnego morza, a cały ten zespół tworzy „multiwersum” – „wieloświat”. Nie ma powodu, aby inne wszechświaty miały takie same prawa fizyki jak nasz. Chociaż nie możemy obserwować innych wszechświatów bezpośrednio, jeden z nich mógłby zderzyć się z naszym. Naukowcy zasugerowali nawet, że „zimna plama” w CMB jest śladem takiej kolizji.

Teoria grawitacji to pomyłka

Grawitacja nie jest w stanie wyjaśnić niektórych obserwacji astronomicznych. Jeśli mierzymy prędkość gwiazd na obrzeżach galaktyki, to poruszają się one zbyt szybko, aby utrzymać się na orbicie, jeśli jedyną rzeczą, która je powstrzymuje, jest grawitacyjne przyciąganie widocznej galaktyki. Podobnie, gromady galaktyk wydają się utrzymywane razem przez siłę silniejszą niż ta, którą można wytłumaczyć grawitacją materii widzialnej. Istnieją dwa możliwe rozwiązania. Standardowe – preferowane przez większość naukowców – jest takie, że wszechświat zawiera niewidoczną ciemną materię, która zapewnia brakującą grawitację. Alternatywą dla niej jest twierdzenie, że nasza teoria grawitacji jest błędna i powinna zostać zastąpiona przez coś, co nazywa się Zmodyfikowaną

Dynamiką Newtonowską (MOND), co naukowcy zaproponowali w 2002 roku w czasopiśmie „Annual Review of Astronomy and Astrophysics”.

Superpłynna czasoprzestrzeń

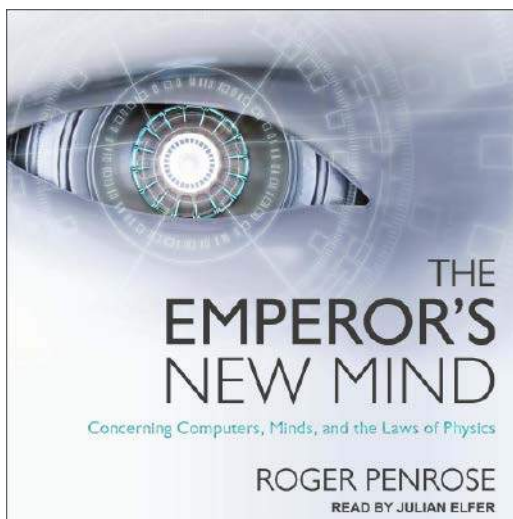
Według niektórych teorii, takich jak ta zaproponowana przez Stefano Liberatiego z International School for Advanced Studies i Lucę Maccione z Uniwersytetu Ludwiga Maximiliana w czasopiśmie „Physics Review Letters”, czasoprzestrzeń nie jest tylko abstrakcyjnym układem odniesienia zawierającym obiekty fizyczne, takie jak gwiazdy i galaktyki, ale to fizyczna substancja sama w sobie, analogiczna do oceanu wody. Tak jak woda składa się z niezliczonych cząsteczek, tak czasoprzestrzeń, zgodnie z tą teorią, składa się z mikroskopijnych cząsteczek na głębszym poziomie rzeczywistości, niż mogą sięgnąć nasze instrumenty. Teoria ta wizualizuje czasoprzestrzeń jako superciecz o zerowej lepkości. Dziwną właściwością takich płynów jest to, że nie można ich wprawić w ruch obrotowy w sposób całościowy, jak to się dzieje w przypadku zwykłej cieczy, gdy się ją miesza. Rozpadają się one na maleńkie wiry – które w przypadku superpłynnej czasoprzestrzeni mogą być nasionami, z których powstają galaktyki.

Wyczytać świadomość z mózgowych fraktali

Z intrygującej, ale co tu dużo mówić, dość trudnej dziedziny topologii przejdźmy do świata teorii znacznie bardziej sensacyjnie brzmiących, które na swój sposób są jeszcze trudniejsze. Tajemnica świadomości, według Rogera Penrose'a, laureata Nagrody Nobla z fizyki w 2020 r., zostanie rozwiązana dopiero wtedy, gdy zrozumiemy, w jaki sposób struktury mózgu mogą wykorzystać właściwości mechaniki kwantowej. Penrose, który otrzymał Nobla za pracę nad naturą czarnych dziur, interesował się świadomością od czasu, gdy był studentem Cambridge. Jest autorem wielu książek na temat świadomości, w szczególności „Nowy umysł cesarza” (5), i uważa, że jest ona tak złożona, że nie może być wyjaśniona przez nasze obecne rozumienie fizyki i biologii.

Po opublikowaniu książki znanej pod angielskim tytułem „The Emperor's New Mind” Penrose otrzymał list od Stuarta Hameroffa, profesora anestezjologii na Uniwersytecie Arizony. W liście tym Hameroff opisał małe struktury w mózgu zwane mikrotubulami, które jego zdaniem były zdolne do generowania świadomości poprzez sięganie do świata kwantowego. Hameroff uważał, że znieczulenie może działać poprzez specyficzne ukierunkowanie świadomości poprzez działanie na neuronalne mikrotubule. Po napisaniu listu spotkał Penrose'a w 1992 roku i przez następne dwa lata rozwijali radykalne idee dotyczące świadomości, które były sprzeczne z myśleniem większości neurologów, i nadal są. Penrose i Hameroff wierzą, że ludzki mózg działa bardziej jak komputer kwantowy niż jakikolwiek klasyczny komputer. Sugerują, że te efekty kwantowe powstają w mikrotubulach, które następnie działają jako łącznik mózgu ze światem kwantowym.

Mikrotubule mają kluczowe znaczenie dla podziału komórki, dzieląc chromosomy idealnie na dwie części. Gdyby mikrotubule nie działały, chromosomy mogłyby być dzielone nierównomiernie na trzy lub cztery, a nie na dwa, wywołując w ten sposób raka. Centralna rola, jaką mikrotubule odgrywały w podziale komórek, doprowadziła Hameroffa do spekulacji, że były one kontrolowane przez jakąś formę naturalnej informatyki. W swojej książce „Ultimate Computing” (1987) Hameroff dowodzi, że mikrotubule mają wystarczającą moc obliczeniową, by wytwarzać myśli. Przekonuje również, że to one, a nie neurony są najbardziej podstawowymi jednostkami przetwarzania informacji w mózgu. Penrose zgodził się z Hameroffem, że mikrotubule mogą utrzymywać „kwantową spójność” potrzebną dla złożonej myśli.



5. Okładka książki „The Emperors New Mind” Rogera Penrose’a

Doprowadziło to Penrose'a i Hameroffa do stworzenia teorii zwanej orkiestrową redukcją, lub OR. Według niej obszary mózgu, w których pojawia się świadomość, muszą być tak skonstruowane, aby mogły pomieścić niezliczoną ilość kwantowych możliwości naraz, zgodnie z zasadami mechaniki kwantowej, pozwalając jednocześnie na kontrolowaną redukcję tych nieskończonych możliwości bez niszczenia systemu kwantowego. Uważają, że mikrotubule są najlepiej nadającymi się do tego znanymi obecnie strukturami w mózgu, w których procesy kwantowe mogą zachodzić w stabilny sposób i być wykorzystywane do generowania naszego świadomego doświadczenia. Zgodzili się zarazem, że świadomość może być ostatecznie odnaleziona w wielu miejscach w mózgu, nie tylko w mikrotubulach.

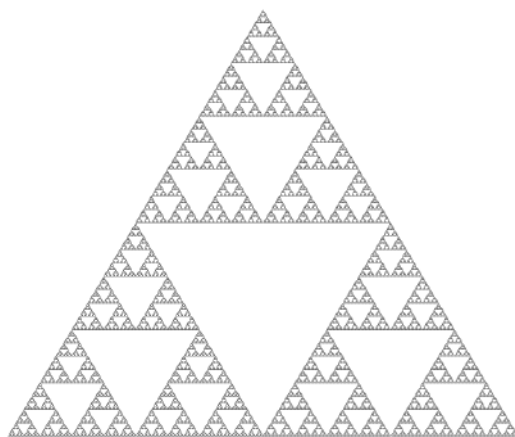
Jednym z głównych zarzutów wobec kwantowej teorii świadomości Penrose'a-Hameroffa jest to, że nie ma sposobu, aby zmierzyć, czy procesy kwantowe zachodzą w mikrotubulach lub innych częściach mózgu. Penrose przyjmuje tę krytykę, ale wierzy, że takie pomiary staną się możliwe w dłuższej perspektywie czasu. Hameroff ma już plany, by sprawdzić, czy wewnątrz mikrotubuli istnieją stany kwantowe. Jeśli uda mu się to udowodnić, jego następnym krokiem będzie sprawdzenie, czy takie stany znikają pod wpływem znieczulenia.

Techniki skanowania mózgu, takie jak PET i MRI, mimo że generalnie zostały znacznie udoskonalone, są jednak, zdaniem Penrose'a, mało lub zupełnie nieprzydatne w badaniach nad świadomością. Mogą one, w jego ocenie, monitorować przepływ krwi

i aktywność w obszarach mózgu, ale nie mogą powiedzieć, czy ta aktywność wiąże się ze świadomą myślą. Do tego potrzebne jest coś innego. Niektórzy naukowcy uważają, że jednym ze sposobów pomiaru myśli jest obserwacja fal mózgowych. Niektóre dowody sugerują, że fale mózgowo, oscylujące z częstotliwością około 40 herców, mogą być skorelowane ze świadomością. Dużą trudnością przy próbie zmierzenia procesów kwantowych w mózgu, wskazuje Penrose, jest to, że takie efekty są niszczone, gdy są obserwowane lub wchodzą w kontakt ze światem zewnętrznym.

Teoria świadomości kwantowej Penrose'a-Hameroffa opiera się na twierdzeniu, że mikrotubule mają strukturę fraktalną i właśnie taka struktura umożliwia zachodzenie procesów kwantowych. W matematyce fraktale pojawiają się jako piękne wzory, które powtarzają się w nieskończoność, tworząc coś, co wydaje się niemożliwe – strukturę, która ma skończoną powierzchnię, ale nieskończony obwód. Wzory fraktalne często obserwuje się w przyrodzie. Wystarczy spojrzeć na strukturę kwiatu kalafiora czy gałęzie paproci, gdzie podstawowy kształt powtarza się w kółko, ale w coraz mniejszej skali. Jest to kluczowa cecha fraktali. To samo dzieje się, gdy spojrzymy do wnętrza naszego ciała. Według zwolenników tej teorii, np. struktura naszych płuc jest fraktalna, podobnie jak naczynia krwionośne w naszym układzie krążenia. Fraktale występują również w dziełach sztuki Eschera i Jacksona Pollocka, a także są wykorzystywane od dziesięcioleci w technologii, np. w projektowaniu anten.

Nie jesteśmy jeszcze w stanie zmierzyć zachowania fraktali kwantowych w mózgu, jeśli w ogóle istnieją. Jednak bada się te wzory w inny sposób. W eksperymentach z wykorzystaniem skaningowego mikroskopu tunelowego (STM), naukowcy z uniwersytetu w Utrechcie starali się tworzyć wzory fraktali kwantowych z elektronów. Kiedy następnie zmierzili funkcję falową elektronów, która opisuje ich stan kwantowy, okazało się, że one również żyły w wymiarze fraktalnym podyktowanym przez fizyczny wzór, który powstał. W tym przypadku wzorzec, którego użyto w skali kwantowej, był trójkątem Sierpińskiego (6), który jest kształtem gdzieś pomiędzy jednowymiarowym a dwuwymiarowym. W kolejnych badaniach przeprowadzonych z badaczami z Uniwersytetu Jiaotong w Szanghaju Holendrzy wstrzykiwali fotony do sztucznego chipu, który został zaprojektowany na kształt trójkąta Sierpińskiego. Obserwowali, jak rozprzestrzeniają się one w całej jego fraktalnej strukturze w procesie zwanym transportem kwantowym. Następnie powtórzyli ten



6. Trójkąt Sierpińskiego

eksperyment na dwóch innych strukturach fraktalnych, obu w kształcie kwadratów, a nie trójkątów.

Obserwacje z tych eksperymentów ujawniają, że rozprzestrzenianie się światła po fraktalu rządzi się innymi prawami w przypadku kwantowym niż w przypadku klasycznym. Tak zdobyta nowa wiedza o fraktalach kwantowych może dać naukowcom podstawy do eksperymentalnego przetestowania teorii świadomości kwantowej. Jeśli pewnego dnia zostaną wykonane pomiary kwantowe w ludzkim mózgu, będzie można je porównać z wynikami holendersko-chińskiego zespołu, aby ostatecznie rozstrzygnąć, czy świadomość jest zjawiskiem klasycznym, czy kwantowym.

Wszechświat wysniony i symulowany?

Koncepcja kwantowej natury świadomości kojarzy się z różnego rodzaju kosmologicznymi poglądami na Wszechświat jako jednorodną świadomość a przynajmniej symulacja. Syntezą takiego myślenia wydaje się hipoteza opublikowana niedawno w pracy uczonych z Instytutu Badań nad Kwantową Grawitacją (Quantum Gravity Research Institute) z siedzibą w Los Angeles, że Wszechświat symuluje sam siebie do istnienia, czyli zakłada istnienie czegoś takiego jak wszechogarniająca panświadomość.

Znany jest z głośnej hipotezy przyjmującej, iż cała nasza egzystencja może być tylko produktem bardzo wyrafinowanych symulacji komputerowych, jest Nick Bostrom, filozof. Tym razem sprawą zajęli się fizycy. Wszechświat jest „dziwną pętlą” – twierdzą w pracy zatytułowanej „The Self-Simulation Hypothesis Interpretation of Quantum Mechanics”. Przyjmują oni hipotezę symulacji Bostroma, która utrzymuje, że symulacja jest dziełem niezwykle zaawansowanych form życia, które stworzyłyby niesamowitą technologię

niezbędną do skomponowania wszystkiego w naszym świecie. Jednak sądzą, iż bardziej efektywne jest myślenie o Wszechświecie jako o „umysłowej autosymulacji”. Wiążą oni ten pomysł z mechaniką kwantową, postrzegając wszechświat jako jeden z wielu możliwych modeli kwantowej grawitacji.

Oryginalna hipoteza Bostroma jest materialistyczna, postrzegając Wszechświat jako z natury fizyczny. Możemy według niej być nawet częścią symulacji naszych własnych potomków, „postludzi”. Nawet sam proces ewolucji może być tylko mechanizmem, za pomocą którego istoty z przyszłości testują niezliczone sy, celowo przesuwając ludzi przez kolejne poziomy biologicznego i technologicznego rozwoju. W ten sposób generują one również rzekomą informację lub historię naszego świata. Ostatecznie nie zauważylibyśmy różnicy. Fizycy z Los Angeles odrzucają rzeczywistość fizyczną na rzecz czystej myśli. Jak mogłaby powstać sama symulacja, która kojarzy się z „Matrixem”? Zawsze tam była, mówią naukowcy. Zgodnie z tą ideą czas w ogóle nie istnieje. Zamiast tego wszechogarniająca myśl, która jest naszą rzeczywistością, oferuje zagnieżdżone pozory hierarchicznego porządku, pełnego „podmyśli”. Dalej jest sugestia, że ludzie sami są takimi „emergentnymi submyślami” i doświadczają oraz odnajdują znaczenie w świecie poprzez inne submyśli. Ezoteryczne? Zdaniem uczonych, my sami w snach przeprowadzamy takie podsymulacje, więc dlaczego nie mogłoby to się dziać na wyższym poziomie? Wskazują oni szczególnie na przejrzyste

rodzaje snów, w którym śniący jest świadomy bycia we śnie, jako przypadki bardzo dokładnych symulacji stworzonych przez twój umysł, które mogą być niemożliwe do odróżnienia od jakiegokolwiek innej rzeczywistości (7).

Pewien fizyk twierdzi, że informacja cyfrowa powinna być uważana za nową formę materii. Oznacza to, że Ziemia zmierza w kierunku kryzysu, który fizyk z uniwersytetu w Portsmouth, Melvin Vopson, nazywa „katastrofą informacyjną”, zgodnie z pracą, którą opublikował w czasopiśmie „AIP Advances”. Vopson twierdzi, że zgodnie z jego obliczeniami, do roku 2500 informacje cyfrowe będą stanowiły połowę masy Ziemi, co sprawi, że ich wzrost będzie niemożliwy do utrzymania.

Vopson przytacza ramy teoretyczne, które nazywa zasadą równoważności masa-energia-informacja, która łączy szereg odrębnych teorii fizycznych. Po pierwsze, istnieje ogólna teoria względności Einsteina, która łączy masę z energią. Następnie jest teoria fizyka Rolfa Laundera, że istnieje podstawowy koszt energetyczny związany z przetwarzaniem informacji, którą Vopson łączy, aby argumentować, że bity cyfrowe mają masę i powinny być uważane za materię obok ciał stałych, cieczy, gazów i plazmy. „Chociaż informacja przejawia się w wielu formatach, w tym informacji analogowej, informacji zakodowanej w biologicznym DNA i informacji cyfrowej, najbardziej fundamentalną formą jest binarny bit cyfrowy, ponieważ może on z powodzeniem reprezentować lub powielać wszystkie istniejące formy informacji.

7. Wszechświat jako symulacja lub system wyobrażeń





8. Wizualizacja rozpadu bozonu Higgsa

Vopson wysuwa też kilka wyjątkowo daleko idących twierdzeń, takich jak to, że ciemna materia utrzymująca galaktyki razem może być w jakiś sposób zbudowana z informacji. Ale poza tym wszystkim zwraca on uwagę na ważny, często pomijany problem: fakt, że stale rosnąca infrastruktura cyfrowa może pewnego dnia wymagać zupełnie niezrównoważonej ilości energii, której nasza planeta po prostu nie będzie w stanie dostarczyć. „Każdy powinien uznać to za interesujące, ponieważ prognozy pokazują, że w najbliższej przyszłości będziemy produkować tak wiele treści cyfrowych, że liczba wyprodukowanych bitów będzie równa wszystkim atomom na Ziemi” – powiedział Vopson w rozmowie z ZME. „Pytanie brzmi więc: gdzie przechowamy te informacje? Jak to zasilimy? To sygnał alarmowy dla przemysłu big data, gigantów internetowych, firm high-tech, badań energetycznych i środowiskowych. Nazywam to niewidzialnym kryzysem, ponieważ dziś jest to rzeczywiście niewidoczny problem, ale prognozy pokazują inną historię” – dodał.

Odpylniliśmy nieco, więc czas powrócić do fizyki na poziomie bardziej realnym, do eksperymentów i teorii, które mamy już jako tako opanowane. Tu ciągle się coś dzieje, np. na początku 2021 r. w detektorach ATLAS i CMS w Wielkim Zderzaczu Hadronów w CERN w Szwajcarii odkryto pierwsze dowody rzadkiego rozpadu bozonu Higgsa (8) na jeden foton i dwa leptony, parę elektronów lub foton i parę mionów o przeciwnym ładunku. Korzystając z Modelu Standardowego, naukowcy są w stanie przewidzieć, na jakie cząstki elementarne może rozpadać się bozon Higgsa, przy czym dość „powszechnym” rozpadem są dwa fotony. Te mniej typowe rozpady bozonu Higgsa są zdaniem naukowców tropem nowej fizyki, która wykracza poza Model Standardowy. A o wykraczaniu poza Model Standardowy słyszymy w ostatnich latach bardzo

często, choć na razie nikt nie zdołał na dobre poza niego wykroczyć.

W eksperymentach LHCb w CERN odkryto latem tego roku nowego tetrakwark oznaczonego T_{cc}^+ , który zawiera dwa ciężkie kwarki i dwa lekkie antykwarki. To pierwszy taki przypadek znany fizykom. Nowa cząstka zawiera dwa kwarki powabne oraz antykwark górny i antykwark dolny. W ostatnich latach odkryto kilka tetrakwarków (w tym jeden z dwoma kwarkami powabnymi i dwoma antykwarkami powabnymi), ale ten jest pierwszym, który zawiera dwa kwarki powabne, bez antykwarków powabnych dla ich zrównoważenia. Cząstki zawierające kwark powabny i antykwark powabny mają „ukryty powab” – liczba kwantowa dla całej cząstki sumuje się do zera, tak jak dodatni i ujemny ładunek elektryczny. Tutaj liczba kwantowa sumuje się do dwóch, a więc cząstka ma podwójny powab. T_{cc}^+ jest pierwszą znaną cząstką, która należy do klasy tetrakwarków z dwoma ciężkimi kwarkami i dwoma lekkimi antykwarkami. Odkrycie otwiera drogę do poszukiwania cięższych cząstek tego samego typu, z jednym lub dwoma kwarkami powabnymi zastąpionymi przez kwarki dolne. Wszystko to sprawia, że prosty wcześniej obraz cząstek tworzących hadrony wciąż się komplikuje, a to, co uznawane było za egzotyczne, staje się coraz bardziej normalne. Komplikująca się rzeczywistość na tym poziomie każe schodzić głębiej w poszukiwaniu czegoś mniejszego i bardziej podstawowego. Zatem potencjalnie te odkrycia prowadzą do nowego definiowania materii.

Wysiłki badacze nie ustają również na obszarze kosmologii, czyli badania Wszechświata jako całości. W pracy opublikowanej w lipcu na serwerze w arXiv, zespół astrofizyków i kosmologów z uniwersytetów w Ulm i w Lyonie zbadał światło pozostałe po Wielkim Wybuchu, znane jako kosmiczne tło mikrofalowe tło (Cosmic Microwave Background, CMB), ustalając,

że Wszechświat może jednak nie być płaski, jak wielu naukowców obecnie uważa, a raczej jest gigantycznym toroidem (9).

Pomysł kształtu torusa dla naszego Uniwersum powstał w latach 80. Jednak w ciągu ostatnich dekad lat naukowcy przeważnie głosili, iż nasz Wszechświat jest geometrycznie płaski i rozszerza się w tej postaci. Jednak, jak zauważają autorzy nowej pracy, w danych dotyczących CMB było coś, co nie do końca się zgadzało. „Widmo [CMB] (...) ma też duże luki”, piszą. Innymi słowy, wydaje się, że w CMB brakuje sygnałów, które byłyby obecne, gdyby Wszechświat był naprawdę, taki jak przewiduje kosmologiczny model standardowy. Proponowane wyjaśnienie brzmi, że Wszechświat w rzeczywistości „wielokrotnie się łączy” w tych brakujących punktach, co oznacza, że jego topologia jest zakrzywiona w taki sposób właśnie jak w toroidzie.

Toroidalna struktura Wszechświata przywołuje omawiane wyżej koncepcje „napędów warp” (bo skoro krzywizna jest naturą wszechrzecz, to zakrzywienie czasoprzestrzeni w celu wyprzedzenia światła nie powinno brzmieć egzotyczne) a także obiektów topologicznych (w końcu toroid wszechświatowy to właśnie taki obiekt) a nawet teorie symulacji, bo coś o skończonej geometrii w sposób naturalny wydaje się bardziej możliwe do symulowania niż nieskończona płaszczyzna.

Fizyka coraz chętniej obraca się w kręgu pojęć, które kiedyś siałaby wśród uczonych zgorszenie, w świecie cząstek, które nie mają masy, a o ich



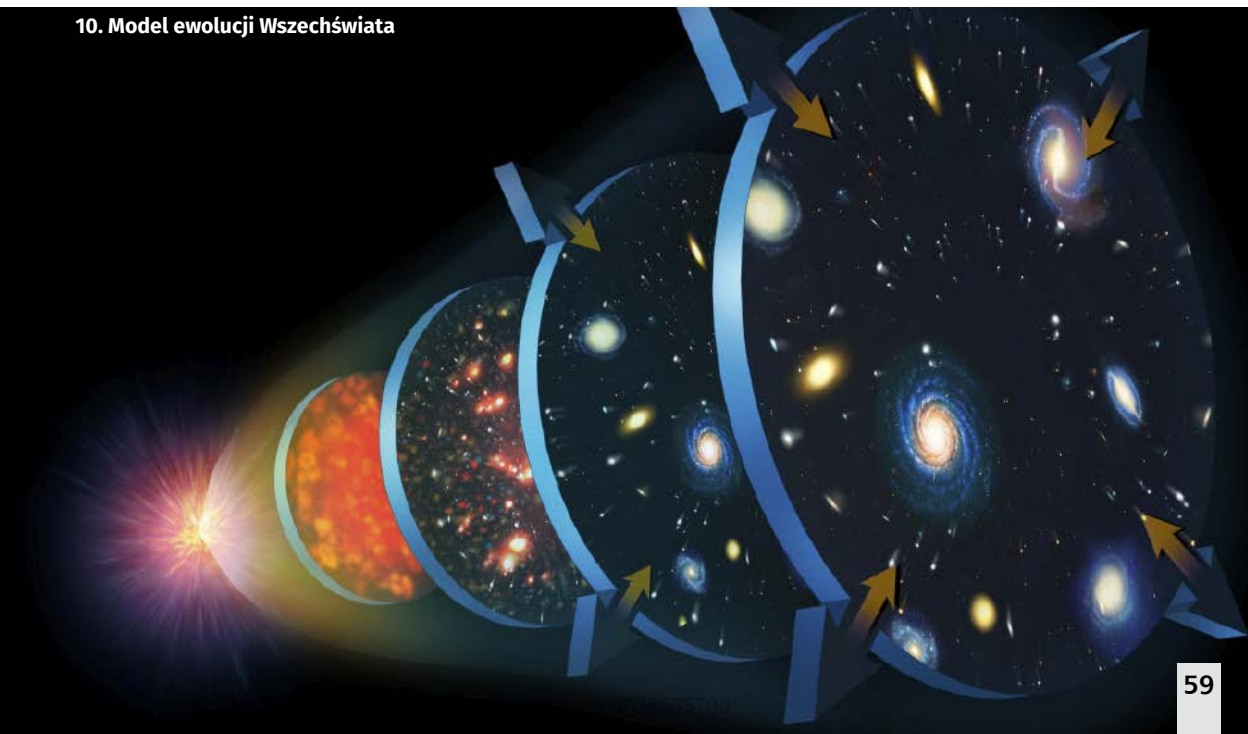
9. Wszechświat jako wielki toroid

właściwościach decydują kształty, topologie, geometria. Zaczyna łączyć światy, których łączenie było dawniej nie do pomyślenia, np. mechanikę kwantową z biologią a nawet psychologią. Czy coś z tego w sensie poznawczych, naukowym i w końcu praktycznym wyniknie? Być może nic, ale z modeli wcześniej utrwalonych w fizyce, np. Modelu Standardowego cząstek też tak po prawdzie nie wynika w praktyce za wiele poza samą wiedzę. Wielkich teorii kosmologicznych ewolucji Wszechświata są wręcz niemożliwe do zweryfikowania.

Rozkołysanie się badań fizycznych po tych wszystkich obszarach, które nie zawsze ściśle kojarzą się z utartymi ścieżkami znanej nauki, świadczy o tym, że fizyka poszukuje nowych dróg, bo ta, którą zmierziała przez dekady, znalazła się, mówiąc najogólniej, w impasie. ■

Mirosław Usidus

10. Model ewolucji Wszechświata



**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 d03866570d

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. Piotr Szulczewski**

CV: Piotr (Peter) Szulczewski

Data i miejsce urodzenia: 1981, Warszawa

Adres zamieszkania: 931 Linda Flora Dr, Bel Air, Kalifornia, USA

Obywatelstwo: kanadyjskie, polskie

Stan cywilny: nieznan

Majątek: 1,8 mld dolarów

Kontakt: @sfpiotr

Edukacja: Uniwersytet Waterloo w Kanadzie

Doświadczenie zawodowe: 2007–2009

– praca w Google, w koreańskim oddziale; 2009

– do dziś – założyciel, szef i główny programista firmy ContextLogic, która opracowała aplikacje i platformę e-commerce Wish

Zainteresowania: psy

Konsekwentny chłopak z Warszawy – **Piotr Szulczewski**

Wygrał stypendium na najlepszej kanadyjskiej uczelni, staż w Google, mógł przebierać w ofertach pracy, ale wybrał własną drogę. Stworzył własny startup i największą mobilną platformę handlową – Wish. Poznajcie historię Piotra (Petera) Szulczewskiego (**1**), który swoją aplikacją podbija świat.

Unika mediów i pytań o prywatność. Dlatego niewiele można opowiedzieć o jego życiu w okresie, zanim zwrócił na siebie uwagę całej Doliny Krzemowej. Uchodzący w medialnych relacjach za skromnego Peter Szulczewski jest z pochodzenia warszawiakiem. Rocznik 1981, zdążył poznać PRL i życie w blokach na Tarchominie.

Miał zaledwie 11 lat, gdy z rodzicami wyjechał do Kanady. Tam ukończył studia na Wydziale Matematyki i Informatyki Uniwersytetu Waterloo w Ontario, uznanym za najlepszą kanadyjską uczelnię w dziedzinie nauk ścisłych. Podczas studiów poznał Danny'ego Zhanga (2), który był najpierw jego przyjacielem a potem partnerem biznesowym. Obaj byli stypendystami Uniwersytetu Waterloo.

Potomkowi chińskich imigrantów tak naprawdę marzyła się kariera piłkarska. Wolał grać z Peterem w piłkę, niż kodować, ale Szulczewskiego ciągnęło do komputera i miał zawsze mnóstwo świetnych pomysłów. Zhang nie otrzymał ostatecznie oferty z żadnego liczącego się klubu piłkarskiego. Połączyli siły, a pierwsze zawodowe kroki stawiali od razu w najważniejszych firmach branży IT.

Szulczewski zaczynał w ATI Technologies Inc., u kanadyjskiego producenta m.in. kart graficznych. Kolejni jego pracodawcy byli już z Doliny Krzemowej, gdzie programował dla Microsoft i Google. Dla Google napisał algorytm dobierający najlepsze i najczęściej wyszukiwane hasła dla reklamodawców. Kod automatycznie dobierał reklamę w popularne słowa kluczowe, których nie uwzględnił zamawiający kampanię menedżer. Dzięki usłudze reklamodawcy zyskiwali więcej odsłon i szans na transakcję, a przychody Google wzrosły, według Szulczewskiego, o ok. 100 milionów dolarów rocznie.

Sukces przyniósł kolejne wyzwanie – w 2007 roku Szulczewski pracował nad optymalizacją stron Google pod gusta koreańskich użytkowników. I otrzymał cenną lekcję od Koreańczyków, którzy nie chcieli tego, co, zdaniem gigantów z Doliny Krzemowej, powinni chcieć, np. białych, ascetycznych stron Google. Szulczewski stworzył nowy projekt pod gusta



2. Szulczewski z Dannyem Zhangiem

i oczekiwania tamtejszych użytkowników. Uczył się myśleć jak klienci, dla których tworzył. A po dwóch latach odszedł z firmy. Ponoć zmęczył go szklany sufit w korporacji, gdzie każdy projekt musiał przejść długą proces od pomysłu do realizacji.

Zaraz za Amazonem i Alibabą

Z oszczędnościami, pozwalającymi na uruchomienie własnej firmy, zasiadł do programowania. Po pół roku miał mechanizm, rozpoznający zainteresowania użytkownika na podstawie jego zachowania w Internecie i dobierający na podstawie tego odpowiednie reklamy. W ten sposób powstał innowacyjny program mobilnej sieci reklamowej, która mogła konkurować z Google AdSense. Był maj 2011 roku. Nowatorski projekt pozyskał 1,7 miliona dolarów na inwestycje i zaangażował dyrektora generalnego Yelp, Jeremy'ego Stoppelmanna. Szulczewski nie zapominał też o starym druhu i zaprosił do współpracy kumpla z uczelni, Zhanga, wówczas pracującego dla z YellowPages.com.

Pojawili się kupcy na nowy produkt. Wśród zainteresowanych zakupem był Facebook, ale Szulczewski zrezygnował z jego oferty w wysokości dwudziestu milionów dolarów za ContextLogic. Wspólnie ze Zhangiem woleli sami dopracować silnik, z którego wyewoluowała mobilna platforma handlowa Wish, jak dotąd najcenniejsze dzieło Szulczewskiego. Pomysł był prosty – samouczący się program i aplikacja, w której użytkownicy dodają swoje zakupowe życzenia, np. koszyk do roweru albo wędkę do łowienia ryb, perfumy itp.

Aplikację w krótkim czasie zainstalowano na dziesiątkach tysięcy telefonów komórkowych. Jednym z najchętniej wybieranych produktów okazały się liczniki rowerowe. Z czasem aplikacja wyszukiwała i pokazywała użytkownikom najlepsze oferty towarów, o których marzyli. Wszystko działało się szybko i wygodnie, bo w smartfonie. Klientami Wish okazały się przede wszystkim kobiety, a oferowane produkty pochodziły głównie od sprzedawców z Chin. Azjatyccy sprzedawcy docenili aplikację. Nie musieli bowiem praktycznie nic robić – wrzucali swoją ofertę, a Wish pokazywał ją potencjalnym klientom.

Na początku twórcy platformy rezygnowali z marży od kupców, pod warunkiem że oferta była zamieszczana z promocyjną niższą o 10–20% ceną. I tak pod bokiem potężnych firm, jak Walmart, Amazon, Alibaba-Taobao itp., pojawił się nowy konkurent – Wish. Szulczewski i Zhang świetnie rozumieli, że nie będzie jednak łatwo pokonać amerykańskich gigantów sprzedażowych. Celowali więc w grupę użytkowników, niewidocznych dla potentatów z Doliny Krzemowej. Chodziło o klientów z mniej wypchanym portfelem, dla których ważniejsza jest cena niż szybka dostawa w ładnym opakowaniu. A takich klientów tylko w USA jest w ocenie Szulczewskiego dużo: „41 procent amerykańskich gospodarstw domowych nie ma płynności większej niż 400 dolarów”, przekonywał inwestorów, dodając, że o klientach z Europy mają jeszcze bardziej mylne wyobrażenie.

W ciągu dekady Wish stał się w światowym handlu internetowym graczem numer trzy, za Amazonem i Alibabą-Taobao. Statystyki wskazywały, że największe grono użytkowników Wish stanowią mieszkańcy Florydy, Teksasu i Środkowego Zachodu USA. Aż 80 proc. z nich po pierwszych zakupach wracało, aby dokonać kolejnej transakcji. W 2017 roku Wish była najczęściej pobieraną aplikacją e-commerce w Stanach Zjednoczonych, zaś ok. 80 proc. klientów Wish wracało na kolejne zakupy. Za pomocą aplikacji Wish kupują również użytkownicy z Grecji, Finlandii, Danii, Kostaryki, Chile, Brazylii i Kanady. I znów Szulczewski otrzymał propozycję sprzedaży Wish, tym razem od Amazona. Jednak do transakcji nie doszło.

Wish jest reklamowany przez wielu znanych sportowców. Ma podpisaną umowę ze znanym klubem koszykarskim Los Angeles Lakers (3). Podczas Mistrzostw Świata w 2018 roku appkę reklamowały gwiazdy futbolu: Neymar, Paul Pogba, Tim Howard, Gareth Bale, Robin van Persie, Claudio Bravo i Gianluigi Buffon. Efektem był wzrost liczby



3. Koszulka Lakersów z logo aplikacji Wish

użytkowników. W 2018 roku Wish okazał się najczęściej pobieraną aplikacją e-commerce na świecie. Twórcom platformy przyniosło to podwojenie przychodów do 1,9 miliarda dolarów.

Bogactwo i życie wśród gwiazd

Peter oprócz tego, że jest zdolnym programistą, ma wybitną żylkę biznesową. W 2020 roku jego firma zadebiutowała na nowojorskim parkiecie, a inwestorzy wycenili Wish na prawie dziewięć miliardów dolarów. Mając prawie jedną piątą akcji, chłopak rodem z Warszawy stał się miliarderem z fortuną wartą 1,7 mld dolarów. W rankingu magazynu „Forbes” znajduje się na 1833. miejscu listy miliarderów w 2021 roku.

Jego firma ma siedzibę na najwyższych piętrach wieżowca przy Sansome Street w San Francisco. Niedawno media ujawniły, że Peter Szulczewski nabył wartą 15,3 miliona dolarów nowoczesną rezydencję w luksusowej enklawie Bel-Air u podnóża gór Santa Monica w Los Angeles. Z rezydencji rozciąga się widok na winnice Ruperta Murdocha, a sąsiadami amerykańskiego miliardera z polskimi korzeniami są m.in. Beyonce i Jay-Z.

Jak wielu miliarderów Szulczewski angażuje się w działania charytatywne – razem ze Zhangiem są fundatorami stypendiów Wish dla studentów ich macierzystej uczelni, Uniwersytetu Waterloo. Na stronie uczelni Szulczewski pisze młodszym kolegom z branży IT m.in.: „konsekwencja jest najbardziej niedocenianą cnotą w przedsiębiorczości”. ■

Miroslaw Usidus

Web 3.0 ponownie, ale znów inaczej

Łańcuchy, które mają nas wyzwolić

Zaraz po wejściu w obieg pojęcia Web 2.0, w drugiej połowie pierwszej dekady XXI wieku, pojawiło się od razu pojęcie trzeciej wersji internetu (1), rozumianej wówczas jako „sieć semantyczna”. Po latach „trójka” ponownie wchodzi w modę jako szlagwort, tym razem jednak Web 3.0 jest nieco inaczej rozumiany.

Nowy sens tego pojęcia proponuje założyciel infrastruktury łańcuchów blokowych Polkadot i współtwórca kryptowaluty Ethereum, Gavin Wood. Jak łatwo się domyślić po tym, kim jest proponent, w nowej odsłonie Web 3.0 tym razem ma mieć coś wspólnego z blockchainem i kryptowalutami. Sam Wood określa nową sieć jako bardziej otwartą i bezpieczną. Web 3.0 miałby być rządzony nie w sposób scentralizowany, przez garstkę autorytetów i jak ma to w coraz większym stopniu w praktyce miejsce – przez monopolistów Big Tech, ale raczej przez demokratyczną i samorządną społeczność internetową.

„Obecny internet coraz silniej sprowadza się do danych, które użytkownicy generują”, mówi Wood w podcaście pt. „Third Web” nagrany w 2019 roku. Jak uważa, dziś startupy z Doliny Krzemowej są finansowane na podstawie ich umiejętności skutecznego zbierania danych. Na niektórych platformach prawie każde działanie podejmowane przez użytkownika jest rejestrowane. „Może to służyć tylko ukierunkowanym reklamom, ale dane mogą być też użyte do innych rzeczy”, przestrzega Wood. „Do przewidywania poglądów i zachowań ludzi, w tym także wyników wyborów”. Ostatecznie prowadzi to do pełnej kontroli w totalitarnym stylu, wnioskuje Wood.

W zamian proponuje otwarty, pozbawiony inwizacji, wolny i demokratyczny internet, w którym decydują użytkownicy sieci, a nie wielkie korporacje.



1. Generacje internetu

Koronnym projektem wspieranym przez Web3 Foundation Wooda jest Polkadot (2), organizacja non-profit z siedzibą w Szwajcarii. Polkadot jest protokołem zdecentralizowanym, opartym na technice łańcuchów blokowych (3), który pozwala na połączenie blockchain z innymi rozwiązaniami w celu wymiany informacji i transakcji w całkowicie bezpieczny sposób. Łączy łańcuchy blokowe, zarówno publiczne, jak i prywatne oraz inne technologie. Zaprojektowano go w czterech poziomach: główny łańcuch blokowy, zwany Relay Chain (łańcuch przekaźnikowy), który łączy różne łańcuchy blokowe i ułatwia wymianę między nimi, parachains (proste łańcuchy blokowe), które składają się na sieć Polkadot, parathreads lub parachains typu pay-per-use i w końcu „mosty”, czyli łączniki niezależnych łańcuchów blokowych. Sieć Polkadot ma poprawiać interoperacyjność, zwiększać skalowalność i bezpieczeństwo hostowanych łańcuchów blokowych. W ciągu niespełna roku Polkadot uruchomiła ponad 350 aplikacji.

Głównym łańcuchem blokowym Polkadot jest łańcuch przekaźnikowy. Łączy on różne parachains i ułatwia wymianę danych, aktywów i transakcje.

Proste łańcuchy parachain biegą równolegle do głównego łańcucha blokowego Polkadot lub do łańcucha przekąźnikowego. Mogą być bardzo różne od siebie, w strukturze, systemie rządzenia, tokenów, itp. Parachains pozwalają również na równoległe transakcje i sprawiają, że Polkadot jest systemem skalowalnym i bezpiecznym.

Zdaniem Wooda, ten system można przenieść na sieć rozumianą bardziej ogólnie niż tylko zarządzanie kryptowalutami. Powstaje internet, w którym to użytkownicy indywidualnie i zbiorowo mają pełną kontrolę nad wszystkim, co dzieje się w systemie.

Od prostego odczytu stron do „tokenomiki”

Web 1.0 był pierwszą implementacją sieci. Jak się przyjmuje, trwał od 1989 do 2005 roku. Można tę wersję zdefiniować jako sieć połączeń informacyjnych. Według twórcy World Wide Web, Tima Bernersa-Lee, była to wówczas sieć „tylko do odczytu”. Zapewniała ona bardzo niewielką interakcję, gdzie można było wspólnie wymieniać informacje, ale nie było prawdziwej możliwości interakcji ze stroną WWW. W przestrzeni informacyjnej obiekty zainteresowania określano jako zasoby identyfikowane przez Uniform Resource Identifiers (URI; URIs). Wszystko było statyczne. Można było poczytać i nic więcej. Był to model biblioteki.

Druga generacja internetu, znana jako Web 2.0, została po raz pierwszy zdefiniowana przez Dale’a Dougherty’ego w 2004 roku jako sieć typu read-write. Strony w technologii Web 2.0 pozwalały na gromadzenie i zarządzanie globalnymi grupami



2. Gavin Wood i logo Polkadot

o wspólnych zainteresowaniach, a medium oferowało społeczną interakcję. Web 2.0 to rewolucja biznesowa w branży komputerowej spowodowana przejściem do internetu jako platformy. W tej fazie użytkownicy zaczęli tworzyć treści na platformach takich jak YouTube, Facebook czy Twitter. Ta odsłona internetu miała charakter społecznościowy i nastawiony na współpracę, ale to zwykle miało swoją cenę. Wadą tego partycypacyjnego internetu, która została uświadomiona z pewnym opóźnieniem, było to, że tworząc treści, użytkownicy udostępniali również informacje i dane osobowe firmom kontrolującym te platformy.

W tym samym czasie, gdy kształtował się Web 2.0, zaczęły się prognozy dotyczące Web 3.0. Kilkanaście lat temu powszechne było przekonanie, że będzie to tzw. sieć semantyczna. Opisy publikowane w około 2008 roku sugerowały nadejście oprogramowania intuicyjnego i inteligentnego, które wyszukiwałoby informacje dostosowane do nas znacznie lepiej, niż to wynikało ze znanych już wtedy mechanizmów personalizacyjnych. Web 3.0 miał być trzecią generacją usług internetowych, stron i aplikacji, która

3. Modelowe przedstawienie techniki łańcuchów blokowych



koncentrowałyby się na wykorzystaniu maszynowego uczenia i rozumienia danych. Ostatecznym celem Web 3.0, według wyobrażeń z drugiej połowy pierwszej dekady XXI wieku, miało być stworzenie bardziej inteligentnych, połączonych i otwartych stron internetowych. Po latach wydaje się, że cele te zostały i są realizowane, choć termin „semantic web” wyszedł z powszechnego użycia.

Dzisiejsze definiowanie trzeciej wersji internetu, opierającego się na zdecentralizowanych sieciach kryptowalutowych Bitcoin czy Ethereum, niekoniecznie jest sprzeczne z dawnymi przewidywaniami internetu semantycznego, ale kładzie nacisk na co innego, na prywatność, bezpieczeństwo i demokrację. Za kluczową innowację ostatniej dekady uznaje się stworzenie platform, których nie kontroluje żaden pojedynczy podmiot, a mimo to każdy może im zaufać. Dzieje się tak dlatego, że każdy użytkownik i operator tych sieci musi przestrzegać tego samego zestawu twardo zakodowanych zasad, znanych jako protokoły konsensusu. Drugą innowacją jest to, że sieci te pozwalają na transfer wartości lub pieniędzy pomiędzy kontami. Te dwie rzeczy – decentralizacja i internetowe pieniądze – są kluczem do współczesnego rozumienia Web 3.0.

Twórcy sieci kryptowalutowych, może nie wszyscy, ale na pewno takie postacie jak Gavin Wood, mieli świadomość, do czego zmierza ich praca. Jedną z najpopularniejszych bibliotek programistycznych używanych do pisania kodu Ethereum nosi nazwę web3.js.

Poza koncentracją na ochronie danych ważny w nowym nurcie Web3.0 jest aspekt finansowy, ekonomika nowego internetu. Pieniądże w nowej sieci, zamiast polegać na tradycyjnych platformach finansowych, które są powiązane z rządami i ograniczone granicami, są kontrolowane przez posiadaczy,

w sposób dowolny, globalny, niekontrolowany. Oznacza to również, że tokeny i kryptowaluty mogą być wykorzystywane do projektowania zupełnie nowych modeli biznesowych i gospodarek internetowych. Coraz częściej obszar ten nazywa się tokenomiką. Wczesnym i jeszcze stosunkowo skromnym przykładem jest sieć reklamy w zdecentralizowanej sieci, która nie musi polegać na sprzedaży danych użytkownikom reklamodawcom, ale polega na wynagradzaniu użytkowników tokenem za oglądanie reklam. Tego typu aplikacja Web 3.0 jest rozwijana w środowisku przeglądarki Brave i ekosystemie finansowym Basic Attention Token (BAT).

Aby wizja Web 3.0 stała się rzeczywistością aplikacji tego rodzaju i wszelkich innych wywodzących się z tego środowiska, musi ich używać znacznie więcej ludzi. Aby tak się stało, aplikacje te muszą być znacznie bardziej czytelne, zrozumiałe dla osób, które nie wywodzą się z programistycznych kręgów. W tej chwili nie można powiedzieć, aby tokenomika była jasna w rozumieniu masowym.

Chętnie cytowany „ojciec WWW”, Tim Berners-Lee, zauważył kiedyś, że Web 3.0 jest w pewnym sensie powrotem do Web 1.0. Dlatego, że nie jest w nim potrzebne żadne pozwolenie od „centralnego organu”, by cokolwiek publikować, zamieszczać, cokolwiek robić, nie ma węzła kontroli, nie ma pojedynczego punktu nadzorowania i... nie ma wyłącznika.

Z tym nowym demokratycznym, wolnym, nie-nadzorowanym Web 3.0 jest tylko jeden problem. Na razie korzystają z niego i chcą go użytkować tylko ograniczone kręgi. Większości użytkowników zdaje się wystarczać wygodny i łatwy w użyciu Web 2.0 w swojej doprowadzonej obecnie do wysokiego poziomu doskonałości technicznej formie. ■

Mirosław Usidus

Kurza mama. Jak ogarnąć kury, dzieci i życie na wsi

Aleksandra Wołkowska

Wydawnictwo Kobiectwo, cena: 44,90 zł

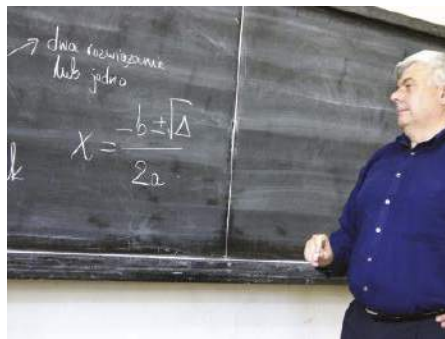
Pełna chata, szczęśliwe kurki i życie w zgodzie z naturą, czyli tak, #jakiubimy! Dzień zaczyna od wystrojenia się w sukienkę. Olewa makijaż i rozczochraną grzywkę. Sprawdza, czy każdy z pięciu synów zjadł śniadanie, i z uśmiechem na twarzy wyrusza do swoich „dziewczynek”. Ola Wołkowska – pani domu, który jest w wiecznym remoncie, żona stolarza, mama gromadki dzieci i opiekunka wesołego kurnika, z którego codziennie nadaje relację na swoim Instagramie, zyskując coraz większe grono obserwatorów. Nigdy nie planowała hodowli kur, bez której teraz nie wyobraża sobie życia. W swojej książce z niespożytą energią i poczuciem humoru dzieli się zdobytą wiedzą i przekonuje, że życie poza wielkim miastem jest pełne uroku. Podpowiada między innymi, od czego zacząć swoją przygodę z założeniem kurnika, ile kosztuje utrzymanie niosek i kiedy można cieszyć się pierwszymi jajkami. Pokazuje również, że perfekcyjna pani domu to ta, która chodzi w ubłoconych kałozach, a w salonie co rusz następuje na kłoczek Lego.



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 (przynaję, 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”.



Matematyka i rower

Być może Czytelnicy pamiętają, że dwa poprzednie odcinki poświęcone były liczbie 21, a sprowokowało mnie do tego pewne zadanie z matury w obecnym, dwudziestym pierwszym roku dwudziestego pierwszego wieku.

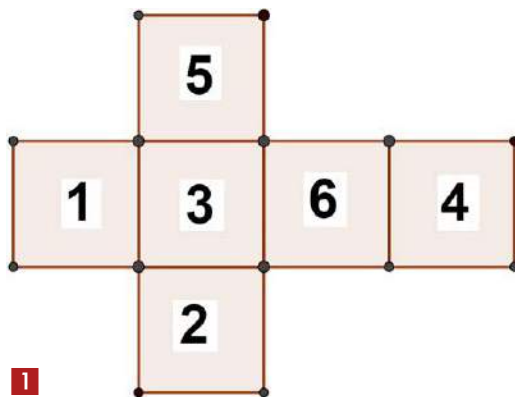
Na kostce do gry liczby rozmieszczone są tak, by suma na przeciwległych ściankach była zawsze równa 7. Dodajmy: $1+2+3+4+5+6=21$. Takie liczby (będące sumami kolejnych liczb naturalnych) nazywamy trójkątnymi, a więc 21 jest liczbą trójkątną (1).

Po inne szczegóły odsyłam do poprzednich numerów, a jeszcze tylko trzy ciekawostki sportowe, dotyczące liczby 21. W siatkówce plażowej set polega na zdobyciu 21 punktów, a koszykówce 3×3 zdobycie 21 punktów przez jedną z drużyn oznacza koniec meczu przed czasem. W Tour de France jest zwykle 21 etapów.

No, właśnie – rower jest wdzięcznym gadżetem w szkolnych zadaniach matematycznych. Zacznę od zadania, sprawiającego kłopoty uczniom i studentom, nawet studentom informatyki, z którymi od lat mam do czynienia. Nieco złośliwie powiem, że nie było studenta... no, *pewnego wydziału* na UW, który by w pełni zrozumiał rozwiązanie. „Jak mogę pamiętać coś z matematyki, kiedy maturę zdawałem półtora roku temu?”

Zadanie 1. Wybrałem się na dwugodzinną okrężną przejażdżkę rowerową. Przez pierwszą godzinę pedałowałem dzielnie z prędkością 20 km/h, ale drugą wlokłem się tylko z prędkością 10 km/h. Cóż, jestem starszym panem. Jaka była moja średnia na całej trasie?

Odpowiedź wydaje się – i jest – bardzo prosta. Średnia to średnia. Dodajemy i dzielimy przez 2; $\frac{20+10}{2}=15$. Piętnaście kilometrów na godzinę.



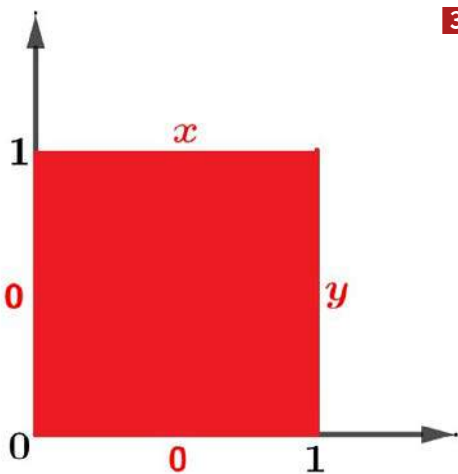
Następnego dnia wybrałem się na przejażdżkę po tym samym lesie, ale z wybranym punktem docelowym. Chciałem dojechać do pewnej znanej mi ładnej polany. Było do niej 20 kilometrów. Jak poprzednio, w tamtą stronę udało mi się utrzymać prędkość 20 km/h, ale zmęczyłem się tak, że w drodze powrotnej moja średnia spadła do 10 km/h. Jaka była moja prędkość na całej trasie?

Jak to jaka? – dziwili się zawsze (!) uczniowie i studenci. Przecież to to samo. Pół drogi z taką prędkością, pół z taką. Średnia to średnia. Dodajemy i dzielimy przez dwa, bo przecież „połowy” są równe. Piętnaście!

Hm. Spójrzmy. Moja wycieczka drugiego dnia trwała trzy godziny (godzinę w jedną i dwie

$$H(a,b) = \frac{1}{\frac{1}{a} + \frac{1}{b}}$$
$$\frac{1}{\frac{1}{a} + \frac{1}{b}} = \frac{2ab}{a+b}$$
$$\frac{400}{30} = \frac{40}{3} = 13\frac{1}{3}.$$

The diagram illustrates a redshifted signal in a curved spacetime geometry. It features two concentric circles on the left, representing different gravitational potentials. The inner circle has a center labeled A and a point K on its circumference. The outer circle has a point M on its circumference. A dashed blue line connects A and M , with a blue '4' at A and a blue '2' at M . A red line connects K and M . On the right, there is a smaller circle with a center labeled B and points L and N on its circumference. A dashed blue line connects B and N , with a blue '1' at B . Red lines connect K to L and M to N . Green arrows indicate the direction of signal propagation from the left towards the right. A green arrow at the bottom is labeled with a green '?'. A red box with the number '2' is in the top right corner.



3

uproszczenia, że nie jest ona zmienna w czasie. Inaczej mówiąc, na obwodzie rozkład temperatury jest taki sam. Jaka jest temperatura w punktach wnętrza kwadratu? Na rysunku 3 lewa i dolna krawędź jest utrzymywana w temperaturze 0, a na pozostałych dwóch bokach rośnie proporcjonalnie.

Napiszę stosowne równanie, zdając sobie sprawę, że może być niezrozumiałe. Dwadzieścia pięć lat temu zrozumiałby to każdy uczeń liceum. Równanie to jest najsłynniejszym równaniem różniczkowym cząstkowym. Nazywane jest równaniem Laplace'a, na cześć markiza Pierre'a Simona de LaPlace (1749–1827). Kilka słów o tym wybitnym uczonym, nazywanym francuskim Newtonem. Urodził się w niezamożnej rodzinie, ale poznano się na jego zdolnościach do matematyki i fizyki i wobec tego zdobył porządne wykształcenie. Wniósł potem olbrzymi wkład do matematyki, mechaniki i fizyki. Wystarczy wspomnieć, że obliczył (tak, obliczył!) masę Saturna. Otrzymał wynik, różniący się od współczesnych pomiarów o jeden procent!

Współpracował z chemikiem Lavoisierem, który mocno zaangażował się w rewolucję francuską i po dojściu do władzy kolejnej frakcji rewolucyjnej został zgilotynowany, a wyrok uzasadniono, że „rewolucja nie potrzebuje uczonych”. Może Laplace był zbyt wielki, może bardziej ostrożny. Został w 1806 roku (a więc za Napoleona) księciem, a w 1813 (czyli już po restauracji Bourbonów) markizem.

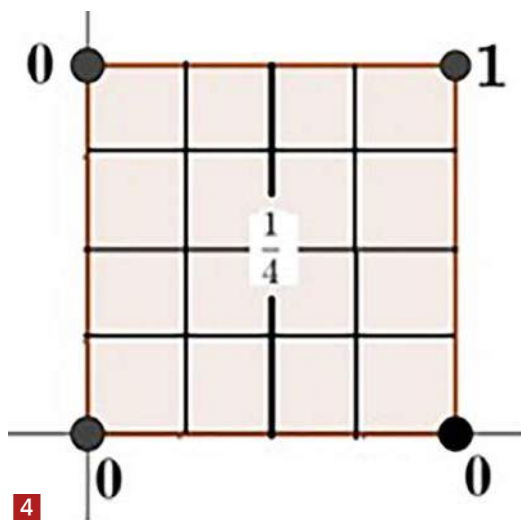
W fizyce i filozofii znany jest tak zwany demon Laplace'a. Jest to hipotetyczna istota, która zna położenia i prędkości wszystkich punktów materialnych świata. Ponieważ wszystko opisane jest i tak układem równań różniczkowych, taka istota może przewidzieć całą przyszłość. Wystarczy, by mogła rozwiązać

ten układ. Tak czy owak – argumentował Laplace – przyszłość świata jest zdeterminowana. Co ma być, to będzie. Nie ma możliwości zmiany, bo wszystko to układ równań różniczkowych. Procesy chemiczne (a więc i biologiczne) też. Nie wiem, czy to dobrze, czy źle, ale współczesna mechanika kwantowa powiada, że nie można znać jednocześnie położenia i pędu cząstki, a więc nic nie jest zdeterminowane. Ale demon Laplace'a odżywa w postaci poglądu, że może cały nasz Wszechświat to po prostu wielki komputer, który sam oblicza swoją przyszłość... Na pewno jednak Anna Kiesenhofer, kręcąc pedałami do mety olimpijskiej, nie myślała o swoich równaniach cząstkowych ani o tym, że według Laplace'a jej wygrana była i tak zapisana w Księdze Przeznaczenia (tzn. w układzie równań różniczkowych).

Wchodząc na kolejne piętro dygresji, wspomnę kilku dawnych mistrzów sportu, którzy mieli porządne wyższe wykształcenie. „Dawnych” – bo teraz dyplom zdobywa się nieporównanie łatwiej..., ale to już temat na inny artykuł. Było wśród tych sportowców wielu lekarzy. Był nim na przykład Roger Bannister (1929–2018, pierwszy człowiek, który przebiegł milę poniżej 4 minut, a było to 6 maja 1954; obecny rekord jest zaledwie o 16 sekund lepszy, 3.43,13 i jest niepobity od 1999 roku). Lekarzem był Stefan Lewandowski (1930–2007, najlepszy średniodystansowiec w Polsce w latach pięćdziesiątych XX wieku)... i inny Lewandowski (Zbigniew, 1930–2010), który był pierwszym Polakiem, który skoczył 2 metry wzwyż (12 maja 1957). Tak, rekord przedwojenny należał do Zbigniewa Pławczyka (1,96). Widocznie *Lewandowski* to takie sportowe nazwisko. Pierwsza zdobywczyni złotego medalu dla Polski po wojnie, Elżbieta Duńska-Krzesińska (1934–2015, skok w dal, Melbourne 1956, 6 m 35 cm), była dentystką. Uzupełniń tę listę bokserem. Kazimierz Paździor (1935–2010, złoty medal w Rzymie w 1960 roku) był magistrem ergonomii. Przyjmujemy powszechnie, że matura sprzed 1939 roku to jakby magisterium z lat mniej więcej 1950–65 i habilitacja dzisiaj. Uniwersyteckie wykształcenie ekonomiczne miała też najsłynniejsza polska lekkoatletka, Irena Szewińska (1946–2018) – widywałem ją na uczelni, bo studiowaliśmy w tych samych latach. Na ogół biegała po schodach. Była młodsza ode mnie o półtora miesiąca.

Wielu matematyków było dobrymi alpinistami, ale nie mam informacji, by w klasycznych sportach osiągnęli wybitne wyniki. Dopiero Anna Kiesenhofer. W prasie pisano potem, że sobie wszystko wyliczyła...

No, to wreszcie trochę matematyki, wróćmy do rysunku 3. Proszę się nie stresować, gdy przez chwilę

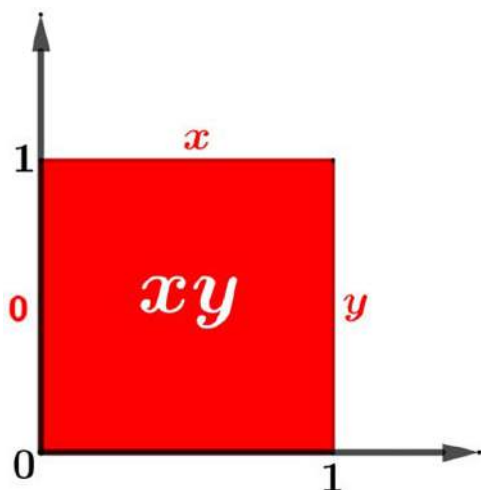


nie będą Państwo rozumieć. Po chwili będzie znów nietrudno. A równanie Laplace'a nawet wygląda ładnie. Po lewej stronie mamy właśnie laplasjan, operator Laplace'a.

$$\frac{\partial^2 u}{\partial x^2} + \frac{\partial^2 u}{\partial y^2} = 0$$

Taki wzór jest zrozumiały dla każdego z Czytelników, znających rachunek różniczkowy. Dla nich również pozostanie kwestią prostego sprawdzenia, że seria rozwiązań to funkcje $1, x, y, x^2 - y^2, 2xy, x^3 - 3xy^2, 3x^2y - y^3$, a ogólnym rozwiązaniem wielomianowym są: część rzeczywista i część urojona wyrażenia dla kolejnych $n=0, 1, 2, 3, 4, \dots$

Ale można to równanie rozwiązać i tak. Temperatura w środku kwadratu jest średnią arytmetyczną



5. Rozwiązanie równania Laplace'a z podanymi warunkami brzegowymi

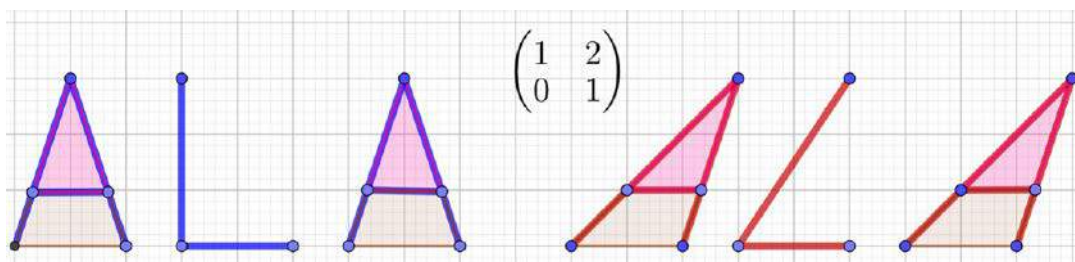
temperatur w wierzchołkach. Zatem w środku kwadratu (3×4) mamy temperaturę $\frac{1}{4}$, średnia liczb 0, 0, 0, 1. Dzieliąc dalej kwadrat na mniejsze oczka, wyznaczymy ją w każdym punkcie gęstej sieci – ale to tak, jakby w każdym punkcie.

Dojdziemy do wniosku, że rozwiązaniem zadania, opisywanego rysunkiem 2, jest funkcja $u=xy$ (5). Dla innych warunków brzegowych rozwiązanie może być bardzo skomplikowane, ale da się je otrzymać z wielomianów, które napisałem powyżej. Potem sprawa się bardzo komplikuje, wchodzi analiza harmoniczna, szeregi Fouriera i w ogóle matematyka z wysokiej półki, a wszystko ma zdecydowany kontekst praktyczny – fizyczny.

W pracach Anny Kiesenhofer pojawia się często geometria symplektyczna. Samo słowo „symplektyczny” jest ciekawe. Utworzył je jeden z najinteligentniejszych matematyków swoich czasów, Hermann Weyl (1885–1955). W słowie „complex”, co po łacinie może znaczyć „splacony razem” (mamy przecież w naszym kraju „szpitale zespolone”), zamieniamy łacińskie „co-” na greckie „sym-” i mamy sympleks, po angielsku symplectic, co wróciło do naszego języka jako „symplektyczny”. O co jednak chodzi? Można to przybliżyć tak. W tradycyjnej geometrii ważne są odległości punktów, kąty między liniami i pola figur. Przesunięcie i obrót nie zmienia cech figury: zachowuje ona ten sam obwód, te same kąty i takie samo pole. Ale są transformacje, które zmieniają i odległości, i kąty – ale nie zmieniają pola.

Na rysunku 6 widzimy odwzorowanie zachowujące pola figur. To zwykła kursywa, *italiki*. W stosunku do prostego druku górna część liter jest przesunięta o dwie jednostki (stąd liczba 2 w tablicy), dół się nie zmienia, a środkowe części przesuwają się proporcjonalnie.

W 1875 roku matematyk niemiecki Felix Klein sformułował tak zwany program z Erlangen. Narzucił on nam nowy topos geometrii. *Topos* to mniej więcej sposób, w jaki mamy myśleć o czymś, przez jakie okulary mamy spoglądać na świat. Od czasów Kleina mówimy, że geometria to badanie niezmienników grupy przekształceń. Przesunięcia i obroty zachowują długości, przekształcenia o wyznaczniku 1 (cokolwiek to znaczy) nie zmieniają pól, jeszcze inne nie zmieniają, powiedzmy, kątów. Znamy takie odwzorowania. Przypomnijmy sobie mapy w rzucie Mercatora. Ceną za wierność jest deformacja okolic podbiegunowych – Grenlandia ma rozmiar Afryki, a Antarktyda ciągnie się przez całą mapę. Ale i taka mapa ma swoje zalety. W geometrii symplektycznej nie zmienia się jeszcze coś innego (a mianowicie



6. Pisanie kursywą ma coś wspólnego z geometrią symplektyczną!

forma różniczkowa skośnie symetryczna), ale to już temat na oddzielny artykuł.

Medal olimpijski jest okrągły jak koło opisane nierównością $x^2 + y^2 \leq r^2$, suma kwadratów nie większa niż

kwadrat promienia. W geometrii symplektycznej często występuje różnica kwadratów $x^2 - y^2$, a pani Annie Kiesenhofer potrzebny jest jeszcze jeden medal, żeby obydwa razem przypominały jej koła rowerowe. ■

REKLAMA

KITY AVT PRZEDSTAWIA

AVTEDU

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

Poznaj całą serię

#AVTEDU #NaukaLutowania #KityAVT

KITY
AVT



Obrońca stali

Uzupełnieniem cyklu artykułów dotyczących korozji i sposobów walki z nią jest bliższe przedstawienie metalu będącego wiernym obrońcą stali. Z przedwakacyjnych odcinków wiesz już, że chroni on stopy żelaza, nim sam nie ulegnie całkowitemu zniszczeniu. Dzisiaj przekonasz się, jakie właściwości cynku to umożliwiają.

Tysiące lat temu produkowano stop zawierający cynk, nie znając tego metalu. Przez setki lat wytopiano ołów z rudy zawierającej pokaźny dodatek cynku i zupełnie nie zdawano sobie sprawy z obecności domieszki. W obecnych czasach takie sytuacje są niemożliwe, ale przez wieki cynk skutecznie ukrywał się przed okiem chemików i metalurgów.

Od historii...

Jeszcze przed kilkoma wiekami metalurzy pracowali metodą prób i błędów, a rudy oceniano „na oko” (parafrazując poetę, brakowało odpowiedniego „szkiełka”), stąd i zdarzające się niespodzianki. W starożytności jedną z nich było otrzymanie mosiądzu (1). Podczas wytopu znanej od niepamiętnych czasów miedzi do jej rudy dostał się minerał, zwany ówczesznie kadmią, a obecnie **galmanem** (jest to węglan cynku ZnCO₃). W procesie wytopu oba metale łączyły się ze sobą, dając żółtej barwy stop. Ze względu na podobieństwo do złota, używano go jako imitacji najcenniejszego metalu.

W późniejszych czasach mosiądz wytwarzano, ogrzewając w zamkniętym piecu miedź z galmanem, ale nadal nie zdawano sobie sprawy z obecności

cynku w stopie (mosiądz uważany był za odmianę miedzi). Zaobserwowano jednak gromadzenie się białego nalotu na ścianach pieców, a na ich dnie – szarego pyłu. Obecnie wiemy, że nalotem był tlenek cynku ZnO, a pyłem – cynk metaliczny.

Cynk otrzymano po raz pierwszy w początkach naszej ery w Chinach lub Indiach. Stało się tak z powodu odmiennej konstrukcji pieców do wytopu i osobiwej właściwości cynku. Do redukcji jego tlenku potrzebna jest temperatura przekraczająca 1200°C, ale cynk wrze już przy nieco ponad 900°C, co oznacza, że od razu po otrzymaniu przechodzi w stan pary. Jeżeli temperatura ścian pieca nie przekracza 420°C, cynk zestala się z pominięciem fazy ciekłej w szary proszek, natomiast przy wyższych temperaturach skrapla w ciecz, która krzepnie w bryłę srebrzystego metalu. Aż do połowy XVIII wieku cała światowa produkcja cynku pochodziła z Dalekiego Wschodu. W Europie cynk pojawił się jako orientalna ciekawostka w roku 1595, a technologię produkcji poznano w roku 1735. Nazwa **cynk** pochodzi z języka niemieckiego, ale jej znaczenie nie jest do końca jasne.

Rudy cynku i ołowiu często występują razem, tak jest i w Zagłębiu Górnośląskim. Już w średniowieczu na tamtych terenach wytopiano ołów, ale żadne kroniki nie wspominają o cynku. Powodem był fakt, że ołów wytopiano w piecach otwartych, więc cynk ulatniał się do atmosfery. W roku 1798 w górnośląskich hutach (w wiekach XVIII i XIX rejon ten był największym cynkowym zagłębiem świata) opracowano technologię produkcji cynku, która z niewielkimi modyfikacjami jest stosowana do dziś. W następnym stuleciu użyto cynku do produkcji ogniw galwanicznych, jako składnika wielu stopów, a także do ochrony stali przed korozją. Spowodowało to wzrost zapotrzebowania na cynk, a w konsekwencji poszukiwanie nowych złóż oraz budowę hut tego metalu.

...do współczesności

Źródłem cynku jest galman lub minerały siarczkowe. Pierwszy etap przeróbki to prażenie rud w celu

1. Złotm mosiądzu to cenny surowiec wtórny



otrzymania tlenku. Ten poddaje się redukcji koksem w piecach, w których pary cynku skraplają się w odbieralnikach, zwanych mufłami (stąd nazwa: **piec muflowy** (2)). Surowy cynk jest oczyszczany przez destylację, przy okazji otrzymuje się kadm stanowiący domieszkę w rudach. Obecnie coraz częściej stosowany jest inny sposób produkcji: rudy cynku rozpuszcza się w kwasie siarkowym i prowadzi elektrolizę roztworu. Metoda ta pozwala otrzymać od razu bardzo czysty metal i jest przyjazna dla środowiska (nie ma emisji dwutlenku siarki), ale kosztowna. Cynk należy do metali produkowanych w największych ilościach. W roku 2020 wytworzono około 12 milionów ton cynku (4. miejsce po żelazie, aluminium i miedzi). Cynk zajmuje 29. miejsce na liście rozpowszechnienia pierwiastków, a powierzchniowa warstwa Ziemi zawiera 0,004% tego metalu.

Około połowy światowej produkcji zużywanej jest do pokrywania wyrobów stalowych w celu zabezpieczenia ich przed korozją, prawie całą resztę pochłaniają stopy. Nadal powszechnie stosowany jest mosiądz. Tombak to również stop miedzi i cynku, od wieków doskonale imitujący złoto. Stop cynku, miedzi i niklu to z kolei zamiennik srebra (nowe srebro). Ze stopu z gliną (zinal) produkuje się części samochodowe i armaturę łazienkową.

Cynk nadal jest też używany do konstrukcji ogniw jednorazowych. Zastosowanie to traci na znaczeniu (coraz powszechniejsze są akumulatory), ale dla Ciebie jest nader istotne. Z tego źródła łatwo uzyskasz bardzo czysty cynk do doświadczeń (patrz: twoje źródło cynku).

Ze związków cynku w największych ilościach otrzymywany jest tlenek, od dawna będący składnikiem farmaceutyków, wypełniaczem przy produkcji tworzyw sztucznych oraz białym pigmentem (biel cynkowa). Chlorek stosuje się do produkcji węgla aktywnego, konserwacji drewna oraz lutowania metali (usuwa tworzące się na ich powierzchni tlenki). Do impregnacji i zabezpieczenia drewna przed pleśnią stosowany jest siarczan cynku.

Cynk to również jeden z mikroelementów, aktywujący liczne enzymy. Przyspiesza gojenie się ran i wspomaga układ odpornościowy.

Analityka bez fajerwerków

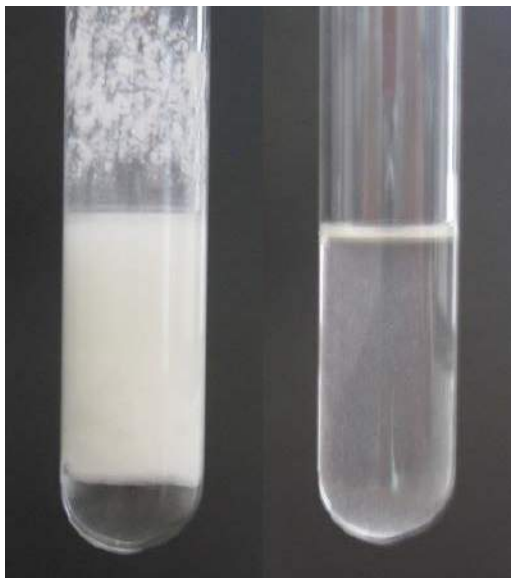
Kation Zn^{2+} nie zapewni ci obserwacji barwnych przemian, wodorotlenek i sole mają biały kolor (o ile anion nie jest zabarwiony (3)). Jak w przypadku wielu metali ciężkich, fosforan i węglan są nierozpuszczalne w wodzie. Wodorotlenki wytrącają z roztworów soli cynku biały osad, łatwo rozpuszczalny



2. Dawnej konstrukcji piec muflowy do wytopu cynku (Górny Śląsk, 1930)



3. Związki cynku mają zazwyczaj białą barwę (Wikimedia/Ondřej Mangl)



4. Biały osad wodorotlenku cynku (po lewej) jest rozpuszczalny w nadmiarze zasady (po prawej)

w nadmiarze odczynnika. To cecha wyróżniająca **metal amfoteryczny**, czyli taki, którego związki (i on sam również) reaguje zarówno z kwasami, jak i zasadami (4). W przypadku reakcji z roztworem amoniaku także dochodzi do wytrącenia białego osadu wodorotlenku, który jednak szybko rozpuszcza się w nadmiarze odczynnika. Ta ostatnia reakcja pozwala odróżnić cynk od glinu. Kationy Al^{3+} , podobnie jak sole cynku, dają biały osad z wodorotlenkami, rozpuszczalny w ich nadmiarze. Natomiast osad $Al(OH)_3$ powstały w reakcji z roztworem amoniaku jest tylko nieznacznie rozpuszczalny po dodaniu większej ilości odczynnika.

Próby analityczne pozostawiam ci do samodzielnego wykonania, a teraz pomyśl...

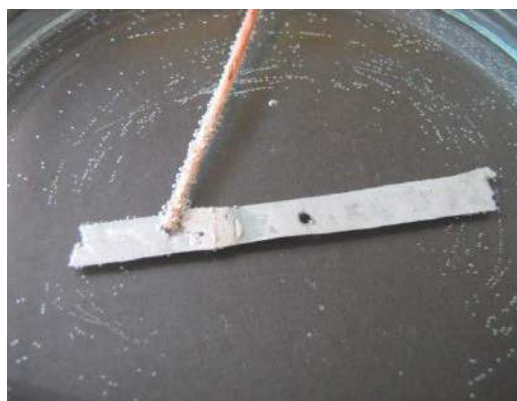
...jak rozpuścić cynk?

– To proste, wrzucam cynk do kwasu, a ten – jako metal nieszlachetny – wkrótce się w nim rozpuści.
– odpowiesz zapewne. Cynk to rzeczywiście metal aktywny, wypierający wodór z kwasów i powinien szybko utworzyć odpowiednią sól. W rzeczywistości jednak reakcja nie przebiega tak gładko.

Cynkową blaszkę wrzuc do zlewki z 5–10% kwasem solnym. Na powierzchni metalu powoli zaczynają się wydzielać pęcherzyki gazu. Reakcja przebiega bardzo opornie, nie takiej szybkości się spodziewałeś. Dodaj kilka kropli roztworu siarczanu(VI) miedzi(II) $CuSO_4$. Po chwili reakcja wyraźnie przyspiesza. Zauważ, że na powierzchni cynku osadziła się metaliczna miedź (blaszka pociemniała). To cynk – metal aktywniejszy – wyparł miedź z roztworu jej soli. Dlaczego jednak wzrosła szybkość reakcji?

Patent na długie życie

Na szalkę Petriego (jeśli jej nie masz, użyj głębszego spodka lub dużej zakrętki od słoika) nalej 5–10% kwasu solnego i wrzuc blaszkę cynkową. Ilość



5. Cynk w roztworze HCl – wodór wydziela się głównie na miedzianym drucie



6. Cynk w roztworze NaOH – wodór wydziela się głównie na stalowym gwóźdźu

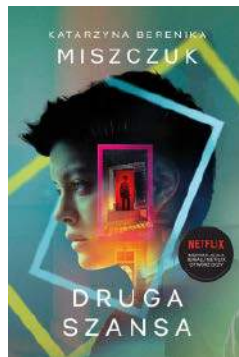
roztworu dobierz tak, aby blaszka całkowicie znajdowała się pod jego powierzchnią. Po wykonaniu poprzedniej próby zapewne nie dziwi cię już powolne wydzielanie wodoru. Do roztworu włóż kawałek miedzianego drutu, ale nie stykaj metali ze sobą. Nic się nie zmieniło: miedź nie reaguje z kwasem

Druga szansa

Katarzyna Berenika Miszczuk

Wydawnictwo Uroboros, liczba stron: 336, cena: 39,99 zł

Czy największe koszmary rodzą się rzeczywiście w naszych umysłach? A może najgorsze czai się tuż za plecami? Julia budzi się w tajemniczym szpitalu. Nie poznaje własnego odbicia w lustrze, nie pamięta, jak się tu znalazła. Z czasem dowiaduje się, że cała jej rodzina zginęła w pożarze. Jedynie Julii udało się przeżyć, choć na skutek odniesionych obrażeń straciła pamięć. Nazwa ośrodka, Druga Szansa, powinna dawać pacjentom nadzieję... Co jednak myśleć o kobiecie, która wróży Julii śmierć, pojawiającej się nagle nieznanym dziewczynie i szeptach rozbrzmiewających dokoła? Dzieje się tu coś dziwnego i trudnego do wyjaśnienia. Julia jest coraz bardziej zagubiona i przerażona, a leczenie nie przynosi pożądanych efektów. Co zrobi w sytuacji, w której traci zaufanie nawet do samej siebie?



Twoje źródło cynku...

...to jednorazowe ogniwa Leclanchego. Załóż rękawice i nożycami rozetnij powłokę. Operację wykonaj na tacy, aby nie pobrudzić wszystkiego zawartością baterii. Wyprostuj blachę z cynkowego kubeczka i oczyść ją papierem ściernym. Tym sposobem dysponujesz bardzo czystym metalicznym cynkiem. Pozostałości po rozbiórce wrzuć do pojemnika przeznaczonego na zużyte baterie. Pamiętaj, aby użyć ogniwa z oznaczeniem zaczynającym się literą R (np. R6, R20). Nie rozbieraj akumulatorów ani ogniwa alkalicznych – nie znajdziesz w nich cynkowej blachy.



7. Twoje źródło cynku: ogniwo R14 (z lewej po zdjęciu osłony widoczny jest cynkowy kubek)

solnym, a cynk – opornie. Szklaną bagietką prześnij drut tak, aby zetknął się z cynkiem. Różnicę widać od razu: gaz wydziela się intensywnie, ale na powierzchni miedzianego drucika (5). Jednak miedź przecież nie reaguje z roztworem HCl (brak niebieskiego zabarwienia, charakterystycznego dla soli miedzi), a ponadto – co stwierdzisz po pewnym czasie – ubywa cynku. Czary?

Rozłącz metale – reakcja cynku z kwasem znowu przebiega bardzo powoli. Dotknij teraz blaszki cynkowej grafitowym pręcikiem (np. z ołówka). Reakcja przyspiesza, a wodór wydziela się na graficie. Znowu czary?

Przygotuj 5% roztwór wodorotlenku sodu NaOH i nalej go na szalkę. Wrzuć cynkową blaszkę. Cynk jest metalem amfoterycznym i reaguje także z zasadami, ale wodór znowu wydziela się bardzo powoli. Do roztworu włóż stalowy gwóźdź. Żelazo nie ulega działaniu zasad, zatem nie zauważysz zmian, ale gdy zetniesz metale, na jego powierzchni gwałtownie zacznie wydzielać się gaz (6).

Po raz trzeci zapewne pomyślałeś o czarach, lecz to tylko do głosu doszło zjawisko powszechne w elektrochemii – **nad napięcie wydzielania wodoru**. Tym terminem określa się wzrost napięcia

potrzebnego do wydzielenia wodoru na powierzchni metalu ponad wartość teoretyczną. Jest ono zależne od wielu czynników, ale najważniejszy to rodzaj metalu. Cynk akurat należy do metali o wysokim nad napięciu (największe ma rtęć), co powoduje powolną reakcję z kwasami i zasadami. Po zetknięciu cynku z innym metalem lub grafitem obie substancje uzyskują ten sam potencjał względem roztworu, a wodór wydziela się w miejscu o niższym nad napięciu. W doświadczeniu z rozpuszczaniem cynku z użyciem soli miedzi zadziałał ten sam mechanizm: wodór wydzielal się na metalicznej miedzi osadzonej na powierzchni cynku. Zauważ, że w każdym przypadku powstawały ogniwa galwaniczne.

Jaki jest zatem patent cynku na długie życie? Po pierwsze wysoka wartość nad napięcia, po drugie – czystość, dzięki której nie tworzą się lokalne ogniwa ułatwiające rozpuszczanie metalu. O ile pierwsza cechę można nazwać wrodzoną, o tyle druga musi zostać nabyta w procesie otrzymywania i oczyszczania cynku. Trud i poniesione koszty opłacają się jednak sowicie. Dzięki nim cynk chroni przed zniszczeniem znaczne ilości produkowanej stali. ■

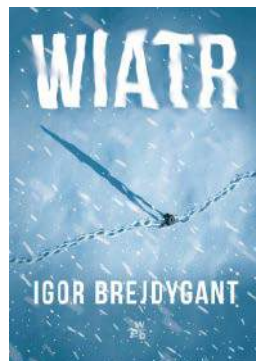
Krzysztof Orliński

Wiatr

Igor Brejdygant

Wydawnictwo Wydawnictwo W.A.B., liczba stron: 416, cena: 34,99 zł

Wczesna wiosna. W Tatrach leży jeszcze sporo śniegu, wieje halny. Dwójka ludzi wyrusza z Morskiego Oka na przełęcz Szpiglasową. Nagle spod grani Miedzianego schodzi lawina. Akcja ratunkowa TOPR nie przynosi rezultatu. Ciała tych dwojga zostały najprawdopodobniej z potężną siłą wtłoczone pod łód stawu. Świadkiem tragedii był Wawrzyniec, który jest przekonany, że chwilę przed zejściem lawiny, usłyszał głuchy dźwięk eksplozji. To punkt wyjścia historii, w której surowa przyroda Tatr, mroczny góralski folklor i środowisko TOPRowców są zaledwie tłem dla śniegiem i wiatrem osnutej, tajemniczej historii, w której mistrz kryminału brawurowo przekracza gatunkowe granice.





dr inż. Jan Sobótka
– nauczyciel akademicki, licencjo-
nowany instruktor i sędzia szachowy

W dniach 13-22 sierpnia 2021 roku został rozegrany w Polanicy-Zdroju 57. Międzynarodowy Festiwal Szachowy im. Akiby Rubinsteina. Od 2013 regularnie uczestniczą we wszystkich kolejnych Festiwalach, które tradycyjnie organizowane są w Polanicy Zdroju. Jest to największy i najstarszy turniej szachowy w Polsce, a w tym roku dodatkowo rozgrywany był turniej kołowy z udziałem silnych arcymistrzów z kilku krajów Europy. W turniejach szachowych Festiwalu uczestniczyło ponad 600 zawodniczek i zawodników z Polski, Belgii, Czech, Francji, Gruzji, Holandii, Litwy, Niemiec, Rosji i Ukrainy. Na zakończenie Festiwalu, na terenie parku Szachowego, odbyło się artystyczne wydarzenie w ramach projektu „Teatr w Parku – Żywe szachy wg poematu Jana Kochanowskiego”.

57. Międzynarodowy Festiwal Szachowy im. Akiby Rubinsteina

Festiwal w Polanicy-Zdroju jest najstarszą szachową imprezą w Polsce.

Akiba Rubinstein (1) urodził się 12 grudnia 1882 roku w Stawiskach koło Łomży, w rodzinie żydowskiego nauczyciela. W 1912 roku w zwyciężył w pięciu najważniejszych turniejach europejskich, czego do tej pory nie dokonał jeszcze żaden z szachistów. W 1926 r. wyjechał na stałe z Polski i osiadł w Belgii. Występował na pierwszej szachownicy w reprezentacji Polski, która zdobyła złoty (Hamburg, 1930) i srebrny (Praga, 1931) medal Olimpiady Szachowej. Zmarł 15 marca 1961 r. w Antwerpii. Jeszcze za jego życia w 1950 r. Międzynarodowa Federacja Szachowa przyznała mu za wcześniejsze osiągnięcia – tytuł arcymistrza.

Pierwszy turniej szachowy im. Akiby Rubinsteina zorganizowany został w Polanicy-Zdroju w roku 1963. Znakomita organizacja, wspaniała atmosfera memoriałów i piękno uzdrowiska sprawiły, że Polanica-Zdrój gościła wielu wybitnych szachistów, min. prezydentów FIDE – Holendra Machgielisa (Maxa) Euwe i Filipińczyka Florencio Campomane-sa, mistrza świata Michaiła Tała, w turniejach grali



1. Akiba Rubinstein, 1907/1908

mistrzowie świata Wasilij Smysłow, Anatolij Karpow i Weselin Topałow a także mistrzynie świata Maja Cziburdanidze, Nona Gaprindaszwili i Zsuzsa Polgar.

W arcymistrzowskim turnieju kołowym 57. Międzynarodowego Festiwalu Szachowego im. Akiby Rubinsteina wystąpiło 8 arcymistrzów i dwóch młodych polskich mistrzów międzynarodowych. Dzień później do rywalizacji przystąpili zawodnicy startujący w pozostałych grupach (2).

W tegorocznym głównym turnieju Festiwalu zwyciężył 19-letni arcymistrz z Ukrainy Kirill Shevchenko (3) z dorobkiem 6 pkt z dziewięciu partii. Shevchenko po pięciu rundach był czwarty, ale w kolejnym meczu pokonał prowadzącego Davida Navarę z Czech i objął prowadzenie. Reprezentant Ukrainy nie oddał już pozycji lidera, a za nim cały czas plasował się Navara, który ukończył turniej z takim samym dorobkiem jak zwycięzca. Najlepszy z Polaków Kacper Piorun zajął trzecie miejsce z wynikiem 5,5 pkt. Powody do radości mają również dwaj młodzi polscy szachiści. Szymon Gumularz (4) dzięki kolejnej trzeciej i ostatniej już normie został arcymistrzem, a z pierwszą normą arcymistrzowską i doskonałym 4 miejscem wyjechał z Polanicy Paweł Teclaf (5). Szymon Gumularz jest 56 Polakiem, który



3. Kirill Shevchenko, zwycięzca Turnieju Arcymistrzowskiego, źródło: Rubinstein Memorial: Shevchenko and Navara tie for first | ChessBase



4. Szymon Gumularz wygrywając w ostatniej rundzie Turnieju Arcymistrzowskiego zdobył ostatnią – trzecią, normę arcymistrzowską, źródło: <https://www.facebook.com/DZSzach/>



2. Banner 57. Międzynarodowego Festiwalu Szachowego im. Akiby Rubinsteina

Tabela 1. 57. Międzynarodowy Festiwal Szachowy im. Akiby Rubinsteina – Wyniki Turnieju Arcymistrzowskiego

Miejsce	Tytuł	Nazwisko Imię	Ranking	1	2	3	4	5	6	7	8	9	10	Punkty	TB1	TB2	TB3
1	GM	Shevchenko Kirill	2619	*	1	½	0	½	1	½	1	1	½	6,0	1,0	4	25,25
2	GM	Navara David	2675	0	*	1	½	½	1	½	½	1	1	6,0	0,0	4	24,00
3	GM	Piorun Kacper	2608	½	0	*	½	½	½	1	½	1	1	5,5	0,5	3	21,25
4	IM	Teclaf Paweł	2520	1	½	½	*	½	½	½	1	½	½	5,5	0,5	2	24,50
5	GM	Tomczak Jacek	2604	½	½	½	½	*	1	0	1	½	½	5,0	0,0	2	22,00
6	IM	Gumularz Szymon	2512	0	0	½	½	0	*	1	1	½	1	4,5	0,0	3	16,50
7	GM	Bluebaum Matthias	2674	½	½	0	½	1	0	*	0	½	1	4,0	0,0	2	17,25
8	GM	Stocek Jiri	2604	0	½	½	0	0	0	1	*	½	1	3,5	0,0	2	13,25
9	GM	Ovsejevitsch Sergei	2579	0	0	0	½	½	½	½	½	*	½	3,0	0,0	0	12,25
10	GM	Krasenkow Michal	2591	½	0	0	½	½	0	0	0	½	*	2,0	0,0	0	9,75



Tabela 2. Czołówka turnieju OPEN A (ELO >= 2000), 60 zawodników w tym 7 arcymistrzów

Pozycja	Kraj	Tytuł	Nazwisko, Imię	Klub	Ranking	Punkty
1.	Polska	GM	Nasuta, Grzegorz	MUKS „Stoczek 45” Białystok	2518+7	7.5
2.	Polska	IM	Kosakowski, Jakub	MUKS „Stoczek 45” Białystok	2429+14	7.0
3.	Ukraina	GM	Omelja, Artem		2499+8	6.5
4.	Czechy	GM	Zilka, Stepan		2577-3	6.5
5.	Litwa	GM	Rozentalis, Eduardas		2523-1	6.5
6.	Polska	I++	Majek, Kacper	UKS „Caissa” Warszawa	2011+89	6.0
7.	Polska	FM	Kucza, Karol	OTSz Ostrów Wielkopolski	2337+11	6.0
8.	Francja		Durand-Le Ludec, Brendan-Bubok		2197+59	5.5
9.	Gruzja	GM	Jobava, Baadur		2603-9	5.5

5. Paweł Teclaf, źródło: <https://bit.ly/3Dc6HWR>

uzyskał tytuł arcymistrza. Wielkie gratulacje dla obu naszych młodych zawodników!

A to partia ostatniej rundy, która zadecydowała o zdobyciu przez polskiego szachistę trzeciej i ostatniej normy na tytuł arcymistrza.

Matthias Bluebaum – Szymon Gumularz, 9 runda

1. d4 Sf6 2. Sf3 d5 3. Gf4 c5 4. e3 Sc6 5. Sbd2 e6 6. c3 Ge7 7. Se5 Sd7 8. S:d7 G:d7 9. Gd3 O-O 10. O-O c4 11. Gc2 f5 12. h3 Ge8 13. He2 po 13. Sf3 nastąpiłoby 13...Gh5 i czarne mogłyby wymienić swojego „złego” gońca za białego skoczka 13...b5 14. a3 a5 15. b3 c:b3 16. S:b3 Hb6 17. Wf1 Wc8 18. a4 b4 19. Sd2 Hd8 20. c4 Gd6 21. G:d6 H:d6 22. Gb3 Se7 23. Wc1 Gd7 24. Wc2 Ha6 25. Sf3 Wc7

Tabela 3. Czołówka turnieju OPEN B (ELO 1700–2200), 76 zawodników

Pozycja	Kategoria	Nazwisko, Imię	Klub	Ranking	Punkty
1.	I	Kołodziejcki, Piotr	KSz „Miedź” Legnica	1932+116	7.0
2.	I	Wieczorek, Mikołaj	MUKS MDK Śródmieście Wrocław	1824+152	6.5
3.	K	Bętkowski, Andrzej	TKKF „Drogowiec” Kraków	2057	6.5
4.	K	Luba, Łukasz	UKS „Giecek” Radków	2155-12	6.5
5.	K	Hasterok, Mateusz	Strzelec Strzelce Opolskie	2115-6	6.5
6.	I	Łukasiewicz, Aleksander	WKS Kopernik	1903+32	6.5
7.	K	Liberadzki, Sławomir	KSz Skoczek Trojański	2049-3	6.5

Tabela 4. Czołówka turnieju OPEN C (SENIORZY), 76 zawodników

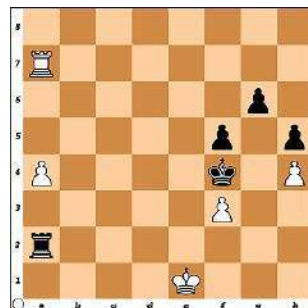
Pozycja	Kraj	Tytuł/ Kategoria	Nazwisko, Imię	Klub	Ranking	Punkty
1.	Ukraina	FM	Shilov, Sergej	SzSON Zagłębie Dąbrowa Górnicza	2043+21	7.5
2.	Polska	I+	Kulon, Bogusław	KS AZS Wroclavia	2117-2	7.0
3.	Polska	K	Kwiecień, Zbigniew		1952+17	6.5
4.	Polska	I++	Gromczak, Dariusz	Kudowianka Kudowa Zdrój	1885+35	6.5
5.	Polska	K	Mazalon, Mariusz	UKS MDK Nr 1 – Hetman Bydgoszcz	2014-1	6.5
6.	Ukraina	IM	Marusenko, Petr		2077-2	6.5
7.	Polska	K+	Zowada, Kazimierz	MCKiS Jaworzno	2029-3	6.5



6. Matthias Bluebaum – Szymon Gumularz, pozycja po 35. Wcc2



7. Matthias Bluebaum – Szymon Gumularz, pozycja po 45...e4



8. Matthias Bluebaum – Szymon Gumularz, pozycja po 56...Wa2

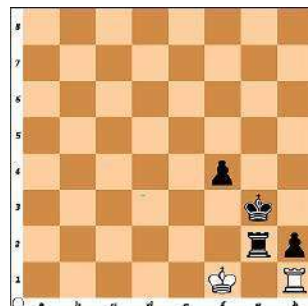
26. Se5 Wfc8 27. Wac1 Ge8 28. Sd3 Hd6 29. Sc5 Gf7 30. Hd3 g6 31. Wd2 Kg7 32. f3 d:c4 lepsze byłoby 32...e5 33. G:c4 Sd5 34. Hb3 Sc3 35. Wcc2 (diagram 6) 35...W:c5 słuszną decyzją o poświęceniu jakości 36. d:c5 H:c5 37. Ga6 H:e3+ czarne mają dwa piony za jakość, świetnie umocowanego skoczka i dużą przewagę pozycyjną 38. Kh1 Wc7 39. Gc4 He1+ 40. Kh2 He5+ 41. Kh1 We7 dość pasywne posunięcie naszego nowego arcymistrza, lepsze było 41...Hf4 i nie można 42. G:e6 ze względu na 42...We7 42. Hb2 Hc5 43. Hb3 e5 44. G:f7 W:f7 45. Wd3 e4 (diagram 7) 46. Wd:c3 białe oddają jakość ponieważ po odejściu wieży grozi zarówno 46...e3 jak i 46...e:f3 46...b:c3 47. H:c3+ H:c3 48. W:c3 powstała końcówka z pionem więcej u czarnych 48...Kf6 49. Kg1 e:f3 50. g:f3 Wd7 51. Wc5 Kg5 dużo lepsze niż pasywne 51...Wa2 52. Kf2 Wd4 wymusza bicie czarnego piona 53. W:a5 Kf4 54. Wa7 h5 55. h4 Wd2+ 56. Ke1 Wa2 (diagram 8), czarna wieża stoi za pionem, a biały król jest odcięty od swoich pionów 57. Wg7 Ke3 precyzyjna realizacja przewagi przez polskiego zawodnika 58. Kd1 Wa1+ 59. Kc2 W:a4 60. W:g6 W:h4 61. Wg1 f4 62. Kd1 Wh2 63. We1+ K:f3 64. Wf1+ Kg3 65. Wg1+ Wg2 66. Wh1 h4 67. Ke1 h3 68. Kf1 h2 (diagram 9) 0-1

Bardzo silną obsadę miał również turniej OPEN A (ELO >= 2100), gdzie wystąpiło 7 arcymistrzów (10). Zwyciężyli dwaj Polacy Grzegorz Nasuta (11) z wynikiem 7,5 pkt. przed Jakubem Kosakowskim 7 pkt.

W turnieju OPEN B dla zawodników z ELO 1700–2200 zakończył się niespodziewanie zwycięstwem dwóch piętnastolatków: Piotra Kołodziejkiego 7 pkt. przed Mikołajem Wiczorkiem 6,5 pkt.

W tradycyjnym już turnieju seniorów (mężczyźni od 60 lat, kobiety od 50 lat) jednym z najsilniejszych obsadzonych turniejów seniorów w Polsce zwyciężył Sergej Shilov z Ukrainy przed Bogusławem Kulonem.

W turnieju OPEN D (ELO <= 2000), w którym uczestniczyło 127 zawodników zwyciężył Mateusz Budnik przed Janem Kukorowskim i Przemysławem



9. Matthias Bluebaum – Szymon Gumularz, pozycja w której Bluebaum poddał partię

10. Najlepsi w turnieju OPEN A (ELO >= 2000), źródło: <https://bit.ly/3D5CPLv>





11. Grzegorz Nasuta – zwycięzca turnieju OPEN A fot. Krzysztof Szeląg/Wikimedia Commons

Szmidtem. Zwycięzcą turnieju OPEN E (ELO $\leq$ 1600) z udziałem 165 uczestników został Jakub Liśkiewicz przed Edgarem Scholtes z Niemiec i Igorem Jankowskim. Dzieci do lat 12 grały w 3 grupach. W silnie obsadzonej grupie F (75 zawodników) zwyciężył Sebastian Baliś, drugi był Wiktor Galla a trzeci Krzysztof Chęś. W grupie G1 (22 zawodników) pierwszy był Jan Gach, drugi Tymon Kamiński i trzeci Mateusz Dorskocz a w grupie G2 (21 zawodników)

12. Symultana z arcymistrzem Zbigniewem Paklezą rozgrywana na 24 szachownicach źródło: <https://bit.ly/3Fcjp4T>



zwyciężył Mateusz Dorskocz przed Nikodemem Górzyńskim i Tymonem Kamińskim.

Festiwalowi szachowemu towarzyszyło wiele imprez. Były to m.in. dwa turnieje szachów błyskawicznych. W pierwszym uczestniczyło 85 zawodników, zwyciężył mł. Piotr Goluch przed Mateuszem Budnikiem. W drugim, z udziałem 112 zawodników, zwyciężył mł. Jakub Kosakowski przed arcymistrzem Grzegorzem Nasutą. W turnieju szachów szybkich (83 zawodników) zwyciężył Tobias Kuegel z Niemiec przed Pawłem Kuligiem. W turnieju szachów Fischera (18 zawodników) zwyciężył Sebastian Basiaga przed Przemysławem Szmidtem.

Dużym zainteresowaniem ze strony zawodników i kibiców cieszyła się symultana z arcymistrzem szachowym Zbigniewem Paklezą (12). Symultana zakończyła się zwycięstwem Arcymistrza we wszystkich 24 partiach.

Na zakończenie Festiwalu zgromadzona w polanickim Parku Szachowym publiczność miała okazję obejrzeć spektakl „Żywe Szachy” (13). W plenerowym widowisku na podstawie poematu Jana Kochanowskiego „Szachy”, które zostało zaprezentowane na szachownicy, wystąpiły dzieci z różnych klubów szachowych z całej Polski. Pierwowzorem dedykowanego hrabiemu Janowi Tarnowskiemu utworu był włoski poemat „Gra szachowa” z 1527 roku, którego autorem był Marco Girolamo Vida. Spektakl wyreżyserowała Aleksandra Paprocka-Paszczyk, a przedstawienie było również transmitowane online. Po spektaklu odbyła się dekoracja zawodników uczestniczących w turnieju arcymistrzowskim Festiwalu (14). W Parku Szachowym prezentowana była również wystawa fotograficzna upamiętniająca historię Międzynarodowego Festiwalu Szachowego im. Akiby

13. Spektakl Żywe Szachy, fot. Jan Sobótko





14. Uczestnicy Turnieju Arcymistrzowskiego, fot. Jan Sobótka

Rubinsteina i przybliżającą postać najwybitniejszego polskiego szachisty.

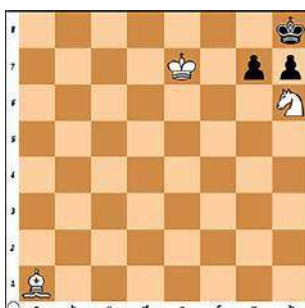
Ja też startowałem, poprawiłem mój wynik rankingowy i jestem zadowolony z rezultatów. Podobnie jak wielu uczestników tegorocznej imprezy,

zarezerwowałem już zakwaterowanie w Polanicy-Zdroju na następny rok, aby uczestniczyć w kolejnej edycji Festiwalu, tradycyjnie w drugiej połowie sierpnia oraz cieszyć się urokami i zdrowotnymi walorami tego pięknego kurortu. ■

Zadania do samodzielnego rozwiązania



Zadanie 1
15. Rodolfo Bozzo – Dag Mad-
sen, Gausdal Int. 1992
Mat w 3 posunięciach



Zadanie 2
16. Aleksander W. Galicki 1900
Mat w 3 posunięciach

Rozwiązanie zadań z MT 9/2021

Zadanie 1

A. Ferrantl 1859

Mat w 3 posunięciach

Rozwiązanie: 1.Sa8 Kd6 2.Kd4 Kc6
3.Hd5#

Zadanie 2

I. Kos 1876

Mat w 3 posunięciach

Rozwiązanie: 1.f8W Ke7 2.c8W Kd7
3.Wc7#

Serce na zakręcie. Rozdroża. Tom 2

K.A. Figaro

Wydawnictwo Lipstick Books, liczba stron: 384, cena: 39,99 zł

Małgorzata Piłat próbuje poskładać swoje życie osobiste. Widmo minionego romansu odciska niezatarte piętno. Pożądanie i namiętność gryzą się z życiem podporządkowanym jednej osobie. Jakiego wyboru dokona Małgorzata? Na szali jest dwóch mężczyzn, tak różnych od siebie jak ogień i woda. A może życie, które potrafi być nieprzewidywalne, samo rzuci jej podpowiedź? Walka, jaką będzie musiała stoczyć, nie będzie lekka, przyniesie ból i wewnętrzne rozterki. Ale przecież nikt nie mówi, że będzie łatwo...





1. Szwalnia w amerykańskim Winder w stanie Georgia w 1953 roku

Era tanich dóbr

Coraz droższa droga w przyszłość

Era dobrobytu dla tych wybrańców, którzy mieli okazję żyć w tamtych czasach w bogatych w krajach zachodnich, w praktyce oznaczała, że wszystko było relatywnie tanie (w stosunku do przychodów zwykłych ludzi). Samochody, elektronika, klimatyzatory, meble, ubrania, buty. Wszystko, co było „dobrem konsumpcyjnym”. Obecnie wielu ekspertów uważa, że era taniości dobiega końca.

Ponieważ głównym motorem gospodarki w krajach kapitalistycznych, w których panował wolny rynek i obfitość dóbr, była konsumpcja, cała gospodarka była skonstruowana w taki sposób, aby napędzać masową konsumpcję. To dlatego rozwijano metody masowej produkcji (1), dzięki której koszt jednostkowy dóbr spadał i można było je tanio sprzedawać, mnożąc obniżki i wyprzedaże, by sprzedać jeszcze

więcej, jeszcze taniej (2). Działo się to kosztem jakości, ale z punktu widzenia celu gospodarki, czyli sprzedawania jak największej liczby produktów, nie była to wada, ale w pewnym sensie zaleta.

Oczywiście nakręcony tak mechanizm znalazł w końcu bariery kosztowe. Nie można było bez końca obniżać kosztów masowej produkcji (mówimy o czasach przed automatyzacją i robotyzacją), ponieważ



2. Wyprzedaże

pracownicy w fabrykach nie godzili się na przeniesienie na nich kosztów tanizny. Sami byli konsumentami i chcieli konsumować więcej. Producenci więc, szukając kogoś, kto zapłaci za taniość produktów na rynkach zachodnich, przenieśli dużą część produkcji do krajów, w których można było utrzymać albo jeszcze dodatkowo obniżyć koszty produkcji.

Za tanie produkty płaciło też środowisko naturalne zanieczyszczane nie tylko przez zakłady wytwórcze ale również przez rosnącą ilość odpadów pozostałych po tanich produktach. Zwłaszcza tworzywo sztuczne było i wciąż jest ogromnym obciążeniem dla środowiska. Plastik to symbol masowych dóbr i ich opakowań, ale również rzeczywistych kosztów konsumpcji, jeszcze tak naprawdę do końca niepoliczonych, bo skutków nasycenia środowiska tworzywami sztucznymi jeszcze w całości nie poznaliśmy, czyli nie poznaliśmy prawdziwej pełnej ceny tego wszystkiego, co wydawało nam się tak tanie.

Braki i niedobory

Wydawało się, że ta epoka taniej masowej obfitości trwać będzie bez końca, zwłaszcza że przemysł rozwinął techniki automatyzacji i robotyzacji, eliminując potrzebę siły roboczej, zarówno tej drogiej, w krajach rozwiniętych, jak i tej taniej, w krajach rozwijających się.

Jednak zobaczyliśmy koniec tego świata. Właściwie wszędzie na świecie ceny obecnie rosną. Dzieje się tak po pierwsze dlatego, że gospodarka obfitości zaczyna przeobrażać się w świat braków i niedoborów. Brakuje surowców, których ceny rosną, brakuje też wysoko przetworzonych półproduktów, czyli np. krzemowej elektroniki nie tylko do komputerów, ale również do samochodów.

Świat znaga się też z niedoborem, który wcześniej w epoce taniej obfitości był mało znany – niedoborem wykwalifikowanej siły roboczej. Pracownicy w krajach Trzeciego Świata nie są już tak tani i przede wszystkim nie tak dostępni w dowolnych ilościach. Maleje liczba miejsc, w których można tanio eksploatować wielkie zasoby ludzkie do pracy.

Kolejne drożący zasób to energia. Ze względu na ograniczenia i opłaty za emisje gazów cieplarnianych nakładane przez bogate kraje na tradycyjne źródła, paliwa kopalne, węgiel i ropę, droższą kilowatogodziny energii elektrycznej, paliwo na stacjach, ogrzewanie i chłodzenie klimatyzatorami. Energia to obok pensji pracowników podstawowy koszt produkcji. Jeśli drożeje energia, to drożeją wszystkie produkty i usługi. Trwale drogie źródła energii oznaczają, że wszystko inne będzie dalej drogie.

Chiny przestają być tanią fabryką świata

Symbolem taniochy były produkty pochodzące z Chin. Zmierzch epoki niskokosztowej produkcji w Chinach sygnalizowano już dekadę temu. Po latach bycia tanią „fabryką świata” wytwarzającą pracochłonne produkty, takie jak tekstylia (3), zabawki, odzież, obuwie i meble dla konsumentów na całym świecie, Państwo Środka weszło na kolejny etap rozwoju gospodarczego i technologicznego. Ten sam próg przekraczały już zresztą inne azjatyckie państwa, najpierw Japonia, potem Korea Południowa, Tajwan, Hongkong. Jeszcze w 2013 roku Międzynarodowy Fundusz Walutowy pisał w swoich analizach, że Chiny nie osiągnęły tzw. punktu zwrotnego Lewisa, w którym gospodarka przechodzi od gospodarki obfitującej w siłę roboczą do takiej, w której występuje jej niedobór. Wtedy w głębi tego wielkiego kraju były jeszcze duże rezerwy tańszej niż na wybrzeżach siły roboczej. Dziś już i te zasoby się wyczerpują w tym sensie, że pracownicy z prowincji już też nie są tacy tani jak byli wcześniej i nie da się eksploatować tych rezerw, utrzymując niskie koszty.

Chiny chcą teraz skupić się na produkcji towarów z wyższej półki, oprzeć się na konsumpcji krajowej, aby napędzać swoją gospodarkę, a pracę przy produkcji tanich, pracochłonnych towarów powierzyć innym. Jest to strategia, o której mówią przywódcy tego kraju od co najmniej kilku lat. Tylko komu powierzyć? Odpowiedź nie jest wcale taka prosta.

Najbardziej oczywistymi następcami Chin wydają się kolejne wschodzące gospodarki Azji, a mianowicie



3. Tekstylny zakład produkcyjny w Chinach

Indie, Bangladesz, Kambodża, Indonezja, Myanmar, Pakistan, Sri Lanka i Wietnam. Jednak tylko w niektórych z tych krajów w ostatnim czasie rośnie produkcja eksportowa. Dotyczy to przede wszystkim Wietnamu i Bangladeszu. Jednak nie są to pod względem przede wszystkim potencjału ludnościowego kraje, które mogą nawet w grupie zastąpić potencjał Chin.

Możliwe, że produkcja wymagająca dużych nakładów pracy ludzkiej może pozostać w Chinach, przechodząc wielkie przemiany, przede wszystkim w kierunku automatyzacji. Chiny są jednym ze światowych liderów we wdrażaniu robotów przemysłowych ale dotyczy to przede wszystkim sektorów takich jak produkcja samochodów i elektroniki. Chińczycy nie wykazali wiele motywacji w implementowaniu automatyzacji produkcji tanich towarów.

Jednym z powodów są techniczne ograniczenia. Miękkie, giętkie materiały, takie jak tkanina, bywają, o czym świadczy problem tapicerki w zrobotyzowanej fabryce Tesli, opisywany swego czasu w MT, trudne dla robotów. Trudno zautomatyzować takie prace jak choćby nawijanie sznurówek do tramppek. Są to rodzaje zadań, w których wciąż prawie nie da się zastąpić ludzkiej siły roboczej.

Mimo że liderzy łańcucha dostaw aktywnie pracują nad rozszerzeniem źródeł zaopatrzenia poza Chiny, a firmy z takich branż jak moda badają możliwości produkcji bliżej swoich klientów końcowych w Europie i Stanach Zjednoczonych, wiele z nich nadal uważa, że opuszczenie Chin jest trudne i nieproporcjonalnie drogie. Firmy te często ostatnio tworzą modele hybrydowe, np. przenosząc część produkcji do tańszego Wietnamu, pozostawiając jednak zasadnicze ośrodki wytwórstwa w Chinach. Tak czy owak

wpływa to na ceny dóbr konsumpcyjnych, które nie mogą być już tak tanie jak były wcześniej.

Będzie drożej, ale czy lepiej?

Spojrzenie na rosnące koszty surowców, narzuty ekologiczne, szzybujące w górę koszty energii i w końcu koniec rajów taniej produkcji takich jak Chiny, nasuwa myśl, że epoka tanich produktów nieuchronnie się kończy. Można to ująć inaczej: świat coraz mniej stać na to, aby wszystkie te dobra, które dotychczas kosztowały niewiele, nadal miały niskie ceny. Zgodnie z tokiem myślenia przyjmującym, że prawdziwa cena tanich produktów jest ukryta przed osobą, która kupuje za niską cenę i ktoś inny zawsze płaci resztę ceny (ale ten ktoś jest daleko w sensie geograficznym lub czasowym, tzn. w przyszłości zapłaci resztę należności, np. przyszłe pokolenia), nie ma już prawie chętnych do płacenia reszty ceny.

No dobrze, można przyjąć, że era taniochy była epizodem w historii ludzkiej cywilizacji. Jeśli ktoś zna relacje cenowe i kosztowe z dawnych epok, wartość pieniądza i zarobki od starożytności po wiek dziewiętnasty, to wie, że w dawnych czasach wszystko było znacznie droższe, od żywności po wysoko przetworzone produkty, narzędzia, ozdoby, buty i ubrania.

Jednak, jeśli mija era niskich cen i wracają dawne relacje kosztowe, to rodzi się pytanie, czy wraz z wysokimi cenami dóbr konsumpcyjnych wróci jakość, z jakiej były znane dawne wyroby, najczęściej powstające drogą rękodzieła i będące solidnymi produktami rzemieślniczymi. Czy jednak tak się nie stanie i będziemy płacić drogo za masowy produkt o jakości szybko przemijającej. To byłoby jednak niezbyt fair. ■

Mirosław Usidus

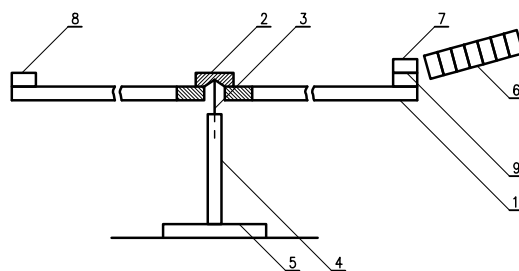
Materiały diamagnetyczne mają niezwykłą właściwość. Niezależnie od bieguna magnesu, który do nich zbliżamy, są od niego odpychane. To zachowanie stanowi rezultat oddziaływania zewnętrznego pola magnetycznego z poszczególnymi elektronami w atomach. Dlatego diamagnetyzm jest powszechną właściwością, maskowaną jednak przez znacznie silniejsze właściwości para- lub ferromagnetyczne. Mimo to, diamagnetyki można wykryć w prostych doświadczeniach, wykonalnych w warunkach domowych.

Te niezwykłe diamagnetyki

Diamagnetyki można wykryć w domu

W szkolnych podręcznikach fizyki znajduje się zwykle informacja, że pod względem właściwości magnetycznych materiały dzielą się na trzy podstawowe grupy. Są to: diamagnetyki, paramagnetyki i ferromagnetyki. Tę ostatnią grupę dzieli się jeszcze na ferromagnetyki magnetycznie miękkie, używane do wytwarzania rdzeni różnego rodzaju uzwojeń oraz magnetycznie twarde, z których są produkowane magnesy trwałe. W bardziej szczegółowych rozważaniach wspomina się o dość rzadko spotykanych: antyferromagnetykach, ferrimagnetykach i metamagnetykach. Wykrywanie ferromagnetyków nie stanowi problemu, ponieważ te materiały silnie oddziałują z zewnętrznym polem magnetycznym. Najbardziej znane ferromagnetyki to: żelazo, jego stopy i spieki oraz nikiel i kobalt.

Materiały diamagnetyczne mają niezwykłą właściwość. Niezależnie od bieguna magnesu, który do nich zbliżamy, są od niego odpychane. Diamagnetyki, w przeciwieństwie do ferromagnetyków, bardzo słabo oddziałują z zewnętrznym polem magnetycznym. Właściwości tej grupy materiałów nie wynikają z uporządkowania momentów magnetycznych całych atomów, ale z oddziaływania zewnętrznego pola magnetycznego z poszczególnymi elektronami w atomach. Dlatego diamagnetyzm jest powszechną właściwością, często jednak maskowaną przez znacznie silniejsze właściwości para- lub ferromagnetyczne. Mimo bardzo słabego oddziaływania z polem magnetycznym, przynajmniej niektóre diamagnetyki można wykryć w prostych i pomysłowych doświadczeniach, dających się wykonać w warunkach domowych. Kilka takich doświadczeń jest opisanych w tym artykule.



1. Wykrywanie właściwości magnetycznych przy użyciu obrotowej dźwigni; 1 – cienka listewka z balsy o długości ok. 1 m, 5 – łożysko oporowe, 3 – igła, 4 – słupek, 5 – podstawka, 6 – magnesy neodymowe, 7 – badany materiał, 8 – przeciwi ciężar, 9 – nakładka kompensacyjna

Listewką z balsy

Już Archimedes w starożytnej Grecji wypowiedział słynne zdanie: „Dajcie mi punkt podparcia, a poruszę Ziemię”. W ten poglądowy sposób wyraził, jak duży zysk na sile można uzyskać za pomocą odpowiednio długiej dźwigni. W układzie doświadczalnym pokazanym na **rysunku 1** też skorzystamy z tej możliwości. Funkcję dźwigni spełnia tutaj cienka i długa listewka z balsy 1 – drewna o bardzo małej gęstości i znacznej wytrzymałości mechanicznej. Długość zastosowanej listewki wynosiła 1 m, a wymiary przekroju poprzecznego 8×2 mm. Taką listewkę można kupić w sklepach z artykułami dla modelarzy lub plastyków. Dłuższy bok przekroju poprzecznego listewki jest skierowany poziomo. W połowie długości listewki ostrożnie wywiercono otwór o średnicy 3 mm i przyklejono do niej od góry łożysko oporowe 2 klejem epoksydowym. Rolę łożyska odgrywał

krążek o średnicy 5 mm i wysokości 4 mm, odcięty z grubego grafitu twardego ołówka dla plastyków. W krążku od dołu nawiercono końcem wiertła stożkowe zgłębienie. Takie rozwiązanie umożliwia obrót listewki w płaszczyźnie poziomej z bardzo małym tarciem. Zgłębienie opiera się na ostrzu igły 3, osadzonej w pionowym słupku 4, przymocowanym do podstawki 5. Podstawka i słupek powinny być wykonane z materiału nieferromagnetycznego, np. z tworzywa sztucznego.

Jako łożyska oporowego, zamiast krążka z grafitu, można użyć metalowej końcówki wyjętej z wkładu od długopisu. Końcówkę należy opłukać denaturatem z resztek tuszu i wypchnąć z niej kulę za pomocą igły. Tak przygotowaną końcówkę przykleja się od góry do listewki nad dopasowanym do niej otworem i opiera listewkę na ostrzu igły 3, osadzonej w słupku 4. Końce listewki 1 mogą nieco ugiąć się pod własnym ciężarem ku dołowi, ale jest to korzystne, ponieważ obniża położenie środka masy ruchomej części układu i poprawia jego równowagę. Przed rozpoczęciem doświadczeń listewka powinna być nieruchoma. Dlatego układ należy umieścić w miejscu osłoniętym od przewiewów powietrza i unikać szybkich ruchów np. ręki.

Do końca listewki zbliżamy dowolny biegun możliwie silnego magnesu trwałego 6 i obserwujemy jej zachowanie. Prostim sposobem na uzyskanie takiego magnesu jest złożenie, skierowanymi ku sobie biegunami różnoimiennymi, kilku magnesów neodymowych w kształcie walca o średnicy i wysokości ok. 1 cm, tak żeby utworzyły one pręt o długości 6–8 cm. Należy zachować przy tym należyta ostrożność, ponieważ magnesy neodymowe są kruche, a ponadto mogą boleśnie przycisnąć skórę palców. Najlepiej byłoby, żeby koniec listewki 1 nie oddziaływał ze zbliżonym magnesem. Wtedy można od razu przystąpić do doświadczeń. Na końcu listewki 1 układamy badany materiał 7, np. skrawek folii aluminiowej lub kawałek druczika miedzianego. Spowoduje to przechylenie obciążonego tym materiałem końca listewki 1 w dół, dlatego na jej przeciwnym końcu należy położyć odpowiedni przeciwważ 8, żeby uzyskać równowagę. Po tym ponownie zbliżamy dowolny biegun magnesu i obserwujemy zachowanie się końca listewki 1. Jeżeli badany materiał 7 ma właściwości diamagnetyczne, to będzie on wypychany z silniejszego pola magnetycznego w pobliżu bieguna magnesu 6 i koniec listewki 1 obróci się w kierunku od magnesu 6. Odwrotna sytuacja nastąpi w przypadku materiałów paramagnetycznych.

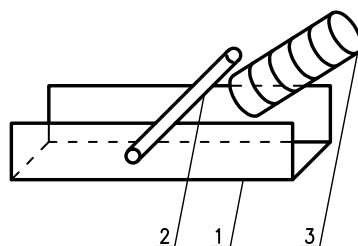
Dla lepszej orientacji, jakich właściwości magnetycznych powinniśmy spodziewać się po wybranym

materiale 7, warto skorzystać z odpowiednich tablic, np. W. Mizerski, „Tablice fizyczne i astronomiczne”, Wydawnictwo Adamantan, Warszawa 2013, str. 187, albo z tablic dostępnych w internecie. Materiały diamagnetyczne mają ujemną wartość podstawowego parametru, który nazywa się podatnością magnetyczną. Czasem zamiast podatności jest podawana przenikalność magnetyczna względna i dla diamagnetyków jest ona nieco mniejsza niż jeden. W przypadku paramagnetyków sytuacja jest odwrotna. Im większe wartości bezwzględne podatności magnetycznej, tym silniejsze właściwości dia- lub paramagnetyczne wykazuje dany materiał i tym wyraźniej będzie oddziaływał z magne sem.

Oczywiście, gdybyśmy na końcu listewki 1 położyli kawałek materiału ferromagnetycznego, np. stalowej blaszki, to oddziaływanie byłoby bardzo silne i blaszka, zlatując z listewki, zostałaby mocno przyciągnięta do bieguna magnesu 6. Gdyby początkowo, po zbliżeniu bieguna magnesu 6, koniec listewki 1 był do niego przyciągany, wskazywałoby to na jego właściwości paramagnetyczne. Wówczas na tym końcu należałoby najpierw położyć nakładkę kompensacyjną 9 w postaci odpowiednio dobranego kawałka materiału diamagnetycznego i zrównoważyć listewkę 1 stosownym przeciwważem 8. Pomocne będą w tym wspomniane tablice. Dopiero po uzyskaniu „obojętności magnetycznej” końca listewki 1 możemy przystąpić do badania wybranych materiałów.

Można z przeciekami

Jeżeli chcemy badać materiały w kształcie przecików, np. kawałki cienkich drutów lub grafity do ołówków, to można posłużyć się bardzo prostym układem pokazanym na **rysunku 2**. W tym celu z kawałka kartonu o rozmiarach ok. 6×3 cm wyginamy rynienkę 1 o przekroju ceownika i kładziemy ją poziomo na stole. Na górnych krawędziach pionowych ramion takiej rynienki układamy poprzecznie badany przecik 2 i zbliżamy do niego dowolny biegun magnesu 3.



2. Wykrywanie właściwości magnetycznych przecików; 1 – rynienka z kartonu, 2 – przecik z badanego materiału, 3 – magnesy neodymowe

Pręciki z materiałów diamagnetycznych będą toczyły się lub przesuwali w kierunku od magnesu, natomiast pręciki z materiałów paramagnetycznych zachowają się odwrotnie. W opisanym układzie zostały zmniejszone opory ruchu, dzięki zastąpieniu tarcia poślizgowego przez tacie toczone i zminimalizowaniu trących się powierzchni.

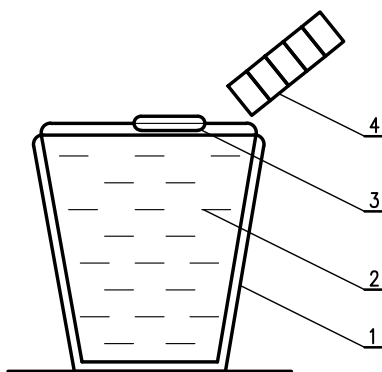
Diamagnetyczna ślizgawka

Inny sposób na wykrywanie dia- lub paramagnetyków polega na wykorzystaniu poślizgu cienkiego płátka badanego materiału po błonie powierzchniowej cieczy. Układ doświadczalny do tego celu też jest niezwykle prosty i został pokazany na **rysunku 3**. Szklankę 1 napełniamy ostrożnie wodą 2, tak żeby woda nie wylewała się, ale utworzyła menisk wypukły ponad górną krawędzią szklanki. Na powierzchni tego menisku kładziemy płatek badanego materiału 3, np. skrawek folii aluminiowej lub tworzywa sztucznego, o rozmiarach kilku mm. Materiał należy kłaść powoli, trzymając go równolegle do powierzchni wody, tak żeby jego krawędź nie przecięła błony powierzchniowej. Wtedy na wodzie można położyć materiał o gęstości większej niż gęstość wody i on nie zatoni, lecz będzie utrzymywał się w zagłębieniu błony powierzchniowej. Do płátka tak przygotowanego materiału zbliżamy dowolny biegun magnesu 4. W przypadku diamagnetyków zaobserwujemy ich poślizg po powierzchni wody w kierunku od magnesu. W przypadku paramagnetyków ruch będzie odbywał się w przeciwną stronę.

Opisane doświadczenie szczególnie efektownie przebiega z płatkami grafitu pirolitycznego. Poślizgowi grafitu sprzyja też fakt, że nie jest on zwilżany przez wodę. Grafit pirolityczny otrzymuje się przez rozkład węglowodorów w wysokiej temperaturze (pirolizę) i osadzanie (depozycję) uwolnionych w ten sposób atomów węgla w postaci równoległych warstw. Cały proces przebiega bardzo powoli, a uzyskany w ten sposób materiał jest anizotropowy i ma bardzo dużą, ujemną podatność magnetyczną w kierunku prostopadłym do warstw atomowych. Ponadto grafit pirolityczny wykazuje duży stosunek podatności magnetycznej do gęstości. Dzięki temu dobrze nadaje się do doświadczeń z lewitacją diamagnetyczną. Płatki grafitu pirolitycznego o rozmiarach kilku mm można kupić przez internet.

Lewitujący grafit

Omówione wcześniej właściwości grafitu pirolitycznego pozwalają zademonstrować lewitację diamagnetyczną. W tym celu należy zestawić układ złożony

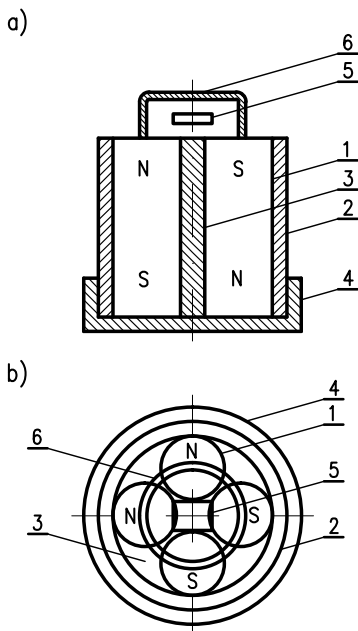


3. Wykrywanie właściwości magnetycznych z wykorzystaniem napięcia powierzchniowego; 1 – szklanka, 2 – woda, 3 – płatek badanego materiału (korzystnie grafitu pirolitycznego), 4 – magnesy neodymowe

z czterech walcowych magnesów neodymowych, umieszczonych obok siebie w narożnikach kwadratu. Magnesy powinny być zwrócone na przemian biegunami różnoimiennymi w tę samą stronę. Jest to tzw. układ Hallbacha. Dla płátka grafitu o rozmiarach 5×5 mm optymalne okazały się magnesy neodymowe o średnicy 8 mm i długości 20 mm. Taki układ magnesów ustawiamy pionowo i nad jego środkiem umieszczamy ostrożnie płatek grafitu pirolitycznego. W obszarze umieszczenia płátka pole magnetyczne wykazuje największą niejednorodność i dużą wartość indukcji magnetycznej. Dlatego płatek jest wypychany z tego obszaru ku górze do miejsca, w którym wartość indukcji pola jest mniejsza i siła odpychania równoważy ciężar płátka. Dzięki temu płatek lewituje nad układem magnesów.

Jeżeli będziemy ostrożnie poruszali w niewielkim zakresie magnesami, to płatek nadal będzie lewitował. Gdy jednak zakres ruchu będzie zbyt duży albo ruch za szybki, wtedy płatek wskutek bezwładności nie wróci do obszaru o najkorzystniejszych parametrach pola magnetycznego i spadnie. Chcąc pokazać to efektowne doświadczenie bezpośrednio dużej grupie osób, warto zbudować prosty przyrząd, pokazany na **rysunku 4** i **fotografii 5**, który każdy uczestnik doświadczeń może bezpiecznie wziąć do ręki bez obawy o zgubienie lub pokruszenie płátka. Cztery magnesy neodymowe 1 zostały złożone we wcześniej opisany układ i umieszczone w przezroczystej plastikowej rurce 2. Wolne przestrzenie między magnesami można wypełnić klejem epoksydowym 3. Od dołu na rurkę 2 jest nałożona plastikowa podstawka 4. W wykonanym modelu był to korek do zamykania fiolek z lekami. Płatek

grafitu pirolitycznego 5 jest umieszczony od góry nad środkiem układu magnesów i zabezpieczony przezroczystą osłoną 6, ograniczającą jego ruch do obszaru o optymalnych parametrach pola magnetycznego. Najprościej wykorzystać w tym celu



4. Budowa przyrządu do pokazu lewitacji diamagnetycznej płytki grafitu pirolitycznego: widok: a) z boku oraz b) z góry; 1 – magnes neodymowy, 2 – przezroczysta rurka z tworzywa sztucznego, 3 – klej epoksydowy, 4 – podstawa, 5 – płytek grafitu pirolitycznego, 6 – przezroczysta osłona

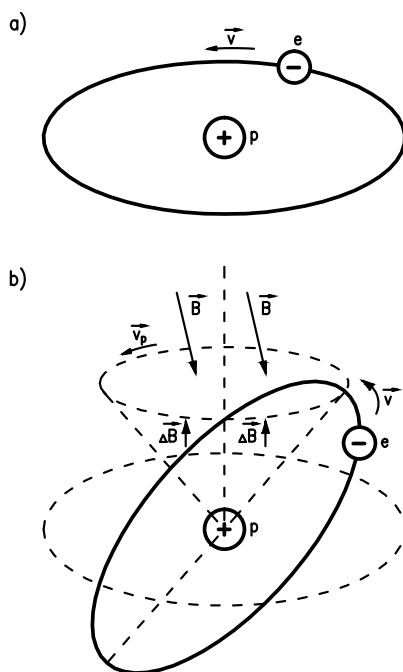


5. Wygląd zewnętrzny przyrządu do pokazu lewitacji diamagnetycznej płytki grafitu pirolitycznego

miseczkę z wytłoczki, tzw. blistra, do tabletek. Miseczkę należy odciąć, przykryć nią lewitujący płatek grafitu i ostrożnie przykleić jej dolny brzeg do górnej powierzchni magnesów, używając przezroczystego kleju epoksydowego.

Zrozumieć diamagnetyzm

Jak wspomniano we wstępie, za zjawisko diamagnetyzmu odpowiedzialne jest oddziaływanie zewnętrznego pola magnetycznego z elektronami w atomach. Dokładne wyjaśnienie szczegółów tego zjawiska wymaga zastosowania mechaniki kwantowej, ale jego podstawy można zrozumieć, wykorzystując prawa fizyki znane ze szkoły. Rozpatrzmy w tym celu najprostszy w budowie atom wodoru (rysunek 6). Na elektron e , poruszający się w atomie tego pierwiastka, umieszczonego w zewnętrznym polu magnetycznym o indukcji B , działa siła elektrodynamiczna, nazywana też siłą Lorentza i skierowana prostopadle do kierunku ruchu elektronu. Skutkiem tego pojawia się dodatkowy ruch elektronu nazywany



6. Schemat ułatwiający zrozumienie zjawiska diamagnetyzmu z wykorzystaniem fizyki klasycznej na przykładzie atomu wodoru: a), b) sytuacja odpowiednio: przed i po umieszczeniu atomu w zewnętrznym polu magnetycznym; e – elektron, p – proton, v – prędkość elektronu, v_p – prędkość precesji orbity elektronu, B – indukcja zewnętrznego pola magnetycznego, ΔB – indukcja dodatkowego pola magnetycznego

precesją, podobnie jak w przypadku wirującego bąka w polu siły ciężkości po odchyleniu jego osi obrotu od pionu. Z tą precesją można powiązać przepływ pewnego prądu elektrycznego, który wytwarza dodatkowe pole magnetyczne o indukcji ΔB . Zastosowanie reguły lewej dłoni i reguły śruby prawoskrętnej pozwala wykazać, że zwrot wektora ΔB jest przeciwny do zwrotu wektora B . W sumie wygląda to tak, jakby atom wodoru stał się magnesem zwróconym biegunem jednoimiennym w kierunku bieguna źródła, wytwarzającego zewnętrzne pole magnetyczne. Dobrze wiadomo, że przy takiej orientacji biegunów następuje odpychanie.

Właściwości materiałów sklasyfikowanych w tablicach jako diamagnetyczne mogą być zmienione przez ich zanieczyszczenie substancjami para- lub ferromagnetycznymi. Dlatego, badając np. grafity z ołówków o różnej twardości, można czasem zauważyć, że są one przyciągane przez magnes. Jest tak dlatego, że grafit w ołówku składa się z proszku grafitu z dodatkiem substancji wiążącej, nadającej mu pożądaną twardość. Zjawisko lewitacji diamagnetycznej znalazło zastosowanie m.in. w grawimetrach nadprzewodnikowych, umożliwiających pomiary przyspieszenia ziemskiego z bardzo dużą dokładnością, rzędu 10^{-9} cm/s², a także w eksperymentalnych pociągach na tzw. poduszce magnetycznej. Nadprzewodniki poniżej temperatury krytycznej stają się bowiem idealnymi diamagnetykami o przenikalności

magnetycznej względnej wynoszącej -1 . Są też produkowane zabawki edukacyjne o nazwie lewitron, w których obserwuje się lewitację, wykorzystując do tego celu płytki wykonane z diamagnetyków – grafitu lub bizmutu.

Właściwości diamagnetyczne ma też woda i można je efektownie zademonstrować w dużej skali w odpowiednio silnym polu magnetycznym. Zrobili to Andriej Geim i Michael Berry, umieszczając żywą, zieloną żabkę nad otworem magnesu Bittera, wytwarzającym pole magnetyczne o indukcji kilku tesli. Dla porównania, pole magnetyczne w naszych doświadczeniach było kilkaset razy słabsze i zajmowało ok. 1000 razy mniejszą objętość. Ponieważ głównym składnikiem organizmu żaby (podobnie jak i naszego) jest woda, żabka lewitowała nad magnesem. Zwierzątko w dobrej kondycji przeżyło eksperyment, a jego zdjęcie trafiło do internetu (zob. np. <https://bit.ly/3zV2W5W>). Wykonawcy eksperymentu otrzymali za swój wyczyn Nagrodę Ig Nobla, nazywaną też anty-Noblem i przyznaną co roku za najbardziej kuriozalne wyniki badań. W przeciwieństwie do innych naukowców nie obrazili się i przyjęli to niezwykle wyróżnienie. Nieco później (w 2010 r.) Andriej Geim, wspólnie z Konstantinem Novosiołowem, otrzymał już normalną Nagrodę Nobla za odkrycie grafenu – materiału, z którym wciąż wiązane są duże nadzieje na szerokie zastosowanie. ■

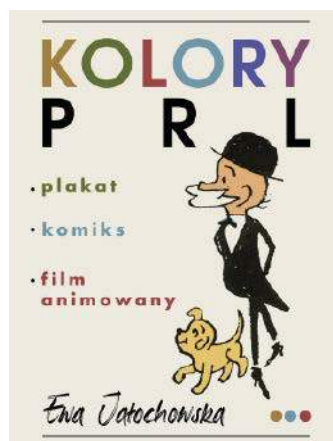
Stanisław Bednarek

Kolory PRL

Ewa Jatochowska

Wydawnictwo MUZA S.A., liczba stron: 397, cena: 59,90 zł

W smutnym peerelowskim pejzażu powstawały dzieła wybitne. Polscy malarze, rysownicy i animatorzy zdobywali światową sławę. W kraju naznaczonym niedostatkiem rozkwitała kolorami sztuka. A że najbliższa ludziom była wtedy sztuka masowa, komiks, plakat i film animowany, ona więc wyrastała na tym zgrzebnym tle – jak maki na betonie. Przeto mowy rok 1945 – oznaczał wprawdzie koniec wojny, ale nie wymazywał z pamięci wojennej traumy. Na murach spalonej stolicy zawisł pierwszy, symboliczny, zamykający ten okres plakat antyhitlerowski „Zniszczymy faszyzm do końca” Eryka Lipińskiego. Ci, którzy ocalili z pożogi, choć poturbowani psychicznie i fizycznie, z nadzieją patrzyli w przyszłość. Ten upragniony „nowy lepszy świat” okazał się jednak pułapką. Dla artystów ponure realia peerelowskie stały się wyzwaniem i inspiracją. W nich narodziła się słynna w świecie polska szkoła plakatu, a międzynarodową sławę zdobywali twórcy polskich animacji. Na kartach książki pojawiają się artyści tej miary co Henryk Tomaszewski, Waldemar Świerzy, Janusz Stanny czy Tadeusz Trepcowski. Uchyła swoje drzwi Pracownia Plakatu Propagandowego i raczkujące studio filmów animowanych w Bielsku-Białej. Śledzimy losy Waleriana Borowczyka i Jan Lenicy – w latach 60. okrzykniętymi twórcami polskiej szkoły animacji. Obserwujemy trud tworzenia pierwszych filmów rysunkowych. Wreszcie i perypetie związane z „najbardziej imperialistycznym” gatunkiem sztuki. A w nim Henryk Chmielewski czy Janusz Christa – królowie komiksu. Przewijają się tu zresztą cała plejada znakomitości: plakatistów, animatorów, twórców komiksów.





Powierzchnia Ziemi wynosi 510 100 000 km². Niespełna 1/3 jest zamieszkiwana przez ludzi. Na początku naszej ery było nas ok. 250 milionów. Szacuje się, że już za rok liczba ludzi na świecie osiągnie 8 miliardów. Zmiany klimatyczne, ekonomiczne i polityczne wpływają na to, że rozmieszczenie ludności na Ziemi nie jest równomierne. Najgęściej zaludnionym miejscem jest Makau. Najluźniej ludziom jest na Grenlandii. Pomimo milionów kilometrów kwadratowych, na których można żyć, ilość dogodnej przestrzeni jest ograniczona. Dlatego też niezbędne jest odpowiednie gospodarowanie przestrzenią, tak by wszystkim nam żyło się komfortowo. Zapraszamy na wydział zajmujący się gospodarowaniem przestrzenią. Zapraszamy na GP.

Gospodarka przestrzenna

Gospodarka przestrzenna znajduje się w ofercie wielu uczelni, w tym: politechnik, uniwersytetów, akademii i prywatnych uczelni wyższych. Można ją studiować dziennie, zaocznie lub wieczorowo, więc dla każdego coś się znajdzie. W zależności od wybranej uczelni, pierwszy etap edukacji, który regulaminowo trwa 3,5 roku, można ukończyć, uzyskując dyplom licencjata lub inżyniera. Studia magisterskie nie powinny trwać dłużej niż 1,5 roku (rzadko zdarzają się wyjątki). Po tym okresie można przystąpić do obrony pracy dyplomowej, co w przypadku sukcesu kończy się uzyskaniem tytułu magistra. O tym, jak sprawnie gospodarować przestrzenią, można się również dowiedzieć na studiach podyplomowych. Ich ukończenie stwarza możliwość przystąpienia do egzaminu na uprawnienia w zakresie planowania przestrzennego. „Podyplomówka”, w zależności od wybranej uczelni, trwa od roku do dwóch lat. Są to studia skierowane do osób, które nie miały do tej pory nic wspólnego z architekturą, urbanistyką i planowaniem przestrzennym. Dokonując wyboru uczelni i trybu studiowania, należy poważnie zastanowić się nad tym, jakiego efektu się oczekuje. Zasadnicze pytanie brzmi: w jakim miejscu widzę się po uzyskaniu dyplomu magistra? Studia licencjackie to w tym przypadku nauka ukierunkowana na gospodarkę przestrzenną w ujęciu społeczno-gospodarczym, ekonomicznym i przyrodniczym. „Inżynierka” wzbogaca treści nauczania o elementy techniczne. Podjęcie dobrej decyzji może okazać się nie lada sztuką.

Dostanie się na GP można uznać za wyzwanie. Nie jest to najbardziej oblegany kierunek, ale niewielka liczba wolnych miejsc na roku wpływa na to,

że konkurencja rośnie. W rekrutacji na rok akademicki 2020/2021 Politechnika Krakowska odnotowała 3,5 kandydata na indeks. Należy zwrócić uwagę na fakt, że już od kilku lat uczelnia systematycznie zmniejsza liczbę miejsc na tym kierunku. W bieżącym roku studentami politechniki Krakowskiej na kierunku gospodarka przestrzenna zostało stu dwudziestu szczęśliwców. Przebrnięcie przez proces rekrutacji wydaje się najtrudniejszym zadaniem. W kolejnych latach można się spokojnie skupić na edukacji. Poziom trudności zwykle określa się tu jako niewygórowany. Tym samym jest czas nie tylko na naukę, ale i na życie studenckie. Łącząc te dwa elementy, edukacja może upływać w przyjemnej atmosferze. Student nie zdąży się obejrzeć, a miną pierwsze lata i trzeba będzie wybrać specjalizację, w jakiej będzie rozwijać się dalsza nauka. W zależności od uczelni pojawiają się różne warianty, wśród których można wybierać na przykład pomiędzy: środowiskowego uwarunkowania gospodarowaniem przestrzenią albo urbanistyką w planowaniu przestrzennym. Wybierając konkretną specjalizację, należy spodziewać się różnego zestawu przedmiotów. Wszystkich bez wyjątku będą obowiązywały treści podstawowe. Wśród nich znajduje się między innymi matematyka. Jest jej tutaj wyjątkowo mało, bo 30 godzin na studiach licencjackich i 60 na inżynierskich. Na licencjacie nie ma fizyki i grafiki inżynierskiej, które na studiach technicznych zajmują po 45 godzin. Ponadto: statystyka, ekonomia, geografia ekonomiczna, rysunek techniczny i planistyczny, socjologia, historia urbanistyki i prawoznawstwo. Studia na tym wydziale



to przede wszystkim projekty, nad którymi spędza się dużo czasu. To rysunki, projektowanie, wielogodzinne planowanie i raz jeszcze dopieszczanie projektów. Podobno nie są one wyjątkowo trudne i skomplikowane, a dla osób, które faktycznie interesują się tematyką, potrafią być nawet miłe i przyjemne. Jedyne problemy, jakie mają z nimi studenci, to to, że są czasochłonne. I na tym kierunku nie mogło zabraknąć przysłowiowej kosi. Absolwenci sugerują z uwagą podejść do geometrii wykreślnej i geodezji. Pozostałe przedmioty kierunkowe nie powinny przysporzyć większych trudności.

Studiowanie na tym kierunku wydaje się stosunkowo łatwe i przyjemne. Należy jednak dodać, że by osiągnąć sukces w pracy zawodowej, wypadałoby już w trakcie studiów inwestować czas w doskonalenie się we własnym zakresie. Kluczowe będzie sprawne poruszanie się po systemach wspierających projektowanie. W branży tej niezbędna jest także dobra znajomość prawa budowlanego i geodezyjnego. Koniec studiów to nie koniec nauki. Absolwent, chcąc zwiększyć zakres swoich kompetencji i kwalifikacji, powinien podjąć starania w celu uzyskania specjalistycznych uprawnień.

Ukończenie studiów nie powinno być wielkim wyzwaniem, tym bardziej dla pilnych uczniów. Trudności czekają na absolwentów chcących podjąć pracę zawodową. Stanowiska, na których może być zatrudniony inżynier gospodarki przestrzennej, są stosunkowo mało sprecyzowane. Można rozwijać karierę naukową, można zostać projektantem, doradcą, managerem przestrzeni. Planowanie swojej kariery zawodowej wymaga nie lada wyobraźni. Oczywiście,

przydadzą się znajomości i kontakty, ale to umiejętność przewidywania i dostosowania do aktualnych warunków stanowią klucz do sukcesu w tej branży. Warto łączyć wiedzę z zakresu gospodarki przestrzennej z umiejętnościami i kwalifikacjami związanymi z, na przykład: budownictwem, prawem, informatyką, ekonomią. Każdy z tych kierunków dobrze współgra z gospodarką przestrzenną, a tym samym skutecznie zwiększa umiejętności i szansę na dobrą i ciekawą pracę. Po ukończeniu GP można także starać się o zdobycie uprawnień do wyceniania nieruchomości lub zostać na uczelni i rozwijać karierę naukową. Jednym z możliwych wariantów jest także otwarcie własnej firmy, natomiast należy dysponować nie lada umiejętnościami nie tylko z zakresu gospodarki przestrzennej, ale i marketingowej, finansowej oraz managerskiej.

Studia na tym wydziale można uznać za proste, łatwe i przyjemne. Absolwent tego kierunku to osoba uśmiechnięta, niezestresowana i zadowolona z siebie. Studiowanie GP to z całą pewnością doskonały pomysł dla osób, które wiedzą, czego chcą od życia i mają plan na siebie. Rzeczywistość po studiach okazuje się już nie tak łaskawa i przyjemna. Poszukiwanie pracy, zdobywanie dodatkowych kwalifikacji, uprawnień i kontaktów wymaga wysiłku i jest czasochłonne. Wszystkie te zabiegi będą kosztowały absolwenta wiele trudu, tak więc jest to kierunek dla osób wytrwałych i zdeterminowanych, by pozostać w tej branży. Gospodarka przestrzenna to kierunek ciekawy, ale godny polecenia jedynie osobom, które wiedzą, że temu właśnie chcą się poświęcić. ■

Michał Pacholski



Szkoła Wynalazców

dozwolone do lat 15

Mielicie zadanie z dziedziny „dyplomacji podstawowej”. Są sprawy, gdzie wezwanie policji i wdrożenie śledztwa załatwia temat, ale są też sprawy wymagające delikatności, w imię zasady: „co się bardziej opłaci”. Zadaniem Waszym było: *Jak uniemożliwić nieuczciwemu artyście zagarnianie dla siebie części wspólnego dochodu z datków publiczności.*

Problem należał do takich właśnie „delikatnych”. Artysta był znany i lubiany, jego występy cieszyły się popularnością, co przekładało się na zysk restauracji. Te jego zyski z nielegalnego uszczuplania wpływów od klientów nie były takie znów duże, no ale... Tak czy owak należało sprawę ukrócić. No i co wy na to?

Mateusz Gawlikowski (4 pkt.) proponuje, żeby artysta chodził po sali w towarzystwie koleżanki, która od niedawna pracuje w restauracji. Taka osoba nie jest jeszcze zżyta z pozostałym personelem, a więc także z artystą i w rezultacie będzie się on obawiał, że coś zauważy i doniesie.

Pomysł niezły, ale krótkoterminowy. Skuteczność będzie zależała od tego jak szybko nowa koleżanka zżyje się z zespołem. W świecie artystycznym jest to niekiedy sprawa nawet paru godzin!

Zbigniew Borek (4 pkt.) zauważył kiedyś, że pracownicy firm pogrzebowych nie mają w ogóle kieszeni, a raczej mają imitacje kłapek itp. Uważa, że to jest właśnie po to, żeby nie było gdzie wkładać wyłudzonych dodatkowych opłat od uczestników pogrzebu. Artysta mógłby występować w kostiumie tak zaprojektowanym, żeby w ogóle nie miał kiszeni.

Dobry pomysł, choć oczywiście nie każdy typ kostiumu będzie pasował do koncepcji „bezkieszniowca” i jednocześnie do treści występu.

Robert Pająk (4 pkt.) proponuje wykonać pojemnik na pieniądze w formie skarbonki, zamkniętej na kluczyk. Moment wkładania pieniędzy do skarbonki będzie dla wszystkich widoczny,

a wyjąć pieniądze będzie mógł tylko szef, bo on ma klucz.

Pomysł raczej typowy: brakuje mu trochę tej „artystycznej iskry”, natomiast niewątpliwie skuteczny.

Dziękuję za udział z konkursu i zapraszam do następnych. Oba kolegom życzę sukcesów w kolejnych zadaniach i sławy „wynalazcy” w rodzinie, gronie przyjaciół i w szkole!

Nowe zadanie:

Nielubiana w szkole (na szczęście nie przez wszystkich!) fizyka ukazuje swoje wszechdobylskie oblicze w bardzo różnych i codziennych sytuacjach. Na przykład taki zwyczajny problem:

Dwaj bracia: Wojtek i Piotrek, zaczęli robić sobie kawę na śniadanie. Zazwyczaj pili kawę z mlekiem. Żeby zrobić taką kawę, trzeba mieszać czarną kawę z mlekiem w odpowiednich proporcjach, zamieszać i wypić. Chłopcy zauważyli, że jeśli mieszają świeżo zaparzoną kawę z mlekiem, to ta mieszanina będzie jeszcze za gorąca. I trzeba będzie poczekać, aż trochę ostygnie. Wojtek uważa, że trzeba najpierw poczekać, aż nieco ostygnie świeża, czarna, gorąca kawa, zmieszać ją z mlekiem i pić! Doszło do zakładu: która metoda będzie szybsza? Spróbujcie pomóc braciom i doradźcie im: jak będzie szybciej. A zatem:

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 zmieszać ją z mlekiem i pić?

Tak naprawdę sprawa jest prosta, ale gdy do głosu dochodzą przyzwyczajenia i przeświadczenia, to nic nie jest proste. Musicie więc pozbyć się nawyków i opinii dziadków i babć, i metodą Leonarda da Vinci wszystko zbadać samodzielnie!

Życzę sukcesów badawczych, prawidłowego wnioskowania i fantazji potrzebnej nawet prawdziwym fizykom! Przypominam o terminie do końca listopada br.

Ranking Szkoły Wynalazców

1. Mateusz Gawlikowski.....(14 pkt.)
2. Marek Miernik(13 pkt.)
3. Mateusz Konarski.....(11 pkt.)
4. Jacek Zieliński(9 pkt.)
5. Zbigniew Borek(5 pkt.)
6. Małgorzata Mickiewicz.....(5 pkt.)
7. Robert Pająk(4 pkt.)

Klub Wynalazców

bez ograniczeń wieku

Mieliście zadanie wymagające znajomości psychologii i obyczajów ludzi biznesu: w jaki sposób malarz wziął za portret pięć razy więcej, niż wynosiła uzgodniona wcześniej cena.

Malarz miał problem: portret musiał być realistyczny albo „pochlebiony”. Model nie miał zbyt sympatycznej twarzy, więc dobry portret musiał wywołać negatywne odczucia oglądających.

Według autentycznej historii malarz poinformował biznesmena, że właśnie zamówił 1000 kopii kserograficznych jego portretu i ma już przygotowaną wstępną umowę z właścicielem strzelnicy, który ma zamiar użyć tych portretów jako tarcz do strzelania (biznesmen był powszechnie nie lubiany). Może jednak odstąpić od umowy ze strzelnicą, ale musi otrzymać rekompensatę za poniesione koszty. To oczywiście ogólny schemat,

bo należało jeszcze zadbać o uniknięcie oskarżenia o „poniewieranie wizerunku”, itd. A jakie propozycje mieli nasi klubowicze?

Zbigniew Przygodzki (4 pkt.). Malarz mógł domalować dolną część postaci i zrobić z całości ośmieszającą karykaturę. Rysy twarzy musiałby jednak trochę skorygować, żeby nie podpaść „pod paragrafy”. Gdyby zapowiedział biznesmenowi, że ma zamiar wystawić tę karykaturę na wystawie – mógłby coś utargować.

Propozycje kolegi gra na urażonej dumie i ambicji. W życiu mogłoby się to jednak różnie skończyć, z pobiciem malarza przez wynajętych osiłków włącznie.

Robert Pająk (3 pkt.). Malarz mógłby namalować nieco uproszczoną kopię i powiesić ją w eksponowanym miejscu. Pod portretem dać napis sławiący biznesowe osiągnięcia modela i przewidywanie objęcia przez niego wysokiego stanowiska.

Byłaby to gra na podlizanie się modelowi, ale kto wie? Molier – jak mało kto znał słabości duszy człowieka i wrażliwość na pochlebstwa, nawet

Ranking Klubu Wynalazców

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

Dzień bez teleranka. Stan wojenny – pomysły na życie

Anna Mieszczanek

Wydawnictwo MUZA S.A., cena: 44,90 zł

13 grudnia 1981 roku zatrzymał na moment rozpędzoną polską rzeczywistość. Przerwał wielkie marzenie ludzi o wolnym kraju. Siłami wojska i służby bezpieczeństwa przywrócił dawny „porządek”. Internował wolność. Chwilowo. Formalnie stan trwał półtora roku, niewiele dłużej niż solidarność „karnawał”. W lipcu '83 został zniesiony. Jednak ostatni więźniowie polityczni siedzieli do połowy 1986 roku.

Goniec z wydawnictwa, reżyser zaraz po filmówce, młody pracownik komitetu dzielnicowego, początkująca dziennikarka, motorniczy, ogrodniczka z MSW, kapitan z wydziału zabójstw, sekretarz organizacji partyjnej w dużej firmie, mały chłopiec, któremu zamknęli ojca, i inni to bohaterowie, których historia zwykle nie zauważa. Każdy z nich stał się częścią reporterskiej opowieści o jednym z najboleśniejszych momentów w dziejach powojennej Polski, kiedy wszystko przesiąknięte było polityką. Ludzi tych, niezależnie od tego, po której stronie barykady stali, łączy trudna historia czasów ich młodości. To bohaterowie szarej codzienności, wciągnięci w wojnę z niewidzialną siłą politycznej przemocy. Jak dostosowywali się do zakazów i nakazów stanu? W jaki sposób zamykały się za nimi ścieżki otwarte przed grudniem '81? Co pozwalało im trwać i tworzyć enklawy wolności? Jakich musieli dokonywać wyborów? Kiedy słabli i gorzknieli? Dzięki czemu znów ogarniała ich nadzieja? Czy ten proces naprawdę się zakończył? Relacje bohaterów, nagrywane na bieżąco, a także w roku 1989 i w latach dwutysięcznych, przeplatają się z dokumentami władz partyjnych czy SB. Pokazują szczegóły tego kolizyjnego kursu, na którym Polacy znaleźli się w latach 80.





dość grubymi nićmi szyte (vide: „Mieszczanin szlachcicem”) i pięknie wykorzystał tę wiedzę w swoich komediach.

Kolegom gratuluję i zapraszam do następnych zadań.

Nowe zadanie:

Zadanie bardzo techniczne i zarazem zaproszenie do skorzystania z metod TRIZ.

W praktyce hydraulicznej znane jest zjawisko „uderzenia wodnego”, które jest niekiedy wykorzystywane jako tzw. taran wodny. W większości przypadków jest to zjawisko szkodliwe. Dochodzi do takiego uderzenia wodnego przy gwałtownym

zamknięciu rurociągu, zwłaszcza o dużej średnicy i sporej prędkości przepływu. Uderzenie wodne może „rozwalić” szwy rur, uszkodzić zawory, narobić sporo jeszcze innych szkód. Jak temu zapobiec? I to jest wasze zadanie:

Zaproponować sposób uniknięcia szkód, do których może dojść w wyniku zjawiska „uderzenia wodnego.”

Należy sięgnąć po podręczniki fizyki – wystarczy licealny, pomyśleć i gotowe!

Wszystkim życzę dobrych pomysłów i skutecznego myślenia. Termin nadsyłania propozycji: koniec listopada br.

Vademecum Młodego Wynalazcy

Niniejszy odcinek VMW rozpoczyna serię, poświęconą „analizie wepolowej” i „systemowi standardów”. Najpierw jednak trochę uwag wstępnych.

Starsi z was, studenci, znają dobrze metodę „na przodka” pomocną przy opracowywaniu prac projektowych. Metoda ma oczywiście różne warianty, poczynawszy od przerysowania „żywcem” ze zmianą danych w tabelce rysunkowej, do studium kilku wariantów podobnych tematycznie do zadanej nam pracy i po analizie opracowanie własnego projektu, w oparciu o zauważone pewne idee, zawarte w „przodkach”. Metody „na przodka”, wyjąwszy wariant pierwszy, nie można oceniać negatywnie. Znane są pojęcia takie: jak „lepsze jest wrogiem dobrego”, ewolucja materii ożywionej: roślin, zwierząt, bakterii itp. Wszystkie bez wyjątku przemiany ewolucyjne obiektów przyrodniczych realizowane są metodą „na przodka”.

Jeśli się bliżej przyjrzeć, to można dojść do wniosku, że Stwórca znał doskonale tę metodę. Stąd całe serie organizmów ideowo podobnych i rozwijających się wg jednej linii „konstrukcyjnej”. Wystarczy przytoczyć serie kotów, głowonogów, ryb, ptaków, itd. W technice serie zunifikowanych konstrukcji są codziennością, co widać doskonale, obserwując rozwój samochodów, samolotów, obrabiarek, sprzętu AGD itd. Jak ma działać wynalazca, który ma ambicję wynaleźć coś nowego, „coś, czego jeszcze nie było”.

Wydawać by się mogło, że w takiej sytuacji „przodka nie ma”! Na szczęście coś niecoś jest. Oczywiście jest sporo technik poszukiwania nowego rozwiązania „czegoś”, czyli znalezienia

systemu, realizującego potrzeby funkcjonalne społeczeństwa lub jego sektorów.

Henryk Altzsuller w latach 70. ub. stulecia zauważył, że niektóre rozwiązania techniczne, a zwłaszcza ich idee, pojawiają się w różnych systemach. Tu można przytoczyć charakterystyczny przykład: ploter – jako urządzenie peryferyjne komputera.

Typowy ploter to urządzenie stosunkowo niewielkie, ale urządzenia pochodne, oparte na idei komputerowego sterowania ruchem dowolnego narzędzia, w układzie 2D lub 3D, to już maszyny o rozmiarach średnio kilkunastometrowych, obejmujące całą gamę maszyn CNC. Rychło zauważono, że takie urządzenie może nie tylko wycinać, wypalać, itp., ale może też tworzyć struktury „drukowane”: ciekłym metalem, tworzywem sztucznym, masą betonową i dało to początek burzliwego rozwoju drukarek 3D, które dziś mogą wybudować nawet dom!

Takich idei – „matek” – poprzez analizę kilkudziesięciu tysięcy opracowań wynalazczych zebrano – w klasycznej postaci – 76. Działo się to w latach 70. ub. stulecia i oczywiście zbiór ten jest ciągle aktualizowany; zmniejszany, powiększany – ciągle gdzieś ktoś coś próbuje ulepszyć.

Żeby jednak móc operować czystymi ideami systemów technicznych, trzeba je jakoś w uproszczeniu przedstawić i nadać reguły ich wykorzystywania. W zasadzie w tym celu Altzsuller opracował symboliczny zapis struktur systemów technicznych, znany jako „analiza wepolowa”. Samo określenie wywodzi się z języka rosyjskiego, w którym „wieszczstwo” to „substancja”, a „pole” też po polsku – pole. Ze

zbitki tych słów powstało określenie „wepole”, oznaczające minimalny system techniczny zdolny do wykonywania „głównej funkcji użytkowej”, dla której system został skonstruowany. Termin „wepole” został przyjęty przez pierwszego tłumacza książki H. Altszullera – prof. A. Góralskiego.

Analiza wepolowa doczekała się rozwiniętego aparatu formalnego, zwanego niekiedy „algebrą wepoli” i dającego spore możliwości. Najlepiej to zrozumieć na paru przykładach. Załóżmy więc, że mamy do czynienia z trzema problemami technicznymi; każdy w zasadzie z innej branży:

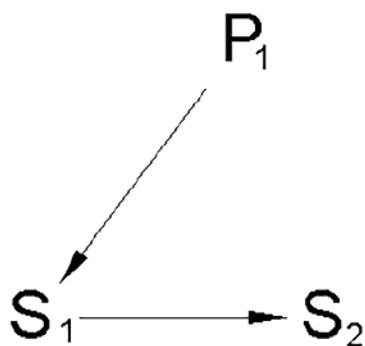
1. Żeby ściągnąć śrubę okrętową z wału, wykorzystuje się stalowe pręty, które kurczą się po uprzednim nagrzanianiu,
2. Dla mikrod dozowania płynnych leków nagrzewa się powietrze wewnątrz ampułki i rosnące ciśnienie wyciska płyn,
3. Dla sprasowania proszku zamkniętego w ogrzanej metalowej tulei, po jej zamknięciu intensywnie się ją chłodzi.

W tych trzech przykładach jest: przedmiot obrabiany przez narzędzie, które jest: chłodzone, podgrzewane, itp. W każdym z tych przykładów mamy trzy elementy: narzędzie, przedmiot i strumień energii w postaci jakiegoś pola – o polach TRIZ nieco później. Jasne jest, że narzędzie i przedmiot działają na siebie za pośrednictwem trzeciego elementu. O ile narzędzie i przedmiot zaliczamy do „substancji”, o tyle energia może mieć bardzo różną postać i ogólnie w TRIZ nazywamy ją „polem”. W rezultacie oprócz znanych z fizyki pól: elektrycznego, magnetycznego, grawitacyjnego i pola słabych i silnych oddziaływań, mamy jeszcze: pole sił mechanicznych, pole ciepła, pole oddziaływań chemicznych, akustycznych i wibracyjnych. Dla wygody pola ujęto w skrót: MATCHEM, obejmujący wszystkie pola ważne w sensie technicznym. Można więc, uogólniając powiedzieć, że w TRIZ „pole” to sparametryzowana przestrzeń dwu- i trójwymiarowa.

Wszystkie przytoczone przykłady możemy przedstawić schematycznie (1):

gdzie: S_1 – substancja – narzędzie, S_2 – substancja – przedmiot, P_1 – pole tu: ciepłne.

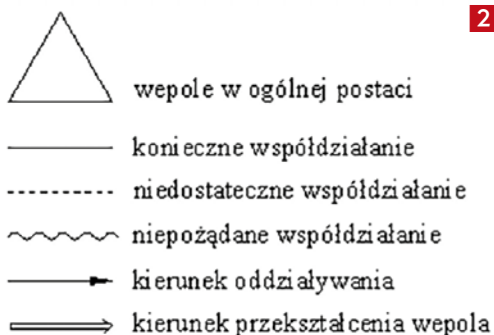
Powyższy schemat jest „minimalnym systemem technicznym, zdolnym do działania. Wystarczy odjąć jeden dowolny element i już systemu nie ma. Taki schemat opisuje całą gromadę systemów: na przykład żona (narzędzie) robi awanturę mężowi (przedmiot), czyli substancja: S_1 (żona), działając za pomocą pola akustycznego P_1 (krzyczy, wrzeszczy) na substancję S_2 (męża).



Ten nieco żartobliwy przykład ilustruje ogólną zasadę: substancją może być wszystko, co ma masę i wymiary, a więc pojedynczy przedmiot, ale także cała fabryka. Dzięki takiemu ujęciu problemów odcinamy się od całej gromady „wektorów inercji” a więc: „my tak nie praktykujemy”, „to się nie uda”, „Amerykanie tak nie robią, to nam też na pewno nie wyjdzie” itp.

Wepola ilustrują „samą esencję” problemu z odcięciem się od jego materialnej postaci. Ogromnie ułatwia to poszukiwanie nowych rozwiązań, które wydają się niemożliwe.

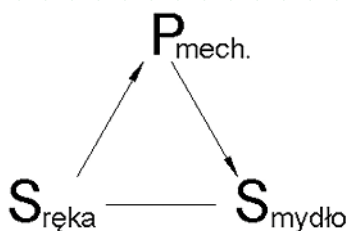
Jakie są zasady zapisywania relacji wepolowych? Oczywiście analiza wepolowa ma swój język, chociaż każdy może sobie stworzyć swój, byle zachować podstawy. Najczęściej posługujemy się symbolami pokazanymi niżej (2).



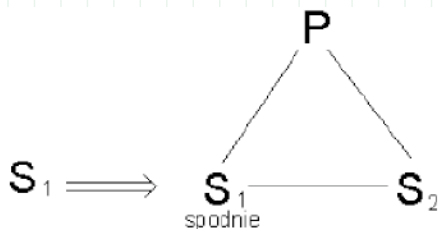
W konkretnych systemach technicznych do podstawowych symboli dodajemy indeksy, ułatwiające orientację w zapisie, zwłaszcza gdy mamy do czynienia z rozwiniętymi schematami.

Przypuśćmy, że chcemy zapisać w wepolach taką oto prostą czynność: „człowiek bierze do ręki mydło”. Zapis wepolowy będzie wyglądał następująco:

Ręka generuje pole oddziaływań mechanicznych, które działają na mydło. Strzałki oznaczają kierunek



3



4

oddziaływać, a linia bez „grotu” oznacza, że konieczne jest działanie (3).

Wepolowy zapis pozwala ujawnić przyczyny powstania problemu, inaczej: pozwala ujawnić niepożądane efekty. Przekształcenia wepolowe podpowiadają twórcy – wynalazcy, co konkretnie należy wprowadzić do systemu dla rozwiązania zadania albo jak trzeba zmienić system, ale nie daje informacji, jaką substancję i jakie pole należy wprowadzić bądź wyprowadzić. Dla otrzymania technicznej odpowiedzi trzeba dobrać substancje i pola. Przy tym zazwyczaj zaleca się skorzystać ze skrótu MATChEM i właśnie w takiej kolejności „przymierzać” różne pola.

„Algebra” wepolowa zawiera szereg zasad, które stosowane w odpowiedni sposób pozwalają ruszyć z miejsca większość zadań. Pierwsza grupa zasad odnosi się do syntezy wepola lub rozbudowy wepola. Widząc słowo „synteza” lub „rozbudowa” wepola, widzimy od razu, że najpierw trzeba się zorientować: czego brakuje w naszym schemacie. I rzeczywiście jest szereg zadań, w których przy próbie zbudowania wepola czegoś brakowało. Ze zdziwieniem i z ulgą (bo wiecie już, co robić) widzicie, że w wepolu brakuje substancji lub pola albo tego i tego. Od razu możecie postawić diagnozę: w wepolu brakuje elementu i należy go dobudować.

Przykład: chcecie wyprasować spodnie, ale jak na złość nie macie żelazka. Co robić? Mamy więc tylko jedną substancję – spodnie. Jest to niepełne wepole. Musimy je dobudować do pełnego, tj. wprowadzić substancję i pole.

S_1 to u nas będąc spodnie. Należy teraz wstępnie określić S_2 i P. Można tu zaproponować cały szereg wariantów, na przykład P – pole ciśnienia. Ciśnienie można zrealizować w różny sposób: zacisnąć spodnie pomiędzy dwie deski, można też użyć... własnego ciała. Sposób znany doskonale w wojsku: spodnie od munduru wyjściowego musiały być wyprasowane, a żelazko „gdzieś się zapodziało”. Wkładało się spodnie pod siennik i rano były gotowe do wyjścia.

Drugi wariant – pole ciepłe. S_2 – dowolny metalowy przedmiot, mający chociaż jedną powierzchnię płaską: np. patelnia, obciążona papierem (zwykle jej dolna powierzchnia nosi ślady używania). Patełnię, można rozgrzać na kuchence i jako tako wyprasować spodnie (bo inaczej warta nie wypuści za bramę!).

Te parę przykładów to wstęp do analizy wepolowej, której poświęcimy trochę miejsca, głównie bazując na przykładach, które najlepiej wprowadzają do tej metody. ■

Prezes Klubu Wynalazców
Instruktor TRIZ
Jan Boratyński

Toksyczna przyjaźń

Polly Phillips

Wydawnictwo MUZA S.A., cena: 42,90 zł

Od lat były nierozłączne, teraz stały się rywalkami. Kiedy dziewczyny z przyjaciółek zmieniają się w konkurentki – nic dobrego z tego nie może wyniknąć. Tak jest i w tym przypadku. Izzy i Becca od lat nierozłączne, jako dorosłe kobiety zaczynają ze sobą rywalizować. Niby wciąż się przyjaźnią i wiedzą o sobie wszystko, ale w tej relacji coraz częściej dochodzi do spięć. I coraz bardziej widać, jak bardzo się różnią. Isabel piękna, zgrabna, bogata, u boku wymarzonego męża wiecie bajkowe niemal życie. Wykwintne przyjęcia, towarzyski blichtr, wszystko na wysokim poziomie. Becca przeciwnie, u niej nic nie jest doskonałe. Aspirująca dziennikarka ma narzeczonego o wątpliwej reputacji, małe mieszkanie i wciąż walczy z nadwagą. Na styku tych dwóch światów iskrzy coraz bardziej. Zgrzyty zdarzają się i podczas towarzyskich spotkań, i podczas wspólnie spędzanych świąt. Bo choć nadal żyć bez siebie nie mogą, dominują teraz w ich kontaktach złośliwości, zawiść oraz brak lojalności. I ciągłe wzajemne pretensje. Gdy Becca na pierwszej stronie popularnego tabloidu czyta informację, którą dopiero co w „największym sekrecie” wyjawiała Izzy, jej wściekłość sięga zenitu... Ma dość tej przyjaźni. Ma dość takiej przyjaciółki...



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.com.pl/ActiveReaderNagrody)

Szymon Miazga mieszka w mieszkaniu z balkonem, usytuowanym od strony południowo-zachodniej. Na tym balkonie, pełniącym funkcję podręcznego magazynu m.in. owoców, ziemniaków, itp. w upalne lato pojawiają się problemy. Otóż nie da się niczego z grupy towarów, psujących się, przetrzymać na balkonie. Oczywiście można kupić lodówkę i postawić na balkonie, ale to już spory wydatek. Szymon proponuje wykorzystanie starej metody chłodzenia przez odprowadzenie wilgoci. Uważa, że można zbudować „łódówkę balkonową” w formie ławy – skrzyni, z wewnętrznym wkładem z blachy ocynkowanej. Do tego wkładu należałoby włożyć również blaszany pojemnik, odizolowany klockami drewnianymi od pojemnika zewnętrznego. Po nakryciu obu pojemników pokrywami (blaszanymi) na wierzch pokrywy zewnętrznej nakładłoby się wielowarstwową „czapę”, której końce zwiisałyby aż do dna. Do przestrzeni pomiędzy obydwojema zbiornikami należałoby wlać wodę i tylko pilnować, żeby nie wyschła całkowicie.

Idea w zasadzie słuszna, ale zważywszy na koszt jednostkowego wykonania „ławy – skrzyni” i obu pojemników, należałoby raczej pomyśleć o wykorzystaniu dużych garnków ocynkowanych lub aluminiowych. W tej wersji pomysł nadaje się raczej na obóz harcerski.

Jan Wojton wpadł na genialny pomysł załatwienia sprawy anteny samochodowej, urywanej przy okazji mycia auta na myjce dotykowej (automatycznej). Zdaniem Janka antenę należy wykonać tak, jak miarkę zwijaną. Na czas mycia antenę by się chowało, a po umyciu wyciągało!

Pomysł z gatunku „sucharów”. Chowanie anteny przed gwałtownymi ruchami szczotek myjki automatycznej doczekało się ponad 3000 patentów. Prócz tego: antena nadrukowana na szybie, antena jednocześnie ogrzewająca szybę tylną, itp. Czy to oznacza, że nic tu już nowego nie da się wymyślić? Życie pokazuje, że nie ma granic dla ludzkiej pomysłowości. No więc myśleć!

Mitosz Kujawa – oglądał kiedyś kotdredę: „4 pory roku”. Składa się ona w zasadzie z dwóch kotdred o różnych grubościach, których można używać oddzielnie lub spi-

Pomysł miesiąca 10/2021

Pomysł lokalizacji różnych przedmiotów w domu podzielał nam na wyobraźnię. Pomysłeliśmy, że mogłaby powstać elektroniczna „mapa domowa” różnego rodzaju przedmiotów, które się łatwo gubią lub są potrzebne na tyle rzadko, że zwykle zapomina się, gdzie je schowaliśmy. W dobie dynamicznego rozwoju IoT – pomysł zdecydowanie do przemyślenia.

Autorem pomysłu jest Artur Gąsowski

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

A oto plon ostatniego miesiąca:

nać je razem, w przypadku niskiej temperatury w po-koju. W rezultacie mamy do dyspozycji trzy warianty kotdredy: najcieńszą – na lato, nieco grubszą na jesień i wiosnę oraz najgrubszą na zimę, po spięciu obu kotdred. Mitosz uważa, że wystarczyłaby jedna kotdreda, ale nadmuchiwana, dzięki czemu można by uzyskać kotdredę o dowolnej – w pewnych granicach – grubości. Izolatorem – tak naprawdę – jest powietrze!

Pomysł na oko niezły, ale żeby kotdreda pneumatyczna „trzymała ciśnienie”, musiałaby być nieprzewiedna, co znacznie pogorszyłoby komfort spania pod nią.

Artur Gąsowski – proponuje „raz na zawsze” załatwić problem gubienia w mieszkaniu różnych potrzebnych drobiazgów, jak: okulary, długopisy, kalkulatory, grzebień, itp. Wobec daleko już posuniętej miniaturyzacji chipów, można by te wszystkie „ginące” obiekty zaopatrzyć w indywidualne chipy lokalizacyjne i na smartfonie lub komputerze, na mapce mieszkania, natychmiast zobaczylibyśmy, gdzie są okulary. Stosowany obecnie sygnał dźwiękowy np. do znajdowania kluczyków samochodowych nie załatwia problemu, bo domowe hałasy: telewizor, dzieci, głośna rozmowa, itp. utrudniają znalezienie zguby.

Pomysł realny; rzeczywiście o wiele skuteczniejszy od zabawy w „ciepło-zimno”, gdy mamy tylko sygnał akustyczny. Ilość otaczających nas gadżetów rośnie, więc i problem z panowaniem nad tą gromadą staje się coraz bardziej dokuczliwy.

Monika Rachwańska po przegodzie z wyrwanym „z korzeniami” pendrive’em, proponuje zaopatrzyć wszystkie komputerowe wtyczki w magnesy, które będą trzymać je we właściwym miejscu, ale ławo „puszczą” przy nieostrożnym szarpnięciu za kabelek. Dotyczy to zwłaszcza pendrive’ów, które nie są przeciętne na stałe wetknięte w port, są dość długie i łatwo je uszkodzić. **Prawdopodobnie niektóre komputery już mają taki system: magnetycznego trzymania pendrive’a. Jest to jednak początek, ale jeśli rzecz się przyjmie, to z uwagi na to, że jest to dobry pomysł, powinien się upowszechnić.**



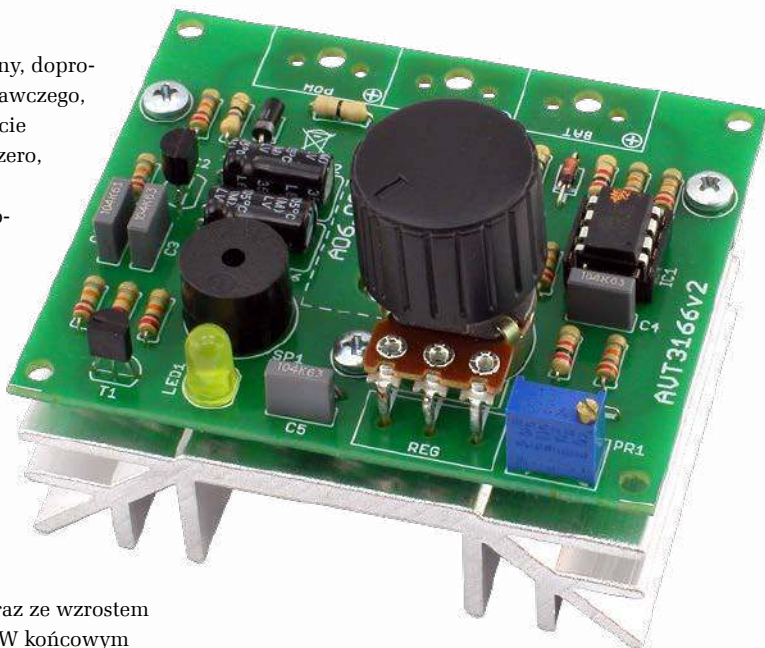
W naszej rubryce „Elektronika dla Ciebie” co miesiąc zachęcamy Cię, drogi Czytelniku, do wykonywania prostych projektów – zabawek, gadżetów itp. Każdy to potrafi. Opis jest zawsze zrozumiały dla nieelektroników, a montaż niemal intuicyjny. A jeśli złapiesz bakcyłę pasji elektronicznej, na co liczymy,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

Regulator do prostownika



Regulator należy traktować jako przystawkę do prostownika. Podstawową funkcją, jaką układ realizuje, jest regulacja prądu ładowania. Regulacja wykonywana jest metodą podobną do regulacji fazowej. Takie rozwiązanie zapewnia dużo mniejsze straty mocy oraz łatwiejsze i elastyczniejsze sterowanie.

Przebieg sinusoidalny, wyprostowany, doprowadzony jest do tranzystora wykonawczego, tranzystor jest otwierany w momencie przejścia przebiegu napięcia przez zero, dzięki temu prąd narasta łagodnie – wraz z przebiegiem sinusoidy. Moment zamknięcia tranzystora jest regulowany, im później to nastąpi, tym większa część przebiegu zostanie przepuszczona i w efekcie będzie płynął większy prąd. Opisany układ nie jest stabilizatorem prądu, nie utrzymuje wartości prądu na stałym poziomie. Pozwala za to ograniczyć początkową wartość prądu, która jest w istocie wartością maksymalną, ponieważ w trakcie ładowania wartość prądu maleje wraz ze wzrostem stopnia naładowania akumulatora. W końcowym etapie prąd ładowania może być dużo mniejszy niż na początku. Wydłuża to czas potrzebny do pełnego naładowania, ale pozwala dokładniej określić



moment zakończenia. Drugą ważną funkcją układu jest kontrola wartości napięcia akumulatora. Dla uzyskania jak najdokładniejszego wyniku pomiar wykonywany jest przy zamkniętym tranzystorze wykonawczym. Taki cykl pomiarowy uruchamiany jest raz na 200 półokresów napięcia zasilającego, czyli co ok. 2 sekund i wtedy nie płynie prąd ładujący (przez ok. 10 ms). Wynik pomiaru nie jest zakłócany prądem ładującym ani pulsowaniem napięcia, ani nawet rezystancją przewodów i połączeń. Jeśli zmierzone napięcie osiągnęło wartość 14,4 V, ładowanie zostaje przerwane, a gdy napięcie spadnie, ładowanie zostaje wznowione. Pod koniec ładowania taki



Wszystkie niezbędne części do tego projektu zawiera kit AVT3166, w cenie 38 zł, dostępny pod adresem: <https://sklep.avt.pl/avt3166.html>

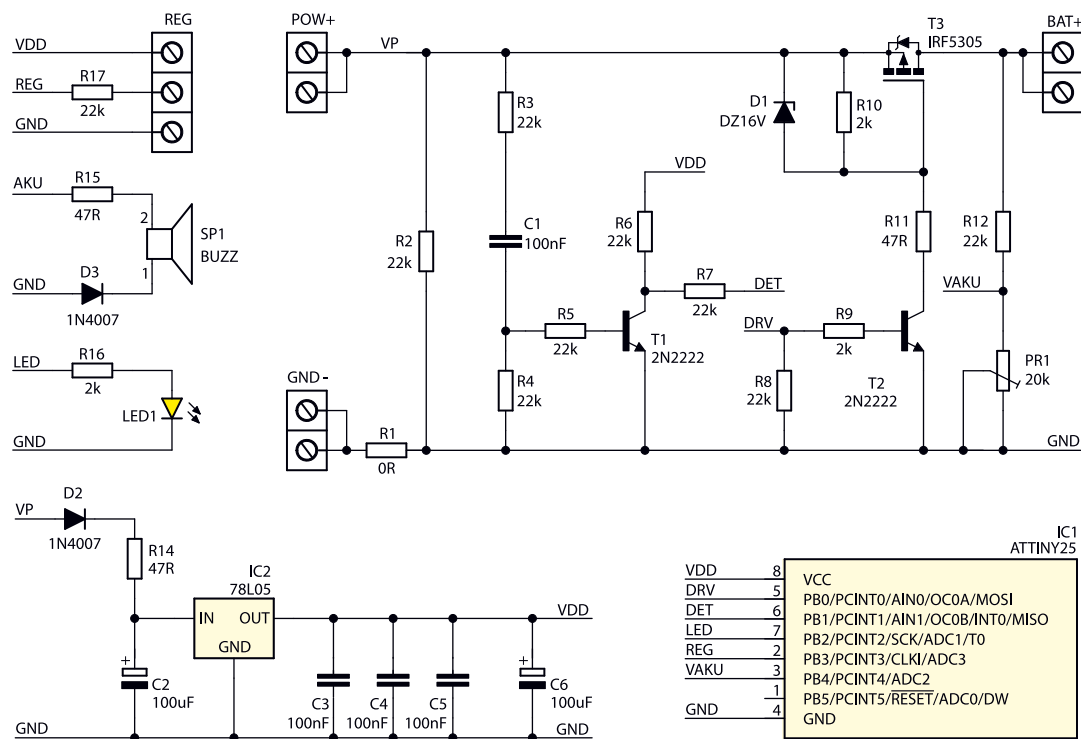


cykl będzie się wielokrotnie powtarzał, ponieważ nawet w pełni naładowany akumulator nie utrzymuje na zaciskach napięcia 14,4 V. Napięcie dosyć szybko spada do wartości ok. 13 V, a potem powinno się ustabilizować w okolicy 12,6 V. Aktualny poziom naładowania sygnalizuje dioda LED. Dioda miga z częstotliwością ok. raz na 2 s, z wypełnieniem zależnym od stopnia naładowania akumulatora. Przy napięciu do ok. 11 V dioda miga z wypełnieniem ok. 5%, im wyższe będzie napięcie, tym dłuższe będzie świecenie diody w każdym cyklu, aż do napięcia 14,4, gdy dioda będzie świeciła światłem ciągłym. W praktyce – nawet po naładowaniu akumulatora dioda może co jakiś czas mignąć, ponieważ wartość napięcia na akumulatorze spada. Będzie to etap tzw. ładowania konserwującego. Dodatkową funkcją układu jest zabezpieczenie przed zwarcieniem. Działanie tej funkcji polega na tym, że dopóki na zaciskach wyjściowych układu nie ma napięcia (nie jest dołączony akumulator), ładowanie nie zostanie załączone. Dopiero dołączenie do wyjścia napięcia o wartości min. 9 V (z akumulatora) odblokuje możliwość ładowania. Stan zacisków wyjściowych sprawdzany jest w każdym półokresie przebiegu napięcia zasilającego, tuż przed załączeniem tranzystora, dlatego nawet

przypadkowe odłączenie przewodów od akumulatora i zwarcie nie spowoduje uszkodzenia układu. Ostatnią funkcją układu jest sygnalizowanie nieprawidłowej biegunowości dołączonego akumulatora. Jeśli do zacisków wyjściowych akumulator dołączymy odwrotnie, to natychmiast odezwie się sygnalizator dźwiękowy. Dla bezpieczeństwa akumulator powinniśmy dołączać, gdy odłączone jest zasilanie prostownika i jeśli nie będzie sygnalizacji dźwiękowej, to możemy podłączyć zasilanie prostownika.

Schemat układu wraz z elementami prostownika widoczny jest na **rysunku 1**. Tranzystor T1 wraz z elementami sąsiadującymi to detektor przejścia przebiegu napięcia przez zero. Tranzystor T2 wraz z elementami sąsiadującymi pracuje jako driver tranzystora wykonawczego T3. Impulsy dodatnie o napięciu 5 V z wyjścia mikrokontrolera zostają zamienione na impulsy masy i otwierają tranzystor wykonawczy T3, natomiast rezystor R10 powoduje jego zamykanie po zakończeniu impulsu. Zbyt wysoka amplituda przebiegu sterującego mogłaby spowodować uszkodzenie obwodu bramki tranzystora MOSFET, dlatego zastosowano diodę Zenera D1. Pozostałe elementy układu to: blok zasilania zbudowany na bazie układu IC2 – 78L05,

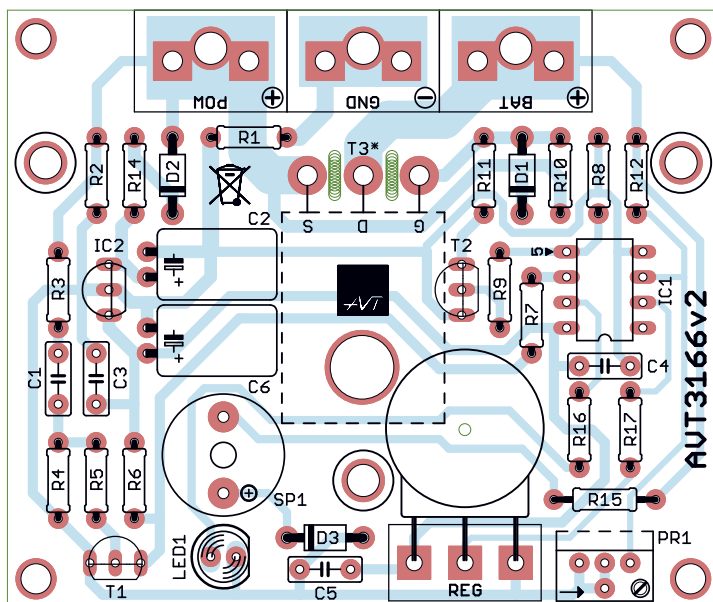
1. Schemat elektryczny układu





mikrokontroler IC1 z zawartym w pamięci programem sterującym, układ sygnalizujący odwrotną biegunowość akumulatora – elementy R15, SP1 i D3, złącze REG służące do dołączenia potencjometru oraz blok pomiaru napięcia – regulowany dzielnik rezystancyjny z elementów R12 i PR1.

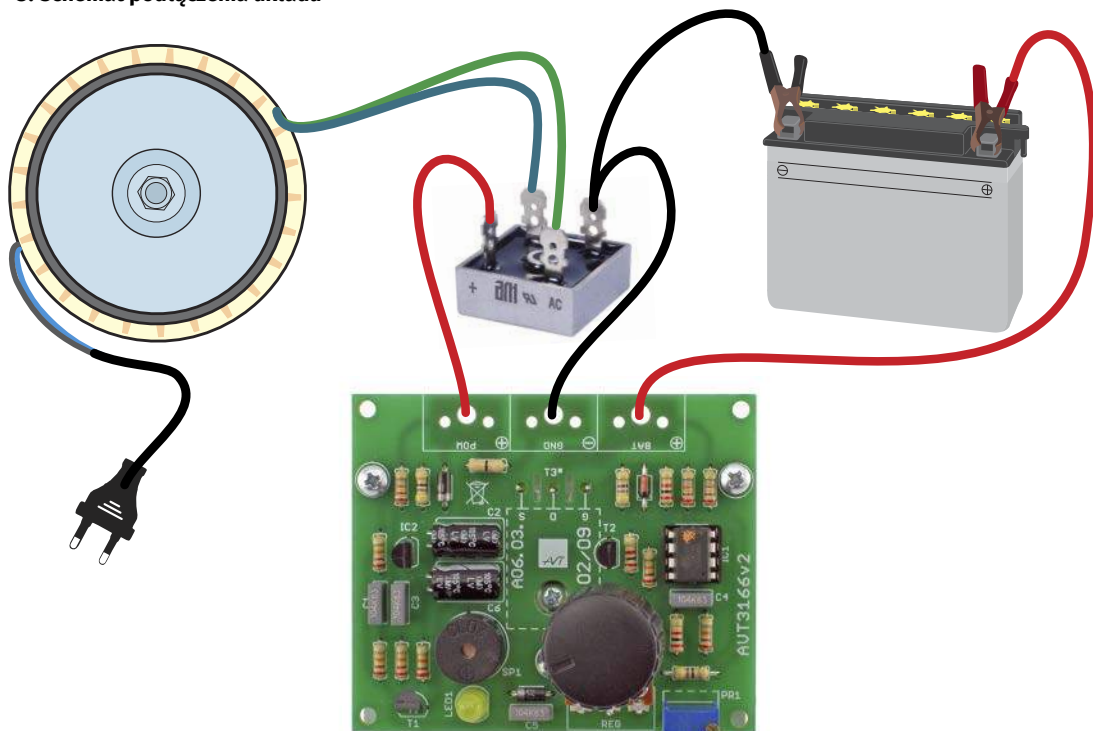
Układ został wykonany na płytce widocznej na **rysunku 2** oraz fotografii otwierającej. Montaż układu wykonujemy według ogólnych zasad. Odkryte ścieżki na płytce należy dodatkowo pociąć. Tranzystor wykonawczy należy zamontować od spodu płytki tak, by jego wkładka radiatorowa była skierowana na zewnątrz, a otwór montażowy pokrywał się z otworem na płytce, jednak nie powinien przylegać do płytki. Radiator należy przykręcić trzema wkrętami, stosując dodatkowo tulejki dystansujące. Na koniec należy przykręcić



2. Schemat montażowy układu

tranzystor do radiatora z zastosowaniem podkładki i tulejki izolującej. Na koniec układ wymaga prostej regulacji opisanej w instrukcji. ■

3. Schemat podłączenia układu





Mechaniczne wahadło

Podstawowym elementem mechanicznego zegara jest wahadło. Wykonując ruch, odmierza krótkie, zawsze jednakowe odcinki czasu. Reszta mechanizmu zegarowego napędza wahadło zegara oraz zlicza liczbę wahaniec i pokazuje upływ czasu za pomocą wskazówek. Proponuję zrobienie modelu mechanicznego wahadła napędzanego grawitacyjnie ciężarkiem.

Było to w słynnej katedrze w Pizie w roku 1584. Galileo Galilei nazywany też Galileuszem patrzył, jak urzeczony na ruch wahającej się lampy oliwnej zawieszony na długim sznurze. Zaczął porównywać czas wahaniec ze swoim pulsem. Zgadzało się, bo każdemu wahaniciu odpowiadała jednakowa liczba uderzeń pulsu. Skonstatował, że wahania lampy przebiegają zawsze w tym samym czasie i mogą być miernikiem upływającego czasu. Dalsze badania nad zjawiskiem wahadła prowadził w domu. Galileusz pierwszy sformułował prawa ruchu wahadła. Jego badania wykazały, że czas trwania wychYLENIA jest zawsze jednakowy, niezależnie od drogi przebytej przez wahadło. Czas zmienia się jednak wtedy, gdy zmieniamy długość sznura, czyli długość ramienia wahadła liczony od osi zawieszenia do końca lampy. Idźmy śladem Galileusza, zróbmy model mechanicznego wahadła z własnym napędem. Nasz model, wiszący na ścianie, nie będzie kompletnym zegarem, ale będzie miał wahadło, zębaty tryb i wychwytowy mechanizm napędowy. Mechanizm ten będzie napędzał wahadło, wykorzystując siłę grawitacji. Po nakręcaniu można będzie przyglądać się, jak pracuje, bo mechanizm nie pozostanie zakryty. Do tego nasz działający model wydaje odgłos przypominający cykanie dawnego ściennego zegara.

Zabieramy się do pracy. Do budowy modelu użyjemy sklejk i drewna, osie będą tkwiły w prawdziwych łożyskach, dzięki czemu nie będzie strat tarcia i mechanizm zadziała możliwie precyzyjnie, a do budowy modelu proponuję użyć całego arsenału elektronarzędzi. I tak wiertarka na kolumnie zapewni nam prostopadłość otworów, a mechaniczna piła włośnicowa



1. Gotowy model mechanicznego wahadła

sprawi, że wycinanie zębów koła zębatego będzie przyjemnością. Szlifierka taśmowa z różnej gradacji papierami ściernymi od 36 do 150 zastąpi pilniki do drewna, a wyrzynarka z laserowym wskaźnikiem upora się błyskawicznie z przycinaniem grubej sklejk dużego koła zamachowego. Jednym słowem, majsterkujemy na bogato.

Działanie mechanizmu: do deski suportu zamocowane jest na osi wahadło. Wahadło zegarowe jest stosowane w zegarach mechanicznych, jak już wiemy, od Galileusza, jako prosty element odmierzający jednakowe odcinki czasu. Zbudujemy je



z prętą zakończonego obciążnikiem nazywanym przez fachowców łożką. Wahadło wykonuje drgania w płaszczyźnie pionowej pod wpływem siły grawitacji. Górna część wahadła, ta powyżej osi, przenosi za pomocą listwy łączącej ruch wahadła na koło zamachowe. Utrzymanie stałej amplitudy ruchu wahadła odbywa się poprzez uzupełnianie strat energii za pomocą mechanicznego wychwyty. Zębaty tryb napędzany jest za pomocą obciążnika zawieszonego na sznurku. Obciążnik zrobiony jest z ciężkiego drewna np. ze szlachetnej dębiny. Aby sznurek obciążony ciężarkiem nie rozkręcił się gwałtownie, z trybem zębatym współpracuje zapadka. Blokując tryb. Zapadka jest podnoszona kształtką umieszczoną na przecie wahadła i pozwala na to, że tryb zębaty może obrócić się tylko o jeden ząb na jeden powrotny ruch wahadła. Widać to na fotografiach. Już widzę radość w momencie, kiedy powiesimy na ścianie nasz mobil, nakręcimy go i uruchomimy.

Narzędzia: papier, ołówek, nożyczki, wiertarka na kolumnie, wiertła piórkowe i otwornica do wiercenia dużych otworów, szlifierka taśmowa, mechaniczna piła włósnicowa, elektryczny strug, ścisiki stolarskie, dremel, wyrzynarka, klej na gorąco w serwowującym go pistolecie, wikolowy klej do drewna.

Materiały: płyta drewniana 500×300×15 milimetrów, sklejka o grubości 12 milimetrów, łąta drewniana, deska o grubości 20 milimetrów lub całowa, pręt drewniany o średnicy 13 i długości 800 milimetrów, listwa 10×10×500, 14 sztuk łożysk kulkowych o wymiarach 7×3×3 milimetra, typ 683ZZ, gdzie średnica zewnętrzna to 7, a średnica wewnętrzna do osi 3 milimetry, sznurek jedwabny lub dratwa szewska, dwa wieszaki do obrazków.

Suport: jest zrobiony z lekkiej deski lipowej o wymiarach 500×350 milimetrów. Do krótszego boku za pomocą gwoździaków mocujemy dwa uchwyty do wieszania obrazków. Starajmy się, żeby były na tej samej wysokości i równoległe do siebie. W przyszłości, gdy będziemy wieszować nasz model, posłużymy się poziomnicą.

Koło zamachowe: dla uniknięcia pomyłek najpierw na papierze rysujemy nasze koło zamachowe o średnicy 200 milimetrów. Dla ozdoby oraz zmniejszenia masy możemy zaprojektować na wzorze a potem wywiercić sześć otworów o średnicy 40 milimetrów. Koła te rysujemy sobie cyrklem co 60° i w jednakowej odległości od środka. Na sklejce o grubości 12 milimetrów przyklejamy po całości płaszczyzny nasz papierowy szablon klejem wikolowym. Gdy klej wyschnie, koło ze sklejki wycinamy wyrzynarką. Pracę bardzo ułatwią uchwyty stolarskie, którymi będziemy mocować naszą przycinaną sklejkę do stołu warsztatowego. Będziemy mieć wtedy

wolne ręce i możliwość skupienia się na prowadzeniu brzeszczotu możliwie blisko linii szablonu. Za pomocą szlifierki taśmowej zaopatrzonej w gruby papier ścierny wyrównujemy linie cięcia do linii szablonu. Widać teraz, jak bardzo potrzebne było naklejenie szablonu na sklejkę. Mamy zgrubnie wycięte koło zamachowe. Teraz pora napunktować osie wszystkich kół zmniejszających masę, co nam potem bardzo ułatwi wiercenie. Punktujemy oczywiście punktakiem i małym młotkiem. Możemy przystąpić do wiercenia. Bardzo istotne jest dokonanie tego wiertarką na statywie. Wiertarka kolumnowa zapewni nam prostopadłość otworów do płaszczyzny sklejki. Centralny otwór proponuję tymczasowo wywiercić wiertłem 5 milimetrów, potem go rozwiernimy, dostosowując do wymiarów zamówionych łożysk. Otwory o średnicy 40 milimetrów bez problemu zrobimy płaskim wiertłem piórowym. Aby koło było naprawdę równe i okrągłe, posłużymy się trikiem zastępującym tokarkę. Do centralnego otworu wkładamy śrubę M5 i mocujemy ją dwiema nakrętkami. Nakrętką i kontrnakrętką. To zapobiegnie odkręceniu się, jak mogłoby się to zdarzyć, gdybyśmy użyli jednej nakrętki. Ta druga mocująca nazywa się kontrnakrętką. Nasze koło mocujemy w wiertarce zamocowanej w uchwycie i papierem ściernym szlifujemy tak, by uzyskać perfekcyjny kształt koła. Szablon z powierzchni zdzieramy papierem ściernym za pomocą szlifierki taśmowej, a brzegi otworów dremelem z tarczką listkową.

Poprzeczka: przenosi ruch wahadła na koło zamachowe. To listewka o przekroju 10×10 milimetrów i długości 220 milimetrów. Na jej końcach wiercimy dwa otwory 7 mm, w które wklejamy klejem wikolowym małe łożyska. Powinny być umieszczone równo z powierzchnią płaszczyzny listewki. Widać to na fotografii. Z innymi łożyskowanymi elementami postąpimy podobnie.

Wahadło: pręt o średnicy 13 i długości 700 milimetrów. W odległości 530 milimetrów będzie umieszczona łożyskowana oś wahadła. Wahadło zakończone jest łożką o średnicy 110 milimetrów. Łóżkę wycinamy z litego drewna wyrzynarką i obrabiamy papierem ściernym do uzyskania odpowiedniego kształtu. W górnej części wiercimy otwór na osi poprzeczki przenoszącej ruch na koło zamachowe. Osiami są śruby M3 o długości 20 milimetrów.

Tryb zębaty: jest wycięty z 12-milimetrowej grubości sklejki. Zaczniemy jednak, a jakże, od szablonu. Na papierze rysujemy koło o średnicy 100 milimetrów i drugie centralnie o średnicy 80 milimetrów. Od środka rysujemy promienie precyzyjnie co 30°. Tak jak widać na rysunku, łączymy prostymi te punkty na obwodzie, tworząc 12 zębów. Koło wycinamy ze sklejki za pomocą



2. Koło zapadkowe i zapadka; **3.** Wycinanie zębów koła za pomocą piły włosowej; **4.** Małutkie łożyska zapewnią znikome tarcie; **5.** Wycinanie krążka ze sklejk; **6.** Montaż mechanizmu obciążnika; **7.** Tak obrabiamy koła ze sklejk; **8.** Koło zamachowe wycinamy wyrzynarką; **9.** Zgrubnie obrabiamy koło szlifierką taśmową; **10.** Po napunktowaniu wiercimy otwór na oś; **11.** Wiercenie otworów 40 milimetrów w kole zamachowym; **12.** Łożyska o średnicy 7 milimetrów; **13.** Koło zamachowe;



14. Zapadka, która ma współpracować z kołem zębatym; **15.** Do klejenia bębna pomocne są ściski stolarskie; **16.** Łożyska wklejamy po jednym z każdej strony płaszczyzny listwy; **17.** Zamontowane łożysko; **18.** Uchwyt pokręta; **19.** Do tego typu prac trzeba użyć wiertarki kolumnowej; **20.** Łezka wahadła obrobiona za pomocą papieru ściernego; **21.** Na sznurku zamocowanym do bębna zawiesznie ciężarek napędu; **22.** Warto sprawdzić położenie deski za pomocą poziomnicy; **23.** Mechanizm napędowy wahadła; **24.** Konstrukcja prowadząca ciężarek;

piły włósnicowej lub wyrzynarki. Całość wyrównujemy i wygładzamy papierem ściernym.

Zapadka: z deseczki o grubości 15 milimetrów wycinamy kształt o długości 170 milimetrów, tak jak to widać na rysunku. Istotny jest ząb współpracujący z trybem. Widzimy to na rysunku. Z jednej strony wiercimy otwór o średnicy 7 milimetrów pod łożyska. Dwa łożyska po jednym z każdej strony klejamy klejem wiskolowym równo z powierzchnią płaszczyzny drewna.

Popychacz: jest przyklejony do ramienia wahadła i składa się z dwóch deseczek o grubości 5 milimetrów lub patyczków od lodów. Jego dokładne położenie i wymiary ustalimy dopiero podczas montażu. Wówczas będziemy mogli określić, w którym miejscu ramienia wahadła powinien być przyklejony i jaką dokładnie muszą mieć długość wystające elementy, aby prawidłowo współpracowały z zapadką.

Bęben napędu: ma kształt walca o wymiarach: grubość 15 i średnica 30 milimetrów; wycinamy go z litego drewna. Na nim będzie nawinięty sznurek.

Pokrywa bębna: koło wycięte ze sklejkii o grubości 10 milimetrów obrobione papierem ściernym. Na powierzchnię pokrywy doklejamy klejem wiskolowym uchwyty pokrętła w kształcie krzyża zrobionego z listewek lub patyczków od lodów.

Obciążnik: dwa jednakowe drewniane elementy o wymiarach 110×100×20 milimetrów. Przy dolnej krawędzi pomiędzy te klocki wklejamy listwę 100×35×10 milimetrów. Po wyschnięciu kleju obrabiamy obciążnik papierem ściernym. Przy górnej krawędzi obciążnika wiercimy dwa otwory o średnicy 3 milimetrów w odległości 70 milimetrów. W otwory wkładamy gwoździe, które będą osiami rolek prowadzących sznurek, na którym zawieszony jest obciążnik. Widzimy to na fotografii.

Suport obciążnika: robimy z listwy o wymiarach 200×35×10 milimetrów. Wiercimy 5 otworów. Dwa z nich posłużą do przymocowania listwy do spodu boku deski suportu. W pozostałych trzech zamocujemy zrobione z nitów rolki prowadzące sznurek od mechanizmu bębna zębatki do obciążnika. Widzimy to na fotografii.

Montaż modelu: zaczynamy od zamocowania na ścianie deski suportu, sprawdzając jej wyrównanie za pomocą poziomnicy. Z prawej strony na swojej osi mocujemy wahadło. Na osi montujemy podkładki dystansowe, tak by nie tarło o deskę suportu. Koło zamachowe powinno mieć swoją oś umieszczoną na wysokości górnej części wahadła. Tu także zastosujemy podkładki dystansowe. Koniec górnej części wahadła łączymy za pomocą poprzeczki z kołem zamachowym. Wyznaczamy miejsce na oś bębna zapadki. Po zamontowaniu tego bębna z łatwością znajdziemy miejsce, gdzie powinna być zamocowana oś zapadki. Ma ona, opadając pod własnym ciężarem, zasekować

się z trybem zębatym a jednocześnie podnosić się wtedy, kiedy bęben z trybem jest ręcznie obracany celem podciągnięcia obciążnika. Nie trzeba dodawać, że zapadka ma precyzyjnie trafiać w zęby koła zapadkowego. Do bębna przyklejamy klejem wiskolowym sznurek o długości około 3 czy 4 metrów, na którym będzie zawieszony obciążnik. Długość sznurka zależy od wysokości, na której zawieszamy model

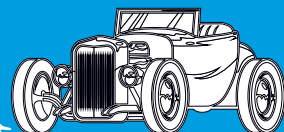
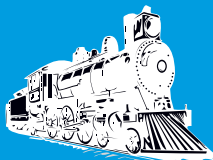
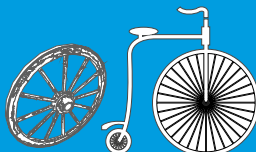
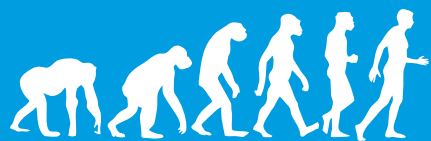
oraz jak długo będzie pracował mechanizm. Najpierw nawijamy na bęben trzy czy cztery zwoje i powlekamy klejem, zwracając uwagę na koniec sznurka. Gdy klej wyschnie, nawijamy resztę sznurka na bęben. Do ramienia wahadła przyklejamy dwie cienkie listewki poruszające mechanizm bębna za każdym razem, gdy wahadło jest wychylone po lewej stronie. Tu także potrzebna jest precyzja, ale klej na gorąco pozwoli nam wyregulować położenie tych popychaczy, zanim stwardnieje. Niższa listewka blokuje tryb zapadki w chwili, gdy górna podnosi zapadkę. Gdy odpowiednio dobierzemy długości tych listewek, ten mechanizm wychwytowy będzie pozwalał na obrót koła zapadki tylko o jeden ząb na każdy ruch powrotny wahadła. Pewnie to się nie uda za pierwszym razem i trzeba będzie przekleić położenie górnej listewki, ale wreszcie musi się to powieść, bo inaczej mechanizm nie zadziała. Zawieszamy nasz obciążnik na sznurku, który jest poprowadzony w rolkach i model jest gotowy.

Zabawa: Nakręcamy mechanizm modelu, obracając uchwyt bębna zapadkowego i poruszamy wahadło. Powinien bez problemu ruszyć. Słychać dźwięk pracującej zapadki, a wahadło porusza się, odmierzając czas. Galileusz, gdyby żył, z pewnością by nas pochwalił za takie ujęcie tematu wahadła, ale pewnie jeszcze bardziej zadziwiłyby go nasze współczesne elektroniczne sterowane internetowo zegarki. ■

Adam Łowicki



25. Działający model mechanicznego wahadła



OGNIWA PALIWOWE

1801

Według przekazów, brytyjski chemik i fizyk, Humphry Davy, historycznie jako pierwszy demonstruje koncepcję ogniwa paliwowego.

1838–42

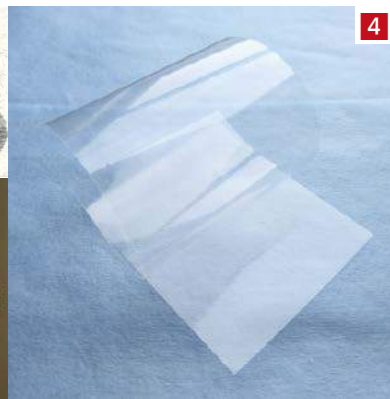
William Grove wpada na pomysł wykorzystania metody pozyskania energii elektrycznej z procesu elektrolitycznego dysocjacji cząsteczek wody na wodór i tlen, opisany wcześniej m.in. przez Williama Nicholsona i Anthony'ego Carlisle'a. To właśnie Williamowi Grove'owi przypisuje się dokonanie pierwszej znanej demonstracji ogniwa paliwowego, urządzenia zwanego przez niego „baterią gazowoltaiczną” w 1842 r., wykorzystującą reakcję łączenia wodoru z tlenem do bezpośredniego wytwarzania prądu elektrycznego. Ogniwo takie, bez części ruchomych, składało się z dwóch platynowych elektrod, umieszczonych w oddzielnych szklanych cylindrach. Jeden z cylindrów napełniony był tlenem, drugi wodorem, a oba zanurzone były w rozcieńczonym kwasie siarkowym (elektrolicie). Ogniwo takie działało bezszumowo, a jego jedyną substancją odpadową jest woda. Inne źródła mówią, że Grove wykorzystał prace niemiecko-szwajcarskiego chemika Christiana Friedricka Schönbeina, który odkrył zasadę działania ogniw wodorowych w 1838 roku i opublikował na ten temat artykuł w „The London and Edinburgh Philosophical Magazine and Journal of Science”.

1889

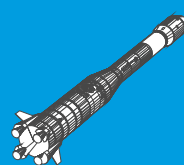
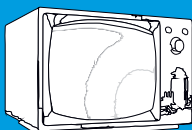
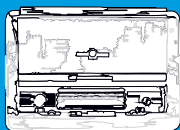
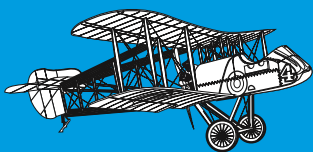
Termin „ogniwo paliwowe” został po raz pierwszy użyty przez Charlesa Langerę i Ludwiga Monda, którzy badali ogniwa paliwowe wykorzystujące jako paliwo gaz węglowy. Przeprowadzili serię eksperymentów z wykorzystaniem gazu pochodzącego z węgla. Używali elektrod wykonanych z cienkiej, perforowanej platyny i mieli wiele trudności z płynnymi elektrolitami. Udało im się uzyskać sześć amperów na stopę kwadratową (powierzchnia elektrody) przy napięciu 0,73 V.

1932–59

Profesor inżynierii z Cambridge Francis Bacon zmodyfikował stary wynalazek Monda i Langerę, opracowując pierwsze alkaliczne ogniwo paliwowe (AFC). Jednak dopiero w 1959 r. Bacon zademonstrował praktyczny system ogniw paliwowych o mocy 5 kW do zasilania spawarki. Mniej więcej w tym samym czasie Harry Karl Ihrig we współpracy z Siłami Powietrznymi USA zamontował zmodyfikowane ogniwo Bacona o mocy 15 kW w ciągniku rolniczym Allis-Chalmers. System ten wykorzystywał wodorotlenek potasu jako elektrolit oraz sprężony wodór i tlen jako reagenty. Opracował następnie szereg pojazdów napędzanych ogniwami paliwowymi, w tym wózek widłowy, wózek golfowy i łódź podwodną.



1. Ilustracja ogniwa Williama Grove'a, 2. Ogniwa paliwowe na pokładzie jednego ze statków wykorzystanych w programie kosmicznym Apollo, 3. Schemat wnętrza samochodu prototypowego Chevrolet Electrován napędzanego ogniwami paliwowymi, 4. Membrana nafionowa



1955–58

W. Thomas Grubb, chemik pracujący dla General Electric Company (GE), zmodyfikował oryginalny projekt ogniwa paliwowego, stosując jako elektrolit membranę jonowymienną z sulfonowanego polistyrenu. Trzy lata później inny chemik GE, Leonard Niedrach, opracował sposób osadzania platyny na membranie, która służyła jako katalizator niezbędnych reakcji utleniania wodoru i redukcji tlenu. Rozwiązanie stało się znane jako „ogniwo paliwowe Grubba–Niedracha”.

lata 60.

Ogniwa paliwowe skonstruowane w latach 50. nie znalazły szerszego praktycznego zastosowania aż do momentu, w którym Stany Zjednoczone zdecydowały się wykorzystać ogniwa z membranami polimerowymi, AFC, jako źródło elektryczności i wody w swoim programie kosmicznym. W ogniwa paliwowe zostały wyposażone takie statki jak np. Gemini 5, zostaje w programie Apollo czy stacja orbitalna Skylab. Do produkcji ogniwi paliwowych stosowano wówczas niezwykle drogie materiały, a do ich działania były potrzebne bardzo wysokie temperatury oraz tlen i wodór o niskim poziomie zanieczyszczenia. Koszt ich wytworzenia sięgał wówczas 100 tysięcy dolarów za kilowat, jednak zdecydowano się na ich użycie, gdyż wodór i tlen wykorzystywano jako paliwo i dzięki temu na statkach kosmicznych były dostępne w dużych ilościach. Dodatkowym atutem ogniwi była produkcja wody pitnej na potrzeby załóg. Firma International Fuel Cells (IFC, później UTC Power) opracowała ogniwo AFC o mocy 1,5 kW, które miało być wykorzystywane w misjach kosmicznych Apollo (2). Ogniwo paliwowe zapewniało energię elektryczną oraz wodę pitną dla astronautów na czas trwania misji. Następnie firma IFC opracowała ogniwo AFC o mocy 12 kW, wykorzystywane do zasilania pokładowego podczas wszystkich lotów promów kosmicznych.

1966

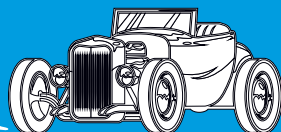
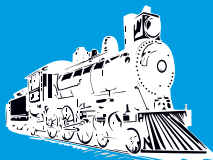
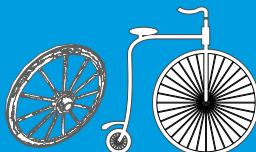
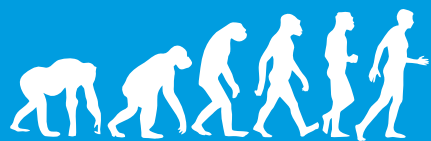
Firma General Motors opracowała pierwszy pojazd drogowy napędzany ogniwami paliwowymi – Chevrolet Electrován (3), który miał ogniwo paliwowe PEM firmy Union Carbide, zasięg prawie 200 km i prędkość maksymalną ponad 100 km/h. Miał tylko dwa miejsca siedzące, ponieważ stos ogniwi paliwowych oraz duże zbiorniki wodoru i tlenu zajmowały tylną część furgonetki. Zbudowano tylko jeden egzemplarz, a projekt uznano za nieopłacalny.

1967–69

Opracowanie przez Walthera Grota z firmy DuPont nafionu, polimeru, który zrewolucjonizował technikę ogniwi paliwowych. Jest to syntetyczny kopolimer tetrafluoroetenu (monomeru teflonu) i perfluorowanego eteru oligowinyloвого zakończonego silnie kwasową resztą sulfonową. Nafion najczęściej formowany jest w postaci cienkiej folii (4) i wykorzystywany jako membrana przewodząca protony (tzw. jonomer), jednocześnie nieprzewodząca elektronów lub anionów.

lata 70.

Obawy o dostępność ropy naftowej wskutek kryzysu paliwowego oraz początek programów ograniczania zanieczyszczenia powietrza doprowadziły do opracowania szeregu demonstracyjnych pojazdów z ogniwami paliwowymi, w tym modeli zasilanych wodorem lub amoniakiem, a także silników spalinowych zasilanych wodorem. W latach 70. kilku niemieckich, japońskich i amerykańskich producentów pojazdów oraz ich partnerów rozpoczęło eksperymenty z pojazdami napędzanymi ogniwami paliwowymi, zwiększając gęstość mocy stosów ogniwi paliwowych z membraną do wymiany protonów (PEMFC) i opracowując systemy magazynowania paliwa wodowego. Finansowanie ze strony amerykańskiego wojska i zakładów energetycznych umożliwiło rozwój techniki ogniwi paliwowych ze stopionym węglanem (MCFC).



lata 80. XX wieku

Kontynuacja prac badawczych ogniw paliwowych w transporcie. Marynarka Wojenna Stanów Zjednoczonych zleciła badania nad wykorzystaniem ogniw paliwowych w łodziach podwodnych, gdzie wysokowydajne, bezemisyjne i niemal bezgłośne działanie oferowało znaczne korzyści operacyjne.

lata 90. XX wieku

Skupienie uwagi na technikach ogniwa paliwowego PEMFC oraz ogniwa paliwowego ze stałym tlenkiem (SOFC), szczególnie w przypadku małych zastosowań stacjonarnych. Ze względu na niższy koszt jednostkowy i większą liczbę potencjalnych rynków postrzegano je jako oferujące bardziej bezpośrednie możliwości komercyjne, na przykład zasilanie awaryjne dla zakładów telekomunikacyjnych i mikrokogeneracji mieszkaniowej. W Niemczech, Japonii i Wielkiej Brytanii zaczęto przeznaczać znaczne fundusze na rozwój technologii PEMFC i SOFC do zastosowań w mikrokogeneracji mieszkaniowej. Polityka rządowa mająca na celu promowanie czystego transportu również przyczyniła się do rozwoju PEMFC w zastosowaniach motoryzacyjnych.

1991

Roger Billings konstruuje pierwszy normalnie funkcjonujący samochód napędzany wodorowymi ogniwami paliwowymi (5). Samochód, nazwany LaserCel 1, został zaprezentowany w Filadelfii. Jest to pierwsze zgłoszone funkcjonalne zastosowanie ogniw paliwowych w małym pojeździe transportowym.

2000

W ramach projektu HyFleet/CUTE (6) w Europie, Chinach i Australii wdrożono dziesiątki autobusów napędzanych ogniwami paliwowymi. Autobusy były, i nadal są, postrzegane jako obiecujące wczesne rynkowe zastosowanie ogniw paliwowych ze względu na połączenie wysokiej wydajności, zerowej emisji i łatwości tankowania, a także ze względu na to, że pojazdy te poruszają się po ustalonych trasach i są regularnie tankowane wodorem w swoich bazach.



5



7



8

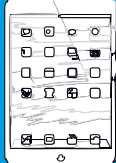
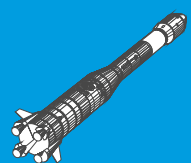
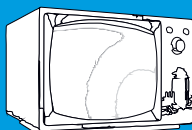
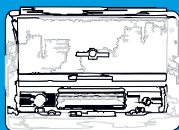
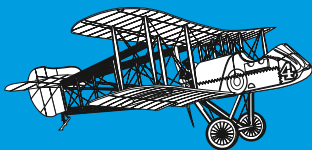


6



9

5. Roger Billings ze swoim pojazdem LaserCel 1, **6.** Autobusy w ramach projektu HyFleetCUTE, **7.** Napędzany ogniwami paliwowymi samolot Boeinga Diamond DA20, **8.** Wodorowe motocykle ENV firmy Intelligent Energy, **9.** Pociąg Coradia iLint Alstomu



I dekada XXI wieku

Tysiące pomocniczych jednostek zasilających PEMFC i DMFC zostało wprowadzonych na rynek w zastosowaniach rekreacyjnych, takich jak łódzie i kampery, a podobnie duża liczba jednostek z mikroogniwami paliwowymi została sprzedana w sektorze przenośnym, w zabawkach i zestawach edukacyjnych. Popyt ze strony wojska zaowocował również setkami przenośnych jednostek zasilających z ogniwami paliwowymi, które trafiły do użytku wojskowego, zapewniając zasilanie urządzeń łączności i nadzoru, zmniejszając zarazem obciążenie żołnierza związane z noszeniem ciężkich zestawów akumulatorów. Zakrojony na szeroką skalę program kogeneracji mieszkaniowej w Japonii przyczynił się do zwiększenia w tym kraju dostaw komercyjnych stacjonarnych urządzeń PEMFC. Urządzenia te zaczęły być instalowane w japońskich domach od 2009 roku, a ich liczba obecnie osiągnęła kilkanaście tysięcy.

2003–08

Koncern Boeing rozpoczyna prace nad wersją demonstracyjną samolotu zasilanego wodorowymi ogniwami paliwowymi PEM (ang. Fuel Cell Demonstrator Airplane), by po pięciu latach, w 2008 roku przeprowadzić lot dwumiejscowej awionetki Diamond DA20 (7).

2005

Firma Intelligent Energy produkuje pierwszy na świecie motocykl ENV (8) skonstruowany pod kątem zasilania ogniwami paliwowymi. Udało się pokonać bariery miniaturyzacyjne, tworząc ogniwa polimerowe zasilane metanolem – DMFC, co pozwala na zastosowanie ich w przenośnym sprzęcie elektronicznym, używanym z dala od źródeł ładowania akumulatorów, np. w komputerach przenośnych – laptopach, czy telefonach komórkowych.

2009

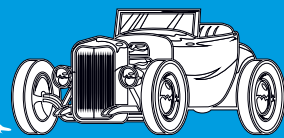
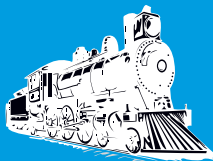
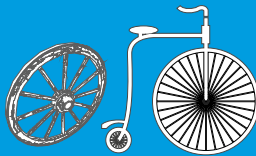
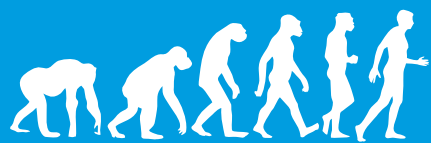
Wprowadzenie przez firmę Toshiba ładowarki Dynario z ogniwem paliwowym. W obliczu limitowanej serii produkcyjnej, liczącej trzy tysiące sztuk, popyt na dynario znacznie przewyższył podaż. Jednocześnie zastosowanie stacjonarnych ogniw paliwowych gwałtownie wzrosło wraz z rozpoczęciem realizacji japońskiego projektu Ene-Farm, a w Ameryce Północnej zaczęto stosować ogniwa paliwowe w zasilaczach awaryjnych (UPS).

2017

W Chinach w trasy ruszają pierwsze na świecie tramwaje na wodór. Firma CRRC Qingdao Sifang tworzy pierwszą partię tramwajów zasilanych ogniwami wodorowymi. Pojazdy znalazły zastosowanie na budowanej w Foshan (południowe Chiny, prowincja Guangdong) linii Gaoming (długość 17,4 km, 20 przystanków). Pierwszy odcinek miał długość nieco ponad 6 km. Tramwaje mogły przewozić 285 pasażerów i rozwijać prędkość 70 km/h. Ogniwa do ich napędu dostarczyła kanadyjska firma Ballard Power Systems, specjalizująca się w produkcji ogniw do celów transportowych.

2018

Pociągi wodorowe zaczynają kursować w Niemczech. Pociągi Alstomu Coradia iLint (9) wyjechały na tory w niemieckiej Dolnej Saksonii. Pierwsza linia obsługiwana przez napędzane wodorem składy liczyła prawie 100 km biegną przez Cuxhaven, Bremerhaven, Bremerförde i Buxtehude, zastępując istniejącą flotę pociągów spalinowych.



WYNALEZKÓW HISTORIĘ ODKRYJ Klasyfikacja ogniw paliwowych

PEM (Proton Exchange Membrane), PEMFC

Ogniwa paliwowe PEM zasilane są czystym wodorem lub reformatem. Membraną ogniwa PEM jest materiał polimerowy np. nafion. Charakterystyczną cechą ogniw PEM jest duża sprawność w produkcji energii elektrycznej, do 65% oraz mała ilość wydzielanego ciepła. Ogniwa PEM są stosowane głównie do napędzania pojazdów oraz do budowy stacjonarnych i przenośnych generatorów energii.

Alkaliczne ogniwo paliwowe AFC (ogniwo paliwowe Bacona)

Jedna z najlepiej rozwiniętych technologii ogniw paliwowych. AFC zużywa wodór i czysty tlen do wytworzenia energii elektrycznej oraz wody i ciepła. Elektrolitem jest stężony roztwór wodorotlenku potasu (KOH). Zaliczane do ogniw wodorowych (wodorowo-tlenowych). Reakcja przebiega w temperaturze od 100 do 250°C. Ogniwa AFC zastosowane zostały na promie kosmicznym Apollo do generacji energii elektrycznej i ciepła. Ogniwa AFC są wrażliwe na wszelkie zanieczyszczenia i wymagają paliwa o dużej czystości, co stanowi przeszkodę w ich komercjalizacji.

PAFC – ogniwa z kwasem fosforowym

Jako elektrolit wykorzystuje się kwas fosforowy. Ogniwa PAFC są stosowane do budowy systemów generacji energii elektrycznej i ciepła. W związku ze swoją wielkością, masą oraz znacznym czasem rozruchu znajdują powszechne zastosowanie w rozwiązaniach do stałej zabudowy. Zaletą tego typu ogniw jest wysoka sprawność, rzędu 80 proc. Temperatura pracy PAFC wynosi 150–200°C. Ciepło uzyskiwane za pomocą tego ogniwa wykorzystuje się w kogeneracji. Dodatkowo para wodna produkowana przez ogniwo może być zamieniana na ciepło. Elektrolitem w ogniwie PAFC jest kwas fosforowy (V).

MCFC – ogniwa paliwowe ze stopionym węglem

Rodzaj wysokotemperaturowego ogniwa paliwowego, pracującego przy $t=600^{\circ}\text{C}$ i wyżej. Ogniwa te powstały w latach sześćdziesiątych XX w. i były bardzo drogie ze względu na elektrody wykonane z metali szlachetnych. W latach siedemdziesiątych XX w. elektrody zaczęto wykonywać z niklu i jego tlenku oraz chromu. Dzięki temu udało się obniżyć nie tylko cenę, ale i zwiększyć moc z 10 mW/cm² do 150 mW/cm². Elektrolitem w ogniwach MCFC jest stopiony

węgiel Li/K. Wysoka temperatura reakcji zachodzącej w ogniwie pozwala na stosowanie szerokiego wachlarza paliw (gaz ziemny, benzyna, wodór, propan).

SOFC – ogniwa paliwowe ze statym tlenkiem

Elektrolitem jest zestalony tlenek, zaś membranę wykonuje się z ceramiki tlenkowej. Pracują w wysokich temperaturach od 650 do 1000°C. Rezultatem wysokiej temperatury reakcji przebiegającej w ogniwie SOFC jest wysoka sprawność w systemach generacji energii elektrycznej i ciepła – nawet 85 proc. Ogniwa SOFC stosuje się w budowie stacjonarnych generatorów energii elektrycznej i ciepła, pracujących w sposób ciągły z jednakowym obciążeniem.

DMFC – paliwowe zasilane bezpośrednio metanolem

Ogniwa zasilane metanolem są to zmodyfikowane ogniwa polimerowe (PEMFC). Metanol jest paliwem łatwym w składowaniu, co w połączeniu z niską temperaturą reakcji (około 80°C) czyni ogniwo DMFC idealnym do zastosowań jako bateria małej mocy. Ogniwa DMFC mają polimerową membranę, taką jak ogniwa PEM. Różnica pomiędzy ogniwem DMFC a ogniwem PEM tkwi w konstrukcji anody, która w ogniwie PEM pozwala na dokonanie wewnętrznego reformingu metanolu i uzyskanie wodoru do zasilania ogniwa. Ogniwo DMFC charakteryzuje niższa sprawność w porównaniu do ogniwa PEM i wynosi 40 proc. Przy obecnym rozwoju technologii ogniwa zasilane bezpośrednio metanolem mogą produkować ograniczone moce, za to mogą magazynować dużo energii na małej przestrzeni, co oznacza, że mogą produkować niewielkie ilości energii elektrycznej w długim okresie (co czyni je użytecznymi np. w telefonach komórkowych, laptopach).

DFAFC – ogniwa zasilane kwasem mrówkowym

Bardzo skuteczną alternatywą dla DMFC jest technologia bezpośredniego ogniwa paliwowego z kwasem mrówkowym o gęstości od pięciu do sześciu razy większej mocy. Tak duże gęstości pomagają w miniaturyzacji ogniw paliwowych, zwłaszcza w urządzeniach przenośnych. Poprawiają także efektywność ogniwa paliwowego bez uszczerbku wydajności netto energii elektrycznej. Zastosowanie kwasu mrówkowego jako paliwa daje takie korzyści, jak mniejsze zużycie paliwa, wyższe stężenie paliwa po stronie anody, dobra kinetyka anody w temperaturze pokojowej i wysoka gęstość mocy. ■

M.U.

*** Pisownia oryginalna ***



PRZEGLĄD TECHNICZNY Rozwój sieci telegraficznej w Stanach Zjednoczonych Ameryki Północnej

Ilość czynnych aparatów telegraficznych, należących do tow. Bell wynosiła w dn. 30 września 1919 r. 7 201 757, inne towarzystwa połączone z siecią międzymiastową tow. Bell posiadały 3 790 568 aparatów; prócz tego było jeszcze czynnych 1 012 000 aparatów, należących do innych towarzystw, nie połączonych z siecią międzymiastową tow. Bell. A zatem ogólna ilość aparatów wynosiła w St. Zj. Am. Pn. 12 004 325, co stanowi 1 aparat na 10 mieszkańców. Z powodu wielkich odległości między ośrodkami przemysłowymi i handlowymi bardzo duże zastosowanie mają w Ameryce telefony międzymiastowe. Choćby pocztą działa bardzo sprawnie, jednak list idzie z New-Yorku do Chicago (1 600 km) dwa dni, a podróż najszybszym pociągami wymaga 20 godzin, tak że dobra komunikacja telefoniczna jest niezbędna. Funkcjonuje ona bardzo dobrze, gdyż np. na połączenie między New-Yorkiem a Chicago bardzo rzadko czeka się dłużej niż 20 minut.

4 października 1921

Dział taboru kolejowego zakładów Kruppa

jest obecnie zatrudniony całkowicie. Zakłady wyrabiają miesięcznie 20 parowozów, a niedawno wypuściły tysiączny, od czasu wznowienia pracy po wojnie, wagon towarowy. Parowozy przeznaczone są dla Rosji. Wykonywane jest również duże zamówienie na obręcze (bandaże) dla Rosji.

4 października 1921

Olbrymi most przez Hudson River w Nowym Jorku

Po dwukrotnym zawaleniu się podczas budowy olbrzymiego mostu w Montreal przez rzekę św. Wawrzyńca i udatnem przeprowadzeniu pod Hudson River tunelu dla kolei Pensylwańskiej do nowej centralnej stacji w środku miasta, wśród inżynierów nowojorskich zapanowało zdanie, że okres budowy mostów w Nowym Jorku minął i że odtąd będą budowane tunele podwodne. Jako argument przeciwko mostom wysuwano konieczność znacznego podnoszenia ich ponad zwierciadło wody dla ułatwienia żeglugi, co doprowadzało do budowy długich rampjazdowych, wymagających w zabudowanej części miasta kosztownego wywłaszczania (zbudowane w pierwszym dziesięciu XX wieku mosty Williamsburg i Manhattan przez East River). Jednakże pogląd ten wobec inicjatywy konstruktorów mostowych się nie utrzymał i dziś rozważany jest projekt mostu przez rzekę Hudson, który pod względem śmiałości budowy i ogromu pozostawia w cieniu wszystko co dotąd na tem polu zdziałano. Projekt obejmuje szereg stacji końcowych i rozrządzających poszczególne koleje w New Jersey, stację centralną osobową w środku miasta na 42 ej ulicy i łączący je most przez rzekę Hudson naprzeciw 59 ej ulicy. Most jest systemu wiszącego o środkowym prześle rozpiętości 988 m i dwóch bocznych po 521 m. Jeźdźnia – dwupiętrowa, 71,5 m szeroka, zawiera na pokładzie górnym miejsce dla 16 linii pojazdów, 2 tory dla omnibusów, 2 dla tramwaju linowego i 2 chodniki, każdy po 4,5 m. Pokład dolny dźwiga 6 par torów kolejowych. Prześwit nad zwierciadłem wody wynosi 47,2 m w środku i po 42,7 w koło wież. Jeźdźnia będą dźwigały dwa łuki wiszące, złożone każdy z dwóch linii potrójnych łańcuchów o 18,3 m jedna nad drugą. Każdy łańcuch składa się z ogniw stalowych, połączonych sworzniami. Dla zabezpieczenia od rdzy każda trójka łańcuchów będzie otoczona powłoką z blachy

miedzianej, tworzącej razem komorę do ułatwienia oględzin łańcuchów. Wieże będą niezwykłej wielkości do 256 m nad źródłem wody i do 50 m poniżej jego na fundamentach o wymiarach 122×61 m. Wieże będą ze stali, obmurowane betonem w celu ochrony od rdzy. Największy spadek na dojazdach wynosi dla kolei 0,025, dla jazdy kołowej 0,0875, na moście spadek, zarówno dla kolei, jak dla jazdy zwykłej 0,025. Przewidywany koszt mostu wynosi 100 000 000 dolarów, a stacji kolejowych z obu stron mostu 115 000 000, razem 215 000 000 dolarów. (...) Względem, przemawiającym na korzyść konstrukcji nadwodnej w tym wypadku jest możliwość ześrodkowania w jednej budowli tak znacznej liczby torów kolejowych i tak szerokiej jezdni i chodników, co przy używanym dotąd systemie tuneli rurowych byłoby zbyt skomplikowane dla konstrukcji podwodnej.

11 października 1921

PRZEGLĄD PRZEMYSŁOWO-HANDLOWY Przemysł emalajerski

Do artykułów wybitnie eksportowych naszego przemysłu należą naczynia emalajowane. Pierwsza fabryka naczyń emalajowanych w Polsce (a jednocześnie w całym państwie Rosyjskim) – „Wulkan” – powstała w Warszawie w r. 1881. Oprócz „Wulkanu” kongresówka posiada 4 fabryki naczyń emalajowanych, które w r. 1912 zatrudniały 5,948 robotników; produkcja roku 1912 wynosiła 860,880 tysięcy pudów. – W Cesarstwie, które stanowiło główny rynek zbytu dla naszych naczyń emalajowanych, istniały tylko 2 fabryki: „Ługański” w Ługańsku i „Werdemüller i S-ka” w Rydze. Trzecia fabryka w Ługawie na Uralu rozpoczęła działalność dopiero przed samym wybuchem wojny. Produkcja pierwszych dwóch fabryk wynosiła około 60,000 pudów rocznie, czyli zaledwie 8% produkcji kongresówki, która prawie całkowicie pracowała dla rynków wschodnich, mianowicie: nasz przemysł emalajerski umieszczał przed wojną w Rosji 90% swej produkcji.

Inne rynki były dotychczas dla naszego przemysłu emalajerskiego zamknięte. Panowały na nich niepodzielnie wyroby zachodnioeuropejskie, z którymi współzawodniczyć nie mogliśmy zarówno wskutek wyższych kosztów produkcji, jak i dlatego, że przemysł zachodni jest doskonale zorganizowany. Ta właśnie dobra organizacja pozwala fabrykom zagranicznym szybko przystosowywać się do każdorazowych wahań koniunktury i zmieniać dowolnie kierunek swego zbytu. Pozatym przemysł zagraniczny korzysta z najrozmaitszych ulg i przywilejów: subwencji, pożyczek długoterminowych, ulg podatkowych itd. – co stwarza stanowczą przewagę po stronie przemysłu zachodnioeuropejskiego. Jest to jednak przewaga natury wyłącznie gospodarczej, gdyż pod względem technicznym nasze fabryki zachodnim nie ustępują, i wyroby nasze jakością swoją z zachodnimi współzawodniczyć mogą, co nam nieraz przynawali nasi konkurenci zagraniczni. Rozmiary naszej produkcji i nieznaczna w stosunku do niej pojemność rynku wewnętrznego nadaje temu przemysłowi charakter wybitnie eksportowy, dlatego jest rzeczą b. ważną dla nas zachowanie rynku rosyjskiego, wchłaniającego przed wojną tak olbrzymi procent naszej produkcji. Z chwilą powstania granicy celnej od strony Rosji ten najważniejszy dla nas rynek może być zachowany tylko przez uzyskanie w traktacie handlowym ulgowej stawki celnej dla wyrobów emalajowanych przywożonych do Rosji z Polski i w Polsce wytwarzanych. Jednocześnie należy odgrodzić się cłem od zachodu, dopóki nie będą wyrównane nasze warunki produkcji z zachodnimi. Dalszym zadaniem naszej polityki gospodarczej w stosunku do przemysłu emalajerskiego powinno być uzyskanie tych rynków, na których dotychczas panowały wyroby zachodnie, np. krajów Bałkańskich.

październik 1921



Głośniki bezprzewodowe, czyli jak ukłęcić stereo z niczego

Na rynku elektroniki popularnej ogromne znaczenie zdobyły tzw. głośniki bezprzewodowe. Wcześniej nazywane głośnikami Bluetooth (kiedy transmisja bezprzewodowa dla tego typu urządzeń ograniczona była w praktyce do BT), wywodzą się z zapomnianych już stacji dokujących, w których miały fizyczne połączenie ze smartfonami, będącymi dla nich głównym źródłem materiału muzycznego.

Faktycznie są więc minisystemami zawierającymi wzmacniacz, głośniki i porcję nowoczesnej techniki cyfrowej służącej komunikacji i sterowaniu. Nazywanie ich głośnikami z jednej strony jest zrozumiałe dla przeciętnego konsumenta, to po prostu urządzenie, które „gra”; z drugiej strony może rodzić konfuzję fakt, że jest czymś zupełnie innym niż zespoły głośnikowe, zwłaszcza te konwencjonalne, pasywne, podłączane do zewnętrznych wzmacniaczy, w ramach tradycyjnych systemów stereofonicznych lub kina domowego, składanych z wielu komponentów (dawniej zwanych „wieżami”). Zamieszanie i nieporozumienie powiększa też rozwój jeszcze innej kategorii – aktywnych zespołów głośnikowych, coraz częściej wyposażonych w transmisję bezprzewodową (lepszą niż Bluetooth), pracujących w parach, w wysokiej klasy systemach.

W tym artykule przybliżymy jednak temat niedroгих, popularnych głośników bezprzewodowych, wykorzystując do tego celu materiał porównawczy, przedstawiony we wrześniowym teście AUDIO – sześciu modeli w cenie ok. 2000 zł. Dla tej kategorii to już co najmniej klasa średnia, propozycje nie tylko młodzieżowe i przenośne, ale też domowe i „dorosłe”.

Głośniki bezprzewodowe można spotkać wszędzie – w sklepach mniej i bardziej ekskluzywnych, w posiadaniu użytkowników młodszych i starszych, mniej i bardziej wymagających, w pomieszczeniach mniejszych i większych, a także poza nimi, na otwartej przestrzeni, w różnych warunkach. Przestały być

nowinką, ale nadal potrafią zaskoczyć oryginalną formą, bogactwem funkcji, a nawet spektakularnym brzmieniem... jak na swoją wielkość. Teoretycznie nie wchodzi w parady regularnym systemom audio ze względu na niższą jakość dźwięku, ewidentną zwłaszcza w dziedzinie stereofonii, bardzo ograniczonej czy wręcz niemożliwej do wykreowania z pojedynczego, małego urządzenia. Mimo to pojawiają się nawet w salonach i tam, gdzie do tej pory niezbędna wydawała się choćby mała, ale wyposażona w parę odseparowanych głośników „wieżyczka”. To ona pada ofiarą wygodniejszego, pod pewnym względem uproszczonego, ale optymalnego sposobu odbioru muzyki, teraz uzupełnionego o nieznane wcześniej metody transmisji sygnału i sterowania. Producenci do niedawna podkreślali obecność funkcji smart, ale stały się one już oczywiste, wrażenia



nie robi też obecność asystentów głosowych. Co więc czeka nas w najbliższej przyszłości? Nawet jeżeli nie pojawią się zupełnie nowe funkcje, to te dostępne już teraz mogą występować w różnych kombinacjach, przy czym „wszystko naraz” wcale nie jest najlepsze dla wszystkich. Dlatego rozwój tego gatunku będzie zmierzał do różnicowania wielkości, formy i funkcji pod kątem odmiennych potrzeb i gustów

Łączy je wspomniany kompromis pod względem stereofonii, ściśle związany z ogólną koncepcją malego, zintegrowanego urządzenia. W zasadzie w takich warunkach... z odległości większej niż metr trudno o coś więcej niż dźwięk monofoniczny, ale producenci nie poddają się i nie rezygnują choćby z namiastki stereofonii, która przez ostatnie pół wieku stała się standardem. Producenci radzą (lub nie radzą) sobie z tym na różne sposoby, które zaraz prześledzimy. Ostatecznie jednak subiektywnie oceniana jakość dźwięku zależy nie tylko od efektu stereofonicznego – mono też da się słuchać, i to czasami z dużą przyjemnością... Ale jak dokładnie brzmią testowane głośniki, możecie przeczytać w AUDIO 9/2021.

Konfiguracje stosowane w głośnikach bezprzewodowych często można zaliczyć do pojennego schematu 2.1, w którym mieści się wiele różnych wariantów. Wyjaśnijmy w szerszej perspektywie, że dla wieloczęściowych instalacji głośnikowych, stereofonicznych bądź wielokanałowych, oznaczenie z kropką oznacza system z subwooferem aktywnym; system 2.1 byłby więc systemem stereofonicznym z parą satelitów (zwykle pasywnych) i subwooferem (aktywnym), przetwarzającym najniższe częstotliwości obydwu kanałów (zmiksowanych do jednego). Oznaczenie takie w zintegrowanym urządzeniu musi więc oznaczać coś innego. Analogia jest taka, że jedynie po kropce wskazuje na wspólny dla obydwu kanałów głośnik niskich częstotliwości, a dwójka przed – niezależne głośniki kanału lewego i prawego dla przetwarzania zakresu średnio-wysokotonowego. Ma to znane uzasadnienie, które przypomnimy w największym skrócie – perspektywę stereofoniczną tworzą głównie średnie i wysokie tony, można więc „pójść na skróty” i zmonofonizować niskie tony bez wielkiego uszczerbku dla efektu, który i tak będzie słaby... ze względu na niewielką odległość między głośnikami obydwu kanałów. Korzyścią z działania wspólnego głośnika niskotonowego jest przede wszystkim oszczędność miejsca, które w głośnikach bezprzewodowych jest deficytowe – a lepszy jest jeden większy głośnik niskotonowy (choć w takich urządzeniach wciąż będzie mały) niż dwa... jeszcze mniejsze (miniaturowe).

Ambitniejsze wśród systemów tego typu, oprócz jednego, wspólnego niskotonowego, mają w każdym kanale układy dwudrożne – złożone z przetworników średnionotonowego i wysokotonowego; w prostszym rozwiązaniu będą tutaj pracować przetworniki średnio-wysokotonowe (szerokopasmowe), spotkaliśmy się też z układem złożonym z jednego, wspólnego dla obydwu kanałów nisko-średnionotonowego i pary wysokotonowych – stereofonia zostaje więc ograniczona już tylko do tego zakresu.

Są również, zwłaszcza wśród większych konstrukcji, systemy z pełną separacją obydwu kanałów – zwykle z układami dwudrożnymi, incydentalnie nawet trójdrożnymi.

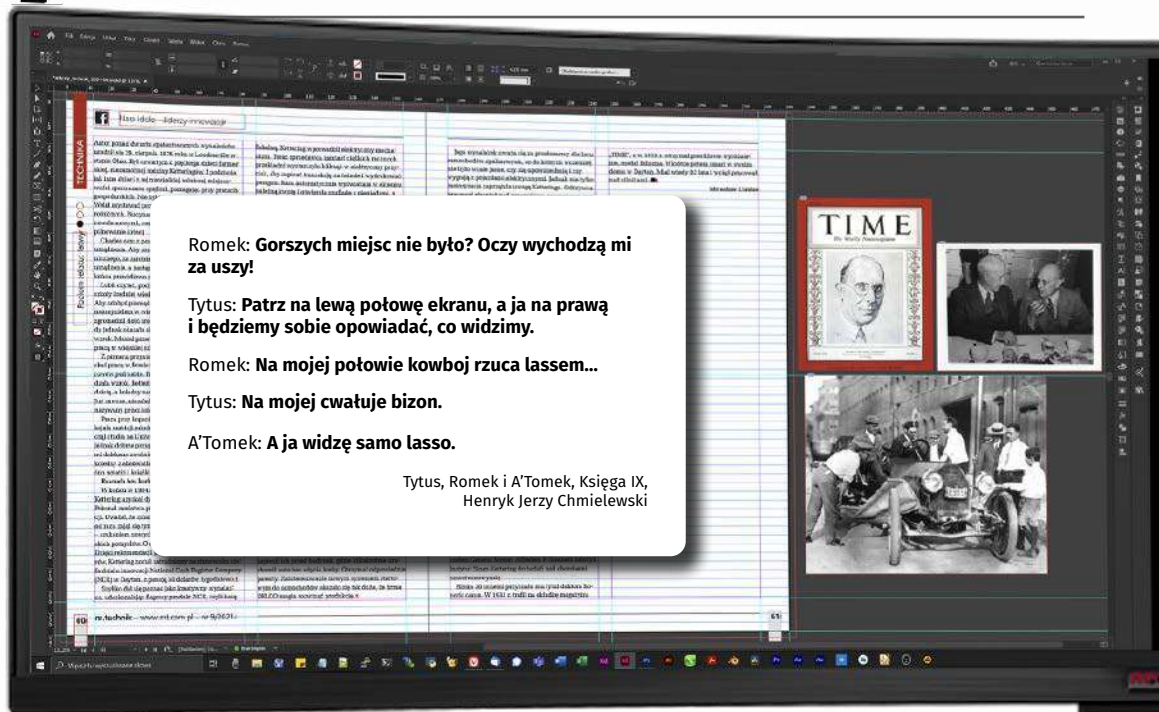
Czasami zastosowanie większej liczby przetworników czy układów dwudrożnych służy skierowaniu ich w różne strony, dla uzyskania lepszego rozpraszania, co zapewni bardziej równomierne nasycenie dźwiękiem pomieszczenia, w przypadku takich urządzeń ważniejsze niż pozory stereofonii dostępne w bardzo ograniczonym obszarze.

Ale nawet wtedy wszystkie przetworniki niskotonowe często łączy wspólny układ rezonansowy obudowy, wykorzystujący membranę bierną. W sytuacji, w której objętość obudowy jest zbyt mała, aby prawidłowo dostroić ją za pomocą otworu i tunelu, pomagamy sobie membraną bierną.

A głośniki bezprzewodowe są właśnie małe, a nawet malutkie. Ich przetworniki również, ale potrafią pracować z zaskakująco dużymi amplitudami, a objętości obudów są miniaturowe; nawet jeżeli udało by się w nich zwinąć i upchnąć tunel o odpowiednich wymiarach, to zajęłby on dużo miejsca – o tyle trzeba by konstrukcję powiększyć, relatywnie znacznie. Membrana bierna wymaga z kolei miejsca na obudowie (a nie w jej wnętrzu, bo jest płaska), z czym jednak projektanci zręcznie sobie radzą, mogąc nadawać membranom biernym w zasadzie dowolny kształt. Ponadto membrana bierna fizycznie (ale nie akustycznie) zamyka obudowę, dzięki czemu nie wpadną do niej... cokolwiek można sobie wyobrazić wpadającego do otworu bas-refleks takiego urządzenia, zwłaszcza przenośnego.

Aktywne systemy głośnikowe – a do nich należą też głośniki bezprzewodowe – korzystają także intensywnie z korekcji elektrycznej, pozwalającej wyrównać charakterystykę i obniżyć dolną częstotliwość graniczną, nie można jednak z takim działaniem przesadzić, aby głośnika niskotonowego... nie wysadzić w powietrze zbyt dużą amplitudą. ■

Andrzej Kisiel



AOC AGON AG493UCX

– gdy oczy wychodzą za uszy

Długo zastanawiałem się nad użyciem powyższego cytatu, ale im bardziej chciałem go pominąć, z różnych względów, tym bardziej on wracał do mnie niczym rzucony bumerang. I prawdę powiedziawszy, tytuł oddaje wszystko, co dotyczy recenzji tego ogromnego monitora. A może nie?

Propozycja rzućenia okiem na te monstrum została przyjęta bez zastanowienia. „Pacuszka” ledwo zmieściła się do samochodu. Lekka i poręczna niestety nie jest. A jaka jest?

Duży może więcej, ale za to jakie ma wymagania

Monitor AOC AGON AG493UCX o przekątnej 49 cali (ok. 125 cm) zajął w zasadzie połowę mego biurka 160×80 cm, na którym na co dzień mieszczą się bez problemu dwa laptopy 15" i 17" oraz dwa 27" monitory. Ba! Pozostaje sporo miejsca na kilka innych przedmiotów, takich jak chociażby ręce 😊.

Tu z racji sporej krzywizny – promień 1800R, połowa blatu została przystońnięta niczym planeta przez gwiazdę śmierci. Rozdzielczość 5120×1440 px przy proporcjach obrazu 32:9 przytłacza, ale z gwiazdką, o której później. Jasność podświetlanej matrycy diodami WLED typu VA to 550 nitów, czas reakcji GTG 4 ms, a MPRT 1 ms; kąty oglądania 178/178, przy

16,7 mln kolorów. To już niemal wartości podstawowe, ale częstotliwość odświeżania na poziomie 120 Hz, tego parametru nie znajdziemy jeszcze przez jakiś czas w innych monitorach. Gracze dobrze wiedzą, że to robi różnicę.

Podłączenie nie sprawia problemów, wszelkie przewody zostały dołączone i każdy nawet bez czytania instrukcji sobie poradzi. Mamy do dyspozycji złącza: HDMI 2.0 ×2, DisplayPort 1.4 ×2, USB-C 3.0 Gen 1, USB 3.2 (Gen 1) ×3, wyjście słuchawkowe.

Waga bez pudła to zaledwie 14,4 kg, gwarancja przez 3 lata, MTBF wynosi 50 000 h, ale z wyłączeniem matrycy. Dodatkowo w pudle znajdziemy pilot do zdalnego sterowania. W chwili testowania sprzętu można było go nabyć za nieco ponad 4 tys. zł.

Gramy! ...znaczy pracujemy

Przekornie zaczniemy od pracy, ponieważ poza gramy, do czego głównie został skonstruowany monitor, reklamowany jest jako wybitne narzędzie do biura.



I tu zaczynają się schody – pojawia się tytułowy problem. Jeżeli monitor znajdzie się na wyciągnięcie ręki – czy muszą wyjść za uszy, nie ma innego wyjścia. Oczywiście, po kilku tygodniach można się przyzwyczaić, ale dla mnie było to męczące, niestety. I to jest połowa gwiazdki, o której była wcześniej mowa.

Praca, chociażby przy łamaniu numeru MT, na tym sprzęcie była łatwa, mimo że czasem drażniła z racji innych przyzwyczajeń. Niestety głowa była zmuszona do większego wysiłku przy obracaniu jej z lewej do prawej i na odwrót. Na ekranie mieści się dużo jednocześnie otwartych okien, ale są jak dla mnie zbyt niskie, co sprawia, że trzeba często przewijać zawartość ekranu – na co dzień korzystam z rozdzielczości 4k na każdym z 3 wyświetlaczy i przy przesiadce na niższą ciągle czegoś brakuje.

Oglądamy! ...znaczy, w dalszym ciągu pracujemy

Przytoczone 4k, do którego się człowiek szybko przyzwyczaja, tu rozczarowuje, bo go nie ma! Monitor wyświetla za mało linii poziomych, by móc obejrzeć np. film z Netflixu we właśnie wysokiej rozdzielczości. Nie ma takiej możliwości. I to jest druga połowa wspomnianej wcześniej gwiazdki.

Na YouTube można znaleźć kilka, kilkanaście krótkich filmów w rozdzielczości 5120×1440 w proporcjach 32:9, które po oddaleniu się od monitora, w slangu młodzieży, robią robotę. Szczególnie gdy oglądany obraz będzie w HDR – czy stają się zaczerwienione ze szczęścia.

Gramy! ...a co, służbowo

To tygrysy lubią najbardziej. Karta RTX4000 w połączeniu z tym ekranem rzuca na kolana, chociaż nie przy każdym tytule. Dlaczego?

Otóż monitor rozwija skrzydła przy symulacjach, daje frajdę latając, jeżdżąc etc. Strzelając już

niekoniecznie, przynajmniej ja nie mogłem się przyzwyczaić. Ciągłe „kątem oka” reagowałem na jakiś nieistotny ruch przy krawędzi ekranu i cała zabawa stawała się trudniejsza.

120 Hz odświeżania działa bardzo dobrze, oko się nie męczy, poruszające się elementy nie smużą, są płynne i to wszystko sprawia, że zainwestowane pieniądze zwracają się w ogólnym zadowoleniu.

Czy warto?

Cóż, straszy cena, ale raczej nie wszystkich. Należy mieć na uwadze gabaryt – nie wszędzie on się zmieści. Trzeba uwzględnić odległość od ekranu – musi być większa niż przy standardowym ekranie 24–27”. To są ew. dodatkowe koszty, jakie być może należy ponieść, by w pełni cieszyć się zakupem.

Miałem problemy przy pierwszym podłączeniu dwóch komputerów do monitora, mianowicie macbook pro i dell uparcie chciały wyświetlać obraz w pełnej rozdzielczości 5120×1440 px na każdej połowie monitora – co było komiczne i funkcja PbP nie działała, jak należy. O ile mac po ręcznym poprawieniu parametrów nie sprawił problemów, o tyle della trzeba było zmusić do poprawnej pracy. Wbudowany przetwornik KVM służący do dzielenia jednej klawiatury i myszki między dwoma komputerami również nie bardzo chciał działać, być może znów gryzły się komputery, ale należy pamiętać, że taka opcja również istnieje i jak donoszą inni – działa.

Ogólnie sprzęt godny polecenia. Wydaje mi się, że takich wielkich ekranów na rynku będzie tylko więcej, co cieszy, przede wszystkim dlatego, że młodzi ludzie nie będą musieli mrużyć oczu, by odczytać coś na malutkim ekranie. Poza tym wszechobecna rozrywka na takim sprzęcie zrelaksuje lepiej niż mniejszy i ciemniejszy ekranik. ■

DW

Nie przegap wrześniowego wydania **Elektroniki dla Wszystkich**



W numerze między innymi:

„Maluch” dla malucha

Samodzielna budowa zabawki dla dziecka może zapewnić podwójną satysfakcję. Czy i Ty doświadczysz takiej radości?

Inteligentny dom także dla Ciebie, czyli jest dobrze, ale nie beznadziejnie. Trochę historii

Każdy, kto zainteresowany jest tematyką smart home, koniecznie powinien choć trochę poznać historię inteligentnych domów.

Silniki prądu stałego

Silniki elektryczne prądu stałego nazywane PMDC niesłusznie są lekceważone jako elementy wręcz prymitywne. Nie można w pełni zrozumieć ich właściwości bez dobrego rozumienia, czym jest ich praca w czterech ćwiartkach.

Współczesne neony, czyli znów tajemnicze lampy EL

Na rynku można znaleźć mnóstwo gadżetów, które z reguły przedstawiane są jako lampy lub neony LED. Jednak nie zawsze zawierają diody LED, a często są to po prostu lampy EL.

Redukcja napięcia sieci energetycznej z 240 V na 220 V

Nominalna wartość napięcia sieci energetycznej to 230 V, a w praktyce coraz częściej sięga 240 V. Jak zapewnić długą żywotność starych radioodbiorników, projektowanych na napięcie 220V?

Ponadto w numerze:

- Filozofia sieci. Protokół TCP
- Panorama audio. Co to jest DAC?
- Droga do RRIO, czyli wzmacniacze operacyjne (nie tylko) dla początkujących
- Lampowa „mrygałka”
- Dołączanie przewodów do złączy kołkowych
- Szkoła Konstruktorów:
 - Jak wdrażać dzieci i wnuki w arkaana techniki, a w szczególności elektroniki, logiki oraz programowania
 - Zaproponuj ciekawe, najlepiej nietypowe zastosowanie diod LED

ELPORTAL.pl

EdW możesz zamówić na

www.ulubionykiosk.pl

lub w empikach i wszystkich

większych kioskach z prasą.

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