



nr 11. listopad 2021

e-suplement www.mt.com.pl



Tu przejrzysz
i kupisz ten numer

NEWS 24/7
przeglądaj codziennie
na swoim smartfonie

młody **m.technik**

Ciekawi świata są zawsze młodzi



SKARB NARODÓW

Prawdziwe bogactwa XXI wieku

RAPORT: Gdzie jest woda w kosmosie?
Życia bez niej nie znamy

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 58 59
e-mail: prenumerata@ruch.com.pl

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

RICHARD ARKWRIGHT
Tłacz rewolucji przemysłowej

11. listopad 2021

m.technik

SKARB NARODÓW
Prawdziwe bogactwa XXI wieku?

Temat okładowy
Pieniądze mogą stracić na wartości. Także złoto i diamenty mogą okazać się iluzorycznym bogactwem. Najprawdziwszą i największą wartość ma to, czego nauczyli się ludzie, wiedza i informacje, które zgromadzili, umiając przekuć w wszystko na nowe idee i wynalazki.

Tam skarb twój, gdzie... głowa twoja

Adam Smith w swoim dziele pt. „Badania nad naturą i przyczynami bogactwa narodów” pisze o trzech wyznacznikach i fundamentach, inaczej skarbach narodów – o cennych metalach, złocie i srebrze, o zbożu oraz oczywiście o pracy ludzkiej. Zasoby te nie przestały być cenne i dziś, ale w czasach po rewolucji przemysłowej, w realiach rozwoju tzw. gospodarki wiedzy i Przemysłu 4.0, skłonni jesteśmy upatrywać źródeł bogactwa głębiej, w czymś z jednej strony mniej namacalnym, a z drugiej – bardziej fundamentalnym.

Kto zaprzeczy, że prawdziwymi klejnotami w skarbnicy dóbr narodowych są dziś innowacje i wynalazki. Najbogatsze dziś i najpotężniejsze państwa są zarazem krajami o najwyższej innowacyjności. Zaś kraje o niskim poziomie innowacyjności pozostają uboższymi krewnymi, skazanymi na korzystanie z rozwiązań tych pierwszych (i płacenia za nie niemałych pieniędzy). Wzrost liczby patentów towarzyszy wzrostowi bogactwa. Wprawdzie ta zależność nie zawsze działa w drugą stronę, gdyż kraj może stać się bogaty bez wysokiego poziomu własnych innowacji, mając np. dużo ropy naftowej. Jednak trudno mi znaleźć przykład wysoko innowacyjnego kraju, który byłby biedny.

Poziom innowacyjności wydaje się jasną pochodną wysokiego poziomu wykształcenia społeczeństwa. W proponowanej przez nas w tym numerze MT metaforze określamy edukację, zwłaszcza tę związaną z dominującymi kierunkami rozwoju nauki i techniki, jako pieniądze, swoistą walutę w skarbnicach narodów.

A złoto? Co jest odpowiednikiem złota w realiach współczesnej opartej na przetwarzaniu komputerowym gospodarce? Uznaliśmy, że złotem naszych czasów są dane. Ta analogia może nie od razu i nie każdemu wyda się oczywista, więc może spójrzmy na to, co jest najcenniejszym zasobem firm oferujących dziś najbardziej zaawansowane technologie, będących zarazem liderami przychodów i zasobów. Potęgą grupy nazywanej niekiedy Big Tech – Google, Amazona, Microsoftu, Facebooka z kilkoma innymi gigantami, opiera się na danych, które efektywnie zbierają, przetwarzają i potrafią wykorzystać do pozyskania... kolejnych danych, co oznacza szybkie pomnażanie bogactwa.

Te właśnie bogactwa wydają się dziś bardziej trwałymi podstawami zasobności niż kruszec, drogocenne minerały czy gotówka, które choć czasem zyskują, to jednak również często tracą na wartości.

Mirosław Usidus

Do
50%
taniej
w prenumeracie
dla szkół
i placówek
oświatowych!

Roczna prenumerata drukowana w promocji dla szkół i placówek oświatowych kosztuje 99,90 zł, roczny dostęp online – 57,00 zł.

Szczegóły na
www.ulubionykiosk.pl/prenumerata/szkolna

PRENUMERATA – TO SIĘ OPŁACA!
Szczegóły na str. 24

STAŁY KONKURS

Active Reader

Supernagrody!

Szczegóły na stronie 2

KSIĄŻKI

GRY

PŁYTY

MODELE

NARZĘDZIA

SPRZĘT

AKCESORIA



Czy fundamenty „bogactwa narodów” zmieniły się zasadniczo od czasów Adama Smitha, który opisywał je prawie dwa i pół wieku temu? Nie tak bardzo jak może się wydawać. Już w jego czasach złoto, pieniądze czy kosztowności, traktowano głównie jako symbole, dostrzegając, że prawdziwe skarby są w czym innym. Dziś rozumiemy to o wiele lepiej, doceniając znaczenie wiedzy, innowacji i informacji.



RAPORT



Spis treści

Temat numeru: Skarb narodów.

Prawdziwe bogactwa XXI wieku?

- 26 • Innowacyjność i wynalazki, czyli – diamenty. Sakiewka pełna pomysłów
- 32 • Wysztafcenie to najlepsza waluta w kieszeni. Potęgi klucz
- 39 • Bogactwo narodów wczoraj i dziś. Tam skarb twój, gdzie... głowa twoja
- 42 • Współczesne złoto, czyli dane. Zanim poleje się krew

Technika

- 8 Info Zoom
- 16 Dodaj do obserwowanych
- Horyzonty mgłą spowite
- 17 • Skomplikowana fizyka prawdziwych wojen gwiezdnych. Łatwo i szybko już było – na Ziemi
- 21 • Maszyny tworzą rzeczy, które stawiają pod znakiem pytania o naturę nauki. Symulacja symulacji
- 23 • Dlaczego komputery analogowe wciąż są nam i będą potrzebne? Drzewo jako hybryda analogowo-cyfrowa
- 46 Raport MT: Poszukiwania wody w kosmosie
- 59 Nasi idole – liderzy innowacji: Jak tkął się kapitalizm – Richard Arkwright

m.technik

- 62 e-Technologie: mobilne aplikacje: test aplikacji: Zarządzanie hasłami

Szkoła

- 64 Chemia inna niż w szkole: Na cześć złośliwych duchów (1)
- 68 Matematyka z ludzką twarzą: Wycieczka w n-ty wymiar
- 73 Koniec i co dalej? Motoryzacja, jaką znamy. Zamiast auta w garażu – mobilność
- 76 Edukacja przez szachy: Problem skoczka szachowego
- Na warsztacie
- 82 • Elektronika dla Ciebie: Uniwersalny timer od 0 do 99 minut
- 84 • Latawiec szaszłykowy
- 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
- Odkryj historię wynalazków
- 98 • Internet Rzeczy
- 102 • Klasyfikacja Internetu Rzeczy według zastosowań

Hobby

- 103 Z pasji do motoryzacji: Na trzech kołach
- 108 Akademia audio: Nowoczesne soundbary (1)

Prezentacje

- 110 MT testuje: Duża rzecz za nieduże pieniądze. Smartfon Motorola Moto G9 Power

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

25

Prawdziwe bogactwa XXI wieku

Poszukiwania wody w kosmosie 46

List miesiąca

nagroda: 30 punktów AR

Szczegóły na stronie 2

„Roje robotów”

Co do rojów robotycznych opisanych w niedawnym raporcie „Młodego Technika” (wydanie 09.2021), warto być może uzupełnić podane informacje o jeszcze kilka faktów.

Wyróżnia się dwa główne sposoby działania robotów w roju:

1. Centralna, czy też może po prostu jednorodna, inteligencja

Wszystkie roboty są kontrolowane przez jeden system. Same roboty nie są inteligentne niezależnie, lecz po prostu reagują na polecenia kontrolującej je nadrzędnej sztucznej (lub może nie sztucznej) inteligencji. Odbyna się to analogicznie do zdalnego sterowania robotem przez człowieka. Taki robot nie posiada żadnej własnej inteligencji poza podstawowymi funkcjami generowania ruchu po otrzymaniu jakiegoś sygnału. Za to można mu przekazać, co ma robić, aby wykonywał zadania. Rój jest wtedy po prostu układem nadrzędnej kontroli sprawowanej nad wieloma robotami w tym samym czasie.

2. Inteligencja poszczególnych robotów (Robot-by-Robot Intelligence)

Roboty w takim systemie mają niezależną inteligencję i często wchodzą ze sobą w interakcję. Działają trochę podobnie do mrówek w mrowisku. Mrówki mają zadania, do których wykonania dążą. Wykorzystują przy tym swoiste formy komunikacji z innymi mrówkami, które są zwykle dość ograniczone, jednak mają znaczenie dla funkcjonowania całości roju.

Można wyróżnić jeszcze trzeci model, który jest w zasadzie hybrydą dwóch opisanych wyżej. Roboty zarządzane są w nim przez pewną ogólną nadrzędną inteligencję, określającą zadania. Każdy pojedynczy robot podejmuje owe zadania i używa swojej własnej, indywidualnej inteligencji oraz komunikacji z innymi robotami roju w celu wykonania tych zadań w sposób skoordynowany i sprawny.

Robotyka roju zajmuje się badaniem konstrukcji robotów, ich ciał fizycznego oraz zachowań. Jest inspirowana przez obserwację emergentnych zachowań w społecznościach, co niekiedy bywa nazywane inteligencją roju. Stosunkowo proste indywidualne reguły mogą prowadzić do wyłaniania się zestawów wysoce złożonych zachowań roju. Kluczowym składnikiem jest komunikacja pomiędzy poszczególnymi elementami grupy, budująca i oparta na systemie ciągłego sprzężenia zwrotnego. Zachowanie roju to ciągłe zmiany w jednostkach we współpracy z innymi jednostkami.

W przeciwieństwie do rozproszonych systemów robotycznych, robotyka rojów kładzie nacisk na dużą liczbę robotów i promuje skalowalność, na przykład przez wykorzystanie wyłącznie lokalnej komunikacji. Taka lokalna komunikacja może być osiągnięta przez bezprzewodowe systemy transmisji, choćby za pomocą częstotliwości radiowych lub zakresów podczerwieni.

W robotyce rojów kluczowymi czynnikami są miniaturyzacja, jak też redukcja kosztów. Wynika z tych czynników m.in. nacisk na prostotę pojedynczego członka/elementu/modułu w zespole/roju. Wielozadaniowość – cecha pożądana w indywidualnych konstrukcjach robotycznych, jest czymś w logice roju zbędnym. Wielozadaniowy może być rój, ale nie pojedynczy element tego roju.

Dobrym przykładem zaprojektowanego według powyższych reguł systemu rojów jest LIBOT Robotic System, w którym tanie roboty zbudowano do wykonywania zadań wewnątrz i na zewnątrz pomieszczeń mieszkalnych i biurowych. Roboty są tu wykonane tak, aby miały wystarczające możliwości do użytku wewnątrz budynków przez Wi-Fi, ponieważ czujniki GPS, używane z kolei na zewnątrz, zapewniają słabą komunikację wewnątrz budynków.

Innym przykładem takiego projektu jest mikrorobot, zbudowany w Computer Intelligence Lab na uniwersytecie w Lincoln w Wielkiej Brytanii. Zbudowany na okrągłym podwoziu o średnicy 4 cm stanowi tanią i otwartą platformę do wykorzystania w różnych aplikacjach robotyki rojowej.

Przykładów jest oczywiście więcej na całym świecie. Myślę, że warto się tej dziedzinie przyglądać uważnie, bo ani się nie obejrzymy, a te chmary będą wszędzie wokół nas, wykonując pożyteczne prace, ale być może również irytując nas i napawając lękiem.

Zbigniew Rak z Essen w Niemczech

Ransomware

Dziękuję za zajęcie się w „Młodym Techniku” (numer wrześniowy 2021) problemem wymuszania okupów przez grupy przestępcze w ramach procederu zwanego popularnie ransomware.

Samo pojęcie ransomware nie jest nowe. Koncepcje leżące u podstaw ransomware opartego na kryptografii z kluczem publicznym zostały wymyślone jako skuteczny mechanizm wyłudzenia pieniędzy w latach 90. dwudziestego wieku. Pierwszy wirus ransomware, nazwany trojanem AIDS, został stworzony w 1989 roku. Jednak dopiero w 2006 roku pojawiła się kolejna znacząca odmiana tego rodzaju ataku (Archiveus Trojan).

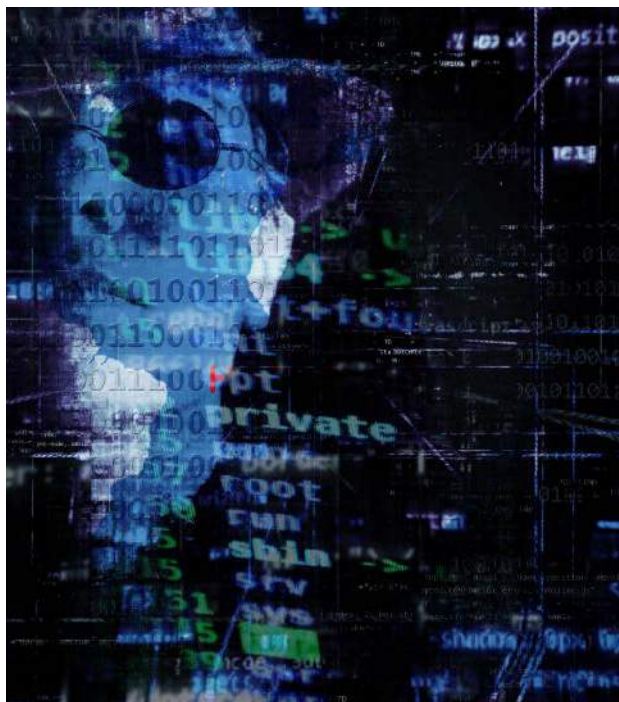
Od tego czasu obserwuje się stały wzrost liczby ataków związanych z oprogramowaniem ransomware. W pierwszym kwartale 2016 r. liczba ataków wzrosła niemal czterokrotnie, osiągając szacunkową wartość miliarda dolarów łącznego okupu wyptaconego w 2016 r., przy czym dziennie dochodziło do czterech tysięcy ataków. Potem liczba ta tylko rosta.

Badania przeprowadzone przez firmę Kaspersky kilka lat temu wykazały, że najbardziej cierpią małe i średnie firmy. 42 procent z nich padło ofiarą ataku ransomware w ciągu 12 miesięcy poprzedzających przeprowadzenie ankiety. Spośród nich jedna na trzy zapłaciła okup, zaś jedna na pięć firm nigdy nie odzyskała swoich plików, mimo że zapłaciła.

Prawdziwy medialny rozgłos ransomware zyskało przy okazji ataku WannaCry, który spowodował awarię systemów w wielu instytucjach systemu ochrony zdrowia, zwłaszcza na terenie Wielkiej Brytanii, oraz kolejnego ataku – Petya, skierowanego głównie przeciw podmiotom gospodarczym i państwowym na terenie Ukrainy.

Do głównych czynników, które odegrały istotną rolę w tym rosnącej popularności ataków tego rodzaju, należy według wszelkiego prawdopodobieństwa zaliczyć pojawienie się silnej infrastruktury zapewniającej anonimowość i ochronę przestępcom, takiej jak sieć Tor, oraz systemy kryptowalutowe, Bitcoin i inne.

Rozwojowi ransomware sprzyja wygodny system płatności w kryptowalutach, trudny do wyśledzenia i namierzenia. Prywatność zapewniana przez wyżej wymienione platformy jest dużą, dodatkową, motywacją do tego typu ataków. Przelewy bitcoin są wysoce anonimowe, co praktycznie uniemożliwia namierzenie sprawców.



Na poziomie oprogramowania wirusy używane w atakach ransomware są stosunkowo proste. Dostępne są na wyspecjalizowanych forach czy też w sieciach dark webu różne zestawy frameworków typu „zrób to sam”. Oferowana jest tam również, o czym MT pisze, ransomware jako usługa zlecona.

Nowe i ulepszone warianty pojawiają się w szybkim tempie. Ataki stają się coraz sprytniejsze, pozwalają przestępcom zarabiać coraz więcej pieniędzy, także z ich sprzedaży jako narzędzi.

Głównym wektorem ataków są złośliwe załączniki do wiadomości e-mail. Jednak z drugiej strony środki ostrożności, pozwalające zneutralizować siłę rażenia takiego ataku, są w sumie dość proste. Zwykle wystarczy posiadanie kopii zapasowych danych i zaktualizowanego oprogramowania. Wówczas groźby szantażysty są bezskuteczne, choć oczywiście tracimy czas, choćby na format dysku i reinstalację oprogramowania. To jednak straty ograniczone w porównaniu do wysokiego okupu jakiego żąda cyberzbir.

Marian Konik, Parczew

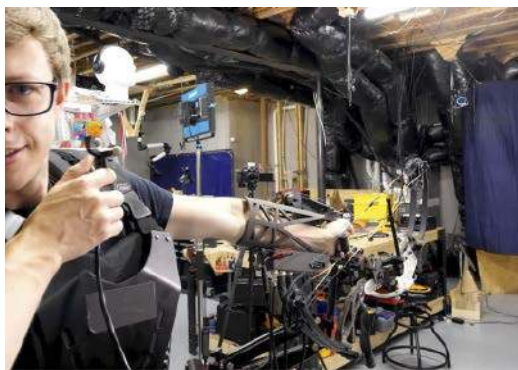
Od Redakcji

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.





WYNALAZKI

Łuk, który nie chybia

Youtuber Shane Wighton prowadzący kanał pt. „Stuff Made Here” skonstruował łuk wyposażony w mechanizm, który sprawia, przynajmniej tak wynika z pokazu udostępnionego w internecie, że strzały nigdy nie chybiają. Co więcej, na demonstracyjnym wideo widać, że Wighton nie musi nawet patrzeć na cel.

Pierwsza wersja łuku z automatycznym celowaniem składała się z dwóch mechanizmów: ręcznego robota, który ustawiał pozycję łuku za pomocą pary silniczków, odpowiadającego za celowanie, oraz drugiego robota, który przytrzymywał i zwalniał naciągniętą cięciwę, obliczając właściwy czas. Układ kamer OptiTrack motion capture w pomieszczeniu, w którym Wighton przeprowadzał testy, łączył się z czujnikami śledzącymi przymocowanymi do łuku i celu. Specjalne oprogramowanie przetwarzało dane z czujników, transferując je do mechanizmu automatycznego celowania.

Celność tej pierwszej wersji rozczarowywała jednak wynalazcę. Zastosował więc dodatkowe ulepszenia wynikające głównie ze specyficznych wymogów strzelania z łuku. Dodał ponadto instalacje do przytrzymywania urządzenia, które przez dodanie wielu elektronicznych elementów stało się zbyt ciężkie do normalnego trzymania. Wyniki ulepszeń widać na filmie udostępnionym przez autora. Efekt sensacji nieco słabnie, gdy zdajemy sobie sprawę, że widoczna tu magiczna celność jest możliwa jedynie w pomieszczeniu wyposażonym w określony zestaw czujników i kamer.



Prezentacja łuku
Shane'a Wightona:
<https://bit.ly/3w1Xazm>

Firma SpaceX z powodzeniem przeprowadziła misję z udziałem czworga, jak to opisywano, „niezawodowych astronautów”. Na orbicie okołoziemskiej w załogowej kapsule Dragon turyście znaleźli się w ramach misji nazwanej Inspiration4. Osoby te to: miliarder Jared Isaacman, który zapłacił za wszystkie cztery miejsca w statku kosmicznym, Sian Proctor, Chris Sembroski i Hayley Arceneaux. Choć nie są to pierwsi turyści w kosmosie, zwraca się uwagę na fakt, że misja ta jest pierwszą od początku do końca komercyjno-turystyczną wyprawą na orbitę. Po trwającej trzy dni wyprawie załoga bezpiecznie wylądowała w kapsule w pobliżu atlantyckich wybrzeży Florydy.

Ciekawostką jest, że przez trzy godziny kapsuła znajdowała się na wysokości orbitalnej 585 km, wyższej niż orbita Międzynarodowej Stacji Kosmicznej czy Kosmicznego Teleskopu Hubble'a i najdalej, na jaką człowiek oddalił się od Ziemi od czasu zakończenia programu księżycowego Apollo w 1972 roku. Załoga, choć składa się z osób niemających niż wspólnego z profesjonalną astronautyką, nie spędzała czasu na relaksie i oglądaniu pięknych widoków przez okna Dragona Resillience. Zaplanowany na trzy dni ich pobytu, na orbicie został dość ściśle zaplanowany i wypełniony zadaniami do wykonania. Wśród nich były m.in. eksperymenty medyczne, telekonferencje z Ziemią oraz praca





TURYSTYKA KOSMICZNA

Kosmiczny rejs turystyczny SpaceX

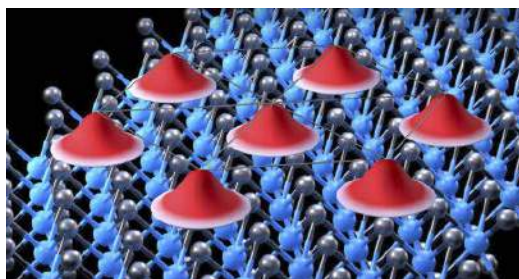


nad nagrywaniem materiału do filmu dokumentalnego dla Netflix.

Nie wszystkie szczegóły misji zostały ujawnione.

Po misji, gdy nie było na żywo żadnych filmów i zdjęć z wnętrza kapsuły podczas lotu, pojawiły się medialne sugestie, że astronauta mieli prawdopodobnie jakiś rodzaj awarii toalety pokładowej. Zmodyfikowany Dragon, wykorzystywany przez turystów wyposażony został w dużą, trójwarstwową kopułę z pleksiglasu, znaną jako Dragon Cupola. Według udostępnionych informacji miejsce to miało być również zapewniającą dyskrecję (zasłonki) toaletą z widokiem.





FIZYKA

Kryształy z elektronów

Zespół naukowców pod kierownictwem Ataca Imamoglu, profesora Instytutu Elektroniki Kwantowej Politechniki w Zurychu, wyprodukował bardzo szczególny kryształ. W przeciwieństwie do zwykłych kryształów składa się on wyłącznie z elektronów. Tym samym udało im się potwierdzić teoretyczne przewidywania, sformułowane prawie dziewięćdziesiąt lat temu. Rezultaty ich badań zostały niedawno opublikowane w czasopiśmie naukowym „Nature”.

W 1934 roku Eugene Wigner, jeden z twórców teorii symetrii w mechanice kwantowej, wskazał, że elektrony w materiale mogą teoretycznie układać się w regularne, kryształopodobne wzory z powodu wzajemnego odpychania elektrycznego. Rozumowanie, na którym się oparł, jest całkiem proste – jeżeli energia

elektrycznego odpychania pomiędzy elektronami jest większa niż ich energia ruchu, to ułożą się w taki sposób, aby ich całkowita energia była jak najmniejsza. Przez kilkadziesiąt lat przewidywania te pozostawały jednak czysto teoretyczne, ponieważ kryształy Wignera mogą powstawać tylko w ekstremalnych warunkach, przy bardzo niskiej temperaturze i bardzo małej liczbie elektronów swobodnych w materiale.

Aby pokonać te przeszkody, Imamoglu i jego współpracownicy zdecydowali się zastosować cienką warstwę półprzewodnikowego diselenku molibdenu, który ma grubość zaledwie jednego atomu i w którym, w związku z tym, elektrony mogą poruszać się tylko w płaszczyźnie. Naukowcy zmieniali liczbę swobodnych elektronów poprzez przyłożenie napięcia do dwóch przezroczystych elektrod grafenowych, pomiędzy którymi umieszczono półprzewodnik. Według przewidywań kryształy Wignera mogły się formować dopiero, gdy cała aparatura jest schłodzona do kilku stopni powyżej zera bezwzględnego, $-273,15^{\circ}\text{C}$. Fizykom udało się uwidocznić regularny układ elektronów dzięki wzbudzaniu za pomocą światła tzw. ekscytonów, czyli par elektronów i „dziur”, które powstają w wyniku brakującego elektronu na jednym z poziomów energetycznych materiału. Częstotliwość światła, przy której powstają ekscytyny, oraz prędkość, z jaką się poruszają, zależą zarówno od właściwości materiału, jak i od interakcji z innymi elektronami w materiale – na przykład z kryształem Wignera.

FIZYKA

Metal, w którym elektrony płyną jak woda

Grupa badaczy z Boston College stworzyła nowy rodzaj metalicznego materiału, mieszaninę niobu z germanem (NbGe_2), w którym ruch elektronów przebiega w taki sam sposób, w jaki woda przepływa w rurze, co diametralnie zmienia dynamikę z cząsteczkowej na płynną. Naukowcy poinformowali o wynikach swoich eksperymentów w czasopiśmie „Nature Communications”.

„Chcieliśmy przetestować ostatnie przewidywania dotyczące „płynu elektronowo-fononowego”, wyjaśniał eksperyment Fazel Tafti, członek zespołu, przypominając, że fonony są drganiami struktury krystalicznej. „Zazwyczaj elektrony są rozpraszane przez fonony, co prowadzi



do zwykłego dyfuzyjnego ruchu elektronów w metalach. Nowa teoria zakłada, że gdy elektrony silnie oddziałują z fononami, będą one tworzyć rodzaj cieczy elektronowo-fononowej, zachowującej się jak płynąca woda”.

Jak dodał, następnym krokiem dla badaczy z jego zespołu będzie poszukiwanie innych materiałów, w których występują podobne efekty pod wpływem oddziaływań elektronowo-fononowych. Uczni chcą też opracować techniki sterowania hydrodynamicznym płynem elektronowym w takich materiałach i konstrukcje nowych urządzeń elektronicznych wykorzystujących ich odkrycia.



ENERGIA SŁONECZNA

Pierwsza na świecie komercyjna instalacja perowskitowych ogniw w Polsce

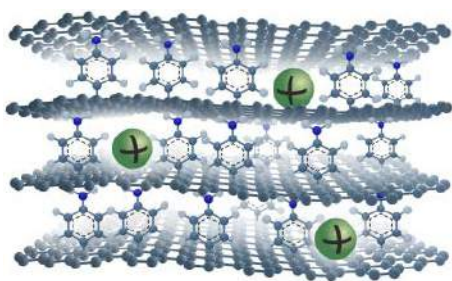


W Polsce, a dokładnie w Lublinie otwarto, jak głoszą komunikaty, pierwszą na świecie komercyjną instalację komercyjną z wykorzystaniem ogniw perowskitowych, które dostarczyła firma Saule Technologies, stworzona przez Olgę Malinkiewicz. Są to tzw. czyli lamele łamacze światła, nie tylko chronią budynek przed przegrzaniem lub też wychłodzeniem, produkując jednocześnie energię z promieniowania słonecznego.

Konstrukcje pokrytych innowacyjnymi ogniwami lameli Saule zaprezentowało niemal już rok temu. W maju 2021 r. otwarta została linia produkcyjna. Teraz zaś łamaczy światła zainstalowano na fasadzie fabryki Aliplast w Lublinie, która zresztą współpracuje z firmą doktor Malinkiewicz. Dzięki zastosowaniu w pionierskiej instalacji systemowi automatyki Animeo marki Somfy, lamele z ogniwami perowskitowymi współpracują z umieszczaną na dachu budynku stacją pogodową. Aktualizowane na bieżąco

dane meteorologiczne oraz funkcja „śledzenia słońca” (suntracking) pozwalają automatycznie zmieniać ustawienie łamaczy względem Słońca.

„Rozpoczynamy nowy etap w rozwoju Saule Technologies – erę komercjalizacji naszych ogniw perowskitowych. Z dumą prezentujemy nasz kompletny produkt – lamele – łamacze światła, działające na fasadzie budynku naszego pierwszego klienta – Aliplast. To dla nas wyjątkowy moment, wieńczący osiem lat ciężkiej pracy całego zespołu Saule Technologies”, mówi współzałożycielka Saule Technologies, Olga Malinkiewicz, cytowana w komunikacie. Firmę założyli w 2014 r. Olga Malinkiewicz, wynalazczyni nowego typu ogniw perowskitowych oraz Piotr Krych i Artur Kupczunas. Saule przez ostatni rok, przy współpracy inżynierów z Korei, Malezji, Wielkiej Brytanii, Singapuru i Japonii, pracowała nad pierwszą na świecie linią produkcyjną drukowanych ogniw perowskitowych.



ENERGIA

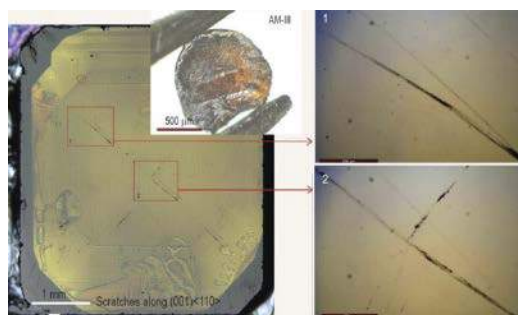
Dzięki grafenowi sód dostaje szansę na zastąpienie litu w bateriach

Zespół badawczy z Uniwersytetu Technologicznego Chalmers w Szwecji zaprezentował rozwiązanie, które pozwala akumulatorom sodowo-jonowym dorównać pojemnością współczesnym akumulatorom litowo-jonowym. Wykorzystują do tego nowy rodzaj grafenu, który ma asymetryczne chemiczne funkcje na przeciwległych powierzchniach i dlatego jest często nazywany grafenem Janusa, na cześć rzymskiego boga Janusa o dwóch twarzach.

„Dodaliśmy cząsteczkę dystansującą po jednej stronie warstwy grafenowej. Kiedy warstwy są ułożone na sobie, molekula tworzy większą przestrzeń między arkuszami grafenu i zapewnia punkt interakcji, co prowadzi do znacznie większej pojemności”, wyjaśnia Jinhua Sun z Wydziału Nauk Przemysłowych i Materiałowych w Chalmers w pracy naukowej, opublikowanej w „Science Advances”.

Zazwyczaj pojemność interkalacji sodu w standardowym graficie wynosi około 35 mAh/g. Jest to mniej niż jedna dziesiąta pojemności jonów litu w graficie. W przypadku zastosowania nowego grafenu pojemność właściwa dla jonów sodu osiąga wartość 332 mAh/g, zbliżając się do wartości typowej dla litu w graficie. Wyniki wykazały również pełną odwracalność i wysoką stabilność cykliczną sodowych ogniw wspomaganych janusowym grafenem.

22,5 kg miedzi, w przewodach elektrycznych i innych częściach, zawiera przeciętnie jeden samochód.



NOWE MATERIAŁY

Szkło twardsze niż diament

Naukowcy z Uniwersytetu Yanshan w Chinach opracowali najtwardszy i najmocniejszy ze znanych do tej pory materiałów szklanych. Według informacji podanych w publikacji na łamach periodyku „National Science Review” może on z łatwością zarysować kryształy diamentu. Zdaniem badaczy, nowy materiał, wstępnie nazwany AM-III – ma „wybitne” właściwości mechaniczne i elektroniczne i dzięki swojej ultrawysokiej wytrzymałości i odporności na ścieranie mógłby znaleźć zastosowanie w ogniwach słonecznych.

Niezwykła twardość AM-III, w porównaniu z siecią krystaliczną diamentu, wynika z pewnego nieładu na poziomie molekularnym, który osiągnięto w materiale. Naukowcy wykorzystali puste w środku cząsteczki węgla zwane fullerunami, a następnie ogrzewali je i zginali przez wiele godzin pod dużym ciśnieniem. Trik polegał na tym, aby nie doprowadzić atomów węgla do zbyt uporządkowanego układu, gdyż w takiej sytuacji szkło uległoby osłabieniu i straciłoby swoje właściwości półprzewodnikowe.

Z analizy materiału wynika, że jego odporność osiągnęła 113 gigapaskali (GPa), podczas gdy naturalny diament w takich samych testach osiąga zwykłe wyniki od 50 do 70. Według naukowców, AM-III ma właściwości regulowanej absorpcji energii porównywalne do półprzewodników powszechnie stosowanych w ogniwach słonecznych, takich jak uwodornione folie z krzemu amorficznego.

1200 – tyle przeciętnie tornad nawiedza w jednym roku tereny Stanów Zjednoczonych.



MARS

Łazikowi NASA udało się pobranie pierwszej próbki

Po pierwszej nieudanej próbie urządzeniom na pokładzie łazika Perseverance udało się pobrać próbkę marsjańskiej skały, która ma zostać zabrana w następnej misji na Ziemię. Przedstawiciele NASA poinformowali o wywierceniu otworu i pobraniu rdzenia skalnego z dużego kamienia nazwanego „Rochette”.

Pierwsze podejście do pobrania próbki, na początku sierpnia 2021 roku, nie powiodło się, ponieważ skała, w którą wwierciło się narzędzie Perseverance, okazała się niespodziewanie bardzo miękka, zaś kruchego materiału nie udało się umieścić skutecznie w zasobniku. Sterujący łazikiem przejechali 455 metrów w inne miejsce, nazwane „Citadelle”, gdzie

znaleziono twardsze skały. W poszukiwaniu dobrych miejsc do pobierania próbek materiału pomaga łazikowi obecnie helikopterek Ingenuity, który po serii testowych lotów stał się dronem rozpoznawczym dla Perseverance.

Próbki te, wraz z co najmniej dwudziestoma innymi, które ma pobrać aparatura Perseverance (na pokładzie łazika znajduje się łącznie 43 zasobniki na próbki), zostaną przewiezione na Ziemię w ramach wspólnej misji NASA i Europejskiej Agencji Kosmicznej, być może już w 2031 roku. Byłyby to pierwsze nietknięte fragmenty innej planety, jakie ludzkość kiedykolwiek przywozła na swoją planetę.

1 800 000 – takiej liczbie zgonów w latach 1971–2009 zapobiegła energia jądrowa, redukując zanieczyszczenia spowodowane spalaniem węgla i ropy, wg badań NASA z 2013 r.

INTERNET MOBILNY

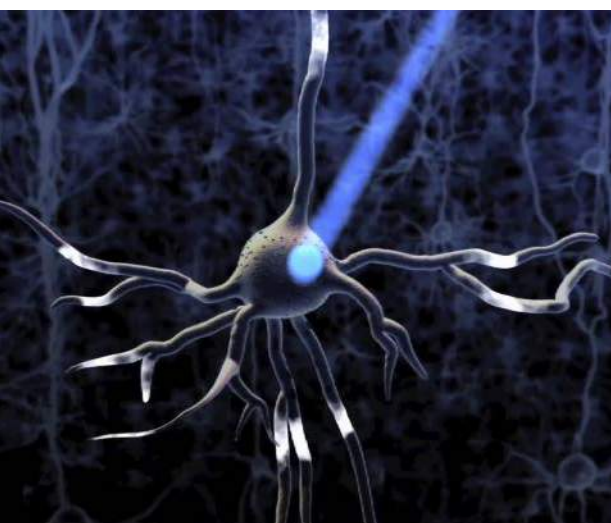
Sieć 6G już w testach

Koreańska firma LG Electronics we współpracy z niemieckim Instytutem Fraunhofera przeprowadziła udany test techniki transferu danych w paśmie, które nazywa 6G THz. Transmisja porcji danych o mocy 15 dBm w zakresie częstotliwości od 155 do 175 GHz na odległość stu metrów była możliwa dzięki nowemu typowi wzmacniacza opracowanego wspólnie przez Koreańczyków i Niemców.

Pojęcie 6G ma w tej chwili charakter nieco umowny, gdyż nie ma jeszcze oficjalnej specyfikacji kolejnego po 5G standardu bezprzewodowej komunikacji. LG chciałoby, aby standaryzacja sieci szóstej generacji została przeprowadzona do połowy tej dekady, by jeszcze przed jej końcem rozpocząć jej komercyjne wdrażanie.



Ogólne założenia są takie, że transfer pomiędzy urządzeniami w sieci 6G będzie znacznie szybszy niż w przypadku 5G. Mówi się o pięćdziesięciokrotności. Sieć piątej generacji oferuje maksymalną prędkość przesyłu danych na poziomie 20 Gb/s. 6G, według niektórych przewidywań, będzie mógł zapewnić prędkość nawet 1 Tb/s. Ponadto ma działać na znacznie wyższych częstotliwościach, od 100 GHz do nawet 3 THz.



BIONIKA

Zdalne sterowanie neuronami dzięki nowemu materiałowi

Technikę zdalnej skutecznej i, co ważne, precyzyjnej stymulacji żywych komórek za pomocą impulsów światła laserowego opracowali badacze z Uniwersytetu Carnegie Mellon. Ich metoda, która, w odróżnieniu od wcześniej badanych, nie wymaga

inwazyjnych implantów, wykorzystuje nowe materiały – płatki węglików/azotków metali przejściowych (MXenes), wyjątkowego dwuwymiarowego nanomateriału odkrytego przez zespół Yury'ego Gogotsiego z Uniwersytetu Drexel.

Zespół rozprowadził płatki na powierzchni tkanki z komórek obwodowego układu nerwowego i naświetlał je krótkimi impulsami światła. W wyniku badań interfejsu między komórkami a materiałami okazało się, że płatki nie są absorbowane przez komórki, a naukowcy mogli dokładnie zmierzyć ilość światła potrzebną do wywołania zmian w komórkach. „Wyjątkowe w materiałach, których używamy w moim laboratorium, jest to, że nie musimy używać impulsów o wysokiej energii, aby uzyskać skuteczną stymulację komórek”, wyjaśniał szef zespołu z Uniwersytetu Carnegie Mellon, Tzahi Cohen-Karni. „Przez świecenie krótkimi impulsami światła na interfejs materiału komórkowego i MXenes, wykazaliśmy, że elektrofizjologia komórki została skutecznie zmieniona”.

Dzięki lepszym metodom stymulacji neuronowej i łatwości produkcji MXenes, badacze mogą skuteczniej stosować fotostymulację na odległość. Jednym z możliwych eksperymentów jest umieszczenie MXenes w sztucznej tkance zaprojektowanej w formie mózgu, a następnie użyć światła do kontrolowania aktywności neuronów i badania roli neuronów w rozwoju mózgu. W medycynie materiał ten mógłby być nawet użyty jako nieinwazyjna metoda leczenia zaburzeń funkcji nerwowych, takich jak drgawki.

AI for General Purpose Robotics

AUTOPILOT CAMERAS

FSD COMPUTER



MULTI-CAM VIDEO
NEURAL NETWORKS

NEURAL NET PLANNING

AUTO-LABELING

SIMULATION & TOOLS

FSD HARDWARE

DOJO TRAINING

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 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.

1 000 000 ofiar śmiertelnych – to jednostka miary zabójczej siły broni jądrowej. Nazwano ją „megadeath”, co przypomina nazwę słynnego metalowego zespołu muzycznego.



KOMPUTERY KWANTOWE

♦ Znany producent procesorów komputerowych, firma AMD, opatentował komputer kwantowy oparty na teleportacji, w którym kwestia skalowalności i stabilności jest rozwiązywana przez zmniejszenie liczby kubitów wymaganych do poszczególnych obliczeń, nie przedstawiając jednak w opisie patentu dokładnego wyjaśnienia mechanizmu teleportacji, co jest oficjalnie tłumaczone obawami przed kopiowaniem pomysłu. ♦ Kwantowe bity zbudowane z wibrujących węglowych nanorurek i par kropek kwantowych mogą być bardzo odporne na zakłócenia zewnętrzne i dekoherencję, zjawiska będące złązą takich układów, zaproponowali uczeni z francuskiego Narodowego Centrum Badań Naukowych współpracujący z kolegami z Hiszpanii i USA jako sposób na zbudowanie trwałych i skalowalnych komputerów kwantowych. ♦

AERONAUTYKA

♦ Firma Dawn Aerospace pomyślnie przeprowadziła w Nowej Zelandii serię pięciu lotów testowych bezzałogowych suborbitalnego samolotu kosmicznego Mk-II Aurora, który wzniósł się w nich na wysokość do 1036 metrów przy wykorzystaniu zastępczych silników odrzutowych, docelowo mających zostać zastąpionymi przez raketowe, a maszyny te po starcie ze zwykłego pasa mają się wznieść na granicę kosmosu – 100 km nad Ziemią. ♦ Thales Alenia Space, konsorcjum francusko-włoskie, podpisało z agencją ESA kontrakt na budowę i wdrożenie europejskiego promu kosmicznego Space Rider, nad którym, w ramach projektu Intermediate eXperimental Vehicle (IXV), prace trwają od kilku lat i który z powodzeniem przeszedł testy przeprowadzone w 2015 roku, gdy IXV wystartował z kosmodromu w Gujanie Francuskiej i osiągnął orbitę ponad 400 km, lądując następnie na Ziemi. ♦ Marynarka Wojenna Stanów Zjednoczonych pracuje nad bezzałogowym samolotem Skydweller, napędzanym energią słoneczną, wykorzystującym oprogramowanie i zmodernizowany sprzęt pochodzący ze znanej jednostki Solar Impulse 2, zdolnym do lotu przez 90 dni bez przerwy, w zastosowaniach ta-

kich jak np. platforma przekąźnikowa lub jednostka rozpoznawcza do eskortowania okrętów. ♦

MOTORYZACJA

♦ Firma Audi zaprezentowała wizję nazywaną „Skysphere Concept”, która polega na tym, że za pomocą naciśnięcia jednego guzika samochód niczym filmowy transformer może radykalnie zmienić wygląd, przekształcając się z jednego typu auta na inny, modyfikując gabaryty, w tym długość, o nawet 25 centymetrów. ♦ Na targach IAA Mobility 2021 w Monachium Mercedes-Benz zademonstrował nową wersję swojego samochodu koncepcyjnego Vision AVTR, po raz pierwszy pokazanego na CES 2020, który, według podawanych w mediach komentarzy, „może czytać w myślach” kierowcy, a mówiąc konkretniej – wykorzystuje urządzenie BCI (brain-computer interface) z elektrodami przymocowanymi do głowy użytkownika z systemem uczących się rozpoznawania intencji kierowcy za pomocą detektorów sygnałów wizualnych w desce rozdzielczej – kierowca musi się na nich skupiać, gdy chce wykonać jakieś zadanie. ♦

NOWE MATERIAŁY

♦ Naukowcy z MIT stworzyli tkaninę elektroniczną z wszytymi w polimerowe włókna setkami krzemowych chipów, zdolną do wyczuwania, przechowywania, analizowania i zapisywania danych na temat aktywności noszącego materiał na sobie. ♦ Specjaliści z Uniwersytetu Technologicznego Nanyang w Singapurze i Kalifornijskiego Instytutu Technologii (Caltech) w Stanach Zjednoczonych opracowali nowy materiał przypominający strukturę kolczugi z drukowanych w 3D ośmiościennych elementów ząbiających się ze sobą, który może być elastyczny jak tkanina, ale może się również usztywnić na żądanie, gdy zostanie owinięty w elastyczną plastikową folię, które jak w próżniowych opakowaniach wypompowane zostanie powietrze, co zmienia go w sztywną strukturę, z której jest 25 razy sztywniejsza lub trudniejsza do zgięcia niż w stanie rozluźnionym. ■

M. U.



1. Wizualizacja projektu Boeing X-20 Dyna-Soar

Skomplikowana fizyka prawdziwych wojen gwiazdnych

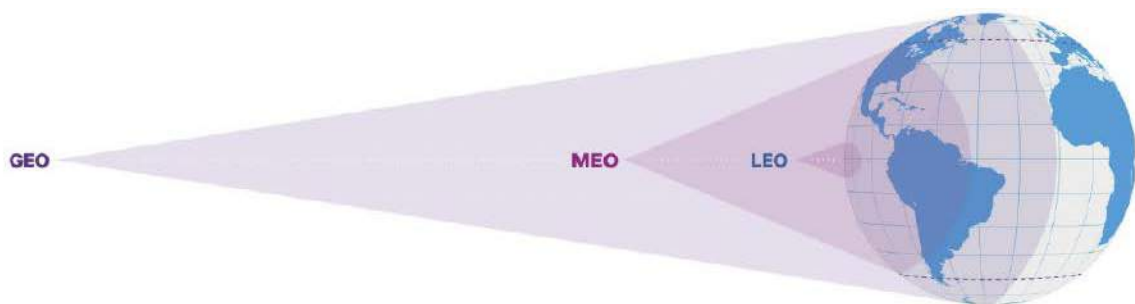
Łatwo i szybko już było – na Ziemi

Popularne filmy przedstawiają wojny w kosmosie tak samo jak wojny na Ziemi albo nawet efektowniej. Piloci gwiazdnych myśliwców manewrują, jak chcą, szybko i nie martwiąc się o zasięg. Transporty wojsk spadają z orbity na powierzchnie planet, z desantami kosmicznych komandosów. Niestety prawdziwa wojna w kosmosie wyglądać będzie zupełnie inaczej...

Na początku ery kosmicznej istniało założenie, że personel wojskowy będzie żyć i pracować w przestrzeni kosmicznej tak jak we wszystkich innych terenach działań militarnych. Siły powietrzne USA dążyły do stworzenia załogowego samolotu kosmicznego w programie Dyna Soar (1). Jednak adaptacja technik, które sprawdzają się dla samolotu, do próżni kosmicznej okazała się poza ówczesnymi możliwościami.

Niestety, ludzie potrzebują w przestrzeni kosmicznej dużo zasobów – jedzenia, wody, a nawet po prostu powietrza i to wszystko musi być wyrzuczone

z Ziemi. Ostatecznie militarne programy załogowe zakończyły się niepowodzeniem. Zamiast tego udoskonalenia w technologii i transmisji danych umożliwiły powstanie satelitów, które spełniają wiele wojskowych funkcji przewidzianych dla wcześniejszych programów załogowych. Obecnie działalność militarna w przestrzeni kosmicznej jest zdominowana przez bezzałogowe satelity, które wpływają na niemal każdy aspekt współczesnych operacji wojskowych. Jednak walki, starcia i bitwy w przestrzeni kosmicznej to zupełnie coś innego.



2. Zakresy orbit niskiej LEO i geostacjonarnej GEO ze średnią orbitą MEO w środku

Powoli, daleko i paliwa mało

Aby opisać, w jaki sposób fizyka ogranicza militarne operacje kosmiczne, trzeba uświadomić sobie pięć kluczowych czynników: satelity poruszają się szybko, poruszają się przewidywalnie, przestrzeń jest bardzo duża, czas ma decydujące znaczenie, a ponadto satelity manewrują relatywnie powoli.

Ruch w przestrzeni kosmicznej niewiele ma wspólnego z lotami w atmosferze ziemskiej. Ponadto nie ma tam możliwości tankowania. Ewentualna wojna kosmiczna prawdopodobnie będzie się toczyć powoli, ponieważ przestrzeń jest duża a statki kosmiczne mogą schodzić z przewidywalnych ścieżek (a tylko zmiana przewidywalnych trajektorii daje pożądanym militarnie efekt zaskoczenia) jedynie przy wielkim nakładzie energii.

Operacje wojenne na Ziemi zazwyczaj toczono są o dominację nad fizycznym miejscem, o kontrolę nad lądem, morzem i przestrzenią powietrzną nad pewną częścią Ziemi, aby zwiększyć swój wpływ na ludzi lub zasoby. Dającą się obecnie wyobrazić wojna w przestrzeni kosmicznej nie ma tym nic wspólnego. Satelity na orbicie nie zajmują i nie dominują w jednym miejscu na przestrzeni czasu. Zapewniają za to określone możliwości, np. komunikację, nawigację i zbieranie danych wywiadowczych ziemskim siłom zbrojnym. Dlatego też, aby „kontrolować przestrzeń kosmiczną” niekoniecznie fizycznie trzeba podbić określone sektory przestrzeni kosmicznej, ale raczej zredukować lub wyeliminować możliwości satelitarne przeciwnika, zapewniając sobie jednocześnie możliwość swobodnego operowania i rozwój własnych zdolności kosmicznych.

Mechanika orbitalna dyktuje, że obiekty na niższych wysokościach zawsze poruszają się szybciej niż te na większych wysokościach. Każda próba przyspieszenia lub zmniejszenia prędkości satelity zawsze będzie prowadzić do zmiany wysokości. Satelity na powszechnie używanych orbitach kołowych poruszają się z prędkością od 3 km/s do 8 km/s

w zależności od ich wysokości nad poziomem morza. Dla porównania, przeciętny pocisk porusza się z prędkością około 0,75 km/s.

Satelita na orbicie zazwyczaj podąża tą samą ścieżką i zmierza w tym samym kierunku. Ścieżki te mogą być kołowe lub eliptyczne. Satelity na orbitach kołowych utrzymują stałą wysokość i prędkość. Na orbicie eliptycznej satelity poruszają się wolniej, gdy jest wyżej i szybciej, gdy obniża wysokość. Zależność pomiędzy wysokością, prędkością i kształtem orbity i kształtem orbity sprawia, że ścieżki satelitów są przewidywalne. Aby zboczyć z wyznaczonej orbity, satelity muszą manewrować, czyli użyć silnika. To istotna różnica w porównaniu z samolotami, które do zmiany kierunku lotu używają w dużym stopniu interakcji z powietrzem. Próżnia kosmiczna nie oferuje takiej możliwości. Ponadto, orbita satelity nie zależy od jego masy, czyli zarówno małe satelity, jak i duże satelity poruszają się z taką samą prędkością dla danej wysokości. To kolejna istotna różnica w porównaniu z warunkami atmosferycznymi, w których masa maszyny latającej ma wielkie znaczenie.

Przestrzeń kosmiczna jest duża. Objętość przestrzeni kosmicznej pomiędzy niską orbitą okołoziemską (LEO) a geostacjonarną (GEO) wynosi około 200 trylionów kilometrów sześciennych (2). Jest to 190 razy więcej niż objętość Ziemi. W dodatku, ponieważ satelity poruszają się szybko, ma dużą bezwładność. W związku z tym zmiana położenia satelity na jego orbicie, znana jako manewr, może wymagać znacznego czasu i energii. Ponieważ przestrzeń kosmiczna jest, jak wielokrotnie już podkreślaliśmy, bardzo duża, a w połączeniu ze ścisłą naturalną zależnością pomiędzy prędkością satelity, wysokością i kierunkiem, zmiana orbity wymaga zarówno ΔV (delta-V – w astrodynamice to wielkość skalarna, która ma wymiar prędkości określająca miarę „wysiłku” niezbędnego do wykonania manewru orbitalnego), jak i czasu. Odbywa się to zazwyczaj poprzez spalanie chemicznych materiałów pędnych.

Ziemiem analogiem do satelity mogłoby tu być manewrowanie pociągiem, który może poruszać się tylko w kierunku wyznaczonym przez jego tory. Jednym z popularnych manewrów kosmicznych jest zmiana płaszczyzny, gdzie płaszczyzna orbity satelity jest przechylona w stosunku do jej pierwotnej orientacji bez zmiany wysokości satelity. Można to porównać do przestawienia pociągu na przecinające się tory bez zmiany jego prędkości. Ponieważ satelity poruszają się szybko i mają duży pęd, potrzeba dużo ΔV , aby wykonać nawet małe zmiany płaszczyzny.

Paliwa na pokładzie jest ograniczona ilość i nieśtetę, jak już wspomniano, na orbicie, w przeciwieństwie do warunków ziemskich, nie ma stacji i jednostek tankujących. Dlatego, nawet dla satelitów z dużym budżetem ΔV , dostępna jest ograniczona liczba manewrów. Ponieważ przestrzeń jest duża, wiele satelitów po prostu nie jest w stanie osiągnąć orbity innych satelitów.

To, że przestrzeń kosmiczna jest duża, oznacza także z punktu widzenia „taktyki”, że starcie kosmiczne nie może być jednocześnie intensywne i długie. Może to być albo krótkie, intensywne użycie dużej ilości ΔV dla dużego efektu, albo długie użycie ΔV dla mniejszych lub trwałych efektów. Ze względu na odległości, planowanie kinetycznego ataku na satelitę wymaga uwzględnienia czasu i ΔV potrzebnych do wykonania misji. Operatorzy satelity atakującego przy obecnych technikach muszą spędzić tygodnie

na ustawianiu satelity do ataku, podczas których mogą zmienić się warunki, które zmieniają potrzebę lub cel ataku. Dodatkowo, jeśli atakujący satelita musi wykonywać kosztowne manewry, może on nie mieć rezerw potrzebnych do odpowiedzi, jeśli cel wykona swój własny manewr, aby uniknąć ataku odwetowego (pamiętajmy o przewidywalności ruchu w kosmosie).

Ustawienie dwóch satelitów na tej samej wysokości i w tym samym tej samej płaszczyźnie jest proste (choć czas- i energochłonne), ale to nie znaczy, że są one w tym samym miejscu. Musi nastąpić fazowanie – rzeczywiste położenie wzdłuż trajektorii orbitalnej dwóch satelitów musi również być takie samo. Ponieważ prędkość i wysokość są ze sobą powiązane, uzyskanie położenia dwóch satelitów w tym samym miejscu to kolejne skomplikowane wyzwanie.

To, co stanowi małą lub bliską odległość jest kwestią oceny, zależy od operatorów satelitów. Na przykład satelita na GEO może być w stanie tolerować 50 km separacji pomiędzy innymi satelitami, ale załogowa stacja kosmiczna może nie pozwolić żadnemu satelicie na zbliżenie się na odległość do 150 km. Szybki atak na LEO wymagałby wydatkowania dużego budżetu ΔV dla atakującego satelity i nie byłby szybki. Oprócz wykonania właściwego manewru fazowania, atakujący i cel muszą znajdować się w tej

3. Bitwa kosmiczna w stylu Gwiezdnych Wojen





samej płaszczyźnie orbity. Każdy manewr zmiany płaszczyzny przez atakującego jest kosztowny.

Pas GEO w wielu analizach uchodzi za nieco „łatwiejsze” miejsce do przeprowadzenia ataku typu kinetycznego (czyli fizycznego uderzenia, zniszczenia satelity przeciwnika). Większość satelitów w GEO znajduje się w tej samej płaszczyźnie, co daje więcej możliwości i celów do ataku. Jednak jest mało prawdopodobne, aby atakujący satelitę, jego ruch pozostanie niezauważony. Wciąż jednak przygotowanie ataków wymaga dużo czasu. Potrzeba dni lub tygodni, aby ustawić broń w odpowiedniej pozycji do ataku.

Wiązki lepsze

Podobnie jak w przypadku wojny naziemnej, wojny lądowej, jednym z celów ataku jest fizyczne zniszczenie celu, znane jako kinetyczny atak uderzeniowy. Istnieją dwa rodzaje kinetycznej broni uderzeniowej w wojnie kosmicznej: naziemne pociski antysatelitarne (ASAT) oraz broń na orbicie (kinetic kill vehicles lub orbitalne ASAT). Ponieważ rakieta wystrzelona z Ziemi ma dużą pojemność ΔV , sama głowica jest umieszczona na właściwej trajektorii przechwycenia i wymaga już niewielkiej ilości materiału pędowego, aby dotrzeć do celu. Czas dotarcia do celu może być krótszy niż 10 minut, jeśli chodzi o cele znajdujące się na LEO i mniej niż 5 godzin w przypadku ataków na GEO. Czasu na reakcję, zwłaszcza przypadku niskiej orbity, jest więc zdecydowanie mniej, gdy strzelamy z Ziemi niż w przypadku operacji, które miałyby naśladować bitwy kosmiczne w Gwiezdnym Wojnach” (3).

Ataki kinetyczne wymagają zbliżenia się do satelity będącego celem ataku. Istnieją jednak również

sposoby, aby atakować z dystansu. Mowa o wojnie elektronicznej i atakach w zakresach fal elektromagnetycznych, które mogą być prowadzone przez satelity, jednostki naziemne lub powietrzne. Wojna elektroniczna to m.in. wykorzystanie częstotliwości radiowych w celu obezwładnienia sygnałów przeciwnika (zagłuszanie) oraz celowe naśladowanie sygnałów przeciwnika w celu wysyłania szkodliwych poleceń lub danych (spoofing). Broń o ukierunkowanej energii wykorzystuje skoncentrowane wiązki w zakresach radiowych (mikrofale o dużej mocy) lub wiązki światła (lasery) do zakłócania pracy satelity. Lasery mogą być użyte do tymczasowego oślepienia czujników optycznych i kamer lub trwale uszkodzić wrażliwe sprzętu pokładowego. Mikrofały o dużej mocy zakłócają działanie elektroniki pokładowej w sposób odwracalny lub trwały. Należy jednak pamiętać tu również o odległościach. Strumienie energetyczne rozpraszają się także w próżni i tracą na dużych odległościach. Jeszcze szybciej tracą moc (przez atmosferę) ataki przeprowadzane z Ziemi. Niemniej z uwagi na opisane wyżej ograniczenia w przeprowadzaniu starć z bliska, działania elektroniczne i elektromagnetyczne są w kosmosie znacznie bardziej obiecujące militarnie.

Warto mieć też świadomość, że kinetyczne zniszczenie satelity przeciwnika może stworzyć odłamki, które następnie zagrażą także stronie atakującej. Niektórzy zastanawiają się, jak wiele starć tego typu uczyni przestrzeń kosmiczną niezdatną do użytku. Dla każdego. Także dla tych, którzy „wygrają” kosmiczną wojnę. ■

Mirosław Usidus

Rosja od kuchni

Witold Szablowski

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

Ta książka to niezwykle połączenie książki kucharskiej i dobrego reportażu. Dowiesz się z niej, jak, kiedy – i po co – kucharz Stalina uczył kucharza Gorbaczowa śpiewać do drożdżowego ciasta. Jak Nina, kucharka z wojny w Afganistanie, zmuszała się, by myśleć o czymś przyjemnym – bo chciała, żeby jej dobry nastrój udzielał się żołnierzom. Jak wokół Czarnobyla, kilka tygodni po katastrofie, zrobiono konkurs na najlepszą stołówkę i kto go wygrał. Ale przede wszystkim zobaczcie, jak jedzenie może służyć propagandzie. W kraju takim jak Związek Radziecki służył jej każdy kotlet, usmażony i podany w każdej stołówce i restauracji od Kaliningradu po krąg polarny i od Kiszyniowa po Władywostok. I to, co jadł pierwszy sekretarz partii komunistycznej, i to, co jadł zwykły obywatel, było polityczne. Rosja jest zaś godną następczynią ZSRR, więc wciąż karmi ludzi propagandą, tak jak robiła to przed laty.



W ostatnich latach wiele osiągnięć naukowych w dziedzinach wymagających obliczeń nie pochodzi ze stosowania bardziej wyrafinowanej matematyki, ale po prostu z większej mocy kalkulacyjnej zastosowanych urządzeń. Okazuje się, że coraz większe możliwości komputerów pozwalają już nie tylko rozbudowywać symulacje, ale niejako symulować układy dokonujące symulacji.

Maszyny tworzą rzeczy, które stawiają pod znakiem pytania o naturę nauki

Symulacja symulacji

Fizyka teoretyczna ma obecnie wiele subdyscyplin zajmujących się symulacjami komputerowymi układów rzeczywistych. Używamy obecnie symulacji do badania procesów formowania się galaktyk i struktur większych niż galaktyki, do obliczania mas cząstek składających się z kwarków, do sprawdzania, co dzieje się podczas zderzeń dużych jąder atomowych i do opisu cykli słonecznych, by wymienić tylko kilka obszarów badań, które są oparte głównie na komputerach.

„Symulacje kwantowe” to systemy składające się z oddziałujących ze sobą, złożonych obiektów, takich jak chmury atomów. Fizycy manipulują oddziaływaniami pomiędzy tymi obiektami tak, aby system przypominał oddziaływanie pomiędzy bardziej fundamentalnymi cząstkami. Na przykład w elektrodynamicie kwantowej naukowcy używają małych obwodów nadprzewodzących do symulowania atomów, a następnie badają, jak te sztuczne atomy oddziałują z fotonami. W laboratorium w Monachium fizycy wykorzystują nadpłynne ultrazimne atomy, aby rozstrzygnąć spór, czy cząstki podobne do cząstki Higgsa mogą istnieć w dwóch wymiarach przestrzeni. Okazało się (w symulacji), że tak.

W niedawnym eksperymencie Raymond Laflamme, fizyk z Instytutu Obliczeń Kwantowych na Uniwersytecie Waterloo w Ontario w Kanadzie, i jego zespół wykorzystali symulację kwantową do zbadania tzw. sieci spinowych, struktur, które w niektórych teoriach stanowią fundamentalną tkankę czasoprzestrzeni. Z kolei Gia Dvali, fizyk z uniwersytetu w Monachium, zaproponował metodę symulacji przetwarzania informacji przez czarne dziury za pomocą ultrazimnych gazów atomowych. Podobny pomysł

jest realizowany w dziedzinie grawitacji, gdzie fizycy używają płynów do naśladowania zachowania cząstek w polach grawitacyjnych. Naukowcy zbadali również gwałtowną ekspansję wczesnego Wszechświata, zwaną inflacją, wykorzystując płynne analogi grawitacji. Ponadto fizycy badali hipotetyczne cząstki fundamentalne, obserwując ich odpowiedniki zwane kwazicząstkami.

Teoria pól dyskretnych, czyli maszyna tworzy rzeczy, których się nie spodziewamy

Algorytm, opracowany przez Honga Qina (1) z Laboratorium Fizyki Plazmy Princeton, stosuje uczenie maszynowe, formę sztucznej inteligencji, która uczy się na podstawie doświadczeń, w celu wypracowania



1. Hong Qin



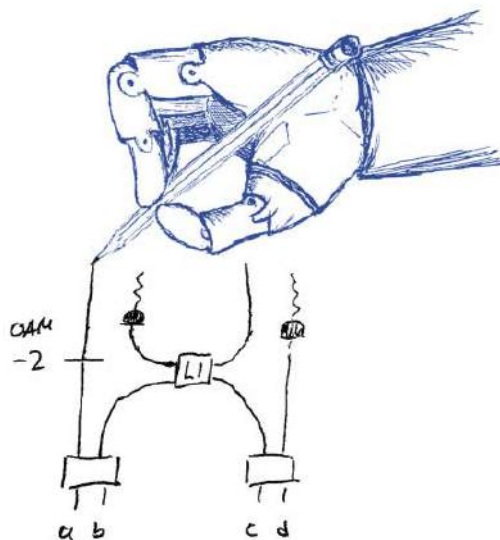
naukowych przewidywań. „Zazwyczaj w fizyce dokonuje się obserwacji, tworzy teorię opartą na tych obserwacjach, a następnie wykorzystuje teorię do przewidywania nowych obserwacji”, wyjaśniał Qin w pracy opisującej algorytm w „Scientific Reports”.

Program, do którego Qin wprowadził dane z wcześniejszych obserwacji orbit Merkurego, Wenus, Ziemi, Marsa, Jowisza i planety karłowatej Ceres, wraz z dodatkowym programem, znanym jako „algorytm służebny”, dokonał dokładnych przewidywań orbit innych planet w Układzie Słonecznym bez wykorzystania wstępnie zaaplikowanych systemowi praw ruchu i grawitacji Newtona. Oznacza to, że algorytm nie miał żadnej początkowej wiedzy o prawach fizyki. Uczył się reguł ruchu ciał niebieskich na podstawie obserwacji. Tworzy więc symulację praw fizyki na podstawie symulacji składającej się z danych, które otrzymał.

Praca Qina opiera się na „teorii pól dyskretnych”. Nazywa on tę teorię rodzajem „powłoki” lub ramy z konfigurowalnymi parametrami, która może być trenowana z pomocą danych obserwacyjnych. Po zakończeniu treningu zamienia się ona w algorytm natury, który komputer „może uruchomić, aby przewidzieć nowe obserwacje”.

Sztuczna inteligencja projektuje eksperymenty z dziedziny fizyki kwantowej wykraczające poza ludzkie wyobrażenia. System uczenia maszynowego, pierwotnie stworzony w celu przyspieszenia obliczeń, dokonuje obecnie szokujących postępów na granicy eksperymentalnej fizyki kwantowej. Fizyk Mario Krenn z uniwersytetu w Wiedniu odkrył pięć lat temu, że MELVIN, algorytm uczenia maszynowego, który zbudował, a którego zadaniem było mieszanie i dopasowywanie elementów składowych standardowych eksperymentów kwantowych i znajdowanie rozwiązań dla nowych problemów, dokonał czegoś niezwykłego. MELVIN pozornie rozwiązał problem tworzenia bardzo złożonych stanów splątanych z udziałem wielu fotonów. Krenn, Anton Zeilinger z Uniwersytetu Wiedeńskiego i ich koledzy nie podali MELVIN-owi reguł potrzebnych do generowania tak złożonych stanów, a mimo to znalazł on sposób. W końcu uczony zdał sobie sprawę, że algorytm odkrył na nowo rodzaj eksperymentalnego układu (2), który został opracowany na początku lat 90. XX wieku. Tamte eksperymenty były jednak o wiele prostsze. MELVIN rozwiązał o wiele bardziej skomplikowaną zagadkę.

„W ciągu kilku godzin program znalazł rozwiązanie, na które my naukowcy, trzech eksperymentatorów i jeden teoretyk, nie mogliśmy wpaść od miesiący”, mówił Krenn w jednym z wywiadów.



2. Eksperyment kwantowy zaprojektowany przez MELVIN-a

Oczekiwania Krenna były takie, że MELVIN znajdzie konfiguracje, które połączą te pary fotonów, tworząc stany splątane o co najwyżej dziewięciu wymiarach. „Ale znalazł jedno rozwiązanie, niezwykle rzadki przypadek, który ma znacznie większe splątanie niż pozostałe stany”, opowiadał.

Podczas wczesnych prób uproszczenia i uogólnienia tego, co odkrył MELVIN, Krenn i jego koledzy zdali sobie sprawę, że rozwiązanie przypomina abstrakcyjne formy matematyczne zwane grafami, które są używane do przedstawiania relacji między parami obiektów. MELVIN najpierw stworzył taki graf, a następnie wykonał na nim operację matematyczną. Operacja ta, zwana „idealnym dopasowaniem”, polega na wygenerowaniu równoważnego grafu, w którym każdy wierzchołek jest połączony tylko z jedną krawędzią.

Eksperymenty wyżej opisane każą zadać pytanie o naturę samej nauki. Czy naukowcy mają teraz już nie rozwijać teorii fizycznych, które wyjaśniają świat, a jedynie po prostu gromadzić dane? Czy teorie nie są czymś fundamentalnym dla fizyki, koniecznym do wyjaśnienia i zrozumienia zjawisk?

Ostatecznie algorytmy mogą na podstawie danych tworzyć alternatywne teorie naukowe, które mogą być podobne do naszych, ale niedokładnie takie same. Mogą też kreować teorie zupełnie nowe, inaczej ujmujące i opisujące świat. Jako że to wszystko obraca się w sferze symulacji, a nie obserwacji fizycznych czy eksperymentów, rodzi się pytania, czy to w ogóle nauka, taka jak ją tradycyjnie rozumiemy. ■

Miroslaw Usidus

W ciągu ostatnich stu lat elektronika przeszła dwie fundamentalne przemiany – przejście z analogowej na cyfrową i przejście z lamp próżniowych na półprzewodnikowe. To, że te przejścia nastąpiły niemal równocześnie, nie oznacza, że są ze sobą nierozzerwalnie związane. Obliczenia cyfrowe były realizowane przy użyciu elementów lamp próżniowych. Przetwarzanie analogowe może być realizowane w półprzewodnikach.

Dlaczego komputery analogowe wciąż nam są i będą potrzebne?

Drzewo jako hybryda analogowo-cyfrowa

Nie ma dokładnego rozróżnienia pomiędzy obliczeniami analogowymi i cyfrowymi. Ogólnie rzecz biorąc, obliczenia cyfrowe zajmują się liczbami całkowitymi, sekwencjami binarnymi, logiką deterministyczną i czasem sprowadzonym do kawałków, porcji, bitów, podczas gdy obliczenia analogowe zajmują się liczbami rzeczywistymi, logiką niedeterministyczną i funkcjami ciągłymi, w tym czasem, który istnieje jako kontinuum w świecie rzeczywistym. Wyobraźmy sobie, że mamy znaleźć środek drogi. Można zmierzyć jej szerokość, cyfrowo obliczyć środek z dokładnością do najbliższego przyrostu. Można też użyć „komputera analogowego”, czyli kawałka sznurka, odwzorowując szerokość drogi na długość sznurka i znajdując środek, bez dzielenia na bity.

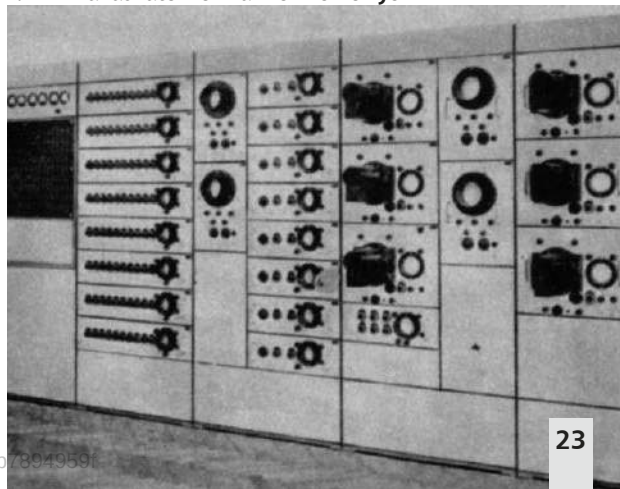
Komputer analogowy to maszyna przetwarzająca sygnał ciągły (analogowy), służąca do rozwiązywania problemów matematycznych i innych (np. zagadnień technicznych, badania zjawisk biologicznych itp.) przez modelowanie (odwzorowywanie) odpowiednich zależności (na ogół ujętych w języku matematyki) za pomocą zjawisk zachodzących w układach mechanicznych, elektrycznych, elektromechanicznych lub elektronicznych. Rozwinięciem komputerów (maszyn) analogowych były komputery analogowo-cyfrowe.

Komputery analogowe do przełomu lat 60. i 70. były znacznie szybsze i tańsze od ówczesnych komputerów cyfrowych. Dobrze sprawdzały się przy

rozwiązywaniu równań różniczkowych i symulacji procesów, jednak wzrost szybkości i spadek ceny komputerów cyfrowych oraz szybki rozwój ich oprogramowania spowodował stopniowy zanik zainteresowania komputerami analogowymi po 1970 r. Pierwszym polskim elektronicznym komputerem analogowym był ARR (Analizator Równań Różniczkowych) Leona Łukaszevicza z 1953 r. Ok. 1954 r. powstał komputer analogowy do rozwiązywania układów równań liniowych ARAL, Analizator Równań Algebraicznych Liniowych (1).

Zdziwilibyśmy się, jak wiele systemów w przyrodzie działa zarówno w reżimie analogowym, jak i cyfrowym. Drzewo integruje szeroki zakres danych wejściowych jako funkcje ciągłe, ale jeśli zetniemy to drzewo, okaże

1. ARR – analizator równań różniczkowych

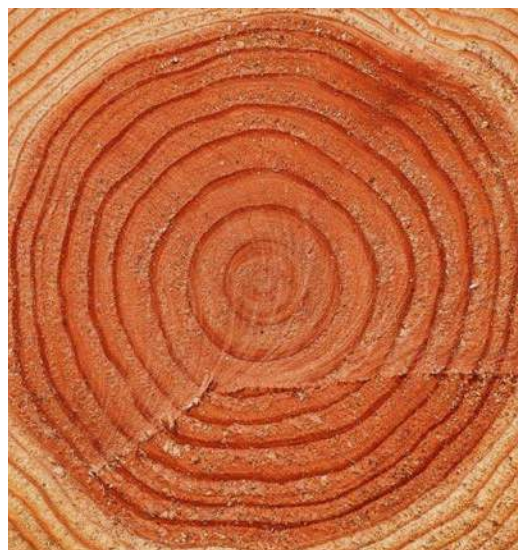




się, że przez cały czas liczyło ono lata na sposób podobny do cyfrowego, na co wskazują słoje przyrostu (2). Natura wykorzystuje kodowanie cyfrowe do przechwywania, replikacji i rekombinacji sekwencji nukleotydów, ale polega na przetwarzaniu analogowym, działającym w systemach nerwowych, w celu zapewnienia rozpoznania i kontroli. System genetyczny w każdej żywej komórce jest komputerem z zapisanymi programami. Mózg już czymś takim nie jest.

Obliczenia cyfrowe, nietolerujące błędów lub niejednoznaczności, zależą od korekcy błędów na każdym etapie. Obliczenia analogowe tolerują błędy. Komputery cyfrowe wykonują transformacje pomiędzy dwoma rodzajami bitów – bitami reprezentującymi różnice w przestrzeni i bitami reprezentującymi różnice w czasie. Przekształcenia pomiędzy tymi dwoma formami informacji, sekwencją i strukturą, są regulowane przez programowanie komputera i dopóki komputery wymagają ludzkich programistów, zachowujemy nad nimi kontrolę.

Zanim przejdziemy do roli przetwarzania analogowego w AI, przypomnijmy trzy prawa sztucznej inteligencji. Pierwsze z nich, znane jako prawo Ashby'ego, od nazwiska cybernetyka W. Rossa Ashby'ego, mówi, że każdy efektywny system sterowania musi być tak złożony, jak system, którym steruje. Drugie prawo, sformułowane przez Johna von Neumanna, stwierdza, że cechą definiującą złożony system jest to, że stanowi on swój własny, najprostszy opis zachowania. Najprostszym kompletnym modelem organizmu jest sam organizm. Próba zredukowania zachowania systemu do jakiegokolwiek formalnego opisu komplikuje sprawę bardziej, a nie mniej. Trzecie prawo stwierdza, że każdy system



2. Stoje w drewnie

wystarczająco prosty, by być zrozumiałym, nie będzie wystarczająco skomplikowany, by zachowywać się inteligentnie, podczas gdy każdy system wystarczająco skomplikowany, by zachowywać się inteligentnie, będzie zbyt skomplikowany, by go zrozumieć.

Każde z tych trzech praw, na swój sposób sugeruje, że rozwoju systemów sztucznej inteligencji nie można oprzeć jedynie na przetwarzaniu cyfrowym. Obliczenia typu analogowego będą nam potrzebne zarówno do konstruowania doskonalszych systemów AI, zrozumienia sposobów ich działania, jak też do kontrolowania ich w sposób efektywny. ■

Mirosław Usidus

Nigdy się nie bałam

Katarzyna Droga

Wydawnictwo Mova, cena: 42,90 zł

Jaki wpływ na przestrzeni lat polskie lekarki miały na historię medycyny? Poznajemy m.in. sylwetkę Anny Tomaszewicz-Dobrowskiej, żyjącej na przełomie XIX i XX wieku, pierwszej kobiety z dyplomem lekarskim na ziemiach polskich, wykształconej w Zurychu, z doświadczeniem zdobytym w Sankt Petersburgu, która nie mogła nostryfikować dyplomu w Polsce ze względu na płeć. Justyny Budzińskiej-Tylińskiej (1867–1936), wykształconej w Paryżu, która była pionierką w dziedzinie propagowania higieny i zdrowia kobiecego. Janiny Bernasiewicz, która podczas otwarcia szpitala psychiatrycznego w Choroszczy w 1930 roku była jedyną kobietą w gronie pracujących tam lekarzy, oraz Bronisławy Dłuskiej, starszej siostry Marii Skłodowskiej-Curie, wykształconej w Paryżu, która wraz z mężem otworzyła w Aninie przy dzisiejszej ulicy Kajki prewentorium przeciwgruźlicze. Kształtując się, praktykując i przekazując swoją wiedzę, udowodniły, że medycyna nie jest dziedziną zarezerwowaną jedynie dla mężczyzn. Pokazały, że także mogą być chirurgami, onkologami, naukowcami, dyrektorami placówek medycznych. Na podstawie dokumentów, rozmów z żyjącymi nestorkami medycyny oraz wspomnień autorka pokazuje determinację, siłę i odwagę Polek w medycynie.





SKARB NARODÓW

Prawdziwe bogactwa XXI wieku



By kraje stały się bogatsze, potrzebna jest nie tylko praca, ale praca coraz bardziej wydajna, ulepszana technicznie i organizacyjnie przez pomysły, na które wpadają ludzie (1). Ważne jest też, aby pomysłów tych ludzi nie lekceważyć, dostrzegać je, doceniać ich potencjał, pomagać wynalazcom i w końcu także chronić te pomysły jako cenną własność intelektualną.

1. Innowacja jako bogactwo

Innowacyjność i wynalazki, czyli – diamenty

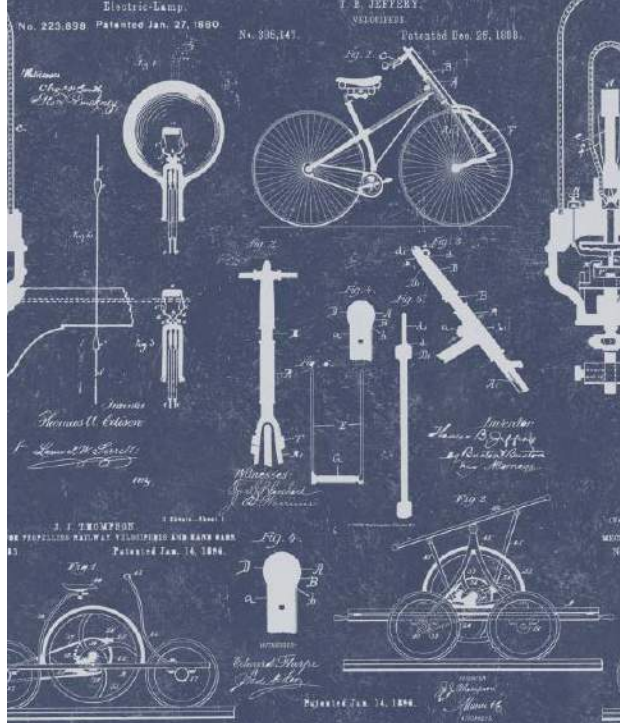
SAKIEWKA PEŁNA POMYSŁÓW

Według danych, publikowanych corocznie przez Światową Organizację Własności Intelektualnej (WIPO) w raporcie World Intellectual Property Indicators, państwa z największą liczbą patentów (2) to albo państwa najludniejsze, których siła w tej dziedzinie wynika z siły wielkich liczb, jak zajmujących w 2019 r. pierwsze miejsce z liczbą patentów ponad 1,4 mln na rok – Chin, lub siódmych na liście Indii z 53 tysiącami patentów w 2019 r., albo państwa najbogatsze – drugie na liście są Stany Zjednoczone (621 tysięcy patentów), trzecia – Japonia

(308 tys.), a jeszcze mocniej znaczenie bogactwa podkreślają pozycja ósma, Australia (prawie 30 tys.) i dziewiąta Kanada (36,5 tys.), państwa z mniejszą liczbą ludności.

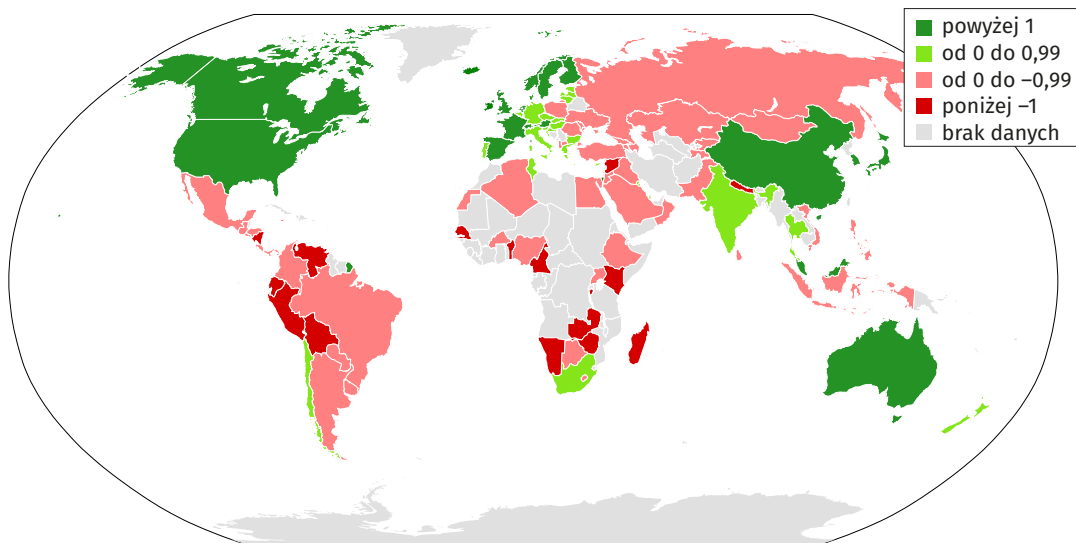
Ogólny poziom liczbowy innowacji, w tym statystykach przeliczanej na liczbę patentów, jest wyznacznikiem potęgi państwa. Jak widać, gdy podzielimy te liczby na liczbę ludności, zbliżamy się do skarbów, jakim innowacje są dla pojedynczych obywateli, choć oczywiście nie w każdym przypadku działa to jako jednoznaczny wskaźnik. Choć przodująca w tym przeliczniku Korea Południowa z 3319 patentami na milion mieszkańców w 2019 r. nie jest krajem ubogim, to jednak nie znalazła się jeszcze w ścisłej elicie najbogatszych pod względem dochodów na głowę. Bardziej to tego modelu pasują: trzecia na liście w tej kategorii Szwajcaria (1122 patenty na mln mieszkańców), piąte w rankingu Niemcy (884), szóste USA (869), siódma Dania (645), ósma Szwecja (601), dziewiąta Finlandia (548) i dziesiąta Holandia (530).

Międzynarodowy Wskaźnik Innowacyjności 2021 (International Innovation Index 2021) jest globalnym wskaźnikiem mierzącym poziom innowacyjności kraju, opracowanym wspólnie przez Boston Consulting Group, Krajowe (amerykańskie) Stowarzyszenie Producentów oraz The Manufacturing Institute (3). Uchodzi za „najbardziej wszechstronny globalny indeks tego rodzaju”. Opiera się na rozległych badaniach zarówno biznesowych rezultatów innowacji, jak i zdolności rządów do zachęcania i wspierania innowacji przez odpowiednią politykę



2. Patenty

w gronie 110 krajów świata i wszystkich pięćdziesięciu stanów USA. Na wartość indeksu składają się m.in. dane o patentach, transferze technologii, wyniki biznesowe, takie jak wydajność pracy i całkowite zwroty z inwestycji w innowacje oraz wpływ innowacji na migrację biznesu i wzrost gospodarczy. Wyniki badań są opublikowane w raporcie „The Innovation Imperative in Manufacturing: How the United States Can Restore Its Edge”. Indeks po raz pierwszy został opublikowany w marcu 2010 roku.



3. Międzynarodowy Wskaźnik Innowacyjności 2021

Pierwsza dziesiątka tego rankingu, czyli kraje o wskaźniku wyższym niż 1, to w kolejności od pierwszego miejsca: Japonia, Korea Południowa, Chiny, USA, Francja, Kanada, Wielka Brytania, Niemcy, Holandia, Australia.

Nie wystarczy mieć pomysły – trzeba je chronić

Gdy próbujemy zdefiniować pojęcie innowacji, to zwykle mówimy o odkrywaniu i rozwijaniu czegoś, co jest nowe, ulepszaniu lub dostosowywaniu czegoś, co już istnieje. Określamy tak rzeczy zarówno całkiem nowe, przełomowe, jak i ulepszenia, uzupełnienia, racjonalizacje rzeczy znanych od dawna. Innowacja bywa również definiowana nieco bardziej ogólnie jako „tworzenie i stosowanie nowej wiedzy w celu ulepszenia świata”. To proces napędzający postęp techniczny i rozumiany szerzej. Banałem jest powiedzenie, że to dzięki procesom innowacji gramofony, radia i telewizory pojawiły się w naszych domach, samoloty wzbily się w przestworza, samochody zrewolucjonizowały komunikację. Kiedy rozglądamy się wokół siebie w naszym codziennym życiu, wszystko, co widzimy, można uznać za produkt innowacji. Komputery, świece, telewizory, satelity, czajniki, samochody, leki, a nawet pismo ręczne. Innowacje nie ograniczają się oczywiście do produktów fizycznie namacalnych. To również rozwój oprogramowania, projekty, doskonalenie metod ich prowadzenia i zarządzania.

Choć, jak wspomniano, są to opisy dość banalne, lecz ten banał pokazuje, że nie są oderwane, abstrakcyjne, pokazuje związek z naszym życiem i cywilizacyjnymi krokami, jakie wykonujemy. Innowacje przekształcają pomysły w wartości gospodarcze i społeczne, we wzrost dobrobytu, produktywności, a ostatecznie nawet w lepsze samopoczucie. Warto pamiętać też o tym, iż jest to mechanizm, dzięki któremu dostosowujemy się do nowych możliwości i wyzwań.

Innowacyjność jest siłą napędową przedsiębiorstw. Dzięki niej firmy mogą konkurować na rynku, tworząc nowe produkty i usługi dla swoich klientów oraz obniżać koszty dzięki poprawie wydajności. Innowacyjne przedsiębiorstwa mają większe szanse na zdobycie większego udziału w istniejących i tworzyć nowe rynki. Wiele popularnych produktów, które są dziś wysoko cenione, jak np. smartfony czy usługi streamingowe, zostało wymyślonych w czasie krótszym niż jedno pokolenie wstecz. Innowacyjne firmy rozwijają się dwa razy szybciej niż firmy, które nie wprowadzają innowacji. Zasada ta przekłada się na całe państwa, choć w tym przypadku jest to nieco

bardziej skomplikowane. Jednak cytowane wcześniej dane wyraźnie pokazują, że innowacyjność idzie w parze z bogactwem i szybkim bogaceniem się.

Jednym z fundamentów bogactwa opartego na innowacyjności jest ochrona owoców tego procesu, wynalazków i rozwiązań. Służy do tego system patentowy i prawo chroniące własność intelektualną. Wynalazcy i stojące za nimi podmioty gospodarcze podejmują zwykle duże ryzyko finansowe. Na drodze do sukcesu doświadczają wielu porażek. Pamiętajmy o „wynalazczej” formule Thomasa Edisona, która brzmiała „1 procent inspiracji, 99 procent potu”. Próby i błędy są dla innowatorów czymś oczywistym. Ślepe zaułki są nieodłącznym elementem innowacji. Czas spędzony nawet w ślepych zaułkach jest inwestycją, która w ogólniejszym rachunku powinna się opłacić. Innowatorzy w przodujących krajach przeznaczają wielkie pieniądze na badania i rozwój. Niektóre z tych inwestycji przeznaczone są na badania stosowane, mające na celu komercjalizację odkryć z zakresu badań podstawowych, odkryć w laboratoriach uniwersyteckich. Gwarancją gratyfikacji dla innowatora za wysiłek, wytrwałość i zainwestowane środki jest system patentowy i prawa dotyczące licencjonowania.

Wynalazcy uczą się na podstawie patentów innych wynalazców i wynajdują nowe lub ulepszone urządzenia, metody, procesy, projekty, funkcje oprogramowania itp. Może to prowadzić do powstania nowych rynków, nowego bogactwa i korzyści dla konsumentów. W ten sposób innowacje w połączeniu z systemem ochrony własności intelektualnej napędzają gospodarkę – powstają całe gałęzie przemysłu, konkurenci wchodzą na nowo utworzone rynki, powstają nowe firmy, miejsca pracy, produkty, bogactwo, uzupełniane są fundusze na badania i rozwój – zaczyna się kolejny cykl innowacji.

Kto inwestuje w badania, ten ma wyniki

Inwestycje w badania i rozwój w całym obszarze OECD wynoszą średnio ok. 2 proc. PKB. Są jednak kraje w których przeznaczają się na nie prawie 5 proc. To Korea Płd. i Izrael. W bogatych krajach zachodnich jest to przeważnie ponad 3 proc. W Polsce ok. 1,3 proc. Wraz ze wzrostem inwestycji w aktywa intelektualne rośnie ich wpływ na gospodarkę.

Kraje takie jak Chiny, Indie i Korea Płd. odniosły już ogromne korzyści z budowania potencjału naukowego i technologicznego, licząc oczywiście na więcej, bo nie są jeszcze w elicie najzasobniejszych. UNESCO podawała, że Korea, która przeznaczająca ogromne środki na badania i rozwój (4), prześcignęła



4. Koreańska innowacyjność

Japonię pod względem wielkości eksportu zaawansowanych technologii już na początku XXI wieku. Jest to tym bardziej godne uwagi, że jeszcze w 1960 r. kraj ten był w dużej mierze gospodarką rolną, odbudowującą się po wyniszczającej wojnie.

Badania i rozwój prowadzone zarówno w sektorze publicznym, jak i prywatnym mają pozytywny wpływ na wydajność w całej gospodarce. Oczywiście wiedza może być zarówno generowana wewnętrznie, jak i pozyskiwana zewnętrznie, poprzez nabywanie licencji na innowacyjne rozwiązania, patentów lub praw autorskich. Kraje o niższym poziomie innowacyjności zwykle skazane są na zakup maszyn i urządzeń, które są nośnikami innowacji generowanych u liderów. Prowadzi to do wprowadzenia lub wdrażania innowacji u nich, co też sprzyja gospodarce i rozwojowi kapitału społecznego. Przykładem może być szkolenie pracowników w zakresie umiejętności wymaganych do opracowania lub wprowadzenia innowacyjnych produktów i procesów z zewnątrz. Nawet więc w przypadku kupowania gotowych innowacji generuje to korzyści gospodarcze dla krajów nabywających. Pamiętajmy, że rosnący poziom wykształcenia też jest skarbem narodów, o czym piszemy w innym artykule w tym wydaniu MT.

Nowszym przykładem podobnego typu są inwestycje przedsiębiorstw w oprogramowanie, które oczywiście również przyczyniają się do wydajności

przedsiębiorstw i wzrostu gospodarczego. Przykładowo w drugiej połowie lat 90. wkład oprogramowania we wzrost wydajności pracy w USA wzrósł ponad trzykrotnie w stosunku do okresu 1974–90. Wkład oprogramowania w amerykańską produktywność pracy w tym okresie był ponad dwukrotnie większy niż sprzętu komunikacyjnego i prawie dwie trzecie większy niż samego sprzętu komputerowego. Według dostępnych danych, w krajach OECD w ostatnich dekadach inwestycje w oprogramowanie generalnie przyczyniały się do wydajności pracy w większym stopniu niż inwestycje w sprzęt komunikacyjny i komputerowy.

Oprogramowanie jest specyficznym produktem. Nowy software to innowacja silnie związana z całą historią innowacji w danej dziedzinie. Współczesne programy, platformy, systemy, frameworki, dostarczone z zewnątrz, są generatorami własnych oryginalnych rozwiązań, na tysiące sposobów. To czysta żywa innowacja, która nieustannie rodzi nowe innowacje, przy czym nie chodzi tylko o nowe wersje, rozszerzenia i moduły przygotowywane przez lokalnych programistów.

Badania prowadzone przez ośrodki naukowo-badawcze, uczelnie i inne instytucje finansowane ze środków publicznych są również istotną częścią procesu innowacji. Charakterystyczne, że to w najbardziej słynących z innowacji a zarazem bogatych krajach najlepiej ułożona jest współpraca podmiotów zajmujących się „badaniami podstawowymi” z przedsiębiorstwami i sektorem rządowym, np. militarnym. Oczywiście i najbardziej znanym przykładem wydajnie działającego systemu są Stany Zjednoczone, ale przepływ ten funkcjonuje dobrze także w innych krajach, oczywiście są to kraje liderujące w dziedzinie innowacji i przekuwania ich na bogactwo.

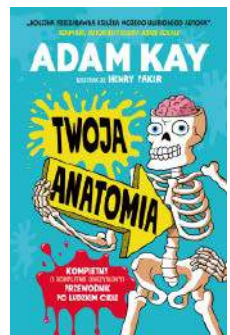
Spójrzmy na Izrael. Dla kraju niewielkiego, geograficznie oddalonego od głównych rynków światowych i niemal pozbawionego zasobów naturalnych,

Twoja anatomia

Adam Kay

Wydawnictwo Insignis, cena: 44,99 zł

Pierwsza książka dla młodych czytelników napisana przez autora głośnych bestsellerów „Będzie bolało” i „Świąteczny dyżur”. Czy zastanawiasz się czasem, jak działa twoje ciało? Ludzkie ciało bywa niesamowite, fascynujące i... bardzo dziwne. Twoje ciało jest dziwne. Moje ciało jest dziwne. A ciało twojego nauczyciela matematyki jest jeszcze dziwniejsze. Ta książka odpowiada na ważne pytania, takie jak: Czy jedzenie kóz z nosa jest bezpieczne? Ile czasu tracisz, siedząc w toalecie? Dlaczego w twoich rzęsach mieszkają obłe żyjątka? Usiądź wygodnie, załóż gumowe rękawiczki i pozwól się zabrać w smrodliwą i lepką podróż po twoich wnętrznościach.





5. Samolot Concorde

innowacje stały się najcenniejszym dobrem narodowym, kluczowym dla jego dobrobytu gospodarczego. Udało się tam stworzyć ekosystem innowacji, w którym rząd zapewnia środki regulacyjne w celu wzmocnienia infrastruktury dla innowacji i skutecznie, za pośrednictwem Izraelskiego Urzędu ds. Innowacji, potrafi skłaniać sektor prywatny do inwestowania w innowacje za pomocą różnych zachęt. Izraelski sektor high-tech jest zwykle chwalony za zwinność, multidyscyplinarne podejście i nieszablonowe myślenie, co pozwoliło wielu firmom i startupom szybko dostosowywać swoje technologie, rozwój i produkty do potrzeb wynikających z nowej sytuacji na świecie.

Czy innowacje gasną?

Na koniec warto wspomnieć o obawach dotyczących ogólnego zmniejszania się poziomu wynalazków i innowacji techniczno-naukowych we współczesnych czasach. Wiek prawdziwych innowacji – nazwijmy go Złotym Czwierćwieczem – trwał od około 1945 do ok. 1970 roku. Jeśli się temu dokładniej przyjrzeć, to trudno oprzeć się wrażeniu, że to w tamtym czasie albo stworzono wszystko, co do dziś definiuje współczesny świat, albo przynajmniej wtedy zasiane zostały nasiona techniki i nauki, która później rozkwitła. Elektronika, komputery i narodziny Internetu. Energia jądrowa. Telewizja. Antybiotyki. Półprzewodniki. Jest tego więcej: zielona rewolucja w rolnictwie, lotnictwo dla mas, tanie, niezawodne

i bezpieczne samochody, pociągi dużych prędkości. Umieściliśmy człowieka na Księżycu, wysłaliśmy sondy do Marsa, pobiliśmy ospę i odkryliśmy podwójną spiralę DNA – klucz życia.

Obecnie postęp wytyczają zorientowane na konsumenta, często banalne ulepszenia w technologii informacyjnej. Amerykański ekonomista Tyler Cowen w eseju „The Great Stagnation” z 2011 roku argumentuje, że przynajmniej w USA, osiągnięto technologiczny stan plateau, płaski etap bez znaczącego rozwoju. Oczywiście, nasze telefony są świetne, ale to nie to samo, co możliwość przelotu przez Atlantyk w ciągu ośmiu godzin lub wyeliminowanie ospy. Jak powiedział kiedyś amerykański guru branży IT, Peter Thiel: „Chcieliśmy latających samochodów, dostaliśmy 140 znaków na Twitterze”.

Ekonomiści opisują ten niezwykle powojenny okres według kryterium wzrostu zamożności. Po drugiej wojnie światowej nadszedł boom. PKB na głowę w USA i Europie gwałtownie wzrósł. Z połiołów Japonii powstały nowy przemysł. Niemcy przeżywały gospodarczy cud. Nawet blok komunistyczny znacząco się wzbogacił. Wzrost ten przypisuje się potężnym powojennym bodźcom, ze strony rządów oraz sprzyjającej wzrostowi i rozwojowi kombinacji niskich cen paliw, wzrostu liczby ludności i wysokich wydatków wojskowych w okresie zimnej wojny. Ale obok tego był ten niezwykle wybuch ludzkiej pomysłowości. Mniej mówi się o tym, być może dlatego, że wydaje się to oczywiste albo

jest to postrzegane jako prosta konsekwencja rozwijającej się gospodarki.

Czy po tym okresie rzeczywiście nic ciekawego nie powstało? Przeglądając świat technologii, ma się wrażenie, że nie. W 1971 roku zwykły samolot pasażerski potrzebował ośmiu godzin na przelot z Londynu do Nowego Jorku i nadal tak jest. Oczywiście, był jeden samolot, również powstały w Złotym Ćwierćwieczu, który mógł odbyć tę podróż w trzy godziny, ale Concorde'a (5) już nie ma. Nasze samochody są szybsze, bezpieczniejsze i zużywają mniej paliwa niż w 1971 roku, ale nie nastąpiła jak na razie radykalna zmiana modelu motoryzacji. Los elektromobilności wciąż nie jest oczywisty, a poza tym – przecież napęd elektryczny to historia jeszcze z XIX wieku.

Od czasu Złotej Epoki nie było nowej rewolucji materiałowej w dziedzinie tworzyw sztucznych, półprzewodników, nowych stopów i materiałów kompozytowych. Po zawrotnych przełomowych odkryciach pierwszej połowy XX wieku fizyka wydaje się kręcić w stanie bezsily, aby zatrzymać się w miejscu.

Dlaczego od lat 70. postęp wyhamował? Kryzys nie wyjaśnia wszystkiego. Zresztą skończył się po kilku latach i świat znów rozwijał się gospodarczo. W „The Great Stagnation” Cowen argumentuje, że postęp stracił impet, ponieważ zebrane zostały wszystkie „nisko wiszące owoce”. Do tych łatwo dostępnych owoców zalicza m.in. uprawę nieużytkowanej wcześniej ziemi, masową edukację i praktyczną „monetyzację” przełomowych odkryć naukowych dokonanych w XIX wieku. Dalszy postęp jest znacznie trudniejszy. Przejście od maszyn śmigłowych z lat 30. do odrzutowców z lat 60. było, jak się okazuje, łatwiejsze niż przejście od dzisiejszych samolotów do czegoś znacznie lepszego.

Niektórzy braku stymulacji do postępu upatrują w fakcie, że bogactwo świata jest obecnie

skoncentrowane w rękach maleńkiej elity. Raport Credit Suisse z października ubiegłego roku wykazał, że najbogatszy 1 procent ludzi posiada połowę światowych aktywów. Bogatych ludzi jest więcej niż kiedyś, ale nie są aktywni we wspieraniu nauki i techniki, skupiając się na konsumpcji. Dlatego że nie potrzeba nowych rozwiązań technicznych, aby się bogacić, jak to było kiedyś. Jak zauważył francuski ekonomista Thomas Piketty w książce „Capital”, pieniądź tworzy pieniądź szybciej, efektywniej niż kiedykolwiek wcześniej w historii. Gdy bogactwo narasta w tak spektakularnym tempie bez konieczności robienia czegokolwiek, nie ma bodźców do inwestowania w innowacje.

W Złotym Ćwierćwieczu duże znaczenie dla rozwoju miało finansowanie publiczne. Jednak z różnych powodów nastąpiło odejście od śmiałych, ale ryzykownych projektów. Przykładem tego procesu jest zahamowanie rozwoju energii nuklearnej po serii „katastrof”, na wyspie Three Mile, w której nikt nie zginął i Czarnobylu, która była specyficzna. Incydenty te spowodowały globalną zapaść w badaniach naukowych, które mogły dać nam bezpieczną, tanią i niskoemisyjną energię. Kryzys klimatyczny jest jedną z cen, jakie płacimy za 40 lat unikania ryzyka.

Do programu Apollo prawie na pewno dziś by nie doszło, dlatego że ryzyko – obliczone na podstawie kilkuprocentowej szansy śmierci astronautów – byłoby nie do przyjęcia. Boeing podejmował ogromne ryzyko, gdy opracował 747, niezwykłą maszynę z lat 60., która w ciągu niespełna pięciu lat przeszła od deski kreślarskiej do lotu. Jego nowoczesny odpowiednik, Airbus A380, po raz pierwszy poleciał w 2005 roku – 15 lat po rozpoczęciu projektu.

Czy owa „innowacyjna stagnacja” wpłynie na rozkład bogactwa na świecie na dłuższą metę? Niewykluczone, ale potrzeba będzie dekad, by zauważyć to zjawisko. ■

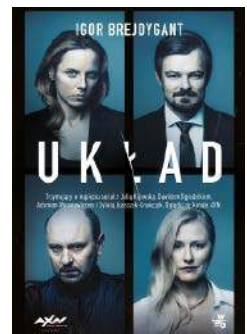
Miroslaw Usidus

Układ

Igor Brejdygant

Wydawnictwo W.A.B., liczba stron: 384, cena: 41,99 zł

Komisarz Monika Brzozowska nie może podnieść się psychicznie po tym, jak w obronie własnej zabiła swojego męża, płatnego zabójcę. Przychodzi jej to z tym większym trudem, że przełożeni odesłali ją na bezterminowy urlop. Niestety pani komisarz przy okazji powraca do zgubnego natogu, który przed laty prawie wyprowadził ją na tamten świat. Tymczasem zaczynają ginąć świadkowie, którzy mieli pomóc policji w doprowadzeniu na tawę oskarżonych szajki pedofilów. Wszystko wskazuje na to, że to właśnie niebezpieczni i wpływowi mężczyźni czerpiący uciechy z zabaw z nieletnimi próbują ukreślić tę sprawę. Z czasem okazuje się, że wszystko jest znacznie bardziej skomplikowane, a bycie katem wcale nie wyklucza bycia ofiarą.





1. Bogactwo w edukacji

Wykształcenie, wiedza, zasoby umiejętności – to część kapitału ludzkiego, która daje się precyzyjnie mierzyć (1). W dodatku tego bogactwa nie można ukraść czy jakoś inaczej utracić, poza szczególnymi przypadkami, które trudno brać pod uwagę. Majątek tego typu stanowi swoisty fundament, na którym opierają się inne skarby, choćby innowacyjność.

Wykształcenie to najlepsza waluta w kieszeni

POTĘGI KLUCZ

Cytowany przez nas wiele razy w tym numerze Adam Smith również pisał o tej części „bogactwa narodów”. Już na pierwszej stronie swojego dzieła stwierdza, że bogactwo narodu regulowane jest

przez „umiejętności, zręczność i rozsądek”, z jakimi praca narodu jest powszechnie wykonywana. Smith zauważa również, że wskaźnik poziomu wykwalifikowania siły roboczej jest ważniejszy niż wskaźnik zatrudnienia. Jak pisze, jeśli dwa narody mają podobną populację i zasoby naturalne, jedyne różnice w produktach, które wytwarzają, zależą od umiejętności ludzi i instytucji rządowych.

Nieco bardziej współczesnym językiem można wyrazić to tak: posiadanie wysoko wykwalifikowanej siły roboczej zwiększa produkcję i bogactwo. Jeśli kraj A zatrudnia większość swojej siły roboczej w górnictwie i rolnictwie, zaś kraj B w inżynierii i robotyce, to oczywiste jest, że kraj B ma znacznie większą wydajność, a zatem wyższy PKB. Kraj B ma dwa godne uwagi atuty. Po pierwsze, jego siła robocza będzie produkować maszyny do automatyzacji produkcji (co bezpośrednio zwiększy produkcję).

Po drugie, przy tej samej wielkości siły roboczej pracownicy z kraju B mają więcej wolnego czasu na dalsze doskonalenie swoich umiejętności i powiększanie wiedzy.

Postulat zdobywania wykształcenia bywa niekiedy uzupełniany o to, by było to wykształcenie sprzyjające rozwojowi gospodarczemu. Ekonomiści znają pojęcie „zwrotu z inwestycji w edukację”. Nie jest on zawsze i dla każdego kraju taki sam. Po pierwsze systemy szkolnictwa różnią się poziomem serwowanych usług. Po drugie – w wielu krajach edukacja jest bardzo droga a korzyści z niej nie w każdym przypadku odpowiadają poniesionym kosztom. Ważne, z gospodarczego punktu widzenia, są też preferowane kierunki kształcenia.

Dziś oznacza to więcej kształcących się w kierunkach technicznych, inżynierskich a także związanych z szeroko pojętą „computer science”, czyli programowaniem, przetwarzaniem danych. W krajach biedniejszych, np. afrykańskich, jest to coraz lepiej rozumiane, dlatego rządy biedniejszych państw, które chcą zmienić swoją pozycję w rankingu bogactwa, prowadzą politykę wspierającą kształcenie w pożądanym z punktu widzenia rozwoju PKB państwa kierunku.

Są też takie inicjatywy jak projekty Microsoftu, którzy poprzez swoją inicjatywę 4Afrika (2) rozwija umiejętności technologiczne w tych krajach. Od 2013 roku 4Afrika zdołała podnieść kwalifikacje blisko miliona młodych ludzi w ponad siedemnastu krajach afrykańskich. Innym przykładem zaangażowania sektora prywatnego jest program Seeds for the Future Africa firmy Huawei, w ramach którego w ciągu pięciu lat tysięcy afrykańskich studentów ma zostać wyedukowanych w zakresie technologii informacyjno-komunikacyjnych.

Dziś, na przykładzie tego, co robi się w Afryce, widać dobrze zrozumienie, w jaki sposób pomnażanie kapitału ludzkiego buduje bogactwo państw. Zamiast wysyłać pieniądze, daje się Afryce cenniejszą walutę. Jest jej zapewne wciąż za mało, ale z pewnością jest trwalszym fundamentem majątku narodowego niż gotówka.

Inwestycje w wykształcenie i szkolenie kreują zasobność także na innych niż państwo poziomach. Na przykład przekładają się na konkretne korzyści dla pracowników w postaci wyższych zarobków, wyższych wskaźników zatrudnienia i wyższego poziomu uczestnictwa w dalszym kształceniu i szkoleniu. Oczywiście w rzeczywistym świecie nie wszystko wygląda tak różowo. W wielu, nawet wysoko rozwiniętych krajach, mamy do czynienia z niedoinwestowaniem kapitału ludzkiego wśród pracowników, ponieważ ani firmy, ani osoby fizyczne nie mają pewności zwrotu z inwestycji. Wreszcie



Microsoft

2. Logo programu 4Afrika

brak rozwiązań, które pozwalałyby na systemowe uznawanie kompetencji nabytych poza formalną edukacją i szkoleniami, oznacza, że znaczna część inwestycji w kapitał ludzki doświadczonych pracowników pozostaje niewidoczna.

Mimo to panuje dość zgodne przekonanie, iż wiedza od dawna uznawana jest za cenny zasób umożliwiający wzrost organizacyjny i trwałą przewagę konkurencyjną, zwłaszcza w przypadku organizacji konkurujących w niepewnym otoczeniu. Wraz z pojawieniem się gospodarki opartej na wiedzy, w której wiedza, kompetencje i związane z nimi wartości niematerialne i prawne są kluczowymi czynnikami przewagi konkurencyjnej, jesteśmy świadkami wielu zmian w charakterze edukacji i wymagań stawianych uczelniom, tak aby stały się one magazynami innowacji, w których źródła talentów są odżywiane i podtrzymywane.

Jak napisał w 2003 roku Arwind Malhotra, znany badacz zagadnień edukacji i przedsiębiorczości, każde społeczeństwo posiada lub kontroluje pewną liczbę zasobów wiedzy, a pomiar poziomu tej magazynowanej wiedzy, zawartej w jednostkach, instytucjach i systemach, jak również potencjału do wzmocnienia istniejących zasobów wiedzy i generowania nowej wiedzy, jest bardzo użyteczny i służy jako cenne narzędzie, podnoszące świadomość i wspierające,

wskazujące braki w dostępnych zasobach wiedzy oraz mobilizujące polityczne wsparcie dla środków zaradczych, które należy podjąć.

Według opinii Komisji Europejskiej z 2006 roku, terminy „aktywa wiedzy”, „aktywa intelektualne” oraz „kapitał intelektualny” są stosowane zamiennie na oznaczenie kombinacji wartości niematerialnych i prawnych oraz działań, które pozwalają organizacji przekształcić pakiet zasobów materialnych, finansowych i ludzkich w system zdolny do tworzenia wartości dla zainteresowanych stron i innowacji organizacyjnych. Jednak w przeciwieństwie do aktywów fizycznych, które mogą mieć ograniczony okres użytkowania ze względu na zużycie, aktywa wiedzy mogą teoretycznie trwać wiecznie. Malhotra uważa, że w przeciwieństwie do danych, które można scharakteryzować jako właściwość rzeczy, wiedza jest właściwością podmiotów, która predysponuje je do działania w określonych okolicznościach.

Edukacja cyfrowych tubylców

Choć wartość wiedzy i edukacji uznawana jest od dekad, a właściwie od stuleci, to jednak podobieństwo do dawniejszych czasów dziś wygląda mocno ogólnie. Zarówno jeśli chodzi o rodzaj najbardziej pożądanej wiedzy, czyli najcenniejszego bogactwa, jak też gdy mówimy o sposobie jej pozyskiwania, edukowania, w ostatnich czasie zachodzą duże zmiany. W założeniu mają sprawić, że skarb ten, rozumiany i pozyskiwany po nowemu, będzie miał jeszcze większą wartość w nowych czasach.

3. Minecraft w szkole

Przekonanie, iż współczesne czasy wymagają nieco innego podejścia, zarówno jeśli chodzi o rozumienie samej wiedzy, nowych wartości, które może przekazywać edukacja, ale również sposobu nauczania, jest dość powszechne. O przykładach nowego rozumienia i podejścia do edukacji pisaliśmy w MT już lata temu. Pisaliśmy m.in. o tym, jak popularna wśród nastolatków gra komputerowa, Minecraft, została wprowadzona do programu nauczania jako jeden przedmiotów szkolnych. Stało się to w zawsze otwartej na edukacyjne eksperymenty Szwecji. Zdaniem nauczycieli pomoże to nauczyć 13-latków samodzielnego myślenia. Zainteresowani edukatorzy mogli ściągnąć z internetu specjalną wersję Minecrafta dla szkół (3).

Uniwersytet stanowy w amerykańskiej Wirginii opracowuje nowe przedmioty nauczania, od przedszkola, po szkołę średnią, wszystkie mają być oparte na wykorzystaniu techniki druku 3D. Naukowcy są przekonani, że technologia ta stanowi przyszłość przemysłu i wytwórczości. Dlatego uważają, że oswajanie dzieci z technologiami trójwymiarowego druku od wczesnego dzieciństwa pozwoli im się potem lepiej odnaleźć, nie tylko zawodowo, w realiach nowej ekonomii. Glen L. Bull, jeden z prowadzących projekt, mówi w rozmowie z „International Herald Tribune” to samo o drukarkach 3D, co od wielu lat mówiło się o komputerach, laptopach, a ostatnio o tabletach, że „każda klasa w szkole powinna być wyposażona w drukarkę trójwymiarową w ciągu najbliższych kilku lat”. Przedmioty wykorzystujące tę technologię uczą kreatywności w tworzeniu z jednej

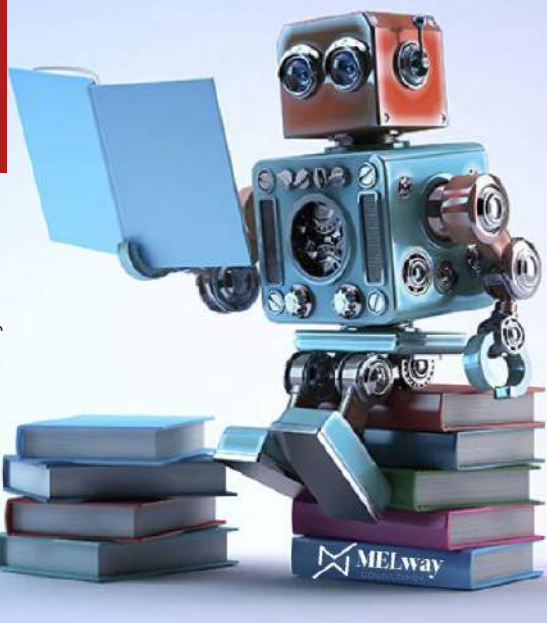


Nowe koncepcje edukacyjne często proponują, aby nowe przedmioty szkolne (o ile w ogóle to mają być przedmioty) wiązały się ściśle z zainteresowaniami uczniów. Chodzi o ty, aby spotkać się uczniem na polu, której jego najbardziej interesuje. Wspomniany Minecraft jako przedmiot jest przykładem tej tendencji. Z biegiem lat pojawiły się pomysły nawet bardziej radykalne, by szkoły wprowadzały takie przedmioty jak „Facebook”, „Lady Gaga”, „Komiksy Manga”, „Deskrolka” itp. Jak łatwo zauważyć, żeby te przedmioty miały sens z punktu widzenia zadań i roli szkoły, nauczyciele i twórcy programów nauczania muszą wykazać się nadzwyczajną elastycznością i talentami kreatywnymi. Przetworzenie przedmiotu „Lady Gaga” na coś użytecznego, wzbogacającego wiedzę i kształtującego umiejętność to przecież nie lada wyzwanie... Co nie znaczy, że się nie da. Za fantazyjnymi często nazwami

Rzeczą, o której się często zapomina, jest fakt, iż spora część życia dzieciom i młodzieży upływa w cyberswiecie, w internecie, w społecznościach, w aplikacjach mobilnych i grach. Oni są „digital natives” (5), czyli rdzennymi mieszkańcami cyberprzestrzeni. A pochodzenia i korzeni, jak wiadomo, nie da się zmienić. I tak zapewne będzie wyglądało ich całe życie, praca, zabawa, kontakty towarzyskie, aktywność zakupowa itd. Tradycyjnie rozumiana szkoła i treści edukacja są śmiesznie odległe od ich świata, więc wartość takiej edukacji spada w z punktów widzenia „cyfrowych tubylców”. Jednak wcale nie musi tak być. Wyobraźmy sobie przedmiot – „przysposobienie do życia w sieci”. Dlaczego nie



eprasa.pl fb7894959f



6. Ucząca się maszyna

nauczyć młodych krytycyzmu i umiejętności selekcji w zalewie informacji, co byłoby prezentem ze starego świata edukacji i wiedzy dla nowego świata? Warto nauczyć młodych ludzi komunikacji według określonych zasad, które sprawiają, że owo „życie w necie” stanie się w wielu jeszcze niezbyt ucywilizowanych sferach nieco znośniejsze oraz, co nie bez znaczenia, bezpieczniejsze.

Od digitalizacji do disintermediacji

W jaki sposób będzie przebiegał proces edukacji „pracowników przyszłości”? Firma Envisioning Technology opracowała swego czasu kompleksowe prognozy dotyczącą przyszłości technologii nauczania. Zauważa w nich m.in., że z jednej strony wykształcenie jest odpowiedzialne za przewidywanie umiejętności związanych z prawdziwym życiem, poprzez przygotowywanie nas do radzenia sobie z ogromną i szybko rosnącą złożonością rzeczywistości, z drugiej jednak – metody edukacji mogą być określone dopiero po praktycznych doświadczeniach. Ta dwoistość jest szczególnie wyraźna, kiedy mowa jest o technologiach, w tym przypadku bowiem szybko postępujące innowacje i nieustanne zmiany są jedynymi pewnikami.

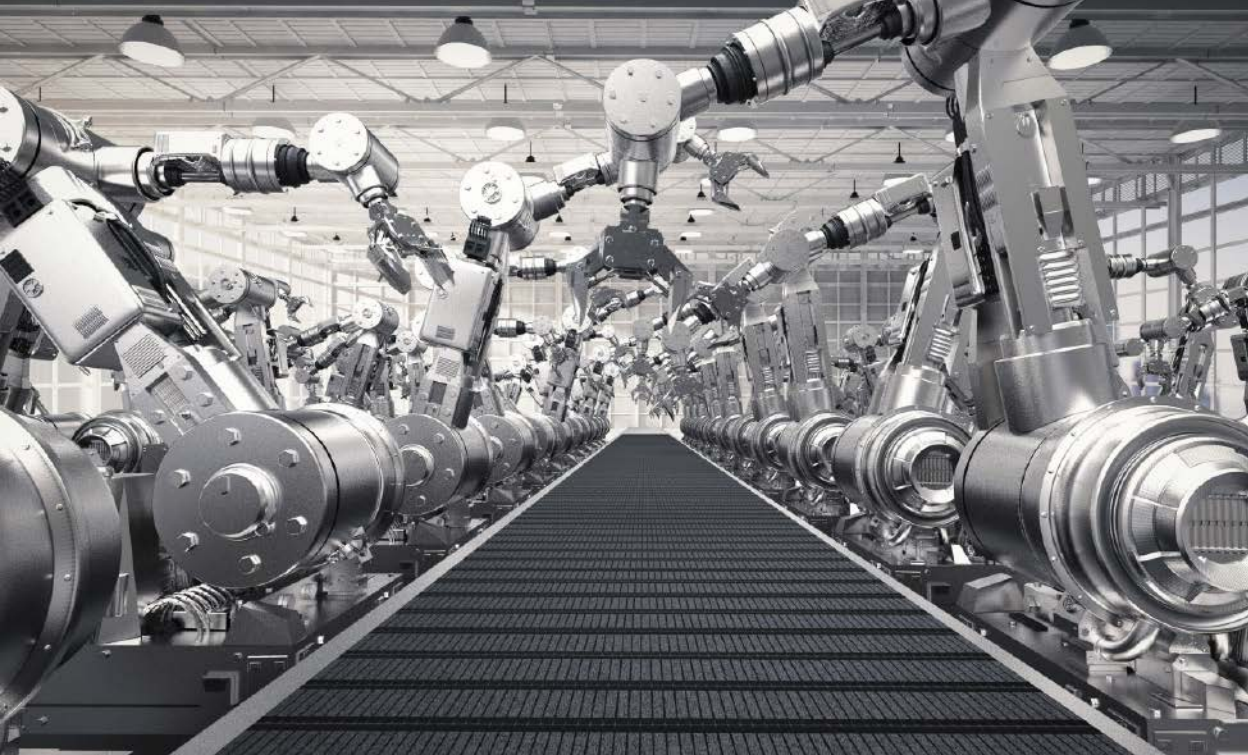
Mapa nowej edukacji, dostarczona przez Envisioning Technology już w 2014 r. – obejmuje perspektywę kolejnych 30 lat (czyli zmiany miałyby już zachodzić). Przedstawia postulowane i konieczne w obecnej sytuacji

zmiany w modelu i procesie wykształcenia. Główne jej punkty to:

- Całkowicie zdigitalizowane klasy szkolne. Nie jest to dokładnie to samo, co znana nam „cyfrowa szkoła”, ale są to pojęcia pokrewne. Całość nauki i procesu dydaktycznego powinna przenieść się do środowiska komputerowo-mobilnego.
- Grywalizacja (Gamification). To pojęcie znane jest także z innych dziedzin, np. ze specyficznych form promocji i marketingu internetowego. Z punktu widzenia edukacyjnego istotne jest stałe ćwiczenie i doskonalenie umiejętności. Elementy rywalizacji w grupie nie są konieczne, bo nie każdy dobrze znosi tego rodzaju stres, ale walka z samym sobą i samodoskonalenie – jak najbardziej.
- Otwarcie informacji. Tworzenie szkolnych i międzyszkolnych społeczności zarządzających i korzystających z wiedzy, która w całkowicie cyfrowych zasobach jest powszechnie dostępna.
- Technologiczna obecność i namacalność technologii na każdym kroku. Chodzi w skrócie o całkowicie skomputeryzowane oraz interaktywne środowisko, w którym przebywają uczniowie/studenci. Także o „Internet Rzeczy”.
- Multimedialne nasycenie miejsc, w których odbywa się edukacja. Obecnie do kreacji wizualnych, technologii wirtualnej rzeczywistości dodalibyśmy Augmented Reality, rozszerzoną rzeczywistość, która przenika i wzbogaca ćwiczenia i lekcje.
- „Disintermediacja”. To chyba najmniej czytelna część wizji Envisioning Technology a jednocześnie najbardziej radykalna. Chodzi w niej o podważenie roli nauczyciela. Nie tylko tej tradycyjnej hierarchicznej relacji, lecz jego lub jej fizycznej obecności w ogóle. Przewiduje się wprowadzenie algorytmów, np. planujących zadania, pytania, ćwiczenia i testy. Coś w rodzaju sztucznej inteligencji zamiast człowieka za katedrą.

Maszynowa przemiana

Jest jeszcze jeden aspekt bogactwa opartego na edukacji, charakterystyczny dla naszych czasów, o którym z oczywistych względów nie pisał ani Smith, ani żaden z teoretyków po nim. Chodzi o wartość gospodarczą wynikającą z postępów nauczania w bardzo specyficznym rozumieniu, mianowicie – nauki maszynowej (6). W rozwoju sztucznej inteligencji, której uczenie maszynowe jest jednym z najważniejszych przejawów, upatruje się obecnie źródeł przyszłego bogactwa i dominacji państw w świecie. Wielu obawia się, że tym razem głównymi centrami wzrostu tego bogactwa będą kraje



7. Zrobotyzowana hala produkcyjna

wschodniej Azji, z Chinami na czele, a nie świat zachodni. Ale to inna kwestia.

W najbardziej ogólnym ujęciu uczenie maszynowe polega na tym, że system komputerowy otrzymuje duże ilości danych, które następnie wykorzystuje do nauki wykonywania określonych zadań, takich jak np. rozumienie mowy czy rozpoznawanie zdjęć. Obecnie najczęściej do tego procesu używa się sieci neuronowych, czyli zainspirowanej przez ludzki mózg sieci połączonych ze sobą warstw elementów, zwanych neuronami, które wprowadzają dane do siebie nawzajem i które można przeszkolić do wykonywania określonych zadań poprzez modyfikację znaczenia przypisywanego danym wejściowym w miarę ich przemieszczania się między warstwami.

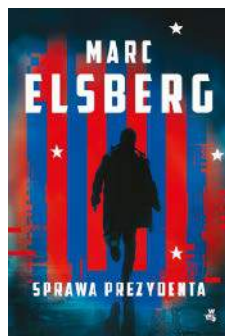
Sieć neuronowa uczy się na przykładach – trzeba jej przedstawić jakąś (w praktyce bardzo wielką) liczbę już rozwiązanych przykładów. Tok uczenia sieci zaczyna się od przypisania różnych lub losowych wag dla wyników. Sprawdzamy, czy odpowiedź sieci jest zgodna z pożądanym wynikiem, a następnie zmieniamy wagi tak, żeby wynik zbliżał się do pożądanego. Mechanizm regulacji wag jest częścią sieci neuronowej, tak zwaną regułą uczenia, i działa automatycznie. Wiele neuronów można połączyć w warstwy, gdzie jedna warstwa przekazuje wyniki kolejnej, aż do najwyższej warstwy, dającej odpowiedź. Struktura sieci (połączenia między neuronami) oraz mechanizm działania każdego neuronu zwykle są określone z góry. Zestaw wag, parametrów określających stopień wpływu

Sprawa prezydenta

Marc Elsberg

Wydawnictwo W.A.B., cena: 49,99 zł

Prawniczka Dana Marin nie wierzyła, że nadejdzie taki dzień: były prezydent Stanów Zjednoczonych w czasie wizyty w Atenach zostaje aresztowany przez grecką policję działającą na zlecenie Międzynarodowego Trybunału Karnego w Hadze. To momentalnie wywołuje dyplomatyczne tornado. W USA toczy się kampania wyborcza i aktualnie urzędujący prezydent nie może sobie pozwolić na jakikolwiek skandal. Biały Dom formułuje groźby pod adresem Międzynarodowego Trybunału Karnego i wszystkich państw Unii Europejskiej. A Dana Marin rozpoczyna walkę ze znacznie silniejszym przeciwnikiem – podobnie jak jej kluczowy świadek, którego zeznania mogą ostatecznie pogrążyć niegdyś najpotężniejszego człowieka na świecie. Amerykańskie tajne służby depczą już sygnaliście po piętach, a oddział specjalny przygotowuje się do siłowego uwolnienia byłego prezydenta, by za wszelką cenę uniemożliwić przekazanie go do Hagi...



poszczególnych wejść na wynik danego neuronu jest podstawową pamięcią sieci.

Uczenie maszynowe dzieli się na nadzorowanie i nienadzorowane z podzbiorem uczenia półnadzorowanego. Wyróżnia się też uczenie ze wzmocnieniem. W pierwszym przypadku zadaniem systemu pod ludzką kontrolą jest nauczenie się przewidywania prawidłowej odpowiedzi na pytanie. Rolę nadzorców w takim systemie wykonują np. pracownicy w trybie online, zatrudnieni na platformach takich jak Amazon Mechanical Turk. Zbiory danych treningowych są ogromne i powiększają się. W dłuższej perspektywie czasowej dostęp do ogromnych zestawów danych może również okazać się mniej ważny niż dostęp do dużych ilości mocy obliczeniowej. W ostatnich latach zaczęto wykorzystywać Generatywne Sieci Konfrontujące lub inaczej Przeciwnastawne (Generative Adversarial Networks – GAN), systemy nauki maszynowej, które mogą wykorzystywać oznaczone dane do generowania zupełnie nowych zasobów danych, na przykład tworząc nowe obrazy z istniejących ilustracji, wykorzystane z kolei do efektywniejszego szkolenia maszyn. Podejście to mogłoby doprowadzić do rozwoju kształcenia półnadzorowanego, w ramach którego systemy mogłyby uczyć się, jak wykonywać zadania, wykorzystując znacznie mniejszą ilość oznaczonych danych, niż jest to konieczne w przypadku systemów szkoleniowych wykorzystujących obecnie naukę pod nadzorem.

Głębokie uczenie (Deep Learning – DL), uchodzące za najbardziej udane podejście w dziedzinie uczenia maszynowego, wykorzystuje „neurony”, które układają się w warstwy i połączenia z innymi „neuronami” – podobnie jak neurony w naszych mózgach. Każda warstwa neuronów wybiera specyficzną cechę do identyfikacji. Kolejne warstwy pracują nad kolejnymi cechami, co prowadzi do nauki dokładnej identyfikacji stworzenia.

Na każdym kroku możemy znaleźć dziś prognozy mówiące o tym, jak w ramach autonomicznej rewolucji roboty wykorzystujące sztuczną inteligencję zastąpią

ludzi w pracy i większości zadań wykonywanych ręcznie lub za pomocą umysłu, które nadal wykonuje człowiek. AI może znacząco podnieść standard życia każdego z nas, automatyzując nie tylko powtarzalne zadania, ale także czynności wymagające podejmowania decyzji.

Z drugiej strony jednym z największych strachów w obliczu rewolucji inteligentnych maszyn jest obawa przed zabranie większości miejsc pracy od prostych, fizycznych, po biurowe i nawet menedżerskie stanowiska. Zdaniem specjalistów z firmy Deloitte, którzy opublikowali na początku sierpnia 2017 r. artykuł na ten temat w serwisie „Quartz”, nie ma powodów do obaw, gdyż automatyzacyjne zmiany doprowadzą do powstania nowych zawodów, których jeszcze nawet nie znamy, o czym wspominał też cytowany wcześniej raport Della. Deloitte przypomina historię z lat 1811–1817, gdy grupa angielskich robotników włókienniczych zagrożonych utratą pracy w związku z projektem zastąpienia ludzi maszynami zaczęła się skupiać wokół legendarnej postaci Neda Ludda i podpalać fabryki, niszczyć maszyny. Historia jednak pokazała, że tuż po redukcji zatrudnienia poprzez pojawienie się maszyny następuje znaczące powiększanie skali produkcji i sprzedaży. Podobne efekty pojawiły się dzięki automatyzacji Forda, która stworzyła podwaliny pod zatrudnienie milionów ludzi na ziemi.

I w końcu sama rewolucja komputerowa potwierdziła tę samą prawidłowość. Dzięki zautomatyzowaniu procesów produkcji (7) i zmniejszeniu ceny jednostkowej komputery stały się polem do popisu dla ludzkiej kreatywności i produktywności. Powstały setki nowych zawodów, nieistniejących w przeszłości w ogóle. Z tych historycznych analogii wynika, jak wnioskuje eksperci, lęk przed zastąpieniem ludzkości przez AI wydaje się całkowicie nieuzasadniony, gdy tak czy inaczej prowadzi do kumulacji większej ilości wiedzy i do zwiększania poziomu edukacji, co, jak wielokrotnie pisaliśmy, zawsze jest korzystne dla narodów jako pomnażanie bogactwa. ■

Mirosław Usidus

Mała Księżniczka

Karolina Lewestam

Wydawnictwo W.A.B., liczba stron: 96, cena: 39,99 zł

Poetycka baśń dla dzieci i dorosłych o tym, co w życiu najważniejsze.

Mała Księżniczka spotyka na swojej drodze zaginioną w tajemniczych okolicznościach legendarną amerykańską pilotkę, Amelię Earhart. Ale zanim dadzą sobie wzajemnie lekcję siostrzeństwa, bohaterka przemierzy kosmos, próbując lepiej zrozumieć swój ukochany kwiat – i napotka przy tym mnóstwo dziwacznych i pięknych postaci. Dzięki niej czytelnicy przeżyją cudowną przygodę, która rozwinie im skrzydła i doda odwagi. No i wreszcie dowiemy się, co naprawdę stało się z Amelią Earhart, która zaginęła podczas pionierskiego lotu nad Pacyfikiem...



Już dawno temu przenikliwi ludzie zrozumieli, że prawdziwe skarby to nie zamknięte w szkatułach kruszce i drogie kamienie. O wiele trwalsze bogactwo rodzi się w połączeniu tego, co ludzie mają w głowach, co wypracują używając zarówno głów, jak i rąk, ze sposobem, w jaki będą współpracować, organizować się w społeczeństwie i w państwie.

Bogactwo narodów wczoraj i dziś

TAM SKARB TWÓJ, GDZIE... GŁOWA TWOJA



1. Adam Smith

„Badania nad naturą i przyczynami bogactwa narodów”, znane również pod skróconą nazwą „Bogactwo narodów”, epokowa praca autorstwa Adama Smitha (1), jest bez wątpienia dziełem, które zmieniło świat. Po dwustu czterdziestu pięciu latach od publikacji jego twierdzenia nie stały się mniej aktualne, jeśli właściwie się je rozumie. Wymagają najwyższej przekładu na język nowych warunków technicznych i społecznych. Wciąż zgadzamy się z fundamentalnymi twierdzeniami Smitha, takimi jak to o konsumpcji jako głównym celu gospodarki. Nie zaprzeczmy, choć minęło prawie ćwierć wieku, tezie, że głównymi cechami życia ekonomicznego są: dążenie do własnego interesu, podział pracy i wolność handlu.

Smith dowodził m.in., że nie ma stałej ilości bogactwa w danym narodzie. Bogactwo, jak przekonywał, powinno być mierzone wielkością handlu towarami i usługami – tym, co dzieje się w zamkowych kuchniach i stajniach, a nie tym, co jest zamknięte w sejfach w zamkowym skarbcu.

Jeśli cechą bogactwa zaś jest zmienność, to naturalnie podobnie jest z jego popularną miarą – pieniądzem. Pieniądze nie mają żadnej wewnętrznej wartości. Choć osiemnastowieczne pieniądze były nadal w większości wykonane z metali szlachetnych, dla Smitha złoto było... warte swojej wagi w złocie, ale jego wartość w odniesieniu do czegokolwiek innego również nie jest stała. W istocie umowność pieniądza

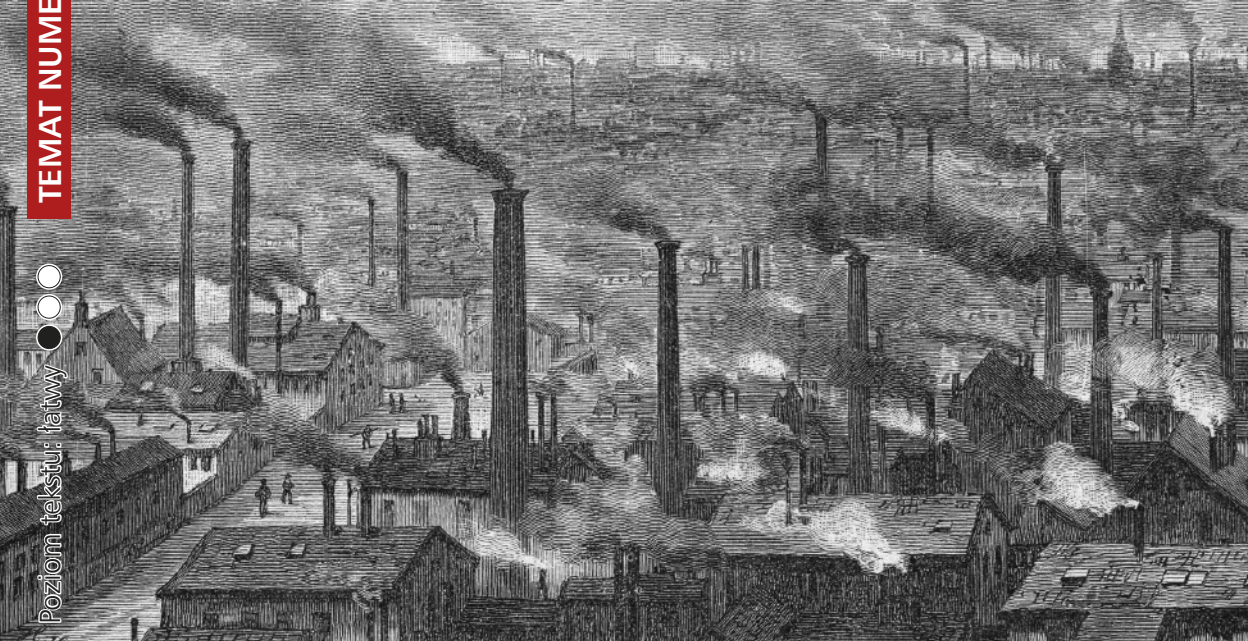
była dobrze rozumiana od czasów klasycznych. Korzystali z tego władcy i fałszerze, przy czym granica nie zawsze była jednoznaczna, zmniejszając ilość kruszcu w monetach.

Jak złoto zaszkodziło rewolucji?

Rozważania Smitha kontynuowali i „aktualizowali” kolejni myśliciele. Na przykład w 1998 roku pojawia się książka Davida Landes, „Bogactwo i nędza narodów”. Tytuł, w którym pobrzmiewa echo książki Adama Smitha, był nieco zwodniczy, bo książka ta jest dziełem historycznym, a nie traktatem ekonomicznym. Jednak można z niej dowiedzieć się wiele na temat rozumienia bogactwa dawniej i dziś.

Landes próbuje m.in. wyjaśnić, dlaczego w niektórych krajach udawało się kumulować zasoby efektywniej niż w innych. Oprócz znanych teorii, że np. Europejczykom sprzyjał umiarkowany, względnie pozbawiony ekstremów klimat i nadmorskie położenie, są tam też poszukiwania źródeł bogactwa w innych sferach.

Jak wiadomo, rewolucja przemysłowa szybciej i dynamiczniej niż gdzie indziej w Europie rozwijała się w Wielkiej Brytanii (2). Dlaczego nie we Francji, która była znacznie większa i bogatsza na starcie? Landes wyjaśnia to tym, że Wielka Brytania edukowała sprawnie i elastycznie swoją siłę roboczą, gromadząc kapitał, także ten ludzki, na bieżąco. Ważne było także zdecentralizowanie władzy w tym kraju.



2. Panorama Londynu z czasów rewolucji przemysłowej

Landes zwraca uwagę na drobiazgi, w których państwo pozostawiało swobodę swoim obywatelom. Suma tych drobiazgów decydowała o przewadze systemu brytyjskiego. We Francji np. jeszcze w latach dwudziestych ubiegłego wieku paryscy urzędnicy zmuszali kierowców wracających do miasta do poddania się pomiarom zawartości zbiorników paliwa za pomocą bagnetowych wskaźników, ponieważ paliwo, które mogło zostać zakupione poza bramami miasta, musiało zostać opodatkowane. Wielka Brytania nie stosowała tylu, aż tak drobiazgowych, administracyjnych szykan.

Już przed wiekami Brytyjczycy zaczęli przenosić się do swoich kolonii, gdzie powstały duże rynki lokalne, które w następstwie zrodziły potrzebę i wynalazek „amerykańskiego” systemu produkcji, czyli np. stosowania części zamiennych. Zanim minęło pół wieku, model ten rozprzestrzenił się w innych krajach, także w Wielkiej Brytanii.

Jak pisał Landes, owe „złoto i kosztowności”, gdy było ich w bród, mogły wpływać hamująco na impulsy rozwoju prawdziwego bogactwa. Na przykład Hiszpanom ogromne zasoby złota i srebra przywiezione z Nowego Świata „zaoszczędziły” konieczności zmian w kluczowym momencie. To klasyczny przykład pokazujący, jak te tradycyjne skarby w nadmiarze utrudniły zdobywanie bogactw w znaczeniu „smithowskim”.

Gdzie indziej w parze z nagromadzeniem „starych” skarbów zdobywanie nowych bogactw utrudniały inne czynniki. Landes pisze, że islam, religia o ogromnej żywotności, nie zdołał jednak zmotywować społeczeństw, w których panował, do nowoczesnego rozwoju. Dlaczego? Landes przygląda się indyjskiemu imperium Mogołów i tureckiemu imperium Osmanów i dowodzi, że w każdym z tych przypadków kluczem do porażki był fakt, że w muzułmańskim dogmacie wiara wystarczała



Mądrość mnicha. Czego nauczyło mnie życie w leśnym klasztorze Björn Natthiko Lindeblad

Wydawnictwo W.A.B., liczba stron: 208, cena: 39,99 zł

Bez pieniędzy, bez alkoholu, bez rodziny i bez wygod współczesnego świata. I co najważniejsze – na własne życzenie. Autor przez 17 lat mieszkał w tajskiej dżungli, gdzie prowadził życie buddyjskiego mnicha. Wyjechał u szczytu kariery w wieku 30 lat, powrócił tuż przed pięćdziesiątką schorowany i mocno zagubiony we współczesności. To, co trzymało go przy życiu, to buddyjskie lekcje, które pobierał przez lata życia w dżungli. „I may be wrong” to autobiografia mnicha zebrana w formie krótkich anegdot i życiowych lekcji, które pozwalają być bardziej „tu i teraz” w pędzącym świecie.

do zbawienia; nie wymagano żadnych innych oznak łaski, w tym doczesnego bogactwa. Zdaniem Landesa, największym błędem islamu była niechęć do korzystania z prasy drukarskiej. Choć krok ten mógł pomóc w zachowaniu ortodoksji, to nawrócone przez islam ziemie stały się intelektualnymi zaściankami. Waluta edukacyjna nie nabierała wartości, nie gromadzono złota, danych a skutkiem tego i diamentów innowacyjności nie było. Tak myśl Landesa zbliża się do metafory proponowanej przez nas w tym wydaniu „Młodego Technika”.

„Ekonomia wiedzy”

Kolejnym autorem, który w tytule i nie tylko nawiązywał do dzieła Smitha oraz do rozważań na temat pochodzenia bogactwa narodów, był David Warsh (3), który wydał w 2006 r. książkę pt. „Wiedza i bogactwo narodów”. Przedstawia ona wielką sprzeczność, która kołaczy w myśli nauki analizującej te problemy od 1776 roku, kiedy to Adam Smith opublikował swoją książkę. Warsh określa to jako walkę między fabryką szpilek a niewidzialną ręką rynku. Z jednej strony Smith podkreślał ogromny wzrost produktywności, jaki można osiągnąć dzięki wynalazkowi podziału pracy, co ilustruje użyty przez niego po raz pierwszy przykład fabryki szpilek, której pracownicy, specjalizując się w wąskich zadaniach, produkują znacznie więcej, niż gdyby pracowali niezależnie. Jednak wąska specjalizacja niesie za sobą zmniejszenie czegoś, co można określić jako „kapitał ludzki”, gdyż sprowadza jednostkę do bardzo wąskiego zakresu umiejętności. Malejący kapitał ludzki nie może prowadzić do kumulowania bogactwa. I mamy problem.

Sprzeczność tę próbuje rozwiązać proponowana przez Warsha i wielu innych, „endogeniczna teoria wzrostu”, która najogólniej polega na tym, że postęp techniczny (rozumiany jako akumulacja wiedzy naukowo-technicznej i kapitału ludzkiego) jest przede wszystkim skutkiem decyzji inwestycyjnych konsumentów i producentów, którzy zawsze postępują



3. David Warsh

racjonalnie. Według tej koncepcji kraj inwestujący w kapitał ludzki nie zetknie się z malejącymi korzyściami (także w sensie społecznym) ze skali. Modele endogeniczne wzrostu kładą nacisk na efekt pośredni polegający na tym, że wzrost inwestycji w jednej firmie przyczynia się do globalnego wzrostu produkcji poprzez zwiększenie ogólnego poziomu wiedzy i kwalifikacji.

David Warsh i wielu innych teoretyków proponuje nowy rodzaj „ekonomii wiedzy”, która opiera się na kumulacji zasobów intelektualnych w ludzkich głowach, repozytoriach, instytucjach i opisanych dokładnie procesach. To jest obecnie prawdziwe bogactwo narodów. To właśnie tę sferę, zdaniem nie tylko tego myśliciela, powinniśmy kultywować i rozwijać. ■

Mirosław Usidus

Ogród na cztery pory roku

Sebastian Kulis

Wydawnictwo Buchmann, liczba stron: 240, cena: 65,00 zł

Co siać w lutym, a co do ziemi włożyć dopiero w kwietniu? Jak zadbać o ogród jesienią i jak przygotować grządki po zimie. No i najważniejsze – co z tymi wszystkimi plonami zrobić. Sebastian Kulis, autor bloga Roślinne Porady, zabiera nas w roczną podróż po swoim ogrodzie. Miesiąc po miesiącu podpowiada co warto wysiać, jak zadbać o ziemię i jak zaplanować kolejne miesiące w ogrodzie, na balkonie i na parapecie, bo jako miejski ogrodnik żadnych wyzwań się nie boi. Co więcej, w swojej książce nie tylko zaraża pasją do ogrodnictwa, ale też równie silnie szerzy idee zero waste. Z Kulisem nic się nie zmarnuje, nawet najędzniejszy pomidorek z parapetowej hodowli. Szpadle w dłoń i do piachu!



Jeśli ktoś nie wierzy w szczerą wartość danych (1), to powinien spojrzeć na Facebooka, na Google i na kilku innych potentatów Big Tech. Ich liczone w miliardach przychody rodzą się nie z czego innego, tylko właśnie danych.

Współczesne złoto, czyli dane

ZANIM POLEJE SIĘ KREW



1. Dane jako złoto

Porównujemy dane do złota, ale to tylko symbol. Równie dobrze można odnosić je do innych tradycyjnych źródeł bogactwa. np. tym, czym sto lat temu była ropa naftowa, tym są teraz dane (2). Tak na to patrzył w 2017 magazyn „The Economist”. Co ciekawe, co zauważa pismo, analogia rozciąga się również na działania władz i regulatorów zaniepokojonych monopolem gigantów technologicznych na dane, podobnie jak ponad wiek temu niepokój wzbudzał monopol na rynku ropy. Dominacja wywołała falę planów rozbicia gigantów technologicznych, tak jak to miało miejsce w przypadku Standard Oil na początku XX wieku.

Jak zauważa „The Economist”, obfitość danych zmienia naturę rynkowej rywalizacji, pozwalając firmom technologicznym na osiąganie wykładniczych korzyści. Im więcej danych gromadzą te firmy, tym większy fundament, na którym mogą ulepszać swoje produkty, co przynosi jeszcze więcej użytkowników, którzy generują jeszcze więcej danych. Niekończąca się spirala generowania bogactwa. np. im więcej danych zbiera Tesla na temat swoich w dużym stopniu autonomicznych samochodów, tym bardziej może poprawić ich wydajność i bezpieczeństwo. Jest to część powodu, dla którego Tesla, która sprzedała tylko 25 tysięcy samochodów w pierwszym kwartale, jest teraz warta więcej niż GM, który sprzedał 2,3 miliona. Dane jednocześnie chronią gigantów technologicznych przed konkurencją, i wzywa do nowego podejścia regulacyjnego.

Kontrola danych przez firmy internetowe daje im nie tylko bajeczne bogactwo, ale również ogromną władzę. Jednocześnie stare sposoby myślenia o konkurencji, opracowane w erze ropy naftowej, wydają się przestarzałe w tym, co zaczęto nazywać „gospodarką opartą na danych”. Potrzebne jest nowe podejście.

Czym zajmuje się General Electric? Czym Siemens? Danymi oczywiście

Pojawienie się najpierw internetu, potem urządzeń mobilnych, głównie smartfonów, sprawiło, że generowana jest ogromna obfitość danych, które są o wiele



2. Dane jako współczesna ropa naftowa

bardziej wartościowe niż wcześniej. Niezależnie od tego, czy idziemy pobiegać, oglądamy telewizję, czy nawet stoimy w korku, praktycznie każda czynność tworzy cyfrowy ślad – kolejne dostawy surowca do destylarni danych. W miarę jak urządzenia, od zegarków po samochody, łączą się coraz powszechniej z internetem. Ilość danych rośnie i będzie rosnąć. Niektórzy szacują, że w pełni autonomicznych samochodów będzie generował 100 gigabajtów na sekundę. Jednocześnie techniki sztucznej inteligencji, takie jak uczenie maszynowe, wydobywają z danych coraz większą wartość. Algorytmy mogą przewidzieć, kiedy klient jest gotowy do zakupu, silnik odrzutowy wymaga serwisowania lub czy dana osoba jest zagrożona chorobą. Znamienne, że starzy giganci przemysłowi, tacy jak General Electric i Siemens, przedstawiają się obecnie również jako firmy zajmujące się danymi.

Ośrodek analityczny CEPS szacuje, że 92 proc. danych świata zachodniego jest obecnie przechowywanych w USA. Amerykańscy giganci technologiczni korzystają z tego od dawna, czerpią korzyści również z efektów sieci. Im więcej użytkowników zarejestruje się na Facebooku, tym bardziej atrakcyjne staje się to dla innych. W przypadku danych mamy do czynienia z dodatkowymi efektami sieciowymi. Zbierając więcej danych, firma ma większe możliwości ulepszania swoich produktów, co przyciąga kolejnych użytkowników, generujących jeszcze więcej danych, i tak dalej. Im więcej danych zbiera Tesla ze swoich samokierujących się samochodów, tym lepiej mogą się one prowadzić.

Dostęp do danych chroni firmy przed rywalami także w inny sposób. Google widzi, czego ludzie szukają, Facebook, co udostępniają. Amazon, co kupują. Są właścicielami sklepów z aplikacjami i systemów operacyjnych, a także wynajmują moc obliczeniową startupom. Mają dobry przegląd działań na swoich rynkach i poza nimi. Widzą, kiedy nowy produkt lub usługa zyskuje na popularności, co pozwala im skopiować go lub po prostu kupić startup, zanim stanie się on zbyt dużym zagrożeniem. Wielu uważa, że zakup WhatsApp przez Facebooka za 22 miliardy dolarów w 2014 roku, aplikacji do przesyłania wiadomości zatrudniającej mniej niż 60 pracowników, należy do tej kategorii przejęć, które eliminują potencjalnych rywali. Poprzez tworzenie barier wejścia i systemów wczesnego ostrzegania dane mogą zdławić konkurencję.

Natura danych sprawia, że środki antymonopolowe z przeszłości stają się mniej użyteczne. Rozbite firmy takiej jak Google na pięć kawałków nie powstrzymałoby efektów sieciowych przed ponownym

ujawnieniem się: z czasem jeden z nich znów stałby się dominujący. Potrzebne jest radykalne przemysłenie – a gdy zarys nowego podejścia zaczyna być widoczny, pojawiają się dwie koncepcje.

Po pierwsze, organy antymonopolowe muszą przesiąść się z ery przemysłowej w XXI wiek. Na przykład, rozważając fuzję, organy te tradycyjnie posługiwały się wielkością, aby określić, kiedy należy interweniować. Obecnie, oceniając skutki transakcji, muszą brać pod uwagę zakres zasobów danych, jakimi dysponują przedsiębiorstwa.

Druga zasada polega na rozluźnieniu kontroli dostawców usług internetowych nad danymi i przekazaniu większej kontroli tym, którzy je dostarczają. Pomocna byłaby większa przejrzystość: przedsiębiorstwa mogłyby zostać zmuszone do ujawnienia konsumentom, jakie informacje przechowują i ile na nich zarabiają. Rządy mogłyby zachęcać do powstawania nowych usług, otwierając więcej własnych skarbów danych lub zarządzając kluczowymi częściami gospodarki opartej na danych jako infrastrukturą publiczną, tak jak to robią Indie ze swoim systemem cyfrowej tożsamości, Aadhaar. Mogłyby również wprowadzić obowiązek udostępniania określonych rodzajów danych za zgodą użytkowników – takie podejście Europa stosuje w usługach finansowych, wymagając od banków udostępniania danych klientom stronom trzecim.

Projekty publicznych, silnie zabezpieczonych repozytoriów prywatnych danych obywateli, odseparowanych od Big Tech, pojawiły się już w niektórych krajach, np. w USA. W Polsce znany jest projekt Jana J. Zygmontowskiego z polskiego projektu SpołTech, który wymyślił termin „wspólnice danych” dla alternatywnego projektu zarządzania naszą prywatnością. Chodzi o przeniesienie komunikacji na platformę zgodną z polskim prawem i podległą jego przepisom oraz danych użytkowników do „wspólnicy”, w której zarządzać nimi będą sami użytkownicy a nie „obce państwo” takie jak Facebook czy Google.

Dobro narodowe, wyznacznik siły i suwerenności

Współczesne rozwiązania informatyczne umożliwiają dostęp do danych na szerszą skalę i w szerszym zakresie niż kiedykolwiek wcześniej. Wartość tych połączonych danych musi szybko przełożyć się na wiedzę i wgląd. Na poziomie strategicznym zdolność Wielkiej Brytanii do wykrywania, rozumienia, przypisywania i działania będzie wynikać z decyzji opartych na danych. Na poziomie operacyjnym przewagę osiągną ci, którzy potrafią opracować



3. Identyfikator w indyjskim systemie Aadhaar

wgląd w dane w czasie rzeczywistym, aby promować dobrobyt, wywierać wpływ i chronić ludzi przed szkodą.

Wymaga to nowych narzędzi wspomagających podejmowanie decyzji oraz nowych umiejętności w zakresie podejmowania wysokiej jakości decyzji w zmiennych, niepewnych, złożonych i niejednoznacznych sytuacjach operacyjnych, przy jednoczesnym zrównoważeniu dużej ilości danych pochodzących z różnych źródeł. Automatyzacja, sztuczna inteligencja (AI) i uczenie maszynowe są obecnie powszechnymi narzędziami do analizy danych, ale nie zostały jeszcze sprawdzone na skalę. Istnieją znaczne możliwości wykorzystania AI do wyciągania wniosków z ogromnych, rozproszonych zbiorów danych, których ludzie nigdy nie byłoby w stanie zidentyfikować przy użyciu tradycyjnych metod analitycznych.

Dane muszą być traktowane jednocześnie jako dobro narodowe. Co do tego dziś już mało kto ma wątpliwości. W przypadku każdego zasobu kluczowe jest zrozumienie jego wartości i właściwe obchodzenie się z nim. Dane mogą poprawić sposób, w jaki działają rządy i organizacje; jeśli są wykorzystywane w sposób inteligentny, mogą zasilać zasoby finansowane ze środków publicznych, z których regularnie korzystamy, czego przykładem jest opieka zdrowotna. Także obrona narodowa i samorządy lokalne mogą być wspierane przez cenne zbiory danych, jeśli będą one wykorzystywane w sposób etyczny, bezpieczny i mądry.

Z uznanego raz faktu, że dane są dobrem narodowym, wynikają pewne konsekwencje, np. takie, że muszą być chronione i potrzeba wręcz „cyfrowego nacjonalizmu”. Otwarcie mówili o tym indyjscy politycy, np. z ruchu Jagaran Manch, wzywając do obronny suwerenności tego kraju w dziedzinie

Małe Licho i babskie sprawy

Marta Kisiel

Wydawnictwo Wilga, liczba stron: 304, cena: 34,99 zł

Czwarta klasa to nie przelewki, o czym Bożydar Antoni Jekiętek przekonuje się na własnej skórze. Również w domu nie ma co liczyć na spokój, kiedy z każdego kąta wytażą krasnoludki i sikają do mleka, Tsadkiel organizuje korki z matmy, Licho zostaje kierownikiem spółdzielni mieszkaniowej, a wujek Konrad usiłuje pracować. Ale prawdziwe zamieszanie zaczyna się, gdy do klasy czwartej be nieoczekiwanie wkracza Jedna Wielka Zmyłka, wywołując wilka z lasu. Czy też, jak w tym przypadku, widmowego Szczęsnego z zaświatów. Krasnoludki, babskie sprawy i tatuś w przypiływie romantyzmu. Nic dziwnego, że pochłonięty Sprawami Bożek nie zauważa, co czai się w szkolnym budynku. A jest to coś straszniejszego niż...



zarządzania danymi generowany przez to ogromne społeczeństwo. W liście do premiera Narendry Modiego, polityk tego ugrupowania, Ashwani Mahajan przekonywał, że dane Hindusów powinny być odpowiednio przechowywane i przestać być wykorzystywane jako karta przetargowa w różnego rodzaju transakcjach handlowych.

„Dane powinny być nie tylko przechowywane w granicach geograficznych Indii, ale muszą być również tutaj przetwarzane”, pisał Mahajan. Oznacza to w praktyce dążenie do budowania indyjskich alternatyw dla usług płatniczych, platform mediów społecznościowych oraz e-commerce. Partia opowiedziała się również za stworzeniem „jak najszybciej” własnego komputerowego systemu operacyjnego Indii.

„The Wall Street Journal” poinformował w czerwcu 2021 r. o planach chińskich władz dotyczące regulacji i wykorzystania zettabajtów danych pochodzących z sektora prywatnego. W czternastym planie pięcioletnim, który został opublikowany w marcu, rząd centralny uznał dane za jeden z zasobów narodowych, które stanowią podstawę gospodarki kraju. Celem jest stworzenie cyfrowego zarządzania i uczynienie wszystkiego, od fabryk po miasta, „inteligentnymi”. Chiński rząd może naciskać na chińskie firmy technologiczne i międzynarodowe przedsiębiorstwa, aby udostępniały dane, nie tylko w odpowiedzi na prośby policji lub śledztwa dotyczące bezpieczeństwa narodowego, ale hurtowo, do dyspozycji rządzących.

Jak napisał „Wall Street Journal”, chiński rząd silnie naciska obecnie chińskich gigantów technologicznych, takich jak Tencent, Alibaba Group i ByteDance, do dzielenia się swoimi danymi, co może być chińską wersją tego, co dzieje się w USA, czyli walki z chińskim Big Techem. Niektórzy widzą w tym sposób na czerpanie zysków z przejęcia danych. Kilka lokalnych rządów prowincji chińskich utworzyło „platformy wymiany danych”, aby ułatwić pozyskiwanie i wykorzystywanie prywatnych danych. Nie trzeba chyba wyjaśniać, że pozycja chińskiego rządu



4. Jack Ma i szef Komunistycznej Partii Chin Xi Jinping

wobec technologicznych potentatów w swoim kraju jest znacznie silniejsza niż w Stanach Zjednoczonych, gdzie wyniki postępowań antymonopolowych nie są takie oczywiste.

Według wielu pogłosek Pekin rozprawia się z firmą Alibaba (tajemnicze zniknięcie na pewien czas jej założyciela Jacka Ma) przede wszystkim na tle dostępu do danych, a nie z innych powodów, które mogą być jedynie pretekstami z zasłoną dymną (4). Niektórzy obserwatorzy branży opisują, że Pekin jest coraz bardziej zmęczony odmową giganta handlu elektronicznego dzielenia się z organami regulacyjnymi ogromnymi ilościami danych o konsumentach. Duże firmy zagraniczne, takie jak Tesla i Apple, ugięły się i w odpowiedzi na zaostżone przepisy dotyczące bezpieczeństwa danych lokalizują ich przechowywanie. Nowa ustawa o bezpieczeństwie danych, która weszła w życie w czerwcu, daje chińskim władzom większe uprawnienia w zakresie ograniczania firm technologicznych i transferu danych do innych krajów. Pozwala również Chinom na wprowadzenie „wzajemnych” ograniczeń w dostępie do danych w odpowiedzi na posunięcia innych krajów.

Niedwuznacznie sugeruje to „wojny o dane” w przyszłości. Jeśli uznajemy je za współczesne bogactwo, złoto czy nawet ropę naftową, to nie jest nic specjalnie nowego i zaskakującego, że o takie zasoby toczą się wojny. ■

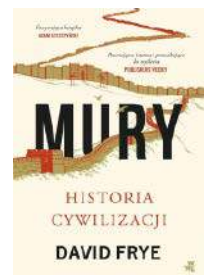
Mirosław Usidus

Mury. Historia Cywilizacji

David Frye

Wydawnictwo W.A.B., liczba stron: 352, cena: 49,99 zł

Przez tysiące lat ludzie budowali mury i atakowali je, podziwiali je i się przeciw nim buntowali. Wielkie mury wyrosły niemal na każdym kontynencie, w Persji, Rzymie, Chinach, Ameryce Środkowej i nie tylko. Towarzyszyły powstawaniu miast, narodów i imperiów. A jednak rzadko pojawiają się w podręcznikach historii. David Frye odkrywa historię murów tych starożytnych i tych całkiem współczesnych, i zadaje pytania, które są zarówno intrygujące, jak i głębokie. Czy mury umożliwiły rozwój cywilizacji? Czy możemy bez nich żyć? Historia murów staje się czymś więcej niż opowieścią o ceglach i kamieniu; staje się opowieścią o tym, kim jesteśmy i jak się takimi staliśmy. To nic innego jak historia cywilizacji.





1. Woda we Wszechświecie

Poszukiwania wody w kosmosie

WszeH₂Oświat

Nasze wyobrażenie o życiu wiąże je ściśle z wodą i to w stanie ciekłym, co, w gruncie rzeczy, oznacza stosunkowo wąski przedział w skali temperatury, mniej niż sto stopni. Choć trzeba mieć na względzie fakt, że temperatura to nie wszystko, miejsc w Układzie Słonecznym i we Wszechświecie, gdzie może występować życiodajna ciecz, odkrywamy coraz więcej...

Woda kosmiczna (1) niejedno ma oblicze. Oprócz zamrzniętych brył, śniegu kometarnego i bardzo słonych solanek, może po prostu unosić się w postaci ogromnych obłoków w przestrzeni. Kilka miesięcy temu międzynarodowy zespół naukowców znalazł w odległości 12 miliardów lat świetlnych chmurę „pary wodnej” wokół kwazaru. Szacuje się, że zawiera co najmniej 140 bilionów razy więcej wody niż wszystkie morza i oceany na Ziemi. Co istotne, uczeni sądzą, że nie są to drobiny lodu, przynajmniej nie w całości, lecz także kropelki. Dlaczego?

Znany jako APM 08279+5255, kwazar ten zawiera czarną dziurę miliardy razy większą od Słońca, produkującą aż bilion razy więcej energii. Czarna dziura

w centrum kwazaru otoczona jest chmurą gazu i pyłu, którą nieustannie pochłania, emitując energię w przestrzeń. W przypadku APM 08279+5255 chmura ta zawiera także ogromne ilości wody. Wyższa temperatura chmur gazu, w porównaniu do innych wód w kosmosie, jest spowodowana promieniowaniem rentgenowskim i promieniowaniem emitowanym przez czarną dziurę. Te dwie rzeczy sprawiają, że woda jest stale podgrzewana. Stanowi też dużą część „posiłku” czarnej dziury. Ze względu na odległość od Ziemi, ta chmura wody w kosmosie znajduje się w odległości 12 miliardów lat świetlnych, co zdaniem uczonych wskazuje, że woda była obecna w naszym Wszechświecie od jego wczesnej fazy istnienia.

Rekordowego odkrycia dokonał międzynarodowy zespół, działający w pobliżu szczytu Mauna Kea na Hawajach. Zespół pracował w Obserwatorium Submilimetrowym Caltech, wykorzystując technologię znaną jako Z-Spec. Urządzenie to funkcjonuje jako spektrograf i jest znacznie bardziej czułe niż inne podobne urządzenia.

Woda to życie... w różnym tego słowa znaczeniu

Jak wspomnieliśmy na początku, naukowcy generalnie uważają obecność wody za kluczowy i konieczny warunek do powstania jakiegokolwiek rodzaju życia podobnego do tego na Ziemi. W związku z tym prawie każda nasza sonda lub lądownik są wyposażone w urządzenia do analizy powierzchni księżycy lub planety w poszukiwaniu cząsteczek wody. To poszukiwanie wody jest jednym z głównych motorów badań Układu Słonecznego. Warto zaznaczyć, że choć istnienie wody wydaje nam się koniecznym do istnienia życia, to rozumowanie nie działa w drugą stronę – istnienie wody nie musi koniecznie oznaczać istnienia życia. Jednak zawsze jej odkrycie budzi zainteresowanie i falę rozważań o konsekwencjach i wynikających z tego możliwościach. Tak było np. z odkryciem wody na Księżycu sprzed paru lat, które zrewolucjonizowało dotychczasowe teorie na temat potencjału tego satelity i po raz kolejny zmieniło spojrzenie na kwestię występowania wody w kosmosie. Wpłynęło również na plany eksploracji kosmosu przez ludzi. Woda w takich miejscach jak Księżyc jest uznawana za niezbędny warunek do ustanowienia tam stałej obecności, bazy, eksploatacji, kolonii ludzkich nawet. Jest kluczowym elementem systemów

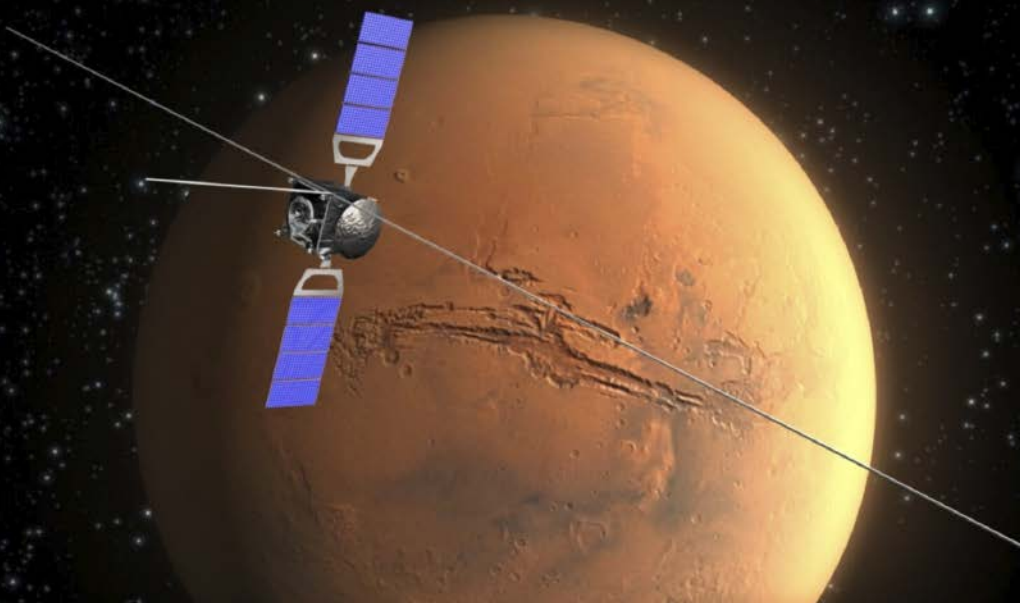
podtrzymywania życia i surowcem do produkcji paliwa raketowego (wodór i tlen).

Ponieważ Księżyc nie posiada żadnej znaczącej atmosfery, woda w stanie ciekłym nie może istnieć na jego powierzchni. Każdy lód wodny wystawiony na działanie promieni słonecznych szybko przekształca się w parę wodną i rozpada na tlen i wodór. Lód wodny byłby stabilny tylko w obszarach Księżyca położonych w cieniu, np. w pobliżu biegunów i w kraterach. Dopiero w sierpniu 2018 roku NASA potwierdziła za pomocą instrumentu Moon Mineralogy Mapper (M3), że na powierzchni Księżyca w pobliżu biegunów znajduje się lód wodny (2). Później Astronomiczne Stratosferyczne Obserwatorium w Podczerwieni NASA (SOFIA) potwierdziło obecność wody na nasłonecznionej powierzchni Księżyca. Odkrycie to wskazuje, że woda może być rozpraszana po powierzchni Księżyca, a nie występować jedynie w zimnych, zacienionych miejscach. Aparatura wykryła cząsteczki wody w Kraterze Claviusa, jednym z największych kraterów widocznych z Ziemi, znajdującym się na półkuli południowej Księżyca. Dane z tego miejsca mówią o występowaniu wody w stężeniach od 100 do 412 części na milion, co można porównać z butelką wody na metr sześcienny gleby. Wyniki zostały opublikowane w „Nature Astronomy”.

Wprawdzie odkryte stężenie wody „gruntowej” na Księżycu jest relatywnie niewysokie – dla porównania, pustynia Sahara ma 100 razy większą ilość wody w gruncie niż stężenie wykryte przez SOFIA na Księżycu. Jednak odkrycie to nasuwa pytania, jak powstaje tam woda i jak utrzymuje się ona w tak surowych warunkach jak księżycowe. SOFIA znajduje się na pokładzie zmodyfikowanego Boeinga 747SP,

2. Obrazy lodu wodnego na biegunach Księżyca wykonane przez Moon Mineralogy Mapper





3. Sonda Mars Express z instrumentem MARSIS na pokładzie

latającego na wysokości 15 tys. metrów ponad 99 proc. pary wodnej zawartej w ziemskiej atmosferze, co pozwala uzyskać wyraźniejszy obraz w podczerwieni. Wykorzystując urządzenie Faint Object infraRed CAmera współpracujące z 106-calowym teleskopem FORCAST, SOFIA zdołała wychwycić specyficzną dla cząsteczek wody, długość fali 6,1 mikrona.

Obserwacje i obliczenia

Chociaż istnieje kilka ciał niebieskich w Układzie Słonecznym, które mają hydrosferę, Ziemia jest jedynym znanym ciałem ze stabilnymi zbiornikami ciekłej wody na swojej powierzchni, z wodą oceaniczną pokrywającą 71 proc. jej powierzchni. Obecność wody na powierzchni Ziemi jest wynikiem jej ciśnienia atmosferycznego i stabilnej orbity w strefie zamieszkiwalnej Słońca.

Wykrycie wody lub znalezienie dowodów na istnienie wody w przeszłości z dystansu jest trudne. Teleskopy optyczne, które rejestrują światło widzialne i dostarczają wizualnych obrazów odległych ciał, dają nam jedynie pewne wskazówki co do jasności i wielkoskalowych kształtów i struktur dużych regionów. Jaśniejsze regiony, szczególnie w pobliżu północnego lub południowego bieguna planety lub księżyca, mogą wskazywać na odbicia zamrożonej wody. Jednakże, gdy jest wystarczająco zimno, lód z dwutlenku węgla również tworzy lśniąca powierzchnię. Zatem same teleskopy optyczne nie mogą potwierdzić obecności wody.

Uczeni przez pomiar odbicia światła od powierzchni mogą pośrednio sprawdzać, czy chodzi o wodę. Różne materiały, w tym woda, absorbują i odbijają różne długości fal świetlnych. Względna intensywność odbicia różnych długości fal jest łącznie nazywana widmem. Charakterystyka widmowa jest znaną od lat i stosowaną często techniką wykrywania składu chemicznego odległych obiektów. Główne metody stosowane obecnie do potwierdzenia tego faktu to spektroskopia absorpcyjna i geochemia. Techniki te okazały się skuteczne do wykrywania atmosferycznej pary wodnej i lodu. Jednak przy użyciu obecnych metod spektroskopii astronomicznej znacznie trudniej jest wykryć wodę w stanie ciekłym na planetach skalistych, zwłaszcza w przypadku wody podpowierzchniowej. Z tego powodu astronomowie, astrobiolodzy i planetolodzy wspomagają się teorią strefy zamieszkiwalnej, danymi o grawitacji i pływach, modelami zróżnicowania wewnątrzplanetarnego i radiometrii do określenia możliwości występowania wody w stanie ciekłym.

Ciekła woda ma odmienną sygnaturę spektroskopii absorpcyjnej w porównaniu z innymi stanami wody z powodu stanu jej wiązań wodorowych. Jednak jej wykrycie może być utrudnione z powodu odległości i przez grube atmosfery na ogromnych odległościach kosmicznych przy użyciu obecnej technologii. Dane ze spektroskopii, choć same w sobie nie potwierdzają ostatecznie i definitywnie obecności wody w stanie ciekłym, to jednak w połączeniu z innymi

obserwacjami na ich podstawie można wnioskować o takiej możliwości. Na przykład, gęstość egzoplanety GJ 1214 b sugerowałaby, że dużą część jej masy stanowi woda, a późniejsze wykrycie przez teleskop Hubble'a obecności pary wodnej silnie sugeruje, że mogą być tam obecne egzotyczne materiały, takie jak „gorący lód” lub „woda w stanie nadciekłym”. Z drugiej strony same inne dane bez spektroskopii też niczego nie przesądzą. np. sezonowe przepływy na nasłonecznionych marsjańskich zboczach, choć silnie sugerują obecność wody w stanie ciekłym, nie wykazały jej na razie w analizie spektroskopowej.

Naukowcy wykrywają ciekłą wodę także za pomocą sygnałów radiowych. Instrument RADAR (radio detection and ranging) sondy Cassini został wykorzystany do wykrycia istnienia warstwy ciekłej wody i amoniaku pod powierzchnią Tytana, co było zgodne z obliczeniami gęstości księżyca. Dane z radaru penetrującego grunt i przenikalności dielektrycznej z instrumentu MARSIS na pokładzie sondy Mars Express (3) wskazują na istnienie 20-kilometrowej stabilnej warstwy ciekłej wody w rejonie Planum Australe na Marsie.

W przypadku jowiszowych księżyców Ganimedesa i Europy, istnienie oceanu podlodowego wywnioskowano z pomiarów pola magnetycznego Jowisza. Naukowcy wykorzystali pomiary grawitacyjne aparatury sondy Cassini do potwierdzenia istnienia oceanu wodnego pod skorupą Enceladusa. Modele pływowe były wykorzystywane do analizowania innych księżyców Układu Słonecznego. Według co najmniej jednego badania grawitacyjnego na danych z Cassini, również księżyc Dione posiada ocean 100 kilometrów pod powierzchnią.

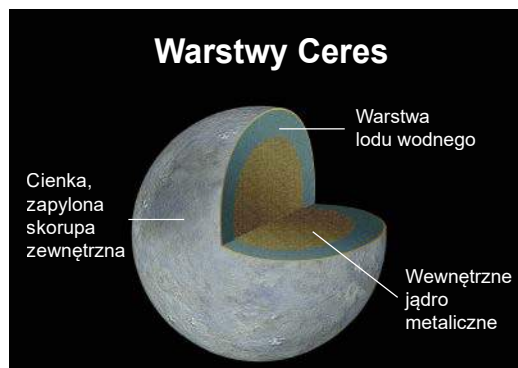
Do obliczania gęstości do określenia składu planet i prawdopodobieństwa występowania wody w stanie ciekłym można wykorzystać również obliczenia gęstości, choć metoda ta nie jest bardzo dokładna, gdyż kombinacja wielu związków i stanów może dawać podobne wyniki. Na obecność podpowierzchniowej warstwy oceanu wskazują obliczenia gęstości księżyców Saturna, Tytana i Enceladusa, wzmacniając wskazówki z innych rodzajów badań.

Takie obliczenia przeprowadza się również dla obiektów poza Układem Słonecznym. Wstępna analiza gęstości egzoplanety 55 Cancri e wskazywała, że składa się ona w 30 proc. z płynu nadkrytycznego, który według Diany Valencia z Massachusetts Institute of Technology może mieć postać słonej wody w stanie nadkrytycznym, jednak dalsze analizy jej tranzytu nie wykryły śladów ani wody, ani wodoru. Mała gęstość GJ 1214 b, drugiej egzoplanety

(po CoRoT-7b), której masa i promień były mniejsze niż masa i promień olbrzymich planet Układu Słonecznego, wskazywała, że jest to prawdopodobnie mieszanina skał i wody, a dalsze obserwacje przy użyciu teleskopu Hubble'a wydają się potwierdzać, że dużą część jej masy stanowi woda, czyli byłby to duży wodny świat.

Kolejnym narzędziem wykorzystywanym przez uczonych szukających wody są modele retencji ciepła i ogrzewania przez rozpad promieniotwórczy. Sugerują one, że mniejsze ciała Układu Słonecznego, Rhea, Titania, Oberon, Tryton, Pluton, Eris, Sedna i Orcus mogą mieć oceany pod litymi lodowymi skorupami, o grubości około 100 km. Szczególnie interesujący w tych przypadkach jest fakt, że modele wskazują, iż warstwy ciekłe są w bezpośrednim kontakcie ze skalistym jądrem, co pozwala na efektywne mieszanie minerałów i soli w wodzie. Jest to więc inny model niż hipotezy dotyczące większych lodowych satelitów, Ganimedesa, Callisto czy Tytana, gdzie warstwy lodu pod wysokim ciśnieniem miałyby znajdować się pod warstwą ciekłej wody. Modele rozpadu radioaktywnego sugerują też, że MOA-2007-BLG-192Lb, mała egzoplaneta krążąca wokół małej gwiazdy, może być tak ciepła jak Ziemia i całkowicie pokryta bardzo głębokim oceanem.

Są hipotezy, że planeta karłowata Ceres (4), największy obiekt w pasie asteroid mógłby mieć co najmniej „wilgotne wnętrze”, na co wskazuje wykryta tam para wodna, być może efekt sublimacji lodu powierzchniowego. Na Ceres znajduje się również góra zwana Ahuna Mons, która jest uważana za kopułę kriowulkaniczną ułatwiającą ruch kriowulkanicznej magmy o wysokiej lepkości, składającej się z lodu wodnego zmiękzonego przez zawartość soli. Uważa się, że globalna warstwa ciekłej wody o grubości wystarczającej do oddzielenia skorupy od płaszcza jest obecna na Tytanie, Europie oraz, z mniejszą



4. Wyobrażenie na temat warstw Ceres



5. Wizja starożytnego Marsa z ciekłą wodą
na powierzchni

pewnością, na Callisto, Ganimesdesie i Trytonie. Inne lodowe księżyce mogą również posiadać wewnętrzne oceany lub kiedyś posiadały wewnętrzne oceany, które teraz zamarły.

Oczywiście do wykrywania wody w sposób najbardziej bezsporny posłużyć możemy się wszelkiego rodzaju sondami, próbnikami, lądownikami i łazikami, które po prostu mogą zbierać próbki z powierzchni planety i analizować je pod kątem składu chemicznego. Statek kosmiczny może również zebrać próbkę i zwrócić ją na Ziemię w celu przeprowadzenia bardziej szczegółowej analizy. Jednak takie misje są bardzo kosztowne i trudne do przeprowadzenia ze względu na konieczność wylądowania na powierzchni ciała pozaziemskiego, ucieczki przed jego grawitacją i powrotu na Ziemię. W przypadku ciał poza Układem Słonecznym to w ogóle na razie niewykonalne.

Układ peten wody

Według oszacowań na podstawie wyżej wymienionych metod badań i analizy opublikowanych pod koniec 2015 roku, wody w stanie ciekłym w Układzie Słonecznym poza Ziemią jest objętościowo 25–50 razy więcej niż wody ziemskiej (1,3 mld km sześciennych). Gdzie ona

jest, a raczej może być, według szacunków, poszlak, pośrednich obserwacji i obliczeń?

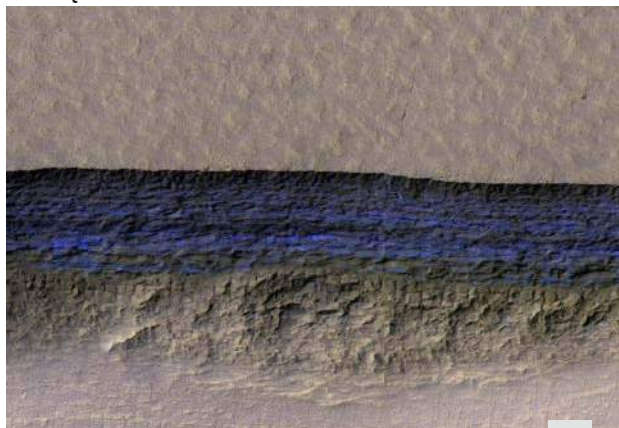
Mars był jednym z pierwszych ciał pozaziemskich, które przyciągnęły ludzką uwagę ze względu na wiarę w ciekłą wodę płynącą po powierzchni i towarzyszące jej życie. Obrazy Marsa z teleskopów

pokazały zmieniającą się jasność na planecie w czasie, co od ok. dwóch wieków wzmacniało wiarę w to, że na Marsie jest mniej więcej tak jak na Ziemi.

Pierwsze sondy kosmiczne w latach 70., Mariner 9, Viking 1, Viking 2 i inne, przyniosły obrazy miejsc, które wyglądały jak suche koryta rzek i kaniony podobne do tych znanych z Ziemi. Naukowcy po dekadach badań wierzą obecnie, że te cechy marsjańskiej powierzchni wskazują na wodną historię Marsa sięgającą miliardów lat w przeszłość (5). Ponadto od początków eksploracji na Marsa wysłano wiele innych satelitów, lądowników i łazików z różnymi instrumentami, które wykryły duże ilości lodu wodnego tuż pod powierzchnią całej planety. Łazik Curiosity wykrył na marsjańskiej powierzchni minerały ilaste, które tworzą się w obecności słodkiej wody, dając kolejne dowody na istnienie w przeszłości ciekłej wody na Marsie. W 2013 roku NASA, na podstawie danych z łazika Curiosity badającego Aeolis Palus w pobliżu Mount Sharp w kraterze Gale, poinformowała, że na Marsie istniało duże słodkowodne jezioro (które mogło być gościnnym środowiskiem dla życia mikrobiologicznego).

W 2015 roku dane sondy orbitalnej Mars Reconnaissance Orbiter zasugerowały, że występujący dość często na marsjańskiej powierzchni wzór ciemnych smug to efekt działania płynnej solanki (wody zmieszanej z solą) tuż pod powierzchnią podczas cieplejszej pory roku na Marsie, która przedostałaby się na powierzchnię na krótki czas. Jednak, jak wspomniano, nie potwierdziły tego badania spektrograficzne, a dodatkowe dane z 2017 roku podały tę hipotezę w wątpliwość, sugerując zamiast tego, że te ciemne smugi są głównie spowodowane ziarnami piasku i pyłu spadającymi kaskadowo z bardzo stromych zboczy na nierównej planecie.

6. Warstwy lodu widoczne w przekroju marsjańskiego gruntu na zdjęciu wykonanym przez satelitę MRO





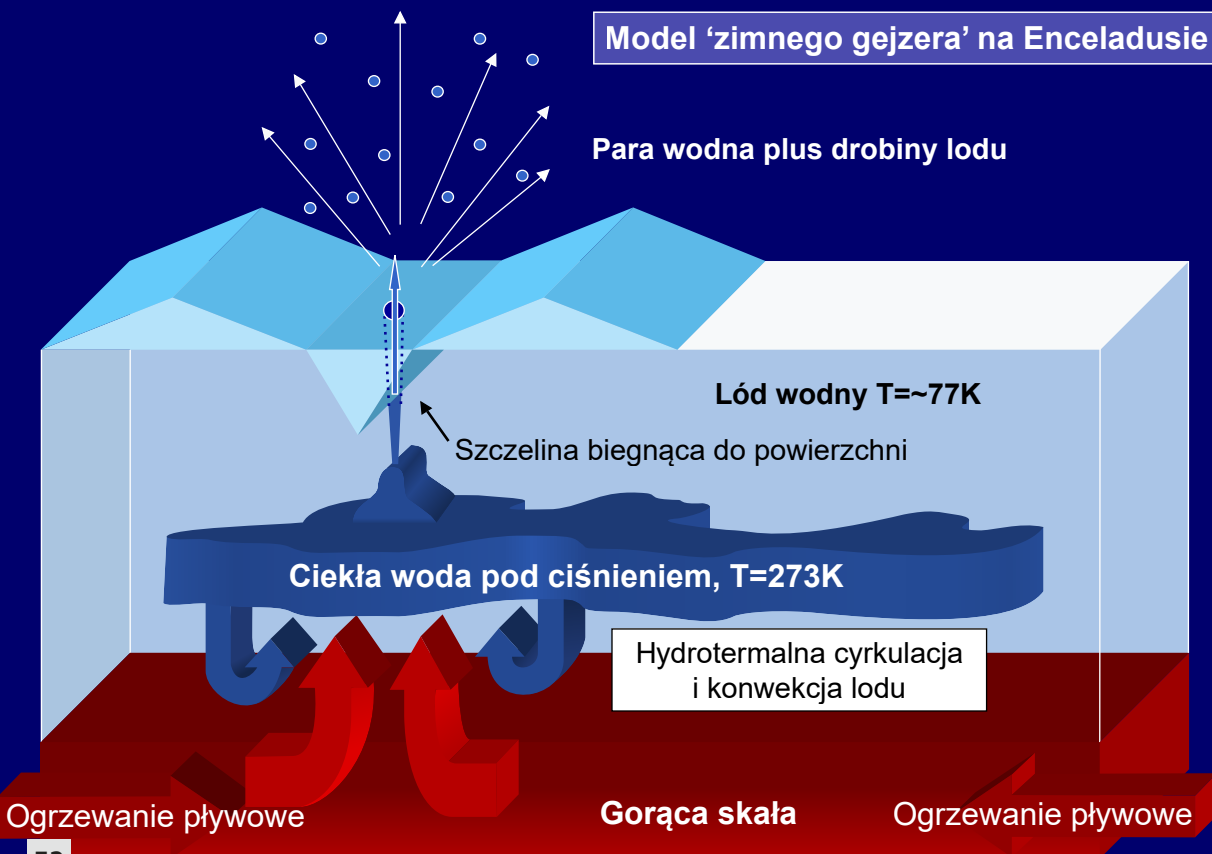
Hipoteza dawnego oceanu na Marsie sugeruje, że prawie jedna trzecia powierzchni Marsa była kiedyś pokryta wodą, a występująca na planecie obecnie woda ma inne postacie (duża jej część znajduje się w czapach lodowych). Przekrój podziemnego lodu na Marsie został ukazany na obrazie satelitarnym MRO jednego ze stromych zboczy, na którym widać jasnoniebieskie warstwy (6). Skarpa opada około 128 metrów z poziomu gruntu w górnej trzeciej części obrazu. Woda na Marsie występuje obecnie prawie wyłącznie w postaci lodu, z niewielką ilością obecną w atmosferze w postaci pary. Woda w stanie ciekłym może występować przejściowo na powierzchni Marsa, ale tylko w określonych warunkach. Nie istnieją duże zbiorniki wody w stanie ciekłym, ponieważ ciśnienie atmosferyczne na powierzchni wynosi średnio tylko 600 paskali – około 0,6 proc. średniego ciśnienia na poziomie morza na Ziemi – a także dlatego, że średnia globalna temperatura jest zbyt niska – 210 K (–63°C), co prowadzi do szybkiego parowania lub zamarzania.

Oszacowania wskazują, że w rejonach podbiegunowych warstwy powierzchniowe składają się

prawie całkowicie z wody pokrytej cienką warstwą drobnego materiału. Potwierdzają to obserwacje instrumentu MARSIS na orbiterze Mars Explorer. Wody tej w południowym rejonie biegunowym ma być wedle obliczeń $1,6 \times 10^6$ km³. Tyle na powierzchni. Ponadto w 2018 roku naukowcy wykorzystujący dane radarowe z orbitera Mars Express Europejskiej Agencji Kosmicznej zidentyfikowali solankowe jezioro ok. 1,6 pod południową pokrywą lodową Marsa. Wysoki poziom soli i niskie temperatury utrudniłyby przetrwanie większości znanych form życia. Jednak nawet na Ziemi odkryliśmy organizmy, które rozwijają się w słonych, choć nie tak zimnych, warunkach.

Jak była o tym wyżej mowa, dane zebrane z obserwacji i badań sugerują, że Europa, księżyc Jowisza, i Enceladus, księżyc Saturna, mają ogromne oceany pod warstwami lodu na powierzchni. Pod powierzchnią tych księżyców materia pozostaje płynna dzięki siłom grawitacyjnym pobliskiej dużej planety i otaczających ją księżyców, które na przemian ściskają i rozciągają podpowierzchniowe oceany, co generuje ciepło poprzez tarcie. Sonda Cassini badająca

7. Wizualizacja hipotetycznego kriowulkanizmu połączonego w oceanem wodnym, który może występować na Enceladusie i podobnych ciałach kosmicznych



Saturna i jego księżyce wykryła gejzery wody wyrzucanej z Enceladusa, wraz innymi substancjami chemicznymi, niewielkimi ilościami soli, azotu, dwutlenku węgla i lotnych węglowodorów które mogą sprzyjać ewentualnemu życiu (7). Kosmiczny Teleskop Hubble'a z daleka i sonda Galileo eksplorująca system Jowisza wykryły podobne wyrzuty pary wodnej z powierzchni Europy. Dane sugerują, że inne księżyce Jowisza i Saturna, w tym duże ciała, Ganimedes i Tytan, również mogą mieć podpowierzchniowe oceany. Taki ocean, jak się podejrzewa, mogłaby mieć także planeta karłowata Pluton.

Szacuje się, że zewnętrzna skorupa stałego lodu na Europie ma grubość około 10–30 km, w tym plastyczną warstwę „ciepłego lodu”, co może oznaczać, że ciekły ocean pod spodem może mieć głębokość około 100 km. Prowadzi to do tego, że objętość oceanów Europy wynosi $3 \times 10^{18} \text{ m}^3$, nieco ponad dwa razy więcej niż objętość oceanów Ziemi. Itsnieją co najmniej dwa modele naukowe hipotetycznego podlodowego oceanu na Europie (8).

Z najnowszych odkryć warto wspomnieć o informacji z lata 2021 r. o tym, że teleskop Hubble'a znalazł pierwsze dowody na obecność pary wodnej na księżycu Jowisza, Ganimedese. Para wodna tworzy się, gdy lód z powierzchni księżycy sublimuje, czyli zamienia się ze stanu stałego w gaz. Domniemany wewnętrzny ocean Ganimedesa znajduje się około 160 km pod skorupą, dlatego też odkryta ostatnio para wodna nie jest wynikiem parowania tego oceanu.

W ramach programu obserwacyjnego mającego na celu wsparcie misji NASA Juno w 2018 roku, Lorenz Roth z KTH w Sztokholmie w Szwecji połączył dane z dwóch instrumentów: spektrografu Hubble'a z 2018 roku oraz archiwalne obrazy z lat 1998–2010. Wbrew pierwotnym interpretacjom danych z 1998 roku, odkrył, że w atmosferze Ganimedesa prawie nie było atomowego tlenu. Oznacza to, że musi istnieć inne wyjaśnienie dla widocznych różnic w obrazach zorzy w zakresie ultrafioletowym. „Dostrzegalne różnice w obrazach UV są bezpośrednio skorelowane z miejscami, gdzie można by się spodziewać wody w atmosferze księżycy. Para wodna, którą zmierzylismy teraz, pochodzi z sublimacji lodu spowodowanej termiczną ucieczką pary wodnej z ciepłych obszarów lodowych”. To odkrycie dodaje przewidywania do nadchodzącej misji ESA (Europejskiej Agencji Kosmicznej), JUICE, co jest skrótem od Jupiter ICy moons Explorer. Planowana do wystartowania w 2022 roku i dotarcia do Jowisza w 2029 roku, spędzi co najmniej trzy lata na szczegółowych obserwacjach Jowisza i trzech jego

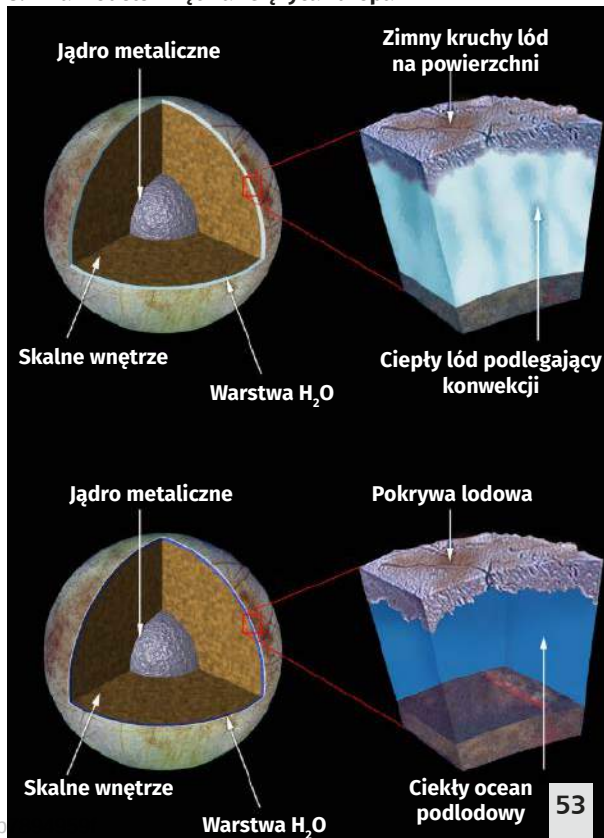
największych księżyców, ze szczególnym uwzględnieniem Ganimedesa jako ciała planetarnego i potencjalnego siedliska.

Obecnie wodę w postaci lodu lub pary wodnej podejrzewa się w miejscach, które wcześniej wolne były od takich podejrzeń. np. ze względu na bliskość Słońca Merkury był uważany za planetę całkowicie „suchą”. Analizy z sondy MESSENGER wskazały na pokłady lodu wodnego w kraterach na biegunie północnym.

Woda jest też częścią składową wielu innych mniejszych obiektów Układu Słonecznego, np. komet zawierających duże ilości lodu wodnego. Niedawno badania pyłu z komety Wild-2 wskazały na obecność ciekłej wody wewnątrz komety w pewnym momencie w przeszłości. Nie jest jeszcze jasne, jakie źródło ciepła (poza Słońcem) mogło spowodować stopienie części lodu wodnego komety.

10 grudnia 2014 r. naukowcy poinformowali, że skład pary wodnej z komety Czuriurow-Gerasimenko, określony przez sondę kosmiczną Rosetta, znacząco różni się od tego, który można znaleźć na Ziemi. Oznacza to, że stosunek deuteru do wodoru w wodzie z komety był trzykrotnie większy niż w przypadku wody ziemskiej. To czyni mniej prawdopodobnymi hipotezy, że woda znaleziona na Ziemi pochodzi głównie z komet.

8. Dwa modele wnętrza księżycy Europa





Wodę i to w ciekawych postaciach znajduje się również na asteroidach. Pierwszą, na której znaleziono wodę, była asteroida 24 Themis. Odkryta na niej również ciecz pod ciśnieniem mniejszym niż atmosferyczne, rozpuszczoną w minerale przez promieniowanie jonizujące. Obecność wody ciekłej stwierdzono na dużej asteroidzie 4 Westa, zaś energia, którą ją topi, pochodzi z uderzeń innych ciał kosmicznych.

Woda w strefach zamieszkiwalnych

Celem obecnych poszukiwań poza Układem Słonecznym jest znalezienie planet o rozmiarach Ziemi w strefie zamieszkiwalnej ich układów planetarnych (zwanej też czasem strefą Żłotowłosej). Kolejnym krokiem jest analiza znalezionych w tej strefie planet. Co udało nam się do tej pory ustalić, jeśli chodzi o ok. 4,5 tysiąca znalezionych i potwierdzonych egzoswiatów?

Myśli kierują się ku najbliższemu układowi gwiazdnemu i odkrytej tam intrygującej egzoplanecie. Jednak mimo że Proxima Centauri b znajduje się w strefie zamieszkiwalnej, możliwość zamieszkania planety została zakwestionowana z powodu kilku czynników. Przede wszystkim faktu, że planeta ta jest zwrócona w stronę gwiazdy zawsze tą samą stroną. Woda w stanie ciekłym może być obecna tylko w najbardziej nasłonecznionych regionach powierzchni planety, w basenach znajdujących się na półkuli planety zwróconej w stronę gwiazdy lub, jeśli planeta jest w rotacji rezonansowej, dwukierunkowo w pasie równikowym. Śladów wody jednak jak do tej pory nie wykryto. Planeta, jako stosunkowo bliska, prawdopodobnie będzie w przyszłości obiektem intensywnych badań.

Bardziej obiecujące pod interesującym nas względem światy odkryto dalej. We wrześniu 2019 r. dwa zespoły naukowiec ogłosiły niezależnie od siebie odkrycie wody w atmosferze egzoplanety K2-18b. Planeta krąży w strefie zamieszkiwalnej swojej gwiazdy. „Jest to jedyna planeta poza Układem Słonecznym, która ma odpowiednią temperaturę i atmosferę, by utrzymać wodę”, twierdził Angelos Tsiaras, astronom z University College London i główny autor publikacji w „Nature Astronomy”. Tsiaras i jego koledzy sugerują, że para wodna może stanowić od jednej setnej procent do połowy atmosfery K2-18b. Ustalenie, ile wody (a także innych gazów, takich jak metan, dwutlenek węgla i amoniak) się tam znajduje, będzie wymagało dalszych obserwacji przy użyciu przyszłych urządzeń kosmicznych, takich jak teleskop Jamesa Webba, który ma zostać wystrzelony w kosmos w grudniu 2021 roku, Atmospheric Remote-Sensing

Infrared Exoplanet Large-Survey (ARIEL), teleskopu Europejskiej Agencji Kosmicznej oraz budowanej generacji megateleskopów naziemnych. K2-18b jest większa od Ziemi i prawie dziewięć razy bardziej masywna, odkryta przez teleskop Keplera w 2015 roku, krąży po 33-dniowej orbicie wokół chłodnego czerwonego karła w gwiazdozbiorze Lwa, oddalonego od nas o około 110 lat świetlnych. Wrócimy do tego obiektu jeszcze we fragmencie dotyczącym zupełnie nowej kategorii egzoplanet wodnych, zwanych „hyceańskimi”.

Gliese 667 Cc to egzoplaneta krążąca w strefie zamieszkiwalnej czerwonej gwiazdy karłowatej Gliese 667 C, należącej do układu potrójnego gwiazd Gliese 667, oddalanej o około 23,62 lat świetlnych w gwiazdozbiorze Skorpiona. W oparciu o obliczenia temperatury ciała czarnego, GJ 667 Cc powinna pochłaniać nieco więcej ogólnego promieniowania elektromagnetycznego niż Ziemia, czyniącego ją nieco cieplejszą niż nasza planeta i w konsekwencji plasując ją nieco bliżej wewnętrznej („gorącej”) krawędzi strefy zamieszkiwalnej niż Ziemię. W 2018 roku Planetary Habitability Laboratory (PHL) uznało Gliese 667 Cc za czwartą najbardziej podobną do Ziemi egzoplanetę znajdującą się w ostrożnie definiowanej strefie zamieszkiwalnej. Jej gwiazda-gospodarz jest czerwonym karłem, o masie około jednej trzeciej masy Słońca. Planeta ma prawdopodobnie jedną stronę stale zwróconą w kierunku gwiazdy, podczas gdy przeciwna strona jest spowita wieczną ciemnością. Jednak pomiędzy tymi dwoma intensywnymi obszarami znajduje się skrawek przestrzeni nadającej się do zamieszkania – zwany terminatorem, gdzie temperatury mogą być odpowiednie (około 0°C) do istnienia wody w stanie ciekłym. Są jednak wyniki badań wskazujące, że planeta ta podlega zbyt silnemu ogrzewaniu pływowemu, by można było myśleć o warunkach sprzyjających istnieniu stabilnej wody w stanie ciekłym.

Super-Ziemia, Kepler-62e, odkryta na orbicie w strefie zamieszkiwalnej Keplera-62, znajduje się około 990 lat świetlnych od Ziemi w gwiazdozbiorze Liry. Kepler-62e może być planetą skalistą lub pokrytą oceanem. Leży w wewnętrznej części strefy zamieszkiwalnej swojej gwiazdy. Kepler-62e okrąża swoją gwiazdę co 122 dni i ma w przybliżeniu o 60 proc. średnicę większą od Ziemi. Biorąc pod uwagę, że ilość energii otrzymywanej przez Keplera-62e z macierzystej gwiazdy jest o 20 proc. większa niż to, co Ziemia otrzymuje od Słońca, możliwe jest, że temperatura powierzchni Keplera-62e może wynosić ponad 350 K (77°C), co wystarczy do wywołania efektu cieplarnianego.



9. Układ planetarny gwiazdy TRAPPIST-1

W tym samym układzie krąży w większym oddaleniu, czyli w odległości 0,718 AU (107 400 000 km), Kepler-62f, jednak wciąż w strefie zamieszkiwalnej gwiazdy Kepler-62. Jej okres orbitalny wynosi około 267,3 dnia. Została wybrana jako jeden z celów do zbadania przez program Search for Extraterrestrial Intelligence (SETI). Biorąc pod uwagę wiek planety (7 ± 4 miliardy lat), napromieniowanie ($0,41 \pm 0,05$ razy większe od ziemskiego) i promień ($1,41 \pm 0,07$ razy większy od ziemskiego), za prawdopodobny uważa się skład skalny (krzemianowo-żelazowy) z dodatkiem prawdopodobnie znacznej ilości wody. Jeśli jej gęstość jest taka sama jak ziemską, jej masa byłaby 1,413 lub 2,80 razy większa od ziemskiej. Chociaż Kepler-62f może być planetą pokrytą oceanem, posiadającą skały i wodę na powierzchni, jest najdalej wysunięta od swojej gwiazdy, więc bez dodatkowej ilości dwutlenku węgla (CO_2), może to być planeta pokryta w całości lodem. Aby Kepler-62f mogła utrzymać klimat podobny do ziemskiego (ze średnią temperaturą około 284–290 K ($11\text{--}17^\circ\text{C}$; $52\text{--}62^\circ\text{F}$), w atmosferze planety musiałoby być obecne co najmniej 5 barów (4,9 atm) dwutlenku węgla. Niska aktywność gwiazdowa pomarańczowych karłów, takich jak Kepler-62, tworzy stosunkowo łagodne środowisko promieniowania dla planet krążących w ich strefach zamieszkiwalnych, zwiększając ich potencjalną możliwość zamieszkania.

Kepler-62f prawdopodobnie dominowałby na liście światów, na które powinniśmy kierować wzrok i najpotężniejsze instrumenty badawcze w poszukiwaniu ciekłej wody (i życia), gdyby kilka lat temu nie odkryto za pomocą Kosmicznego Teleskopu Spitzera wokół gwiazdy TRAPPIST-1 układu aż siedmiu niezwykle interesujących planet (9).

Zacznijmy od TRAPPIST-1e, która jest egzoplanetą o rozmiarach zbliżonych do Ziemi, orbitującą w strefie zamieszkiwalnej ultrachłodnego czerwonego karła TRAPPIST-1, około 40 lat świetlnych od Ziemi w gwiazdozbiornie Wodnika. Potwierdzono również,

że posiada zwartą atmosferę. W listopadzie 2018 r. naukowcy ustalili, że spośród siedmiu egzoplanet w układzie TRAPPIST-1e jest tą, która ma największe szanse na bycie planetą z oceanem, podobną do Ziemi i najbardziej wartą dalszych badań. Pomimo, że cechuje ją obrót synchroniczny, czyli jedna półkula stale zwrócona jest w stronę gwiazdy, podczas gdy druga nie, co może zmniejszyć możliwość zamieszkania planety, ma odpowiednią temperaturę, by na powierzchni mogła zbierać się woda w stanie ciekłym. Tym się różni od innych planet układu – TRAPPIST-1f, g, i h. Ponieważ jest to jedna z najbardziej obiecujących, potencjalnie nadających się do zamieszkania znanych egzoplanet, TRAPPIST-1e będzie jednym z pierwszych celów Teleskopu Kosmicznego Jamesa Webba, który ma przeprowadzić dokładniejsze analizy atmosfery planety w poszukiwaniu chemicznych oznak życia lub biosygnatur.

To, że TRAPPIST-1e jest na czele listy, nie znaczy, że inne planety układu są skreślone. Jej sąsiadka, TRAPPIST-1d, która orbituje na wewnętrznej krawędzi strefy zamieszkiwalnej, jest najmniej masywną planetą układu i prawdopodobnie posiada zwartą, ubogą w wodór atmosferę podobną do Wenus, Ziemi lub Marsa. Otrzymuje zaledwie 4,3 proc. więcej światła słonecznego niż Ziemia. Pod pewnymi względami ta egzoplaneta jest jedną z najbardziej podobnych do Ziemi. Znacznie mniejsza od Ziemi, co może wpływać na jej magnetosferę. Planeta może posiadać wodę w stanie ciekłym i atmosferycznym, nawet wielokrotnie więcej niż Ziemia. Gęsta atmosfera sprzyjać może przekazywaniu ciepła na planecie zawsze zwróconej do swojej gwiazdy jedną stroną. Ostatnie badania przeprowadzone przez Uniwersytet Waszyngtoński wykazały, że TRAPPIST-1d może jednak być podobną do Wenus, niegościnną planetą.

Kolejna z planet układu, TRAPPIST-1f, jest egzoplanetą prawdopodobnie skalistą z masywną wodno-parową otoczką gazową o bardzo wysokim ciśnieniu i temperaturze, orbitującą albo wewnątrz, albo nieco



10. Wizualizacja powierzchni planety typu hyceańskiego

poza strefą zamieszkiwalną swojej gwiazdy macierzystej. Ma promień mniej więcej taki sam jak Ziemia, na poziomie około, ale tylko około dwóch trzecich masy Ziemi. Jest więc mało prawdopodobne, że jest to planeta w pełni skalista, podobna do Ziemi. Symulacje sugerują, że planeta składa się z 20 proc. z wody, co jest znacznie wyższą wartością niż w przypadku Ziemi. Przy tak masywnej otoczce wodnej, ciśnienie i temperatura będą wystarczająco wysokie, aby utrzymać wodę w stanie gazowym, a woda w stanie ciekłym będzie istnieć tylko w chmurach w górnych warstwach atmosfery TRAPPIST-1f.

Kolejna planeta, TRAPPIST-1g, znajduje się w optymistycznie definiowanej strefie zamieszkiwalnej swojej gwiazdy. Jest planetą większą od Ziemi, ale mniej gęstą, co oznacza, że prawdopodobnie zawiera jakąś formę wody. TRAPPIST-1g może posiadać globalny ocean wodny lub wyjątkowo grubą atmosferę parową nad warstwą lodu w stanie nadkrytycznym.

Planety oceaniczne

Wiele razy powyżej padało określenie „świat wodny” lub „planeta pokryta oceanem”. W trakcie badań naukowcy doszli do wniosku, że interesujące z punktu widzenia poszukiwań życia są nie tylko ściśle podobne do Ziemi planety skaliste, ale również planety, gdzie jest woda w postaci ciekłej, półlodowej lub w różnych nadkrytycznych stanach i fazach.

Nazywa się to różnie – świat oceaniczny, wodny, akwaplaneta, a od niedawna „planeta hyceańska” (10). Określa się tak ciała niebieskie zawierające znaczną ilość wody jako hydrosferę na powierzchni lub w oceanie podpowierzchniowym. Termin ten jest również używany czasami dla ciał astronomicznych z oceanem składającym się z innej płynnej

substancji, takiej jak lawa (przypadek księżycy Io), amoniak (w mieszaninie z wodą, jak jest prawdopodobnie w przypadku wewnętrznego oceanu Tytana) lub węglowodory, jak na powierzchni Tytana. W czerwcu 2020 roku naukowcy z NASA poinformowali, że na podstawie badań modelowania matematycznego jest prawdopodobne, że egzoplanety z oceanami są powszechne w Drodze Mlecznej.

Najlepiej poznane wodne światy w Układzie Słonecznym to Callisto, Enceladus, Europa, Ganimedes i Tytan. Europa i Enceladus są uważane za jedne z najbardziej interesujących celów badań ze względu na ich stosunkowo cienką skorupę zewnętrzną i kriowulkanizm. Wiele innych ciał w Układzie Słonecznym kandyduje do miana światów z podpowierzchniowymi oceanami, w tym Ariel, Ceres, Dione, Eris, Mimas, Miranda, Oberon, Pluton i Tryton.

Pozaziemski ocean mógłby być tak głęboki i gęsty, że nawet w wysokich temperaturach ciśnienie zamieniłoby wodę w lód. Ogromne ciśnienie w niższych rejonach takich oceanów mogłoby prowadzić do formowania się płaszcza egzotycznych form lodu, takich jak lód V. Lód taki niekoniecznie byłby tak zimny jak lód konwencjonalny. Jeśli planeta znajduje się wystarczająco blisko swojej gwiazdy, aby woda osiągnęła temperaturę wrzenia, przejdzie ona w stan nadkrytyczny i nie będzie miała dobrze zdefiniowanej powierzchni. Nawet na chłodniejszych planetach, zdominowanych przez wodę, atmosfera może być znacznie grubsza niż ziemska i składać się w dużej mierze z pary wodnej, wytwarzając bardzo silny efekt cieplarniany. Jak wspominaliśmy wyżej istnienie wewnętrznych oceanów może wspierać również energia rozpadu radioaktywnego.

Astronomowie zidentyfikowali w 2021 roku nową klasę planet nadających się do zamieszkania, zwanych planetami „hyceańskimi”. Miałyby to być gorące, pokryte oceanami planety z atmosferami bogatymi w wodór. Wiele kandydatek na planety „hyceańskie”, zidentyfikowanych przez naukowców, jest większych i gorętszych niż Ziemia, ale wciąż ma cechy sprzyjające istnieniu dużych oceanów, które mogłyby wspierać życie mikrobiologiczne podobne do tego, które można znaleźć w niektórych środowiskach wodnych na Ziemi. Planety tego typu pozwalają również na znacznie szerszą strefę zamieszkania niż ściśle rozumiana strefa Żłotowłosej, w porównaniu z planetami podobnymi do Ziemi. Oznacza to, że mogą podtrzymywać życie, nawet jeśli leżą poza jej zakresem.

Zdecydowana większość z tysięcy odkrytych do tej pory egzoplanet to planety o rozmiarach pomiędzy Ziemią a Neptunem, często nazywane „super-Ziemią” lub „mini-Neptunami”. Według założeń mogą to być planety skaliste lub lodowe olbrzymy z atmosferami bogatymi w wodór lub coś pomiędzy. Większość mini-Neptunów ma rozmiar ponad 1,6 razy większy od Ziemi, czyli są mniejsze od Neptuna, ale zbyt duże, by mieć skaliste wnętrza jak Ziemia. Wczesniejsze badania takich planet wykazały, że ciśnienie i temperatura pod ich bogatymi w wodór atmosferami byłyby zbyt wysokie, aby mogło tam powstać życie.

Jednak najnowsze badania mini-Neptuna K2-18b przeprowadzone przez zespół Nikku Madhusudhana z uniwersytetu w Cambridge wykazały, że w pewnych warunkach planety tego typu mogą wspierać życie. Chociaż mogą być nawet 2,6 razy większe od Ziemi i mieć bogatą w wodór atmosferę temperaturze o nawet do 200 stopni Celsjusza, to jednak ich

warunki oceaniczne mogą być podobne do tych, które sprzyjają życiu mikrobiologicznemu w ziemskich oceanach. Kategoria ta przewiduje też podkategorie „ciemnych” światów hyceańskich, które mogą mieć warunki do zamieszkania tylko po energii od swoich gwiazd. Planety tej wielkości dominują w gronie egzoplanet, choć nie zostały zbadane tak szczegółowo jak super-Ziemia. Są prawdopodobnie jednak dość powszechne, co oznacza najbardziej obiecujące miejsce do poszukiwania życia w Galaktyce.

Zespół z Cambridge zidentyfikował sporą próbkę potencjalnych planet „hyceańskich”, wskazując je jako pierwsze kandydatki do szczegółowych badań za pomocą teleskopów nowej generacji, w tym oczywiście Kosmicznego Teleskopu Jamesa Webba (JWST). Wszystkie te planety krążą wokół czerwonych gwiazd karłowatych w odległości 35–150 lat świetlnych od Ziemi.

Woda jest substancją, jak się wydaje, całkiem zwykłą w kosmosie. A zarazem jest czymś niezwykłym, bo tylko w jej obecności powstaje jedyny rodzaj życia, jaki umiemy identyfikować. Potwierdzenie jej wszechobecności we Wszechświecie może oznaczać, że powszechne jest życie o biologii podobnej do naszego, ziemskiego. Ale nie musi. Gdyby tak się okazało, a zapewne szybko o tym się nie przekonamy, sugerowałoby to, że woda nie jest tak istotnym czynnikiem w powstawaniu i utrzymywaniu życia, jak tradycyjnie uważamy, a prawdziwa tajemnica tkwi w czym innym.

Zanim jednakże do tego dojdziemy, minie sporo czasu. Na razie przydałyby się jakiekolwiek twarde dowody na istnienia ciekłej wody gdziekolwiek poza Ziemią. Bo pomimo wszystkiego, co zostało wyżej napisane, takich niezbitych dowodów wcale jeszcze nie ma. ■

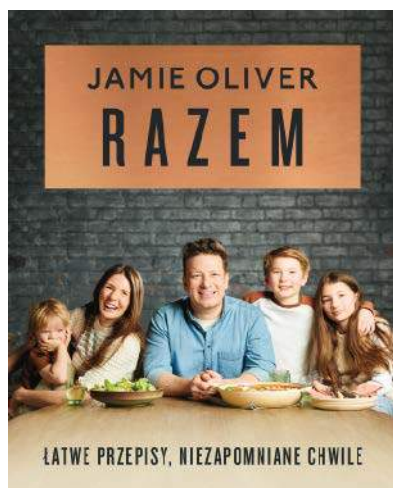
Miroslaw Usidus

Razem

Jamie Oliver

Wydawnictwo Insignis, cena: 79,99 zł

Pełen inspiracji, a przy tym praktyczny zbiór ponad 120 łatwych, lecz zachwycających przepisów opowiada o dodawaniu otuchy, świętowaniu, tworzeniu nowych wspomnień i przede wszystkim o dzieleniu się pysznymi potrawami przy jednym stole. Spotkania z bliskimi nigdy nie wydawały się tak ważne jak dzisiaj. Impreza taco, wyzerka na luzie, a może rodzinny piknik lub kolacja we dwoje? W każdym rozdziale „Razem” znajdziecie specjalnie przygotowany zestaw niesamowitych potraw, które w dużej mierze można przygotować wcześniej, a nie na ostatnią chwilę. Jamie Oliver postawił sobie za cel, byście nie tkwili samotnie w kuchni, lecz cieszyli się jedzeniem wraz ze swoimi gośćmi.



**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 fb7894959f

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. Richard Arkwright

CV: Richard Arkwright

Data i miejsce urodzenia: 23.12.1732
(zm 03.08.1791)

Adres zamieszkania: nie żyje

Obywatelstwo: brytyjskie

Stan cywilny: nie żyje

Majątek: 500 tys. funtów w 1791 roku, co może się równać nawet 5 mld funtów współczesnych

Kontakt: nie żyje

Edukacja: Brak formalnego wykształcenia

Doświadczenie zawodowe: fryzjer, perukarz, wynalazca, przedsiębiorca

Zainteresowania: kwestie społeczne

Jak tkał się kapitalizm – Richard Arkwright

Jeśli ktoś chciałby mieć wzorzec i najbardziej typowy przykład wczesnego, dynamicznego, kapitalisty to Richard Arkwright (1) świetnie by się do tego nadawał. Był zarazem prekursorem rozwiązań w dziedzinie organizacji pracy i życia pracowników, które przyjmowały się znacznie później.

Urodził się w Preston 23 grudnia 1732 roku jako trzynaste, najmłodsze, dziecko krawca Thomasa i jego żony Sarah. Rodzice nie byli w stanie zapewnić synowi edukacji. Czytać i pisać uczyła więc Richarda kuzynka Ellen. A gdy chłopiec podrośł, oddano go na naukę zawodu do fryzjera w pobliskim Bolton.

Richard okazał się zręcznym stylistą fryzur a także cyrulikiem. Z dużą wprawą rwał ludziom bolące zęby, więc zwykle miał kolejkę chętnych na mycie, strzyżenie i opiekę stomatologiczną. Choć w zawodzie wiodło mu się nieźle, to zarobki uznał za mało satysfakcjonujące i postanowił zająć się bardziej dochodową profesją – perukarstwem. A że do wyrobu tej modnej i cenionej wówczas przez mężczyzn ozdoby głowy niezbędne były włosy, to wyruszył w podróż po kraju, by skupować za zgromadzone oszczędności niewieście pukle i warkocze.

Wtedy też ujawnił się inny talent młodego Anglika, który miał zostać w nieodległej przyszłości wielkim wynalazcą, bogaczem i szlachcicem. Podróżując i skupując włosy, Richard poznaje wiele kobiet. Większość z nich to prądky i tkaczki. Richard chętnie słucha opowieści o trudach ich pracy, poznaje używane urządzenia i metodę farbowania tkanin, dzięki której płótna stają się wodoodporne. Fryzjer-cyrulik postanowił wykorzystać ten sposób do barwienia swoich peruk. Z sukcesem. Jego peruki rzeczywiście nie nasiąkały wodą, jak te oferowane przez konkurencję. Niestety, gdy w 1762 roku założył własną firmę w Bolton, moda na peruki zaczęła mijać.

Diabelski hałas epokowego wynalazku

Przedsiębiorczy Anglik musiał wymyślić inny biznes. I teraz bezcenna okazała się wiedza, jaką zdobył

na temat tkactwa. Oszczędności, które mu pozostały, postanowił zainwestować w projekt wydajnej maszyny przędzalniczej. Potrzebował jednak pomocy fachowca w dziedzinie mechaniki i ogólnie kwestii technicznych. Poznał zegarmistrza Johna Kaya, specjalistę od kół zębatach i wszelkich mechanizmów. Wybór nie był przypadkowy, bo Kay wcześniej współpracował z wynalazcą Thomasem Highsem nad podobnym projektem – maszyną włókienniczą, w której palce prądky zastępować miały obracające się wałki. Highs prace nad wynalazkiem ostatecznie zarzucił, co nie przeszkodziło mu potem wystąpić z roszczeniami wobec Kaya i Arkwrighta.

Zanim jednak powstała maszyna warta procesu sądowego, wspólnicy spędzili kilka lat nad przygotowaniem prototypu. W 1766 r. Kay podpisał kontrakt z Richardem Arkwrightem na 21 lat, z honorarium wynoszącym pół gwiney tygodniowo. Wynalazcy od początku wiedzieli, czego chcą – maszyny do przędzenia z napędem mechanicznym, wydajniejszy niż ludzkie mięśnie. Wiosną 1768 r. dwie kobiety zeznały lokalnym władzom, że z sąsiedniego budynku dochodzą łącznie diabelskie dźwięki...

Okazało się, że źródłem hałasu była maszyna włókiennicza (2) Arkwrighta i Kaya. Mechanizm działania urządzenia opierał się na układzie kół zębatach, które napędzały wrzeciona i wałki, rozciągające i skręcające włókna. Urządzenie wymagało napędu konnego i było najszybszym urządzeniem, jakimi wówczas dysponowało włókiennictwo. Przędzarka mogła obsługiwać jednocześnie aż 96 wątków, a wytwarzana nić była mocniejsza niż produkowane innymi metodami.

Arkwright potrzebował teraz wsparcia finansowego, by rozwinąć produkcję. Pierwsze kroki skierował do banku. Tam co prawda pożyczki nie dostał, ale przedstawiono go potentatom branży pończoszniczej. Przedsiębiorca Samuel Need oraz fabrykant i wynalazca Jedediah Strutt docenili wynalazek Arkwrighta i zostali jego partnerami biznesowymi. Dzięki pozyskanym środkom Arkwright mógł uzyskać w 1769 r. patent na swoją ramę wodną, jak nazwano nowe urządzenie.

Pierwszą fabrykę z innowacyjną maszyną włókienniczą uruchomiono w Nottingham. Szybko okazało się, że wydajniejszy niż napęd końskich mięśni jest napęd wodny, lecz lokalizacja fabryki nie gwarantowała do niej dostępu. Inwestorzy musieli

2. Rekonstrukcja maszyny Arkwrighta w Helmsore Mills Textile Museum



przenieść produkcję w inne miejsce. Wybór padł na niewielką miejscowość Cromford, położoną w dolinie rzeki Derwent (3).

W 1770 r. partnerzy zbudowali młyn wodny, który istnieje do dziś i jest wpisany na listę światowego dziedzictwa UNESCO. Kolejny rok zajęło wspólnikom zbudowanie fabryki i zainstalowanie ramy wodnej. Była to pierwsza nowoczesna fabryka włókiennicza w historii. Rama wodna Arkwrighta zyskała rozgłos, co przyniosło wynalazcy zyski, ale też pozwy sądowe roszczących sobie prawa do wynalazku.

Pojawiają się i inne problemy. W okolicy zaczęło brakować rąk do pracy. Aby ściągnąć robotników do fabryki, buduje dla nich pionierskie na owe czasy osiedle robotnicze. Obok domów powstają również szkoła dla dzieci i kościół. Początkowo fabryka zatrudnia ok. 1200 osób, liczba ogromna jak na ówczesne czasy, w których wciąż produkcja tkanin prowadzona była głównie w wiejskich chałupach, a wynajęci chałupnicy raz w tygodniu oddawali swój urobek i odbierali zapłatę. W fabryce Arkwrighta robotnicy byli zatrudnieni na stałe lub okresowe umowy, nie pracowali w trybie zleceń. Czas ich pracy odmierzali zegary. Produkcja szła nieprzerwanie w trzynastogodzinny system pracy, choć zdarzało się, że maszyny działały nawet 24 godziny na dobę.

Przy przędzy pracowały często całe rodziny. Wynalazca zatrudniał również dzieci, które w kolejnych latach stanowiły aż dwie trzecie pracowników. Dziś budzi to kontrowersje, ale wówczas nie było niczym szczególnym. Wprowadził program socjalny i oferował urlopy wypoczynkowe dla robotników. Mogli przez tydzień nie przychodzić do pracy pod warunkiem, że nie opuszczają miasteczka.

Patenty zabrane, ale majątek pozostał

Biznes rozwijał się tak dobrze, że Arkwright wkrótce spłacił swoich współników. Otworzył kolejne własne fabryki m.in. w Manchesterze, Matlock Bath i New Lanark. Otrzymywał także wynagrodzenie autorskie od innych przedsiębiorstw, które korzystały z licencjonowanej ramy wodnej. Udoskonalał nieustannie wynalazek, a w 1775 r. uzyskał tzw. wielki patent na swoją ramę wodną. Dwa lata później w Wirksworth użył silnika parowego Jamesa Watta do pompowania wody, niezbędnej do napędu młyńskiego koła wodnego.

Skrupulatnie strzegł szczegółów technicznych swojego wynalazku. Za ich ujawnienie miała grozić



3. Zakład Arkwrighta w Cromford w dolinie rzeki Derwent

kara śmierci. Mimo to coraz trudniej było mu utrzymać monopol. Przedsiębiorcy z branży w końcu zjednoczyli się, by wytoczyć mu proces o odebranie patentów. Wyciągnięte zostały argumenty najcięższego kalibru. Zeznania przeciwko Arkwrightowi złożyli wspomniany Highs oraz dawny współnik Kay i jego żona. Wszyscy zeznawali zgodnie, że potentat włókienniczy ukraść ich pomysły.

W rezultacie, po kilku latach procesów, w 1885 r. sąd wycofał wszystkie patenty Arkwrighta. Choć utracił ochronę prawną dla swoich rozwiązań, pozostał już potentatem w branży włókienniczej. Jak się ocenia, proporcjonalnie zgromadził fortunę większą niż Thomas Edison czy Alfred Nobel.

Jak na archetypicznego, stereotypowego kapitalistę-parweniusza przystało miał pewne kompleksy. Do końca życia systematycznie korzystał z lekcji poprawnej wymowy. Z okazji objęcia prestiżowej posady szeryfa hrabstwa Derbyshire wydał bankiet godny wizyty członka rodziny monarszej. Jako nowy szeryf przejechał przez miasto w ozdobnej karecie, w otoczeniu kilkudziesięciu gwardzystów i trębaczy. Wydarzenie stało się sensacją opisywaną w lokalnych gazetach.

W 1782 roku kupił za 8864 funtów od Thomasa Halleta Hodgesa posiadłość Willersley Hall w Cromford, by postawić tam zamek. W rezydencji nigdy nie zamieszkał. W 1791 roku prawie już gotowy zamek uległ zniszczeniom podczas pożaru. Gdy uporano się z odbudową latem 1792 roku, Richard Arkwright już nie żył. Zmarł 3 sierpnia w wieku 59 lat. Swoim dzieciom, córce i synowi, zostawił majątek o wartości pół miliona funtów oraz tytuł szlachecki. W 1786 roku za zasługi dla rozwoju przemysłu otrzymał tytuł szlachecki. ■

Mirosław Usidus



Test aplikacji: Zarządzanie hasłami



1Password

Aplikacja szyfruje dane przy użyciu 256-bitowego algorytmu AES, a przy stosowaniu ustawień domyślnych umożliwia przechowywanie bazy na serwerach producenta. Pierwsza publiczna wersja aplikacji została wydana w 2006 roku. Pozwala na kategoryzację danych uwierzytelniających w ramach folderów, oferuje także funkcję generatora haseł oraz możliwość sprawdzania używanych danych pod kątem ich obecności na listach niebezpiecznych haseł.

Do dyspozycji użytkownika pozostawiono rozbudowane opcje synchronizacji danych między urządzeniami (chmura, sieć Wi-Fi, foldery współdzielone) oraz funkcję sporządzania automatycznych kopii zapasowych. Domyślnie baza haseł jest zapisywana na serwerach producenta, przy czym istnieje możliwość zmiany miejsca synchronizacji internetowej (obsługiwane są serwisy Dropbox i iCloud) bądź jej całkowitego wyłączenia.

Dzięki usłudze Watchtower aplikacja informuje o serwisach internetowych, które padły ofiarą ataku Heartbleed, sugerując użytkownikom tych stron odnowienie dotychczasowych haseł. 1Password jest programem typu shareware, udostępniającym 30-dniową wersję testową. Po upływie bezpłatnego okresu próbnego z oprogramowania można korzystać wyłącznie za opłatą.

1Password	
Producent	AgileBits Inc.
Platforma	Android, iOS, Win., Linux
Oceny	
Możliwości	7,5/10
Łatwość obsługi	8,5/10
Ocena ogólna	8/10



Dashlane

Menedżer haseł i sejf umożliwiający zapisywanie i szyfrowanie poufnych danych, m.in. loginów i haseł, a także danych kart płatniczych. Aplikacja jest przeznaczona na komputery z systemami Microsoft Windows i macOS oraz urządzenia mobilne z Androidem i iOS. Program integruje się z przeglądarkami internetowymi i umożliwia automatyczne uzupełnianie formularzy.

Dashlane umożliwia zapisywanie danych w ramach internetowej przestrzeni w chmurze oraz ich synchronizację między różnymi urządzeniami. Do dyspozycji użytkownika postawiono także możliwość przechowywania danych wyłącznie na dysku lokalnym. Oferuje możliwości kategoryzacji różnych danych, zawiera wbudowany generator haseł oraz funkcję ostrzegania przed zagrożeniami internetowymi.

Aplikacja Dashlane jest darmowa do użytku na jednym urządzeniu. Wersja odpłatna oferuje synchronizowaną obsługę wielu urządzeń, a także opcję sporządzania kopii zapasowych oraz dostęp do haseł za pośrednictwem interfejsu internetowego. W wersji płatnej dostępne jest także nieograniczone udostępnianie haseł oraz uwierzytelnianie dwuetapowe U2F.

Dashlane Password Manager	
Producent	Dashlane
Platforma	Android, iOS, Windows
Oceny	
Możliwości	9/10
Łatwość obsługi	9/10
Ocena ogólna	9/10

Smartfony i ich systemy operacyjne, czyli słówko o platformach

Podobnie jak komputer, tak i smartfon, choćby nie wiadomo jak wspinały, to tylko kupka elektronicznego złomu, jeśli brak w nim oprogramowania. Podstawowym oprogramowaniem każdego urządzenia z procesorem, pamięcią i wyświetlaczem jest system operacyjny. To dopiero on decyduje, jakie możliwości ma dane urządzenie i jednocześnie wyznacza jego popularność, mierzoną liczbą dostępnych aplikacji – jako że aplikacje pisane są na określony system operacyjny, a nie „na sprzęt”. Przykładowo, dwa identyczne telefony tej samej firmy mogą być zupełnie różnymi funkcjonalnie urządzeniami, jeśli na jednym producent zainstaluje system Android, a na drugim system Symbian. Aplikacje na Androida nie będą działać na Symbianie i odwrotnie. Najpopularniejsze smartfonowe systemy operacyjne to:

- **iOS** – system firmy Apple (tej od komputerów Macintosh), instalowany w urządzeniach iPhone, iPod Touch, iPad;
- **Android** – system firmy Google, niektórzy twierdzą, że wkrótce podbić cały świat. Rzeczywiście, Android jest coraz częściej instalowany w smartfonach m.in. takich firm, jak Huawei, HTC, LG, Motorola, Samsung, Sony Ericsson, ZTE (a także, co oczywiste, w smartfonach firmy Google);
- **Symbian** – system operacyjny open source (czyli bezpłatny i z tzw. otwartym kodem), obecnie najczęściej spotykany w telefonach firmy Nokia.

Inne, mniej popularne systemy operacyjne dla telefonów komórkowych, to:

- **Bada** – system rozwijany przez firmę Samsung;
- **Windows Phone** – system firmy Microsoft, następca Windows Mobile, czyli po prostu Windows do urządzeń przenośnych;
- **BlackBerry** – system kanadyjskiej firmy Research in Motion, przeznaczony przede wszystkim do zastosowań biznesowych, instalowany w produkowanych przez nią smartfonach z charakterystyczną, pełną klawiaturą QWERTY. Także w niektórych telefonach innych firm (HTC, Motorola, Nokia, Samsung, Sony Ericsson).



LastPass

Ten menedżer haseł umożliwia przechowywanie zaszyfrowanych haseł w ramach konta internetowego. Platforma jest wyposażona w interfejs internetowy, a także oferuje rozszerzenia dla szeregu przeglądarek internetowych oraz aplikacji mobilne. Oferuje wiele ciekawych możliwości m.in. z rozbudowanymi opcjami uwierzytelniania dwuskładnikowego.

Program współpracuje z najpopularniejszymi przeglądarkami internetowymi i umożliwia automatyczne wprowadzanie danych dostępu. Pozwala na synchronizację danych między różnymi przeglądarkami i urządzeniami. Usługa oferuje również sejf na wszelkie dodatkowe informacje i zapiski oraz funkcję generatora haseł.

Podstawowa wersja usługi jest dostępna bezpłatnie. LastPass Premium umożliwia m.in. zaszyfrowane przechowywanie plików, daje dostęp do dodatkowych opcji uwierzytelniania dwuskładnikowego (2FA) oraz zapewnia możliwość awaryjnego dostępu do kolekcji haseł. Płatna wersja dodaje także nieograniczoną synchronizację między wszystkimi urządzeniami, wsparcie dla fizycznych kluczy uwierzytelniania dwuskładnikowego, 1 GB pamięci masowej online, monitorowanie kont dark-web i dostęp do pomocy technicznej.

LastPass	
Producent	LogMeIn, Inc.
Platforma	Windows, Mac, iOS, Android, Linux, Chrome OS
Oceny	
Możliwości	9,5/10
Łatwość obsługi	8,5/10
Ocena ogólna	9/10



Blur

Oprócz typowych funkcji menedżera haseł, Blur pozwala na generowanie jednorazowych numerów kart kredytowych podczas sieciowych transakcji. Za pomocą tej appki można także stworzyć wirtualny numer telefoniczny, by nie trzeba było podawać w sklepie naszego prawdziwego. Program ten jest właściwie kompleksowym rozwiązaniem chroniącym prywatność.

Blur tworzy zaszyfrowane hasła, zabezpiecza dane użytkownika podczas dokonywania płatności online oraz zapisuje i sortuje utworzone hasła. Ułatwia również użytkownikowi śledzenie i zabezpieczanie wszystkich danych osobowych w Internecie. Wysoko oceniany w niezależnych testach.

Za to, że oprócz zwykłego menedżera haseł mamy pakiet chroniący prywatność, trzeba zapłacić więcej niż za inne programy zarządzające hasłami. Blur ma zresztą kilka poziomów opłat. Oprócz zwykłego premium są też opłaty za generowanie numerów kart płatniczych. Dopiero w nieograniczonym planie premium użytkownik „ma wszystko”. Jest też wersja darmowa, ale jej oferta jest dość uboga.

Clear Scan	
Producent	Indy Mobile App
Platforma	Android, iOS, Huawei
Oceny	
Możliwości	8,5/10
Łatwość obsługi	7,5/10
Ocena ogólna	8/10



Keeper

Keeper pozwala na przechowywanie wrażliwych danych takich jak login i hasła do różnego rodzaju usług w bezpiecznym miejscu. Aplikacja szyfruje dane, przez co zapewnia zachowanie bezpieczeństwa nawet na wypadek kradzieży urządzenia. Oferuje również dodatkowe funkcje, takie jak synchronizacja danych w chmurze.

Przed skorzystaniem z aplikacji użytkownik musi założyć konto, podając swój adres e-mail, bądź też załogować się do już istniejącego konta. Kolejnym krokiem jest podanie hasła głównego aplikacji, które odpowiada za szyfrowanie danych. Oprócz loginu i hasła użytkownik może wprowadzić własny komentarz, a także załączyć plik lub obrazek. Dodatkowa funkcja to tworzenie dowolnych pól o dowolnej nazwie i tekstowej zawartości. Keeper ma również zintegrowany generator haseł, umożliwiający szybkie stworzenie bezpiecznego ciągu znaków.

Program oferuje dodatkowe możliwości ochrony – kopię zapasową danych w chmurze i synchronizację między wieloma urządzeniami. Pozwala ponadto na konfigurację weryfikacji dwuetapowej za pomocą wiadomości SMS, połączenia telefonicznego lub generatora kodów. Aplikacja jest udostępniona za darmo na 30 dni. Po tym okresie konieczne jest wykupienie abonamentu na jedno lub nieograniczoną liczbę urządzeń.

Microsoft Lens	
Producent	Microsoft Corp.
Platforma	Android, iOS
Oceny	
Możliwości	8,5/10
Łatwość obsługi	9,5/10
Ocena ogólna	9/10

Na cześć złośliwych duchów (1)

Nazwy pierwiastków pochodzą z różnych źródeł. Odkrywcy często nadawali im imiona kojarzące się z właściwościami (np. kolorem czy też zapachem) lub minerałami, z których je wydzielono. Nazwy wielu z nich wzięto od nazwy kraju, miasta, rzeki czy kontynentu. W imionach kilku uwieczniono też naukowców oraz postacie mitologiczne. Natomiast w nazwach kobaltu i niklu kryją się imiona złośliwych duchów.

Oba pierwiastki to towarzysze żelaza, razem z nim znajdujące się w jednej triadzie (patrz: Trójkami). Podobieństwo kobaltu i niklu do metalu najbardziej rozpowszechnionego na Ziemi powoduje, że poznając ich właściwości, nadal pozostaniesz w tematyce „żelaznym kręgu” (tegoroczne artykuły działu chemicznego traktowały o żelazie, korozji oraz obrońcy stali – cynku).

Złośliwe duchy

Podobnie jak w przypadku innych pierwiastków, związki kobaltu znano od znacznie dawniejszych czasów niż sam metal. Już w starożytności cenione było niebieskie szkło, któremu kolor nadawał pigment otrzymywany z rudy kobaltu (1). Jednak umiejętność wytwarzania barwnika o pięknym, nasyconym kolorze zanikła w Europie, jak wiele innych wynalazków i osiągnięć technologicznych, wraz z upadkiem Cesarstwa Rzymskiego (w Chinach wytwarzano w tym czasie porcelanę zdobioną niebieskim barwnikiem kobaltowym). Europejczycy ponownie poznali związki kobaltu w późnym średniowieczu, gdy Wenecjanie zaczęli stosować je do barwienia szkła. Odkrycie utrzymywano w ścisłej tajemnicy, co spowodowane było chęcią zachowania monopolu na wytwarzanie drogiego niebieskiego szkła. Jednak szpiegostwo przemysłowe nie jest wynalazkiem naszych czasów i z biegiem lat cała Europa poznała tajniki wyrobu kobaltowego barwnika – **smaltu**.

Gdzie jednak kryją się duchy? Otóż do produkcji pigmentu używano

szczególnej rudy. W Górach Kruszcowych (Rudawy leżące obecnie na pograniczu Czech i Niemiec) górnictwo kwitło od bardzo dawna. Saksońscy górnicy czasem znajdowali minerał podobny do rudy srebra, ale nie udawało się z niego otrzymać cennego metalu. Winne były oczywiście złośliwe górskie duchy – koboldy – zamieniające rudę srebra w bezwartościowy kamień. Stąd też nazwa minerału (**kobold** lub **kobolt**), którego reputację poprawiło użycie jako surowca do produkcji niebieskiego pigmentu (ale to jednak nie to co srebro).

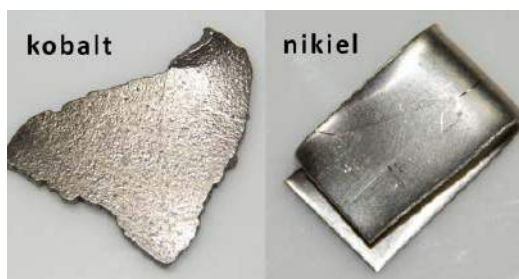
Nikiel ma podobną historię. Zanim poznano go na Starym Kontynencie, już w starożytności wchodził w skład stopów używanych głównie do bicia monet (wtedy również nie znano niklu, a stopy wytwarzano z rudy zawierającej jego domieszkę). Nazwa pierwiastka pochodzi zaś od gwarowego słowa *nicken*, czyli m.in. oszukiwać (samo *nickel* oznaczało natomiast diabła). Rudawscy górnicy często znajdowali czerwonej barwy minerał, z którego mieli nadzieję otrzymać cenną miedź. Jednak, jak w przypadku kobaltu, wytop się nie udawał, za co winą obarczano siły nieczyste. Z tego też powodu rudę nazwano **kupfernicken**, co oznacza oszukańczą miedź lub miedzianego diabła.

Narodziny

Choć rudy kobaltu i niklu znane były już w średniowieczu, zawarte w nich metale świat ujrzał dopiero w wieku XVIII, a odkrywcami stali się Szwedzi (2). W roku 1735 **Georg Brandt** dokładnie zbadał minerał używany do produkcji niebieskiego pigmentu



1. Butelka ze szkła kobaltowego (<https://images-of-elements.com>)



2. Kobalt i nikiel – błyszczące srebrzyste metale
(<https://images-of-elements.com>)

i wydzielił z niego metal (dość zanieczyszczony), któremu nadał nazwę **kobalt**. W końcu stulecia otrzymano już czysty metal, a w dalszych latach kobalt zastosowano jako składnik stopowy. Cały czas doskonalono również technologię produkcji błękitnego barwnika.

Kilkanaście lat później **Axel von Cronstedt** poddał analizie kupfernickel i stwierdził, że faktycznie w minerale nie ma miedzi, ale zawiera on nieznaną składnik. Udało mu się go otrzymać w roku 1751, a srebrzystemu metalowi nadał nazwę **nikiel**. Początkowo odkrycie nie spotkało się z zainteresowaniem, ponieważ zanieczyszczenia powodowały, że nikiel był kruchy. Dopiero otrzymanie czystego metalu ujawniło jego zalety: odporność na korozję i kowalność. W kolejnym stuleciu został użyty jako składnik stopów, m.in. stosowanych – jak przed wiekami – do wyrobu monet.

Złośliwe duchy ukryte w kobalcie i niklu sprawiły kłopot samemu Mendelejewowi. Dziś pierwiastki są ułożone w tablicy zgodnie z rosnącą liczbą atomową (liczbą protonów w jądrze), ale za czasów Mendelejewa o kolejności decydowała masa atomowa. Ponieważ jest ona na ogół proporcjonalna do liczby atomowej, możliwe było ustawienie pierwiastków w poprawnym porządku. Jedynie w dwóch przypadkach kolejność została zaburzona: telluru i jodu oraz właśnie kobaltu i niklu (w końcu XIX wieku pojawiła się jeszcze jedna nieprawidłowa para – odkryty wtedy argon i znany już potas). Kobalt ma większą masę atomową niż nikiel i powinien znaleźć się za nim w szeregu, ale Mendelejew przestawił pierwiastki tak, aby – zgodnie z właściwościami chemicznymi – trafiły do odpowiednich grup. Kobalt wykazuje podobieństwo do leżących pod nim rodu i irydu, np. ze względu na tworzenie wielobarwnych związków, natomiast nikiel do doskonały katalizator przyłączania wodoru (jak pallad i platyna z tej samej rodziny). Wyznaczenie liczby atomowej pierwiastków, czego dokonano w początkach XX wieku, potwierdziło prawidłowość decyzji Mendelejewa.

Zasoby

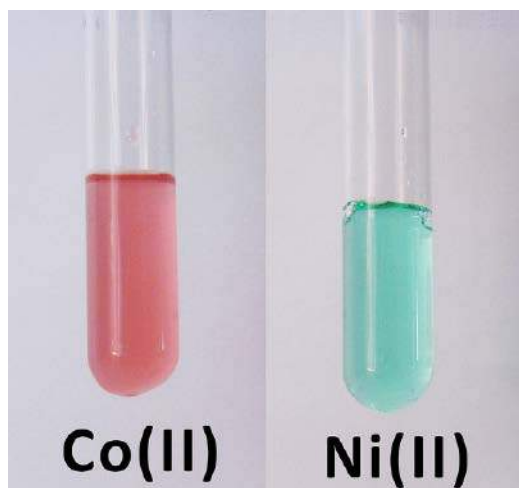
Na liście rozpowszechnienia pierwiastków w powierzchniowej warstwie Ziemi nikiel zajmuje 23. miejsce, stanowiąc 0,01% masy tej części globu. Kobaltu jest mniej – 0,004% i 28. miejsce na liście. Jeżeli potraktujemy Ziemię jako całość, nikiel awansuje aż na 6. miejsce z prawie dwuprocentowym udziałem w masie planety. To oczywiście konsekwencja istnienia żelazno-niklowego jądra Ziemi o średnicy prawie 7000 km. Sytuacja jest odbiciem proporcji w Układzie Słonecznym i Wszechświecie: pomijając wodór i hel, żelazo i nikiel należą do najbardziej rozpowszechnionych pierwiastków. Maksimum trwałości jąder atomowych przypada w pobliżu żelaza, zatem jądra tam leżące dość łatwo się tworzą i trudno ulegają rozpadowi. Nikiel to także typowy składnik meteorytów żelaznych, np. ogromne złoża rud tego metalu w kanadyjskim okręgu Sudbury powstało prawdopodobnie na skutek upadku gigantycznego przybysza z kosmosu około 2 miliardów lat temu. Kobalt jest znacznie rzadszy, ale pierwiastki o nieparzystej liczbie atomowej (kobalt ma 27 protonów w jądrze) są mniej rozpowszechnione niż ich „parzyści” kuzyni (to z kolei efekt mechanizmu nukleosyntezy).

Nikiel i kobalt tworzą własne minerały (kobalt przy tym często towarzyszy niklowi), ale występują również jako domieszki rud swoich sąsiadów z lewej i prawej – żelaza i miedzi. Proces produkcji jest dość skomplikowany i obejmuje szereg operacji chemicznych związanych z wydzieleniem metali z rud, a następnie rozdzieleniem niklu i kobaltu. Czyste metale otrzymywane są najczęściej przez elektrolizę ich soli, a niklu również **metodą Monda** (wykorzystuje się w niej łatwe tworzenie oraz rozkład lotnego połączenia niklu z tlenkiem węgla).

Najwięksi dostawcy niklu to Indonezja i Filipiny, a kobaltu – Demokratyczna Republika Konga. W roku 2020 światowa produkcja wyniosła 2,5 mln ton niklu i 140 tys. ton kobaltu. KGHM podczas otrzymywania miedzi uzyskuje rocznie kilka tysięcy ton siarczanu niklu.

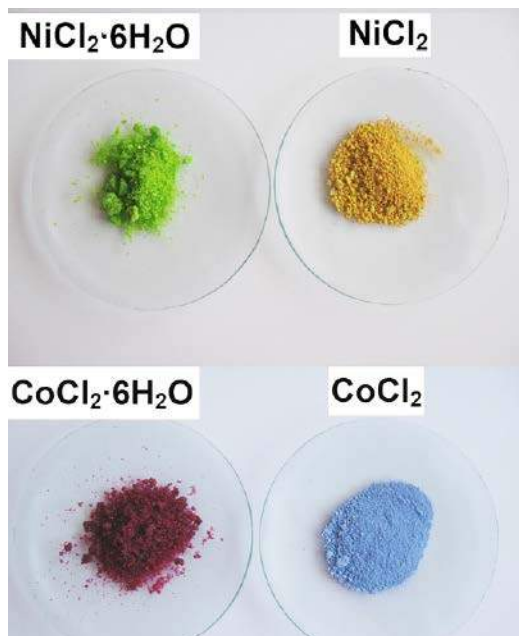
Wodorotlenki

Najłatwiej dostępne związki kobaltu i niklu to chloroki lub siarczany(VI) sprzedawane w postaci soli uwodnionych. W roztworze wodnym trwałe są dwuwartościowe kationy tych metali: sole kobaltu mają w nim różowe zabarwienie, a niklu zielone (3). Podobny kolor mają krystaliczne sole uwodnione, ale bezwodne są innej barwy (tak zwykle dzieje się w przypadku związków metali z grup od 4 do 11) (4).

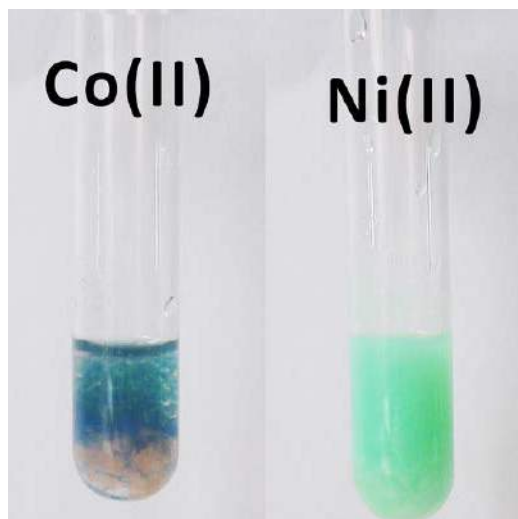


3. Barwy kationów kobaltu(II) i niklu(II) w roztworze wodnym

Pora na doświadczenia. Najpierw wodorotlenki mające duże znaczenie analityczne podczas wykrywania kationów niklu i kobaltu. Sporządź roztwory soli oraz roztwór wodorotlenku sodu NaOH. Po zmieszaniu roztworu soli danego metalu z zasadą zauważysz wytrącanie barwnych osadów: zielonego w roztworze soli niklu i niebieskiego dla kobaltu. W pierwszym przypadku rzeczywiście powstaje wodorotlenek o wzorze $\text{Ni}(\text{OH})_2$, natomiast kobalt

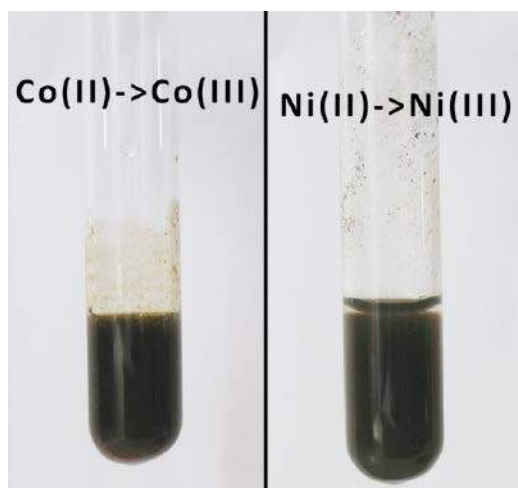


4. Barwy uwodnionych i bezwodnych soli kobaltu i niklu



5. Wytrącone wodorotlenki kobaltu(II) i niklu(II)

tworzy osad soli zasadowej (5). Początkujący analityk może pomylić go z osadem wodorotlenku miedzi, jednak $\text{Cu}(\text{OH})_2$ ma błękitną barwę w przeciwieństwie do wyraźnie niebieskiego związku kobaltu. Ponadto, gdy nieco poczekasz, niebieska barwa zmieni się w brunatnoróżową. To efekt utleniania osadu do związku kobaltu trójwartościowego przez tlen z powietrza rozpuszczony w roztworze (w przypadku wodorotlenku miedzi nic takiego nie zaobserwujesz). Reakcja jest analogiczna do przypadku wodorotlenku żelaza(II), który także łatwo utlenia się do brunatnych połączeń żelaza(III). Teoretycznie przez ogrzewanie niebieskiego osadu mógłbyś uzyskać czerwonoróżowy $\text{Co}(\text{OH})_2$, ale nie będziesz w stanie



6. Efekt działania utleniaczami na wodorotlenki kobaltu(II) i niklu(II)

Trójkami

Mianem **triada**, dziś już historycznym, określa się trzy sąsiadujące ze sobą metale z obecnych grup 8, 9 i 10 (dawniej była to jedna grupa VIII). Wykazują one większe podobieństwo między sobą w szeregu (okresie) niż z pierwiastkami położonymi w tej samej kolumnie (grupie). W ten sposób wyróżniano: żelazowce (żelazo, kobalt, nikiel), **platynowce lekkie** (ruten, rod, pallad) i **platynowce ciężkie** (osm, iryd, platyna). Otrzymane w końcu XX wieku pierwiastki o liczbach atomowych 108, 109 i 110 utworzyły kolejną triadę – **superplatynowce** (has, meitner, darmsztadt).

zapewnić braku dostępu tlenu i zawsze powstanie związek kobaltu(III).

W przypadku wodorotlenku niklu(II) nie zaobserwujesz zmian barwy, ponieważ nie ulega on wpływowi tlenu z powietrza, potrzebny jest silniejszy czynnik utleniający. Dodaj nieco wody utlenionej (3% roztwór H_2O_2) lub kilka kropli perhydrolu (**30% roztwór; ostrożnie, odczynnik żrący!**). Zielony osad

natychmiast przechodzi w czarny $\text{Ni}(\text{OH})_2$. Nadtlenek wodoru spowoduje również szybkie utlenianie niebieskiego osadu związku kobaltu do brunatnego $\text{Co}(\text{OH})_3$ (6).

Przeprowadzone reakcje potwierdzają podobieństwo kobaltu i niklu do żelaza. Sole metali dwuwartościowych reagują z wodorotlenkami, dając trudno rozpuszczalne, barwne osady. W przypadku żelaza i kobaltu następuje utlenianie powstałych związków tlenem z powietrza do połączeń, w których metal jest trójwartościowy (szybciej w przypadku żelaza niż kobaltu). Wodorotlenek niklu(II) jest odporny na działanie powietrza, ale ulega silniejszym utleniaczom, np. nadtlenkowi wodoru (podobnie jak analogiczne połączenia żelaza i kobaltu). Powstałe związki mają ciemniejsze barwy: brunatne w przypadku żelaza i kobaltu oraz czarną dla niklu. Wodorotlenki żelazowców nie są amfoteryczne: rozpuszczają się oczywiście w kwasach, ale nie reagują z nadmiarem dodanej zasady. ■

Krzysztof Orlński

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



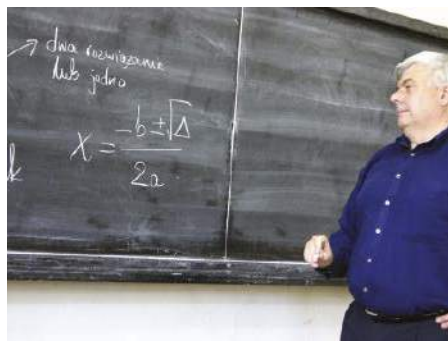
Michał Szurek tak mówi o sobie: „Urodzony w 1946. Ukończyłem UW w 1968 roku i od tego czasu tam pracuję na Wydziale Matematyki, Informatyki i Mechaniki. Specjalność naukowa: geometria algebraiczna. Ostatnio zajmowałem się wiązkami wektorowymi. Co to jest wiązka wektorowa? No, trzeba wektory mocno powiązać 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”.



Wycieczka w n-ty wymiar

W dzisiejszym odcinku Rozmaitości Matematycznych streszczam nie tylko swój wykład dla Akademickiej Telewizji Naukowej ATVN (zbieżność oznaczenia ze znaną stacją telewizyjną jest przypadkowa) – ale także zajęcia dla uczniów, w ramach Krajowego Funduszu na rzecz Dzieci.

Na początku zajęć o przestrzeniach wielowymiarowych pada zawsze pytanie: co jest tym czwartym wymiarem? Może czas? Odpowiadam, że może być nim cokolwiek. Jak to? Zobaczymy. Zrozumiałe, że potrzebna nam będzie pewna wyobraźnia. Oto ładny cytat:

– Najpierw trzeba szeroko rozpostrzeć skrzydła, potem pomachać ze trzy razy na próbę, żeby sprawdzić, czy wszystko w porządku, a potem....

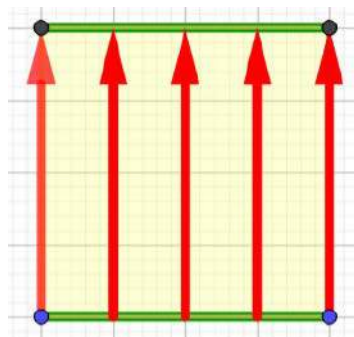
– Ale ja nie mam skrzydeł – przerwał mu łoś.

– Aha – odrzekł zmieszany swoją nieuwagą kruk.

– Ale mam bogatą wyobraźnię. Może wystarczy, jak sobie wyobrazę, że mam skrzydła – dodał łoś bez kompleksów.

Joanna Haręża, *Kruk między innymi*, wydawnictwo AB, Warszawa 2005.

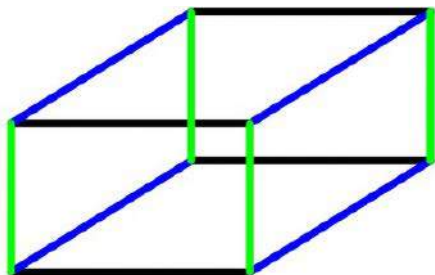
W młodości grywałem w tenisa. Kort trzeba było samodzielnie przygotować, w szczególności wyrównać. Służyła do tego płaska szczotka o szerokości 2 metrów, którą ciągnęło się za sobą. Należało dbać o to, by iść równo i miarowo. Jeżeli przyjmniemy, że szczotka była dwumetrowym odcinkiem, to po przejściu dwóch metrów zostawał na korce wyszczotkowany kwadrat 2 na 2 metry. Można powiedzieć, że kwadrat powstaje przez przeciągnięcie odcinka w dodatkowym, drugim wymiarze. Bez użycia tradycyjnych oznaczeń x, y – zielony odcinek został przesunięty w czerwonym, prostopadłym kierunku.



1. Jak z odcinka powstaje kwadrat

Podobnie zbudujemy sześcian. Trzeba przesunąć kwadrat w dodatkowym, trzecim wymiarze. Jeżeli jednak chcemy to narysować, mamy od razu kłopot. Co prawda nauczono nas w szkole rysunku w perspektywie. Przyjrzyjmy się w telewizji boisku piłkarskiemu. Wiemy, że na środku boiska jest koło (matematycznie: okrąg; przypomnę, że koło to pełna figura, a okrąg jest jej brzegiem), ale na ekranie jest elipsa. To nic, nasz mózg przetwarza ją na równy okrąg. Spójrzmy na **rysunek 2**. Możemy umówić się, że to też jest sześcian, bo to także tylko kwestia perspektywy. Widzimy, że czarno-niebieski kwadrat został przesunięty w prostopadłym, zielonym kierunku.

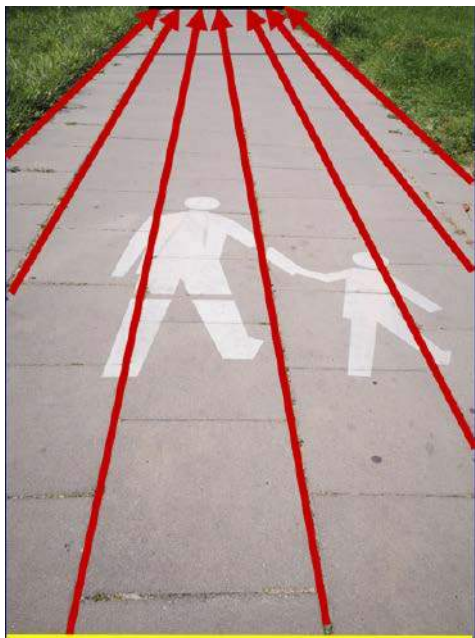
Nie wydaje się to bardzo odkrywcze. Ale to teraz, w XXI wieku. Gdy wiek XIX ledwo przekroczył



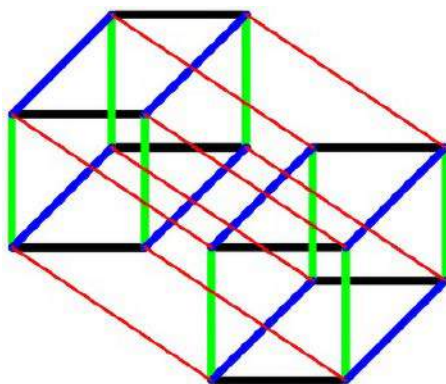
2. Kolorowy sześcian w perspektywie równoległej

połowę swojego stuletniego żywota (czyli w roku 1854), Bernhard Riemann wygłosił w Getyndze swój wykład habilitacyjny, *O hipotezach, które leżą u podstaw geometrii* (oryginalny tytuł niemiecki *Über die Hypothesen, welche der Geometrie zu Grunde liegen*). O samym tym wykładzie piękną książkę napisał ks. prof. Michał Heller, polecam!

Riemann zmierzył się z zagadnieniami przestrzeni o większej liczbie wymiarów. Ówczesnie było to nie do pojęcia, nawet najwybitniejsi uczeni (jak na przykład Karl Friedrich Gauss) nie dopuszczali możliwości rozważania takich przestrzeni, a i samo podejście Riemanna też jest kręte. No, ale łatwo to pisać z dzisiejszego punktu widzenia, gdy wszystko zostało gruntownie opracowane, wyjaśnione, uproszczone i do końca zmatematyzowane. Spójrzmy na **rysunek 3**. Startujemy z dołu rysunku, z żółtej



3. „Spacer Riemanna” – przejście od pewnego wymiaru do większego o jeden

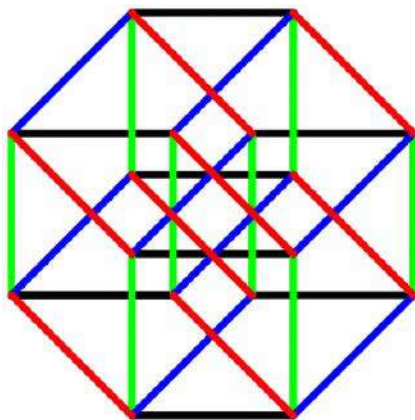


4. Łapiemy czarno-zielono-niebieski sześcian i ciągniemy w czerwonym kierunku

linii i idziemy przed siebie. Wszystkie możliwe drogi zamiatają coś, to jest wymiaru 2 – tak jak przy szczotkowaniu kortu. A więc możliwe drogi od n -wymiarowej figury A do innej figury, B , wypełniają figurę wymiaru $n+1$, o jeden większego.

Wobec tego nic nie stoi na przeszkodzie, by zbudować „coś” wymiaru 4 – a raczej: narysować owo coś. Wybierzmy kierunek, który będzie na rysunku reprezentował czwarty wymiar (czerwony) i przesunijmy sześcian w tym wymiarze. Wynik widzimy na **rysunku 4**, a gdy postaramy się o ładny wykres, możemy dostać coś takiego, jak na **rysunku 5**. Oto przed Wami, drodzy Czytelnicy, bryła wymiaru 4 ..., oj, oj, wcale nie bryła, tylko jej dwuwymiarowy rysunek. Sześcienne klocki możemy wziąć do ręki, naszej bryły... raczej nie. Ale co to szkodzi? Czyż nie mamy bogatej wyobraźni?

Inny pomysł potrzebny jest do stworzenia bryły mającej taki sam charakter, co trójkąt na płaszczyźnie



5. Kostka czterowymiarowa. Po angielsku ma ona oddzielną nazwę: tesseract

i czworościan w przestrzeni. Myślmy o trójkacie równobocznym i czworościanie foremym. Pierwszy z nich powstaje przez połączenie trzech punktów, położonych w tej samej odległości od siebie. Drugi – przez połączenie czterech punktów w przestrzeni, każdy z każdym. W następnym, czwartym wymiarze, jest równie prosto. Łączymy pięć punktów, każdy z każdym. Wszystkie mają być tak samo odległe od siebie wzajemnie. Dlaczego musimy uciec aż do czwartego wymiaru? Dlatego, że w trójwymiarowej przestrzeni takich pięciu punktów nie ma. Cztery tak, ale piątego nie ma już gdzie umieścić.

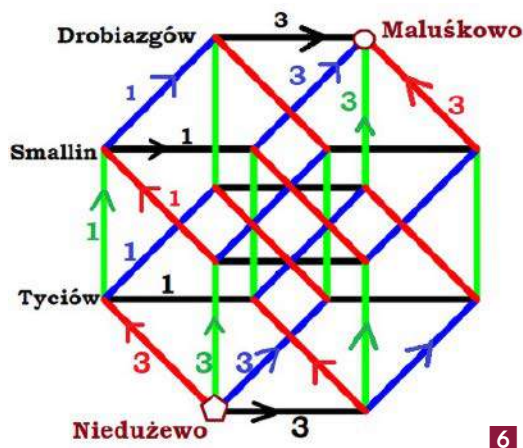
Wszystko da się wyrazić w nieskomplikowany sposób za pomocą współrzędnych. Trójkąt o współrzędnych przestrzennych $(1,0,0)$, $(0,1,0)$, $(0,0,1)$ jest równoboczny – długość każdego z trzech boków wynosi $\sqrt{2}$. W przestrzeni wymiaru cztery mamy 4 współrzędne. Cztery punkty $(1,0,0,0)$, $(0,1,0,0)$, $(0,0,1,0)$, $(0,0,0,1)$ są odległe od siebie wzajemnie też o $\sqrt{2}$. Tworzą zwykły czworościan, piramidkę – tyle że leży ona w przestrzeni wymiaru 4. Pójdźmy wymiar wyżej – albo jeszcze dalej, nawet do wymiaru 100. Wyobraźmy sobie sto punktów. Pierwszy z nich ma współrzędne $(1, 0, \dots, 0)$ – zer jest 99. Drugi ma współrzędne $(0,1,0,\dots,0)$ – po jedynce następuje 98 zer, trzeci $(0,0,1,0,\dots,0)$ – z 97 zerami po jedynce, wreszcie setny ma 99 zer i jedynkę na końcu. Z połączenia tych punktów powstaje bryła 99-wymiarowa o własnościach podobnych do trójkąta i czworościanu.

W wymiarach od 5 w górę jest jeszcze jedna bryła foremna – tworzą ją środki ścian kostki (tego odpowiednika sześcianu). Innych brył foremnych w tych wymiarach nie ma, ale – uwaga – w przestrzeni czterowymiarowej jest ich aż sześć. To też temat na oddzielny artykuł, a właściwie na duży wykład na poziomie uniwersyteckim.

Być może ktoś zaprotestuje: przecież czwarte-go wymiaru nie ma, żyjemy w trzech i wszystko to są tylko igraszki matematyczne. Zgodzę się z tym, że są to igraszki matematyczne. Ale...

Można przecież powiedzieć, że cała muzyka to tylko zabawa dźwiękiem, cała literatura to zabawa słowem, a malarstwo – to tylko kształty i kolory. A jednak wszyscy robimy to dla pewnej przyjemności, dla satysfakcji intelektualnej. Czyż nie tak? Czy nie wzrusza nas obraz, dobra książka, nastrojowa muzyka?

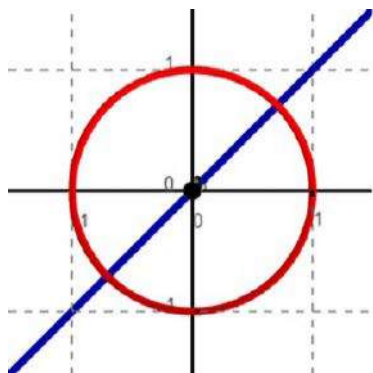
To wszystko byłoby za mało, jeśli chodzi o teorie matematyczne. Na szczęście matematyka jest nauką, którą spotykamy niemal wszędzie. Konkretnie: wszystko to ma zastosowanie, nie tylko w fizyce, ale na przykład w ekonomii. Przestrzenie wielowymiarowe spotykamy wszędzie...



Popatrzmy na rysunek 5 inaczej. Wyobraźmy sobie, że jest to sieć, sieć połączeń. Jakich połączeń? Wszystko jedno. Może to być sieć elektryczna, sieć kolejowa, informatyczna, kanałów wodnych, a nawet np. schemat wyciągów narciarskich. Przyjmijmy tylko jeden z wierzchołków jako początek (start, źródło, input itd.), a inny jako metę, koniec, ujście, output itd. Ja zdecyduję się na sieć kolejową. Lubię zadania „ruszają w dal pociągi dwa”. Są naprawdę kształtujące, a ja sobie dodatkowo wyobrażam, że wszystko dzieje się w jakimś odległym kraju, gdzie naprawdę istotne jest, gdzie i kiedy się one spotkają.

Spójrzmy na **rysunek 6**. Między Niedużewem a Maluśkowem jest gęsta sieć automatycznych połączeń kolejowych (wiem, że w Lille we Francji metro jest całkowicie zautomatyzowane – jeździ bez udziału człowieka; podobno w Warszawie też mogłoby tak być, ale przepisy nie pozwalają!). Każdy odcinek sieci ma przepustowość 1, to znaczy, że można nim wysłać jeden wagon na jedną jednostkę czasu. Jeżeli puścimy trzy wagony, będą one jechać trzy razy wolniej. Jaka jest przepustowość całej sieci?

Spójrzmy. Dyżurny ruchu w Niedużewie wysła 3 wagony w każdym z czterech kierunków, a potem dzielą się one same. Niech jednostką czasu będzie 10 minut. A więc do Tyciowa dojeżdżają po 30 minutach trzy wagony. Dzielą się one po jednym i po dalszych dziesięciu minutach docierają do następnych stacji. W Smallinie spotykają się (ważne – w tym samym momencie) dwa wagony z dwóch różnych kierunków i każdy jedzie dalej. W Drobiazgowie spotykają się trzy wagony (jeden ze Smallina, drugi z południa, zieloną trasą i trzeci z południowego wschodu czerwoną). Znowu: ważne jest, że dojeżdżają w tej samej chwili i wobec tego zgodnie ruszają już do końcowej stacji Maluśkowo.

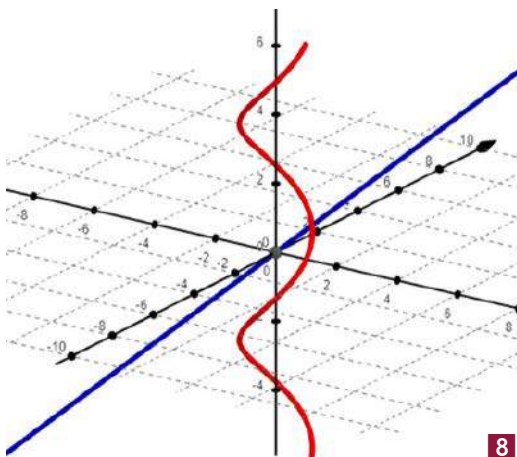


7

Zbierzmy nasze obliczenia. Dwanaście wagonów w czasie $3+1+1+3=8$, nie jest ważne, że to 80 minut. Chwila zastanowienia wystarczy, żeby zrozumieć, że przepustowość to 12 dzielone przez 8, czyli 1,5. W jakich to jednostkach? Trudno powiedzieć. Na wykładach mówię, że to w szurkach. Jeden szurek odpowiada jednemu wagonowi na jednostkowym odcinku.

W dydaktyce matematyki jest pojęcie zadań izomorficznych. To są zadania dotyczące różnych sytuacji, ale mające tę samą treść matematyczną. Nauczyciele szkół podstawowych znają ten problem: jeżeli w klasie jest zadanie o cukierkach, a na klasówce będzie o czekoladkach, to dla części dzieci będzie to inne zadanie! Takiego podejścia nie pozbywają się niekiedy i studenci. Ale dla Czytelników jest zrozumiałe, że zamiast wagonów mogą w sieci przepływać prądy elektryczne, megabajty informacji, woda, narciarze na wyciągu, żołnierze... i tak dalej. Teoria przepływów w sieciach jest ważną gałęzią matematyki obliczeniowej.

W pojęciu wymiaru ważniejsze jest coś innego. Najpierw wyobraźmy sobie pusty plac, po którym jeżdżą samochody (nie wytłumaczę, po co jeżdżą – może po prostu tak sobie). Czy jeżeli dwa samochody znajdują się w tym samym punkcie, to mamy wypadek? Oczywiście, że... nie. Jeżeli przejadą przez to samo miejsce, ale w odstępie minuty, to wszystko będzie w porządku. To znaczy, że kolizja zdarzy się, gdy będą w tym samym punkcie czasoprzestrzeni (x, y, t) . Proste, prawda? Przyjrzyjmy się wykresom przestrzennym. Samochód A jedzie ze stałą prędkością po niebieskiej linii prostej, zaś samochód B jeździ w kółko, jak na **rysunku 7**. Mogą się zderzyć, ale nie muszą. Zróbmy wykresy ich ruchu w czasoprzestrzeni (x, y, t) . Wybrałem takie warunki początkowe, żeby krzywe przestrzenne nie stykały się. Jest bezpiecznie (**rysunek 8**) – chociaż naprawdę widać to dopiero na animacji albo po stosownych rachunkach. Zatem zadanie.



8

Zadanie. Wykazać, że krzywe (t, t, t) i $(\cos t, \sin t, t)$ nie mają punktów wspólnych.

Rozwiązanie jest proste. Ma być $t = \sin t = \cos t$. Pierwsze z tych równań jest spełnione tylko dla $t=0$, ale $\cos 0=1$. Punktu wspólnego nie ma.

Omówmy zatem kwestię bezpieczeństwa. Przyjmijmy, że bezpiecznie jest, jeżeli samochody są odległe o 10 metrów w przestrzeni albo o 10 sekund w czasie. Przypomnijmy sobie, jak mierzy się odległość w przestrzeni. Wystarczy nam wzór uproszczony – odległość punktu (x, y, t) od początku układu $(0,0,0)$ to

$$\sqrt{x^2 + y^2 + t^2} \leq 10$$

Każdy samochód zatem roztacza wokół siebie swego rodzaju „kulę bezpieczeństwa”. Kontroler ruchu może podejrzeć, czy dwie takie kule nie zahaczają o siebie. Wtedy powinien ostrzec kierowców o możliwej kolizji. Nie jest to jeszcze może powód do alarmu – sądzę, że powinien go ogłosić, gdy jeden samochód wejdzie w kulę bezpieczeństwa drugiego.

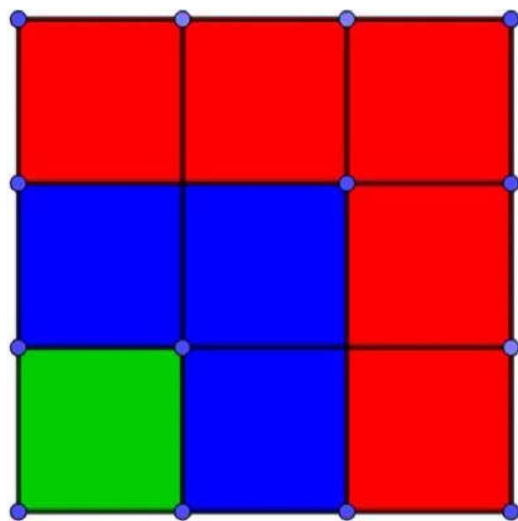
Samochody na pustym placu sobie jakoś poradzają, ale z samolotami w przestrzeni jest już inaczej. „Kula bezpieczeństwa” wokół nich jest już w przestrzeni czterowymiarowej (x, y, z, t) – ale tak naprawdę w przestrzeni o większej liczbie wymiarów, bo jej promień zależy również od prędkości, od pory dnia, od pogody i może od wielu innych parametrów. Pamiętam, że gdy latały jeszcze samoloty Concorde, nad lotniskiem paryskim było o włos od katastrofy: na kursie concorde’a pojawił się mały samolot w odległości 400 metrów. To kilka sekund lotu.

Dochodzimy tu do jednej z możliwości spojrzenia na wymiar, co to jest wymiar? Można powiedzieć, że jest to liczba parametrów potrzebnych do opisanie pewnego zjawiska. Parametry te mogą, ale nie muszą

być niezależne. Jeżeli są zależne, dostajemy do analizy skomplikowany obiekt geometryczny.

Napisałem, że to tylko jedna z możliwości. Inna prowadzi nas do fraktali – figur, które mają wymiar niecałkowity, na przykład nie 1, nie 2, coś pośrodku. Jak to jest możliwe? Punktem wyjścia jest taka oto obserwacja. Jeżeli powiększymy długość boku kwadratu dwukrotnie, to pole tego kwadratu wzrośnie czterokrotnie (rysunek 9), a jeżeli potroimy bok, to pole wzrośnie aż 9 razy. Dla sześciangu idzie to jeszcze szybciej: zwiększenie długości krawędzi dwukrotnie spowoduje ośmiokrotny wzrost objętości, a potrojenie – wzrost 27-krotny. W czwartym wymiarze zależność ta zmienia się wraz z czwartą potęgą, w piątym – z piątą. Z ciekawości przyjrzymy się wymiarowi 10. O ile trzeba zwiększyć krawędź kostki dziesięciowymiarowej, żeby podwoiła się objętość? To proste: $\sqrt[10]{2}$, czyli w przybliżeniu 1,07. Siedmioprocentowy wzrost krawędzi podwaja objętość. Ale to jest także zadanie z ekonomii. Jeżeli inflacja roczna w kraju wynosi 7%, ceny podwajają się po 10 latach.

Urodzony w Polsce Benoit (=Benedykt) Mandelbrot (1924–2010) wprowadził do matematyki figury o ułamkowym wymiarze. Jak to jest możliwe? Właśnie przez analogię z powiększaniem kwadratu i sześciangu. Mówimy, że wymiarem figury jest n , jeżeli można ją podzielić na takie k^n części, że każda z nich jest podobna do wyjściowej w skali 1:k. Taka liczba n nie musi być całkowita. Tak jest właśnie dla fraktali (nazwa od łacińskiego *fractus* = złamany). Fraktale są „modne”, przy kilku okazjach gościły i na naszych łamach. Nie będę więc o nich pisać. Kto nie miał



9. Podwojenie i potrojenie boku kwadratu

okazji się z nimi zetknąć, może w wyszukiwarce wpisać stosowne słowo i wyskoczą tysiące stron. Jednym z ostatnich pomysłów Mandelbrota było zastosowanie fraktali i tak zwanych quasi-fraktali do opisu zjawisk giełdowych. Zastosowania matematyki do gry giełdowej traktuję z dużą podejrzliwością – przypomina mi to odróżnianie banknotów prawdziwych od fałszywych na podstawie popiołu, który powstaje po ich spaleniu.

O ciekawej, „prawdziwej” geometrii figur wymiaru 4 i wyższych – to już innym razem. ■

Michał Szurek

Europejski raj

Alexa Lavenda

Wydawnictwo Lipstick Books, liczba stron: 352, cena: 39,99 zł

Szczęście było tak blisko, na wyciągnięcie ręki... jednak los po raz kolejny zakpił sobie z Patrycji i Maxa i ponownie ich rozdzielił. Czy tym razem Max będzie w stanie odnaleźć Patrycję? Do czego jest zdolny posunąć się George, żeby zatrzymać dziewczynę przy sobie? Dla miłości robi się wszystko, miłość rządzi się własnymi prawami i nie zna żadnych granic... szczególnie miłość niebezpiecznego szaleńca.



Tradycyjny samochód i jego rola jest w sytuacji, w jakiej maszyna do pisania była w latach 80. XX wieku. Wkrótce zostanie zastąpiony przez coś nowego. Maszyny do pisania wydawały się niezastąpione. Dekadę, może półtorej, później, już się nie liczyły. Dość często spotyka się dziś oceny, że przemysłowi samochodowemu zostało tyle samo czasu.

Motoryzacja, jaką znamy

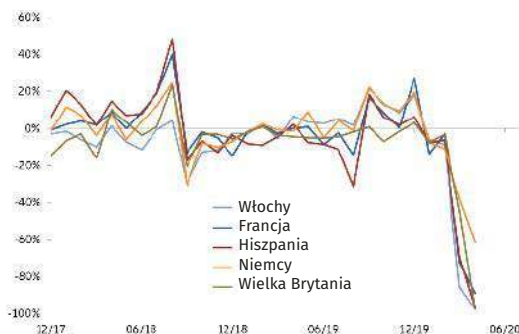
Zamiast auta w garażu – mobilność

Jeśli ktoś uważa, że analogia z maszynami do pisania jest niewyraźna, to może lepsza byłaby historia fotografii, opartej na światłoczułych filmach, wywoływaniu, odbitkach na papierze. W ciągu kilku zaledwie lat wyparte zostało to przez technikę cyfrową, w postaci masowej trafiając ostatecznie do smartfonów będących integratorami wielu innych tradycyjnych urządzeń, od telefonu, przez aparat fotograficzny, po komputer, którego smartfony nie wyparły, ale zastąpiły w wielu prostych zastosowaniach. I tak, jak sądzi wielu, może stać się z samochodami, a raczej określonym modelem korzystania z nich, posiadania, roli w życiu i biznesie.

Procesy te przyspieszyły w ostatnim czasie wskutek pandemii i jej konsekwencji. W Europie prawie natychmiast po jej wybuchu, w kwietniu 2020 spadła gwałtownie liczba rejestracji nowych samochodów (1). W przypadku pięciu największych rynków, które reprezentują 75 proc. rynku UE, spadek wyniósł w kwietniu –84 proc., co oznacza znaczny spadek sprzedaży o 1,6 mln nowych samochodów.

Ekspansja elektryków, z dużym państwowym wsparciem

Jednocześnie widać ekspansję elektrycznych pojazdów, choć w wielu krajach dzieje się to w dużym stopniu za sprawą potężnego wsparcia dotacjami i ulgami wszelkiego rodzaju. Model 3 Tesli w połowie 2020 roku stał się najlepiej sprzedającym pojazdem w Kalifornii we wszystkich kategoriach. Dane opublikowało Kalifornijskie Stowarzyszenie Dealerów Nowych Samochodów, które zapewne zachodziło w głowę, jak to się stało, że najlepiej sprzedającym



Źródło: narodowe organizacje producentów samochodów, Euler Hermes, AZ Research

1. Liczba rejestracji nowych samochodów w Europie

się pojazdem jest akurat ten, który nie jest sprzedawany w sieciach dealerskich.

Potęgą rynku motoryzacyjnego czują oddech nowe go na plecach. Koncern General Motors wprowadza na rynek nowy elektryczny samochód dostawczy, chcąc zapobiec przejściu tego rynku przez Teslę, która zapowiada elektrycznego Cybertrucka. Jednocześnie GM ogłasza plany wprowadzenia do 2023 r. co najmniej dwudziestu całkowicie elektrycznych modeli we wszystkich kategoriach. Ford jeszcze wcześniej podjął rękawicę rzuconą przez Teslę na będącym wcześniej jego domeną rynku półciężarówek i zapowiada również elektryczne modele. Coraz więcej producentów przyspiesza realizację swoich planów, próbując wywalczyć sobie pozycję na rynku.

Pojazdy elektryczne przestają również być domeną bogatych klientów. Chiny produkują obecnie modele, których ceny początkowe wynoszą około tysiąca dolarów i które można kupić za pośrednictwem Alibaby.

Kraje takie jak Hiszpania, Francja czy Niemcy, ogłosiły niedawno nowe plany wspierania sprzedaży pojazdów elektrycznych, zwiększając finansowe zachęty do ich zakupu. W Niemczech zachęty te zbiegają się z podwyżkami podatków od dużych pojazdów z silnikiem Diesla, a także z wprowadzeniem obowiązku instalowania na wszystkich stacjach benzynowych punktów ładowania pojazdów elektrycznych. W Hiszpanii, gigant energetyczny Iberdrola przyspieszył swoje plany rozwoju sieci ładowania pojazdów elektrycznych, koncentrując się również na stacjach benzynowych z punktami szybkiego ładowania i zamierza zainstalować 150 tys. punktów w domach, firmach oraz miastach i miasteczkach w ciągu najbliższych pięciu lat. Wielka Brytania ustaliła datę wprowadzenia zakazu sprzedaży pojazdów napędzanych olejem napędowym i benzyną już na rok 2035.

Niemcy nieprzyzwyczajeni do tego, że są tylko „dość dobrzy”

Przyjrzyjmy się tym zmianom na przykładzie starej potęgi motoryzacyjnej, jaką są Niemcy i presji wywieranej na tamtejszych producentów (2) przez ekspansję elektryków, wspieraną dodatkowo przez państwa. Media zastanawiają się od pewnego czasu, czy niemieckie samochody stracą na atrakcyjności w czasach elektrycznej rewolucji. Trudno zaprzeczyć, że marki takie jak Mercedes, BMW, Audi i Porsche nadal należą do liderów w przemyśle samochodowym. Jednak nic nie jest dane raz na zawsze.

Czy dotrzymują kroku np. Tesli, która zwykle jest podawana jako przykład i symbol nowych czasów? Volkswagen i BMW oferowały kompaktowe samochody z zasięgami ok. 300 km. Audi wprowadziło swój model SUV-a elektrycznego e-tron w 2019 roku, auto wyższej klasy. Porsche zaoferowało Taycaną w 2020 roku, który jest niezwykle drogi a zasięgiem,



2. Niemieckie marki aut

parametrem kluczowym dla tej klasy aut, nie imponuje. Mercedes zapowiedział niedawno rezygnację z modeli spalinyowych do 2030 roku. Wydawałoby się więc, że niemieckie firmy motoryzacyjne jakoś dotrzymują kroku i złapały bakcyli elektrycznej rewolucji.

W szczegółach nie wygląda to już tak wesoło. Dla niemieckich aut powiedzenie, że są „dość dobre”, to za mało. One aspirowały zawsze do miana „najlepsze”. I tu jest problem. Na przykład, Audi e-tron ma pakiet baterii 95 kWh, co daje zasięg ok. 350 km. Tesla Model X zaś oferuje pakiet 100 kWh, co daje prawie 200 km więcej zasięgu.

Można by stwierdzić, że nie minie wiele czasu, zanim niemieckie samochody dogonią Teslę i ponownie staną się samochodami z najwyższej półki. Jednak nikt jeszcze nie prześcignął Tesli, choć wielu się stara. Chevrolet, na przykład, choć pokonuje Teslę ceną, to jeśli chodzi o wydajność, jest daleko w tyle. Ford dogonił Teslę w kategorii zasięgu, a nawet prześcignął, jeśli chodzi o komfort, ale wciąż jest w tyle za wydajnością Tesli i jej szybkością ładowania.

Niemieckie marki mimo pewnych doświadczeń wciąż są nowicjuszami w branży aut elektrycznych. Wciąż muszą sobie radzić z problemami związanymi z łańcuchem dostaw, niewystarczającym zasięgiem i ceną modeli. Może być tak, że Niemcom zabraknie czasu, zwłaszcza gdy Tesla przezwycięży różne bóle rozwojowe, z jakimi się boryka, ustabilizuje jakość produkcji, będzie w stanie sprostać popytowi. Krótko mówiąc, gdy Tesla dojrzeje, powody, by kupować niemieckie auta, mogą zniknąć.

Rewolucja samochodowa = rewolucja ubezpieczeniowa

Mimo niebywałego sukcesu motoryzacji konwencjonalny samochód jest obecnie często opisywany jako jedno z największych marnotrawstw wszech czasów. Samochody, mimo że zostały stworzone do jazdy, stoją zwykle przez 95 proc. czasu, w miastach nawet średnio 97 procent. Do tego gigantycznego marnotrawstwa surowców dochodzi jeszcze marnotrawstwo przestrzeni, zwłaszcza w miastach. Podczas epidemii COVID-19 ruch uliczny w większości miast Europy zmniejszył się o około 80 proc., co oznacza, że typowy samochód był w użyciu tylko przez 1 proc. czasu. Samochody powinny jeździć, a nie stać. Nie ma sensu wykorzystywać jednej z najważniejszych inwestycji gospodarstwa domowego w tak niewielkim stopniu.

Są także inne bilanse, bardzo niekorzystne dla tradycyjnej motoryzacji. Według szacunków Światowej Organizacji Zdrowia, w 2019 roku na całym świecie

w ruchu drogowym zginęło 1,3 miliona osób. Ruch pojazdów autonomicznych radykalnie zminimalizuje liczbę ofiar śmiertelnych i rannych (3).

Ma to przy okazji interesujące konsekwencje dla rynku ubezpieczeń komunikacyjnych. Z raportu firmy Majesco „Rethinking Auto Insurance: From a Transactional Relationship to a Mobility Customer Experience” wynika, że trzeba się przygotowywać do zupełnie nowego rynku, na którym szkody w razie wypadków będą pokrywane nie przez firmy ubezpieczeniowe, a przez producentów pojazdów lub specjalne fundusze tworzone wspólnie przez producentów aut, instytucje odpowiedzialne za infrastrukturę, regulatorów i inne podmioty. Pojedynczy użytkownik, zwłaszcza autonomicznego pojazdu, nie będzie zobowiązany ubezpieczać się indywidualnie a jeśli będzie płacić, to pośrednio, w cenie opłaty za przejazd.

Zanim autonomiczna motoryzacja się upowszechni, w modelach subskrypcyjnych i wynajmach krótkoterminowych, zamiast ubezpieczeń „całościowych” upowszechnić się mogą różnego rodzaju elastyczne rozwiązania krótkoterminowe i częściowe, na przykład wykluczające odpowiedzialność użytkownika za dokonywanie przeglądów, serwisowania itd.

Jednak, jak się uważa, masowe wprowadzenie pojazdów autonomicznych to kwestia wciąż nieokreślonej i raczej odleglejszej niż bliskiej przyszłości. Na razie zmiany widać wyraźniej po stronie modeli korzystania z aut i ich posiadania. W krajach rozwiniętych intensywnie rozwijają się różnorodne formy subskrypcji, ride-sharing i modele krótkoterminowego wynajmu. Młodsze pokolenia są coraz mniej przywiązane do myśli, że muszą mieć auto na własność. To koszty i zwykłe straty, bo wartość samochodu, w którym lokuje się niemałe przecież pieniądze, prawie zawsze spada z czasem.



3. Wizja samochodu autonomicznego

Według badań Automotive Landscape 2025, przeprowadzonych przez Roland Berger Strategy Consultants, do połowy rozpoczętej dekady nastąpi szereg dużych zmian w światowej motoryzacji. Po pierwsze ma nastąpić dramatyczne przesunięcie produkcji i sprzedaży na rynki azjatyckie, na znaczeniu zyskają auta małe i tanie a na atrakcyjności tracą modele droższe jako symbole statusu. Ekosystemy mobilności w dużych aglomeracjach miejskich świata doprowadzić mają do „demotoryzacji”. Zdaniem autorów badania do 2025 r. pojazdy elektryczne będą stanowiły około 10 proc. sprzedaży nowych pojazdów. Udział hybryd osiągnie 40 procent.

Warto jednak pamiętać, że wszystko, o czym mowa wyżej, dotyczy przede wszystkim rynków w krajach rozwiniętych. Nawet jeśli zmiany zajdą tak szybko, jak się przewiduje, to będą widoczne głównie w USA i w Europie. I wtedy powstaną na planecie dwa motoryzacyjne światy, jeden nowy, być może już nie oparty na paliwach kopalnych, być może coraz bardziej autonomiczny i elastyczny w rozwiązaniach mobilnych, i drugi – na innych kontynentach, który nadal będzie wyglądał tak jak tradycyjna motoryzacja. ■

Mirosław Usidus

Książka o przyjaźni

Maciej Marcisz

Wydawnictwo W.A.B., liczba stron: 304, cena: 44,99 zł

Kasia, Michał i Dorota znają się od zawsze. Ich przyjaźń przetrwała kilkanaście lat, przeprowadzki do odległych miast, a nawet coraz bardziej rozjeżdżające się życiowe drogi. Gdy jednak skrywana tajemnica wycieka na grupowym czacie, los relacji zaczyna wisieć na włosku. Czy należy walczyć o marniejącą przyjaźń przez wzgląd na stare dobre czasy? A może pewnym rzeczą należy pozwolić umrzeć? Czy nasi rodzice mieli rację, mówiąc, że w pewnym wieku prawdziwych przyjaciół policzymy na palcach jednej ręki?





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

Problem skoczka szachowego to zadanie polegające na odwiedzeniu skoczkiem (nazywanym też koniem lub konikiem) każdego pola szachownicy dokładnie jeden raz. Jeśli skoczek może po ostatnim ruchu wrócić na pole, z którego zaczynał, to mówimy o zamkniętej ścieżce skoczka szachowego. Jeśli skoczek może obejść wszystkie pola, ale po ostatnim ruchu nie może wrócić na pole startowe, to mówimy o ścieżce otwartej.

Problem skoczka szachowego

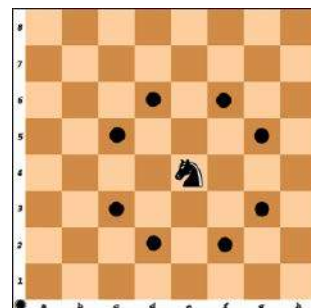
Ruch konika szachowego (1) jest następujący: dwa pola w jednym kierunku i jedno pole w bok w stosunku do pierwotnego kierunku. Ruch tej figury jest najbardziej charakterystyczną cechą łączącą szachy



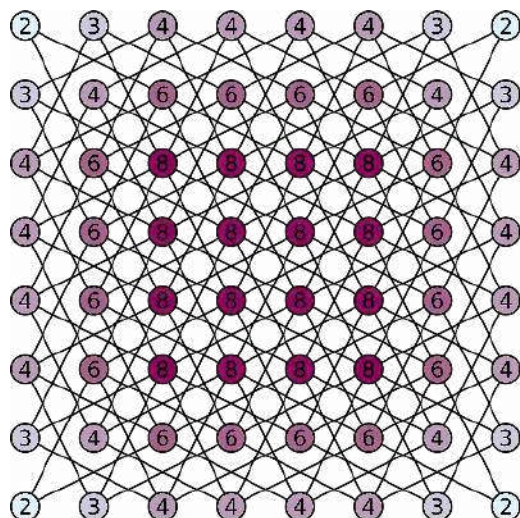
1. Skoczek szachowy

i inne pokrewne gry z ich przodkami, ponieważ występował już w czaturandze i nie uległ zmianie co najmniej od VII wieku n.e.

Poniżej zaznaczono osiem pól, na które może



2. Możliwe ruchy skoczka szachowego z pola e4



3. Wykres przedstawiający liczbę możliwych posunięć skoczka, które można wykonać z wybranej pozycji, na standardowej szachownicy 8x8

1	6	3
4		8
7	2	5

4. Próba rozwiązania problemu dla planszy 3×3

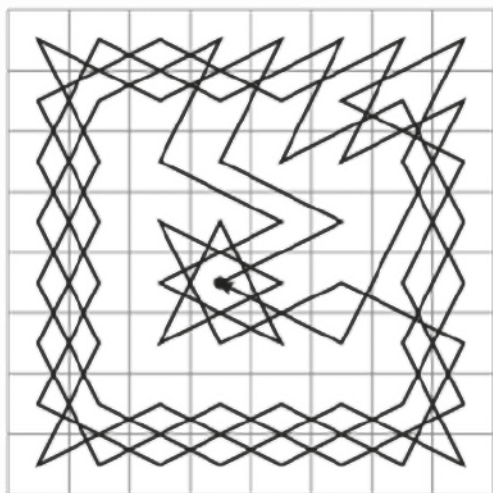
skoczyć skoczek, znajdujący się na polu e4 – jednym z centralnych pól szachownicy (2). Z rogu szachownicy konik może poruszyć się tylko na jedno z dwóch pól. Na wykresie (3) przedstawiona jest liczba możliwych posunięć, które konik może wykonać z wybranej pozycji,

23	8	21	16	25
20	15	24	7	12
9	22	13	4	17
14	19	2	11	6
1	10	5	18	3

5. Przykład otwartej ścieżki na planszy 5×5

Dla planszy o wymiarach 3×3 skoczek może przemierzyć osiem pól, ale na dziewiąte, środkowe nie może wskoczyć (4). Rozpoczynając ze środka takiej planszy, skoczek nie jest w stanie wykonać poprawnego ruchu.

Nie istnieje również rozwiązanie tego problemu na planszy 4×4. Dla wszystkich plansz kwadratowych o boku większym od czterech można znaleźć taką ścieżkę otwartą dla skoczka, żeby odwiedził



6. Zamknięta ścieżka dla szachownicy 8×8

wszystkie pola dokładnie jeden raz. Problem skoczka szachowego na planszy $n \times n$ nazywa się też konikówką rzędu n .

Przykład możliwej ścieżki skoczka na planszy 5×5 przedstawiony jest na **rysunku 5**. Na takiej planszy możliwe są tylko ścieżki otwarte, ponieważ konieczne jest wykonanie nieparzystej liczby posunięć.

Jedno z bardzo wielu rozwiązań zamkniętej ścieżki dla szachownicy 8×8 przedstawione jest na **rysunku 6**. Kolejność wykonywania ruchów przez skoczka: d4 f5 d6 e8 c7 a8 b6 a4 b2 d1 f2 h1 g3 h5 g7 e6 f8 d7 b8 a6 b4 a2 c1 e2 g1 h3 f4 d3 c5 e4 c3 d5 e3 c4 e5 c6 d8 b7 a5 b3 a1 c2 e1 g2 h4 g6 h8 f7 h6 g4 h2 f1 d2 b1 a3 b5 a7 c8 e7 g8 f6 h7 g5 f3 d4.

Historia

Najwcześniejsze wzmianki o problemie skoczka szachowego pochodzą z IX wieku naszej ery. Rudrata, kaszmirski myśliciel i poeta, w sanskryckim traktacie pod nazwą *Kavyalankara* problem skoczka szachowego przedstawił jako skomplikowaną figurę poetycką, gdzie akcenty słów odpowiadają temu, jak przemieszcza się konik po planszy 4×8.

W 1759 roku szwajcarski matematyk Leonhard Euler (1707–1783), jeden z najwybitniejszych matematyków XVIII wieku, przedstawił pierwszą obszerną matematyczną analizę problemu skoczka szachowego na szachownicy 8×8 (7).



7. Leonhard Euler, autor portretu Jakob Emanuel Handmann



1	48	31	50	33	16	63	18
30	51	46	3	62	19	14	35
47	2	49	32	15	34	17	64
52	29	4	45	20	61	36	13
5	44	25	56	9	40	21	60
28	53	8	41	24	57	12	37
43	6	55	26	39	10	59	22
54	27	42	7	58	23	38	11

8. Jedna ze ścieżek skoczka szachowego

Poniżej jedno z możliwych rozwiązań (8). Kolejne ruchy numerowane są kolejnymi liczbami. Euler uzyskiwał rozwiązania, zakrywając pola, na których stanął skoczek, monetami. Zadanie to w literaturze nazywa się problemem szachowym Eulera.

Oprócz matematyki i fizyki prace Eulera dotyczyły również: teorii nawigacji, teorii muzyki, przypływów i odpływów morza, ruchów planet i komet, balistyki, dioptryki i wielu innych. Był jednym z najbardziej płodnych matematyków w historii: w sumie był autorem ok. 900 prac naukowych. Euler był wybitnym uczonym obdarzonym fenomenalną pamięcią. Pamiętał np. *Eneidę* Wergiliusza i mógł całą recytować słowo po słowie, pamiętając, na jakich słowach każda strona się kończy i następna zaczyna. Po utracie wzroku przez ostatnie 11 lat życia dyktował swoje książki i rozprawy. Jedną z ostatnich rozpraw napisał na pod dyktando Eulera, słynne *Kompletne wprowadzenie w algebrę* (*Vollständige Anleitung zur Algebra*) została przetłumaczona na wiele języków i stała się pierwowzorem podręcznika algebry.

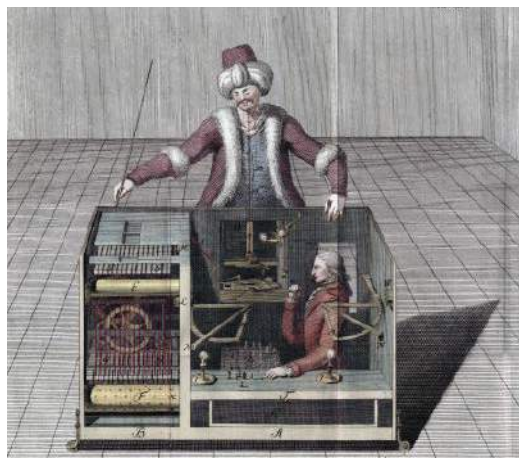
Mechaniczny Turek Kempelena

Pierwszy szachowy automat, zwany Mechanicznym Turkiem, zbudował w 1769 r. Farkas Kempelen (1734–1804), znany również jako Wolfgang von Kempelen. Nie była to prawdziwa maszyna, gdyż w urządzeniu ukryty był światowej klasy szachista. Oprócz gry w szachy automat Kempelena wykonywał bezbłędnie ruchy skoczka szachowego wg zamkniętej ścieżki, rozpoczynając od dowolnie wskazanego pola.

Mechaniczny Turek (9) był to tak naprawdę pseudoautomat w postaci drewnianej figury,



9. Mechaniczny Turek Wolfganga von Kempelena, źródło: Joseph Racknitz, Humboldt University Library



10. Zasada działania Mechanicznego Turka, źródło: Joseph Racknitz, Humboldt University Library

wielkości dorosłego człowieka. Przedstawiał postać wąsatego mężczyzny, ubranego w tradycyjny turecki strój, z turbanem na głowie, siedzącego przy biurku (skrzyni), na którym była ustawiona szachownica, wykonana z kości słoniowej. Mechaniczny Turek trzymał prawą rękę obok szachownicy, lewą zaś wykonywał posunięcia. W czasie gdy przeciwnik zastanawiał się nad swym posunięciem, Mechaniczny Turek unosił do ust fajkę, z której od czasu do czasu wydobywał się dym. Poruszał także głową i oczami. Kempelen przed rozpoczęciem każdego występu automatu pokazywał jego wnętrze, dowodząc w ten sposób, że nie ma tam nikogo. Automat znajdował się na kółkach, niemożliwe było więc wejście do maszyny przez otwór (zapadnię) w podłodze. Pokaz gry odbywał się z reguły w gabinecie

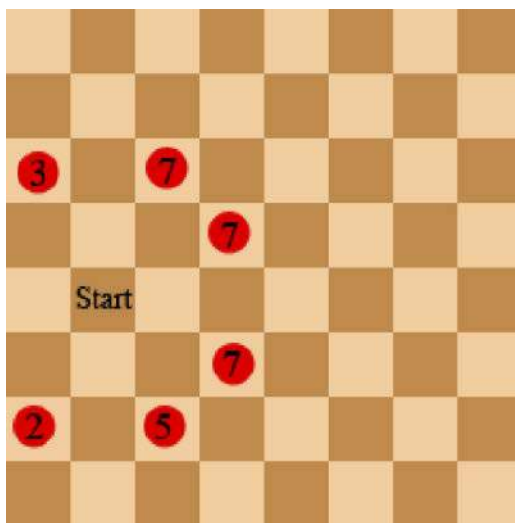
konstruktora. Widzowie stali za balustradą – aby nie dotykać szachownicy, Kempelen podchodził do wynalazku, nakręcał tryby i rozpoczynała się gra. Zazwyczaj automat grał białymi i wykonywał pierwszy ruch. Skinieniami głowy informował przeciwnika o groźbie szacha i szacha królowej przeciwnika. Gdy przeciwnik wykonywał nielegalne posunięcie, Mechaniczny Turek wykonywał przeczący gest głową i ustawiał figurę przeciwnika na poprzednią pozycję. Automat grał agresywnie i zazwyczaj pokonywał przeciwnika w ciągu 30 minut.

Teraz już wiemy, że była to genialna mistyfikacja, a automat poruszany był przez ukrytego pod biurkiem szachistę. Był nim przeważnie jeden z graczy o najwyższej klasie mistrzowskiej, którego Kempelen na czas swych występów „wynajmował”. Tajemnica działania Mechanicznego Turka pozostała niewyjaśniona przez niemal 60 lat, chociaż takie próby podejmowali różni wynalazcy i inżynierowie. Pierwszym z nich był Joseph Racknitz, który uczestniczył w pokazie automatu w Dreźnie w 1784 roku i swoje przemyślenia przedstawił w roku 1789 w książce *Über den Schachspieler des Herrn von Kempelen und dessen Nachbildung*. Książka zawiera szczegółowe szkice Racknitta pokazujące interpretację, jak automat działał (10).

Kiedy Kempelen otwierał kolejno poszczególne drzwiczki ukryty szachista szybko przesunął się w figurę Turka. Sprężyny, śruby i tryby zamontowane były tylko „na pokaz”, a część elementów była wykonana z tektury. Obserwację wykonywanych posunięć na szachownicy umożliwiały ukrytemu szachiście magnesy przymocowane na nitkach pod szachownicą. Kiedy szachista podnosił figurę, magnes opadał, a kiedy stawiał ją z powrotem na szachownicy inny unosił się i „przyklejał” do deski. Szachista siedzący we wnętrzu „automatu” wykonywał swoje ruchy za pomocą pantografu poruszającego ramieniem Mechanicznego Turka. Aby ukryty szachista dobrze widział w ciemnościach skrzyni, podczas pokazu współpracownicy zapalali na biurku dwa jasne kandelabry, niezależnie od warunków oświetlenia panującego we wnętrzu pomieszczenia, w którym odbywał się pokaz.

Reguła Warnsdorffa

Pierwsze algorytmy (sposoby postępowania prowadzące do rozwiązania problemu) wyznaczania ścieżki konika szachowego zaprezentowane zostały przez H.C. von Warnsdorffa w książce *Des Rösselsprungs einfachste und allgemeinste Lösung* w 1823 roku. Zgodnie z regułą Warnsdorffa skoczek



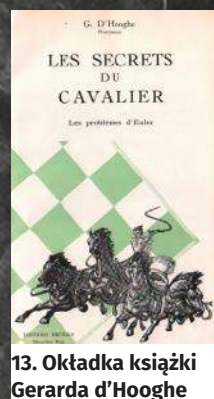
11. Reguła Warnsdorffa – skoczek powinien przejść na pole, z którego będzie miał najmniej ruchów do przodu

zawsze porusza się na pole, z którego ma najmniej wolnych pól (tj. jeszcze nieodwiedzonych), dostępnych na jego następny ruch (11). Ta zasada zapobiega wcześniejszemu odwiedzeniu jednego z dwóch pól, do których skoczek może dotrzeć z narożnika, aby nie mógł później uciec z narożnika. Oznacza to, że automatycznie pierwsze ruchy będą na zewnątrz i stopniowo do wewnątrz.

Zasada nie podaje żadnych instrukcji, co zrobić, jeśli jest kilka pól, z których jest kilka pól, do których można dotrzeć w następnym ruchu. Nawet jeśli istnieje rozwiązanie, zasada Warnsdorffa nie może zagwarantować, że zostanie ono znalezione, a skoczek dla dużych planszy coraz bardziej znajduje się w ślepym zaułku. Nawet na szachownicy ($n=8$) algorytm może zawieść, jeśli wybierze się niewłaściwą z kilku możliwych alternatyw, ale jest to bardzo mało prawdopodobne.

Program komputerowy, który znajduje trasę skoczka szachowego dla dowolnej pozycji wyjściowej, korzystając z reguły Warnsdorffa, został napisany przez Gordona Horsingtona i opublikowany w 1984 roku w książce *Century/Acorn User Book of Computer Puzzles*.

W latach pięćdziesiątych XX wieku farmaceuta Gerard d'Hooghe skonstruował maszynę, która demonstrowała trasę skoczka szachowego z dowolnie określonego pola wyjściowego. Jego tak zwany t'Zeepaard (zeepaard to konik morski w języku niderlandzkim) został pokazany publicznie podczas Olimpiady Szachowej w Lipsku w 1960 roku (12).



13. Okładka książki Gerarda d'Hooghe

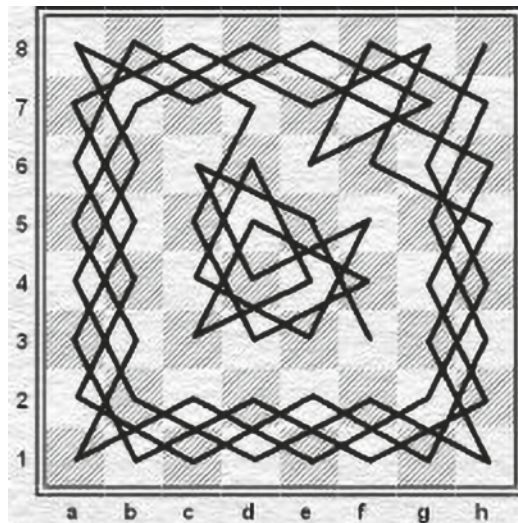
12. Gerard d'Hooghe, wynalazca maszyny prezentującej trasę skoczka szachowego, źródło: <https://bit.ly/31kfuZ8>

Gerard d'Hooghe był honorowym prezesem klubu szachowego i autorem książki *Les Secrets du Cavalier*, *Le Problème d'Euler*, wydanej w 1962 roku, w której opisuje swój wynalazek (13).

Jest szereg strategii prowadzących do znalezienia ścieżek konika szachowego. Przykładowo na rysunkach 14a i 14b widoczne są wyraźnie dwie fazy procesu rozwiązywania problemu ścieżki konika

34	49	22	11	36	39	24	
21	10	35	50	23	12	37	40
48	33	62	57	38	25	2	13
9	20	51	54	63	60	41	26
32	47	58	61	56	53	14	3
19	8	55	52	59	64	27	42
46	31	6	17	44	29	4	15
7	18	45	30	5	16	43	28

14a. Jeden ze sposobów znajdowania ścieżki konika szachowego (strategia „klamry”) opracowany przez francuskiego matematyka Abrahama de Moivre



14b. Ścieżka konika szachowego (strategia „klamry”) opracowana przez francuskiego matematyka Abrahama de Moivre

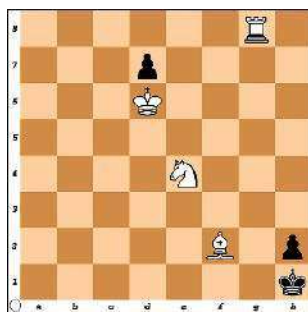
Tabela 1. Liczba możliwych ścieżek skoczka szachowego dla kwadratowych szachownic (źródło: <https://bit.ly/3nGCMck>)

n	wszystkich	zamkniętych
1	1	0
2	0	0
3	0	0
4	0	0
5	1728	0
6	6637920	9862
7	165575218320	0
8	19591828170979904	13267364410532

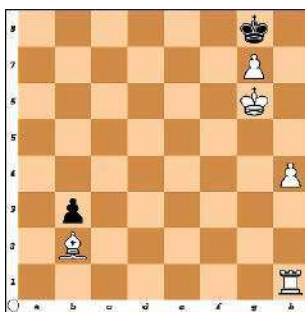
szachowego. W pierwszej kolejności figura przemie-
rza pola, które są najbardziej oddalone od środka
szachownicy, tworząc w ten sposób pewnego rodzaju
„ścianę”, która pozwala skupić się na pozostałym,
mniejszym obszarze do rozwiązania.

Inny przykład to strategia „sektorowa” polegająca
na podziale szachownicy na sektory. Rozwiązywa-
nie problemu tą metodą sprowadza się do odnale-
żenie ścieżki dla każdego z sektorów oraz spojenie
ich w jedną dla całej szachownicy. ■

Zadania do samodzielnego rozwiązania



Zadanie 1
15. W. Czeliznyj 1986
Mat w 3 posunięciach



Zadanie 2
16. W. Simonow 1992
Mat w 3 posunięciach

Rozwiązanie zadań z MT 10/2021

Zadanie 1

Rodolfo Bozzo – Dag Madsen, Gausdal
Int. 1992

Mat w 3 posunięciach

Rozwiązanie: 23.Hg7+ K:g7 24.Wg5+ Kh6
25.Gg7#

Zadanie 2

A. Galicki 1900

Mat w 3 posunięciach

Rozwiązanie: 1.Gf6 g:f6 2.Kf8 f5 3.Sf7#

Pastaman. Sztuka robienia makaronu krok po kroku

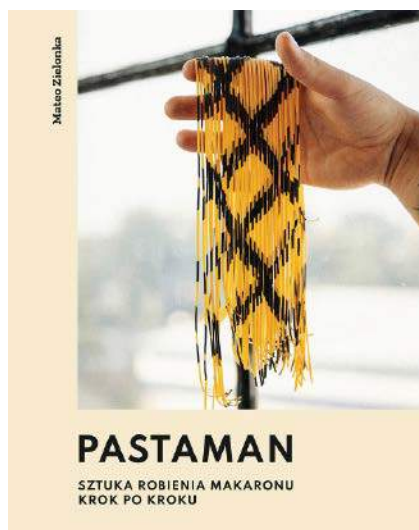
Mateo Zielonka

Wydawnictwo Buchmann, liczba stron: 160, cena: 59,99 zł

Mateo Zielonka, nazywany Pastamanem, potrafi przyrządzić
niezwykłe makaronowe cuda. Makaron w paski, w kropki,
czerwono-zielono-czarny, w każdym możliwym kształcie
– to spełnienie marzeń wielbicieli węglowodanów.

W swojej pierwszej książce Mateo podpowiada, jak zrobić
makaron własnoręcznie. Znalazły się tu przepisy na klasyczne
ciasto jajeczne oraz jego wegańskie odpowiedniki, a także
wskazówki, jak przygotować makaron o rozmaitych kolorach,
uzyskiwanych z naturalnych składników, i formach – od płat-
ków do lasagne przez pappardelle i tagliatelle aż po ravioli,
tortellini i inne makaronowe pierożki. To także 40 pomysłów
na pyszne sosy i dania na kolację, którymi można zachwycić
bliskich oraz gości. Przepisom towarzyszą piękne zdjęcia
i jasne instrukcje.

Niezależnie od tego, czy dopiero zaczynasz swoją przygodę z makaronem, czy jesteś pasjonatem z doświadczeniem, wsłuchaj się w rady mistrza i przygotuj z nim własne makaronowe dzieła sztuki!





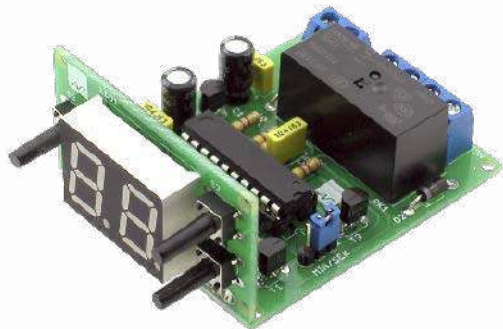
W naszej rubryce „Elektronika dla Ciebie” co miesiąc zachęcamy Cię, drogi Czytelniku, do wykonywania prostych projektów – zabawek, gadżetów itp. Każdy to potrafi. Opis jest zawsze zrozumiały dla nieelektroników, a montaż niemal intuicyjny. A jeśli złapiesz bakcyłę pasji elektronicznej, na co liczymy,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

Uniwersalny timer od 0 do 99 minut



Prosty układ timera przeznaczony do odliczania „w dół” zadanych odcinków czasu w zakresie 0...99 s lub 0...99 min. Wbudowany przekaźnik o znacznej obciążalności prądowej oraz prosta, intuicyjna obsługa kwalifikuje układ do realizacji funkcji czasowych w nieskomplikowanych układach automatyki domowej oraz przemysłowej.

Schemat ideowy timera pokazano na **rysunku 1**. Urządzenie jest przystosowane do zasilania napięciem stałym z zakresu 8...12 V. Dioda prostownicza D1 zabezpiecza układ przed niewłaściwą polaryzacją. Napięcie zasilające jest stabilizowane przez U1, natomiast kondensatory C1... C4 zapewniają odpowiednie jego filtrowanie. Pracą timera steruje mikrokontroler ATtiny26 taktowany wewnętrznym sygnałem zegarowym. Stan pracy jest obrazowany na podwójnym wyświetlaczu siedmiosegmentowym ze wspólną anodą. Katody 2-cyfrowego, multipleksowanego wyświetlacza LED dołączone zostały przez rezystory R5...R12 do portów PA0...PA7 mikrokontrolera. Funkcję kluczy załączających zasilanie wyświetlaczy pełnią tranzystory T1 i T2 sterowane z portów PB3 i PB4. Na potrzeby wprowadzenia nastaw oraz obsługi timera urządzenie zostało wyposażone w 3 przyciski oznaczone S1, S2 i S3. Sygnały z przycisków doprowadzono do portów PB0 i PB1 oraz PB6, poziomem aktywnym jest logiczne „0”. Jako układ wykonawczy zastosowano przekaźnik typu RM84P12 (cewka 12 VDC, styki 8 A/230 VAC).



Aby rozszerzyć funkcjonalność timera, na złączach NC i NO wyprowadzono styki przekaźnika normalnie zwarte i normalnie otwarte.

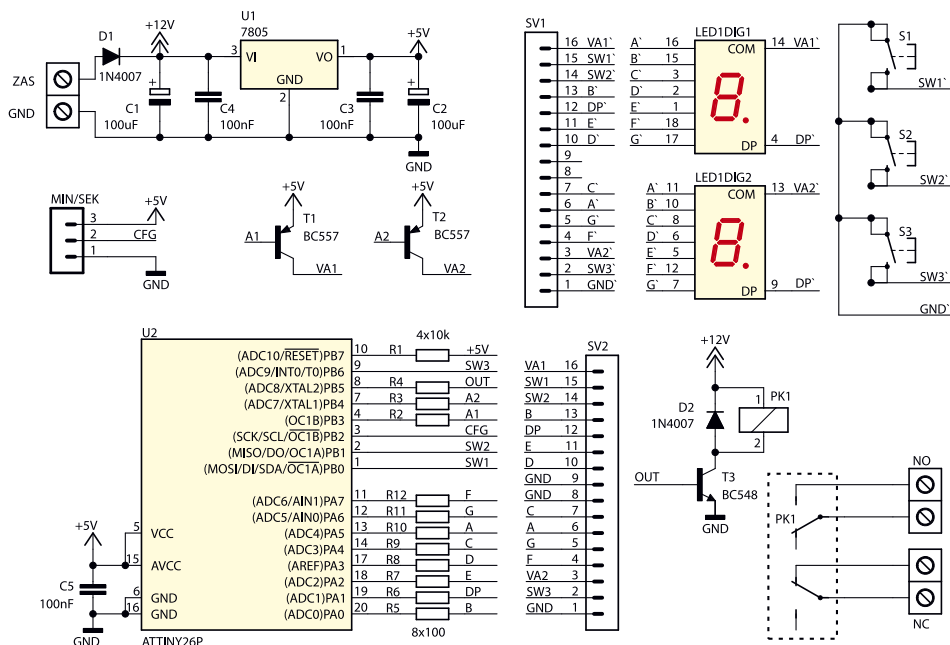
Timer należy zmontować na dwóch płytkach drukowanych. Montaż układu jest typowy i nie powinien przysporzyć problemów, przebiega on w sposób standardowy, zaczynając od elementów najmniej-szych, a kończąc na największych. Po zmontowaniu obydwu płytek należy połączyć je ze sobą za pomocą kątownej listwy szpilek goldpin. Układ zmontowany bezbłędnie będzie działał od razu po włączeniu napięcia zasilającego. Przy sterowaniu obciążeniem o znacznej mocy należy zwrócić uwagę na obciążenie styków przekaźnika oraz ścieżek płytki drukowanej. Aby poprawić ich obciążalność, można dodatkowo pocynować odsłonięte ścieżki lub jeszcze lepiej ułożyć na nich i przylutować drut miedziany.

Do wyboru jednostki czasu, jaka ma być odmierzana (sekundy lub minuty), służy umieszczona na płytce głównej zworka MIN/SEK. Przyciski S1 i S2 służą do zwiększania i zmniejszania wartości,



Wszystkie niezbędne części do tego projektu zawiera kit AVT3200 B, w cenie 46 zł, dostępny pod adresem: <https://sklep.avt.pl/avt3200.html>

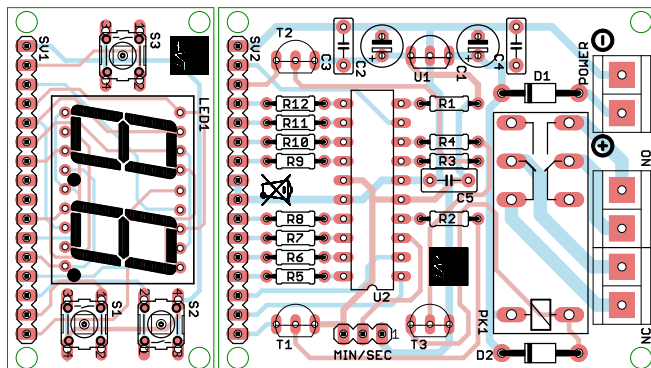




1

natomiast do uruchomienia odliczania służy przycisk S3. Każde przyciśnięcie S2 spowoduje zwiększenie, a przyciśnięcie S1 zmniejszenie wartości. Aby zmiana wartości następowała szybciej, bez potrzeby wielokrotnego przyciskania, należy dany przycisk przytrzymać dłużej. Ustawiona wartość zapisywana jest w pamięci nieulotnej, dzięki temu po ponownym włączeniu zasilania układu nie trzeba jej na nowo wprowadzać. O tym, czy timer pracuje w trybie minutowym, czy sekundowym, informuje kropka przy cyfrze jedności. Jej pojawienie się oznacza, że timer odliczać będzie minuty (zworka w pozycji MIN), natomiast jej wygaszenie (zworka w pozycji SEK) oznacza, iż timer skonfigurowany został do odliczania sekund. Miganie kropki w każdym z trybów sygnalizuje pracę timera. Po uruchomieniu timera w każdej chwili poprzez przyciśnięcie przycisku S3 możliwe jest zatrzymanie

2



odliczania czasu. W tej sytuacji cyfry na wyświetlaczu zaczną migać. W trybie tym timer oczekuje na ponowne, krótkie naciśnięcie przycisku S3 bądź jego dłuższe przytrzymanie, po którym nastąpi powrót urządzenia do początkowej wartości. Używając timer, należy mieć świadomość, że odmierzenie czasu może być obciążone pewną niedokładnością, w szczególności dotyczy to pracy w zakresie minutowym. ■

KIT Y AVT PRZEDSTAWIA

AVTEDU

Zostań mistrzem lutownicy!

Poznaj całą serię

#AVTEDU #NaukaLutowania #KityAVT





Kiedy ma się w domu od lat modelarski warsztat – choćby najmniejszy, zwykle nie ma większego kłopotu ze znalezieniem lub docięciem odpowiedniej listwy na szkolny latawiec. Gorzej, jeśli takiego zaplecza nie ma – a wykonanie modelu latawca to na dodatek zadanie domowe, o którym przypomniało się w niedzielę. Wtedy rozwiązania materiałowe i konstrukcyjne zastosowane w dzisiejszym modelu mogą uratować co najmniej honor, o ocenie już nie wspominając.

Latawiec szaszłykowy

Nie tylko Kite Man, czyli zwiad powietrzny przed wynalezieniem dronów

Latawce płaskie to najstarszy typ statków powietrznych – również osobowych! Wynalezione może nawet około 770 r. p.n.e., choć dalekie od idealnej stabilności i zwykle wyposażone w mało praktyczny ogon – już od zarania swoich dziejów były wykorzystywane również do lotów załogowych – istnieje wiele przekazów pisemnych dotyczących wykorzystania jednopłaszczyznowych latawców załogowych do celów wojskowych (w uniwersum DC zielonym romboidalnym latawcem przemieszcza się Kite Man). W kwestii aparatów cięższych od powietrza w realnym świecie w zasadzie nie było innego wyboru aż do czasu wynalezienia w 1894 roku przez Lawrence Hargrave'a znacznie stabilniejszego latawca skrzynkowego (ang. box kite). Udostępnienie ówczesnym pionierom aeronautyki informacji o doświadczeniach Australijczyka było wyraźnym katalizatorem szybkiego rozwoju lotnictwa na przełomie XIX i XX wieku, co w kilka lat zaowocowało budową pierwszych faktycznie latających samolotów. Wojskowi obserwatorzy nie zrezygnowali jednak szybko z latawców – przesiedli się tylko na znacznie pewniejsze latawce skrzynkowe. Na nich również

(aż do niedawna) montowano często aparaty fotograficzne i kamery filmowe. Dopiero w ostatnich latach zastąpiły je coraz doskonalsze i coraz tańsze drony.

Nie wszędzie jednak można wykorzystywać latawce – nawet zabawkowe. W krajach rządzonych przez fundamentalistów islamskich latawce są zakazane z przyczyn religijnych. Zakaz puszczania latawców obowiązywał w Afganistanie (i prawdopodobnie znów wrócił wraz z odzyskiwaniem władzy przez talibów).



1. Najstarsze wzmianki pisane, dotyczące załogowych lotów latawcowych, wspominają wyczyn chińskiego generała Han Hsina, który w 200 roku p.n.e. wzniósł się latawcem nad mury obleganego miasta, by ocenić konieczną długość tunelu, jaki wykonywali jego żołnierze. Później jeszcze wielu śmiałków szło w jego ślady – w wojsku taka obserwacja była szczególnie cenna, stąd latawce oddziały obserwacyjne działały nawet podczas II wojny światowej. **2.** Współcześnie również lata się latawcami i choć zwykle są to projekty bezszkieletowe typu spadochronowego lub paralotniowego, to latawce płaskie także są używane w tym celu. Płaskim latawcem latał m.in. główny bohater francuskiej komedii „Wielkie wakacje” (oryg. „Les grandes vacances”) z 1967 r. grany przez Louisa de Funèsa. **3.** O zakazie zabawy dzisiejszym modelem w Afganistanie opowiada m.in. książka i nakręcony na jej podstawie film pt. „Chłopiec z latawcem”. Zakaz zniesiono wprawdzie w 2002 roku, pod wpływem działań wojsk koalicji zachodniej, ale sytuacja w regionie znów się diametralnie zmieniła...

Podobne restrykcje z powodów religijnych wprowadzono w Pakistanie – antylatawcowe przepisy wprowadziły w 2001 władze prowincji Pendżab, a podtrzymał je pakistański Sąd Najwyższy w 2005 roku. W oficjalnym uzasadnieniu podano jednak, że latawce mogą ułatwić dokonanie zamachu terrorystycznego.

Cieszymy się zatem szerokimi możliwościami budowy i zabawy latawcami, przystępując do budowy dzisiejszego modelu.

Projekt – nowy, choć na sprawdzonych rozwiązaniach

Dzisiejszy model jest zbliżony wielkością i wielkością detali konstrukcyjnych do opisywanego „Na warsztacie” we wrześniu 2007 roku latawca FEST (polecam jego opis również zobaczyć – link w tabeli...). Jest jednak od swojego starszego brata lżejszy i delikatniejszy – więc lata nawet lepiej od standardowego FEST-a, pokrytego zwykłym papierem – szczególnie przy małych wiatrach. To także dobry projekt do zajęć w pracowniach modelarskich, świetlicach i na zajęciach technicznych. Większość z zastosowanych tu rozwiązań została już pozytywnie zweryfikowana przez kilkanaście lat użytkowania przez setki młodszych i starszych zawodników.

Materiały i narzędzia (oraz ich ekwiwalenty)

Przede wszystkim potrzebny będzie delikatny i lekki papier – optymalnie bibuła gładka (dwa arkusze A3 – najczęściej dostępne w handlu). Może to też być

4. Kadr z Mistrzostw Aeroklubu Wrocławskiego Modeli Latawców z latawcem typu FAST z logo Polski, jakie Ministerstwo Spraw Zagranicznych PR zatwierdziło do użytku w 2002 roku. Jednak ten logotyp nie przyjął się powszechnie – także z powodu pewnych zarzutów profesjonalnych grafików oraz faktu, że mimo wszystko latawce raczej nie są kojarzone ze stałością, z jakim powinno kojarzyć się państwo. Na latawcach jednak wciąż chętnie jest stosowany



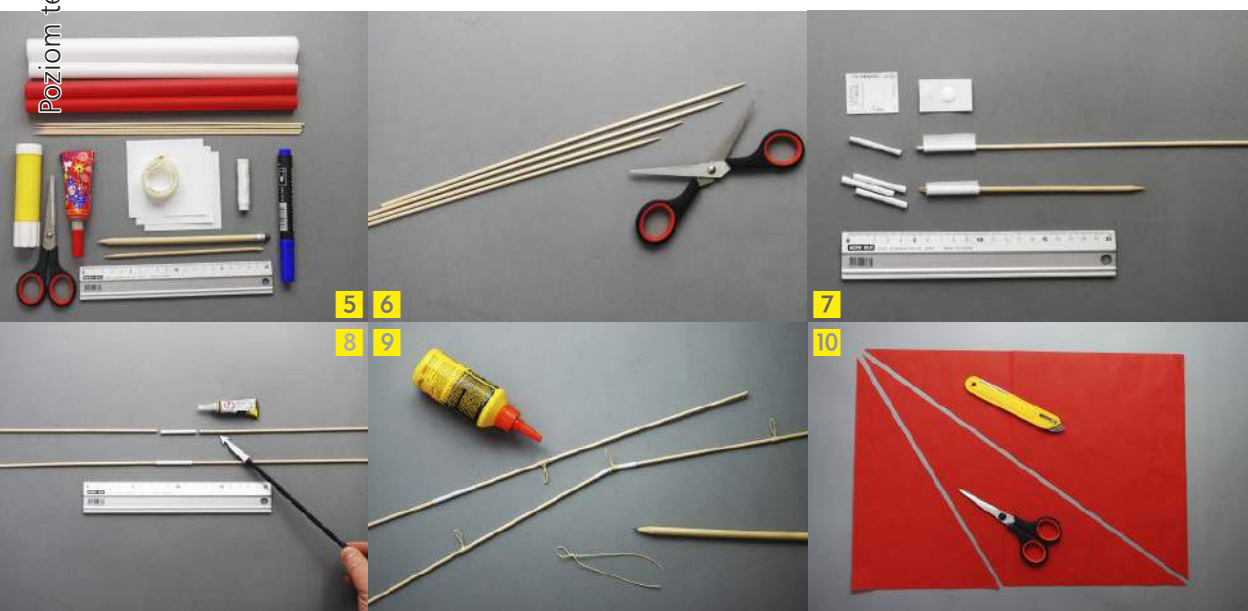


papier, jaki czasem dodawany jest do opakowań butów lub koszul. Jeśli jest pognieciony – nie szkodzi – da się go w miarę dobrze wyprasować żelazkiem. W najgorszym razie może być to nawet papier śniadaniowy (choć będzie cięższy). Drugim niewralgicznym materiałem będą patyczki szaszłykowe – najlepiej o średnicy 3 mm i długości 30 cm – jeśli do dyspozycji będą krótsze, trzeba będzie dostosować wielkość poszycia. Na ogon przyda się bibuła karbowana (krepina), ale od biedy do tego projektu można wykorzystać nawet biały papier toaletowy. Z dodatkowych materiałów potrzebna będzie mocna nić (taka, która nie zerwie się przy pierwszej próbie wytrzymałości). Zwykły papier, klej w sztyfcie, klej POV (typu wiskolowego), cyjanoakrylowy (typu super

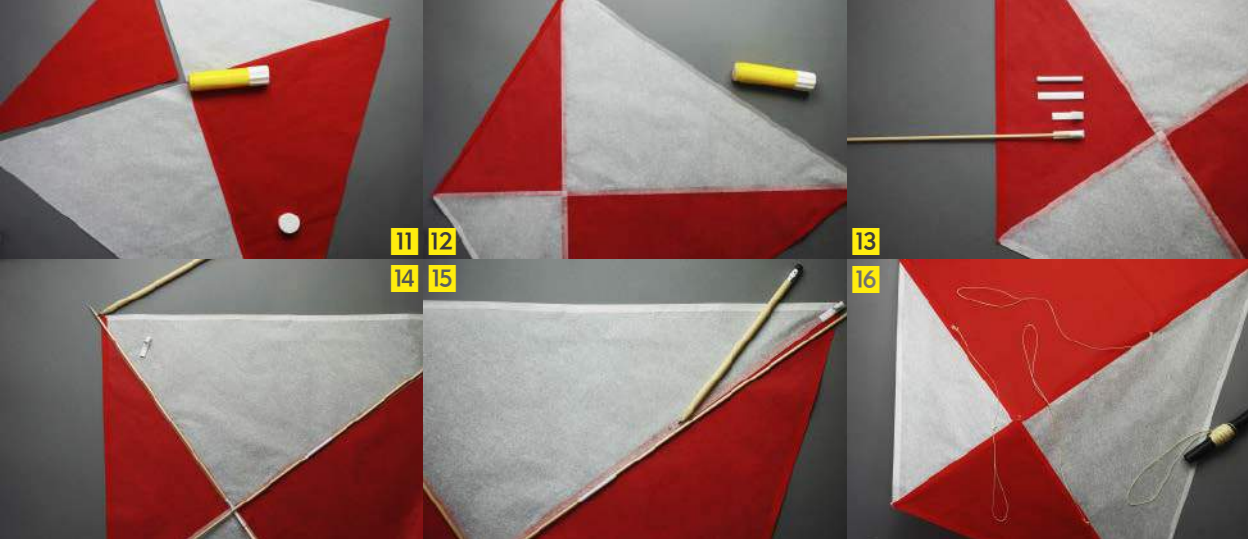
glue), kęs papieru ściernego gradacji 150–200, ołówek, dłuższy liniał, nożyczki i nożyk do tapet w zasadzie wyczerpują listę potrzebnych materiałów i narzędzi. Markera można użyć do zdobienia albo (jeśli już jest wypisany) do wykonania całkiem wygodnego zwijadła holu.

Montaż modelu

Kiedy wszystkie materiały i narzędzia zostały w końcu skompletowane, polecam zacząć od wykonania papierowych mufek (łączników) prętów szaszłykowych i podobnych rurek na kieszenie poszycia. Oba te rodzaje akcesoriów do dzisiejszego modelu wymagają zwinięcia prostokątnych formatek zwykłego papieru (o wymiarach 30×40 mm) w rurki o długości



5. Podstawowe materiały i narzędzia do budowy dzisiejszego modelu – w tym dwa arkusze bibuły gładkiej formatu A3. Do kompletu przyda się jeszcze klej cyjanoakrylowy i bibuła marszczona na prosty, taśmowy ogon. **6.** Zaostrzone końce patyczków trzeba będzie docelowo usunąć – można to zrobić piłą o drobnych ząbkach, nacinając dookoła nożykiem modelarskim – ale można również wykonać najpierw rowek jednym z ostrzy nożyczek (kładąc obrabianą końcówkę pręta i jedno z rozwiedzionych ostrzy nożyczek na krawędź stołu), a następnie delikatnie odtamując zbędną część pręta oraz fazując delikatnie nowy koniec elementu. Łącząc najpierw niezastrzone końce prętów, z docięciem ostrych końcówek można się jednak wstrzymać do czasu przymierzenia szkieletu do poszycia (vide fotografia 14). **7.** Łączniki i kieszenie dla szkieletu to papierowe rurki, ściśle zwinięte i czysto sklejone na odpowiedniej średnicy szablonie. **8.** Prawidłowe wklejenie prętów w mufkę wymaga oznaczenia połowy głębokości łącznika na każdym z prętów. Przy tej wielkości latawca papierowe mufki trzeba koniecznie zaimpregnować klejem C-A (Uwaga! łącznik poziomej belki szkieletu będzie nieco przegięty przed zaimpregnowaniem – vide kolejne zdjęcie). **9.** Na szkielecie znajdują się punkty zaczepu uzdy (może być dwu- lub trzylinkowa – w prezentowanym projekcie będzie możliwość nawet i późniejszego wyboru). Wygodnie jest zrobić pętelki na fragmentach nici, które potem przykleja się do konstrukcji w oznaczonych na rysunku montażowym miejscach – tym razem klejem POV, żeby już nie kleić palców. **10.** Wielkość tego latawca wynika w zasadzie z wielkości formatek folii (A3). Żeby wygodnie było skleić poszycie w szachownicę, przyjęty został taki podział arkusza (podstawa pierwszego trójkąta z lewej niech ma 20 cm). Bibułę należy uważnie wyciąć nożycami albo odpowiednim nożem przy długim liniale.



11. Poszycie skleja się na zakładkę – optymalnie szerokości ok. 5–6 mm. Warto jednak wcześniej rozważyć, która strona ma być prawa, by po tej właśnie ciemniejsze arkusze były klejone na wierzchu. **12.** Po sklejeniu brytów poszycia, najpierw na sucho, a następnie z pomocą kleju w sztyfcie, klei się mankiet szerokości 5–6 mm, wzmacniający zewnętrzne krawędzie bibuły. Tradycyjne wklejanie nici w mankiet nie jest w tym modelu konieczne. **13.** Pora na kieszenie, w które wkładane będą końce prętów szkieletu (szerszy opis ich montażu w tekście). **14.** Z założenia długość szkieletu ma odpowiadać długości przekątnych poszycia, więc do tej długości należy dociąć pręty szkieletu – koniecznie symetrycznie w przypadku poprzecznej, wygiętej części konstrukcji. Końce prętów należy zawsze szfzować, aby nie uszkadzały kieszeni. **15.** Kieszenie przykleja się klejem w sztyfcie, nasadzone na belki szkieletu, do lewej strony naprężonego poszycia. Ponieważ drewniane pręty nie są wklejane w papier, będzie można je wyciągać – co nieraz się jeszcze przyda. Na poszyciu oznacza i przebija się otwory na linki uzdy (przy tak klejonej bibule warto wykorzystać we wszystkich miejscach pas podwójnej warstwy). **16.** W przetknięte przez poszycie pętelki zaczepów przewleka się kolejno: najpierw pętelkę linki uzdy, a następnie przez tę pętelkę koniec tej samej linki i trzeba zaciągnąć. Czynności powtórzyć dla odpowiednio długich linek uzdy (dwóch lub trzech, w zależności od wybranej opcji).

40 mm: minimum dwóch sztuk na pręcie szaszłykowym – oraz min. 4 poprawnie wykonanych sztuk zwiniętych na pręcie (ew. rurce po lizaku lub wiertle) średnicy ok. 4 mm. Skąd te minimalne (a nie dokładne) liczby? Ano stąd, że mało kto rodzi się z wrodzonym talentem do zwijania cienkich rurek z papieru i zapewne pierwsze rurki nie przejdą wewnętrznej kontroli technicznej – nie należy się tym martwić, lecz przyjąć, że to raczej reguła niż wyjątek (to trochę jak z pierwszym naleśnikiem). Do klejenia rurek wygodniej użyć kleju wiskolowego, nakładanego na wewnętrzną stronę końcowej krawędzi zwijanej rurki. Ważne – rurki powinny ściśle przylegać do szablonu (pręta 3 lub 4 mm).

Przedłużanie prętów szaszłykowych z wykorzystaniem właśnie zwiniętych węższych rurek będzie łatwiejsze, jeśli końce prętów delikatnie szfzuje się wkładane w mufkę końcówki oraz oznacza 2 cm – czyli tyle, ile pręta powinno znaleźć się w łączniku. Poprzeczną część szkieletu warto będzie wygiąć, aby uzyskać usztywniające lot podniesione końcówki skrzydeł (usztywnienie) latawca. Można to zrobić (na wzór FEST-a – vide ramka) poprzez wygięcie w łuk sklejonych dobrze prętów – ale można też odpowiednio

przebiegnąć mufkę i utrwalić jej kształt, impregnując papier łącznika klejem cyjanoakrylowym. Druga opcja została zastosowana i z sukcesem przetestowana w opisywanym modelu. Jedną z rzeczy, o których należy jednak pamiętać, to fakt, że okrągłe pręty łatwiej obracają się w kieszeniach niż prostokątne w przekroju listwy – ma to szczególne znaczenie w przypadku poprzecznej belki szkieletu, która miała tendencje do obrotu (szczególnie przy trzypunktowej uździe). Rozwiązaniem jest ciasne związanie skrzyżowania belek lub wręcz jego sklejenie, jeśli nie przewiduje się konieczności demontażu modelu.

Poszycie dla tego dość sporego latawca (jak na wymiary prętów do szaszłyków) można kupić w jednym arkuszu wielkości 50×70 cm lub (częściej) w zrolowanych „wielosztukach” formatu A3. W drugim przypadku można będzie zastosować różnokolorowe arkusze, dobierając kształty i kolory brytów odpowiednio do potrzeb i możliwości. W modelu, którego budowę przedstawia ten artykuł, wykorzystano dwa różnokolorowe arkusze i dość łatwy w wykonaniu, a przy tym chyba całkiem ciekawy podział poszycia w szachownicę. Do klejenia bibuły, ze względu na możliwość przebarwienia, najlepiej używać kleju

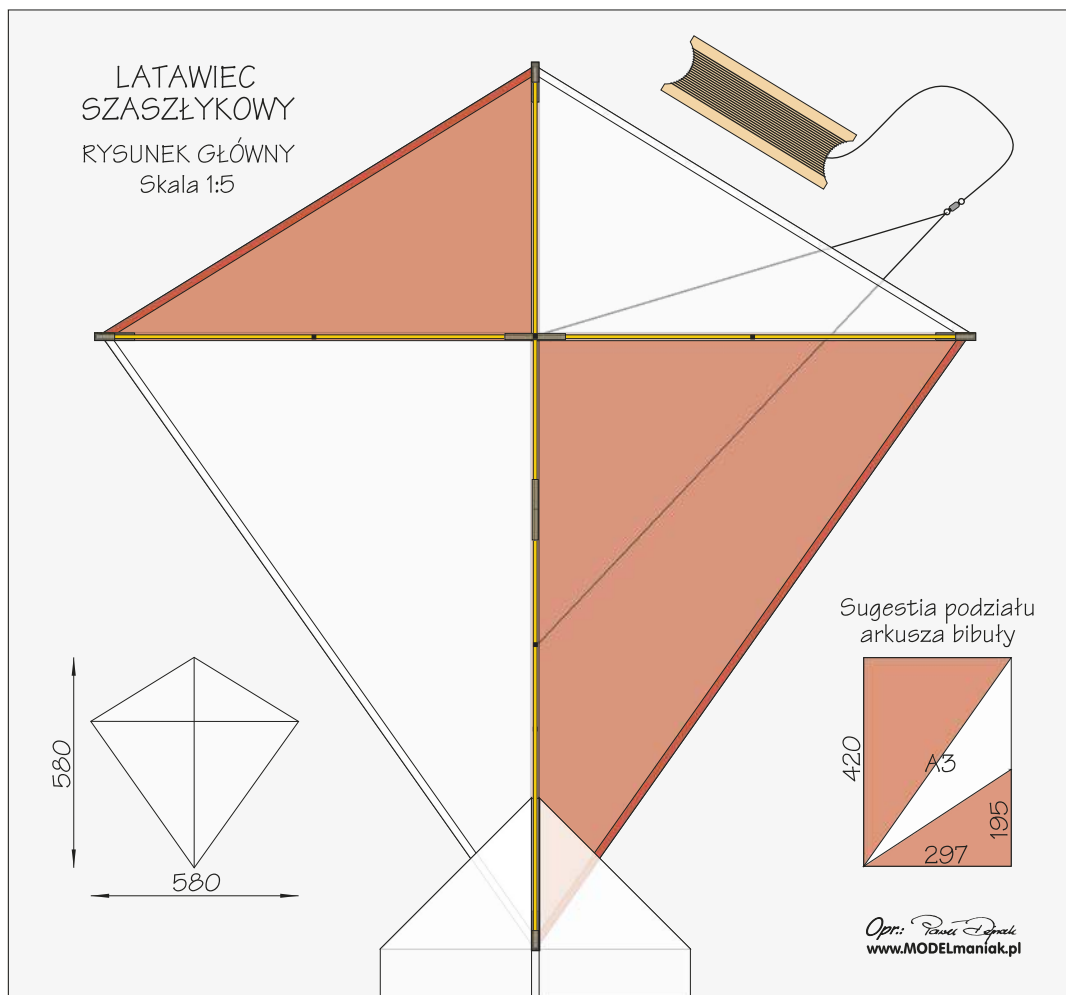


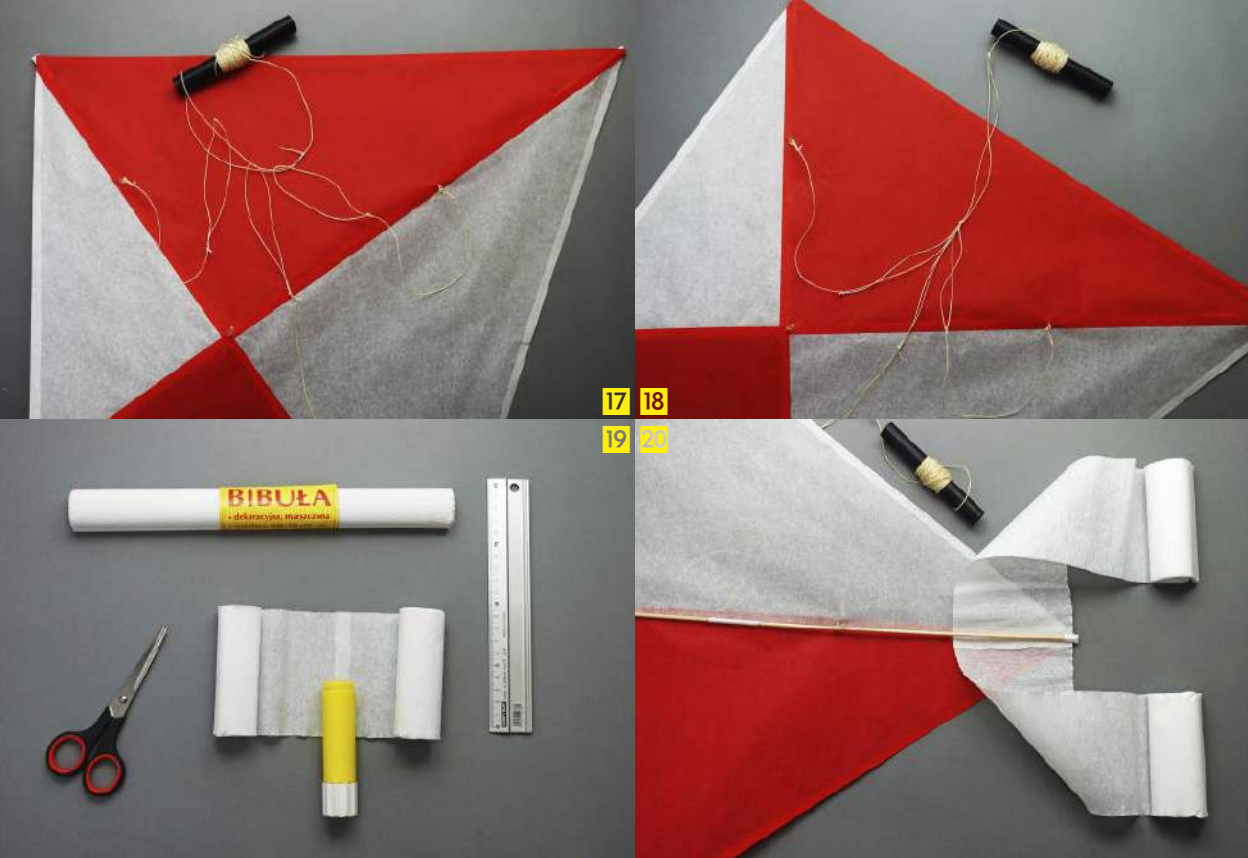
w sztyfcie (można go nawet dosuszać żelazkiem, o ile takie jest na podorędziu). Warto jednak wcześniej dobrze przemyśleć kolejność nakładania brytów (11 i 12). Przy mniejszych wielkościach latawców można nie wzmocniać krawędzi poszycia dodatkowym mankietem, jednak tu warto już go wykonać, naklejając 5–6 mm krawędzi bibuły na lewą (górną docelowo) stronę poszycia. Dodatkowo poszycie można ozdobić, używając przede wszystkim markerów i flamastrów (farby wodne w zasadzie się do tego celu nie nadają). Inną opcją jest naklejenie (raczej niedużych) liter, numerów, zdobień z ciemniejszej bibuły gładkiej lub lekkiej folii.

Kieszenie wykonuje się stosunkowo łatwo, choć należy się liczyć z pierwszymi nieudanymi egzemplarzami. Ich zaletą jest przede wszystkim możliwość wyciągania z nich szkieletu – do transportu,

przechowywania lub ew. remontu czy modernizacji modelu. Wykonuje się je w następującej kolejności:

1. ciasne zwinięcie i sklejenie rurki z prostokąta (30×40 mm) zwykłego papieru (gramatury ok. 100 g/m²) na szablonie ok. 1 mm szerszym niż pręt szkieletu (czyli w tym przypadku na szablonie ok. 4 mm);
2. spłaszczenie papierowej rurki ze szwem (łączeniem) na górze (na środku);
3. zagięcie elementu w 1/3 długości, w taki sposób, aby szef spłaszczonej rurki znalazł się wewnątrz;
4. włożenie do końca krótszej części papierowego elementu sfazowanej końcówki drewnianego pręta szkieletu latawca (trochę jak nogi do końca kaptcia – stąd też wzięło się drugie nieformalne określenie tego elementu we wrocławskich pracowniach MDK im. Kopernika);





17 18

19 20

17. Pętliki na przyholowych końcach linek uzdy nie muszą być aż tak duże jak na zdjęciu (w odróżnieniu od pętli na końcu holu) – tutaj lepiej jednak pokazać zasadę łączenia z linką holu. Pętlę końcówki holu (koniecznie dużą) przewleka się przez wszystkie pętliki linek uzdy, a następnie przez nią przekłada się zwiądadło. Po zaciągnięciu nitki połączenie uzdy z holem jest gotowe do lotu. **18.** Tak bezsuptowo wykonane połączenie można zdemonstrować bez cięcia (a jeśli nawet szybciej będzie ciąć, to wystarczy tylko jedną pętelkę – przy holu). Jeszcze doskonalszym rozwiązaniem byłoby umieszczenie między pętlami linek uzdy a holem wędkarskiego krętlika, przeciwdziałającego skręcaniu się linki holu – jednak w aeroklubowych zawodach latawców wszelkie metalowe elementy latawców są regulaminowo zabronione. **19.** Ogon latawca najszybciej można zrobić z dwóch pełnych taśm bibuły marszczącej – $2 \times (10 \text{ cm} \times 200 \text{ cm})$. Niech zakładka (sklejka) ma 15–20 mm szerokości. Również do klejenia tej bibuły lepiej używać najbardziej suchego z popularnych klejów – w sztyfcie. **20.** Ogonu taśmowego można używać na dwa sposoby – jako pojedynczej taśmy ze wzmocnioną końcówką i z zarobioną pętelką, zamocowaną do pętliki na dolnej końcówki szkieletu – lub (jak na zdjęciu) przez przełożenie taśmy pod pręt ogonowy, przez co latawiec otrzyma dwa ogony połowy długości taśmy. Pierwsze rozwiązanie jest lepsze aerodynamicznie, drugie jest natomiast praktyczniejsze (szczególnie podczas zawodów, w których łatwo o zerwanie ogona na ziemi).

5. przyklejenie „podeszw kapci” założonych na listewki szkieletu do naprężonego poszycia.

Zaczepty uzdy, wykonane w tym latawcu, są charakterystyczne dla większych konstrukcji – ale i tu mają sporo zalet: nie są skomplikowane w wykonaniu, umożliwiają odcięcie linek uzdy (do wymiany lub przełożenia, kiedy latawiec będzie służył jako dekoracja pokoju lub modelarni). Sposób wykonania prostych, pętelkowych zaczepów, mocowanych do szkieletu omotką i klejem pokazuje **fotografia 9**.

Uzda w tym latawcu może być dwulinkowa (pierwsza linka zaczynałaby się na skrzyżowaniu

szkieletu, druga w połowie odległości pomiędzy skrzyżowaniem a ogonem latawca) lub – nieco bardziej stabilizująca latawiec – trójlinkowa (przednie linki w połowie każdego ze skrzydeł, tylna jak w dwupunktowej). W opisywanym latawcu zaczepty umożliwiają zamienne stosowanie obu uzd – przyda się to dla pokazania różnic młodym adeptom modelarstwa lotniczego w prowadzonych przez autora pracowniach. Pewną innowacją jest w tym projekcie wykorzystanie linek uzd, obustronnie zakończonych pętelkami o średnicy minimalnej ok. 10 mm. Umożliwia to łatwe mocowanie i demontaż tych linek,



zarówno do konstrukcji, jak i holu – bez konieczności ponownego wiązania na szkieletcie.

Hol, z możliwie trudnej do zerwania, ale lekkiej nici, zwinięty na rurce o średnicy ok. 15 mm (choćby ze zużytego markera) również ma pętlę – w odróżnieniu jednak od pętelek uzdy, ta powinna umożliwiać swobodne przełożenie przez nią holu. Będzie to potrzebne przy połączeniu z uzdą, co pokazują **fotografie 17 i 18**.

Taśmowy ogon bibułowy nie wymaga komentarza dłuższego niż przy **fotografiach 19–20**. Można też wykonać bardziej pracochłonny i kłopotliwy, choć przez wielu uznawany za bardziej klasyczny ogon z kokardkami (choćby z listków papieru toaletowego, kiedy nie ma bibuły). Wybór opcji pozostawmy wykonawcom – z zastrzeżeniem, że nie warto oszczędzać na długości i gęstości elementu, który odpowiada za stateczność modelu w locie.

Oblatywanie modelu

Podobnie jak w przypadku każdego innego modelu latawca, na miejsce lotu należy wybrać dostępny teren (najlepiej płaski), bez przypadkowych przechodniów, drzew, słupów, linii elektrycznych, dróg kołowych, lotnisk i innych przeszkód stanowiących zagrożenie dla latawca – ale przede wszystkim ludzi i zwierząt. Wielkość i masa własna opisanego latawca pozwala na udane loty zarówno na zewnątrz przy delikatnym i umiarkowanym wietrze, jak i w szkolnej sali gimnastycznej (pod warunkiem stosowania się do zasad bezpieczeństwa i poleceń koordynatora lotów).

Przed startem na zewnątrz należy określić kierunek wiatru (np. podrzucając kilka żdźbeł trawy) i pamiętać, że każdy szanujący się lotnik startuje zawsze pod wiatr, jeśli tylko jest taka możliwość. Przy starcie z dłuższego holu przyda się przeszkolony pomocnik

Warto zobaczyć również

Doskonalwsze wersje modeli latawców opracowane przez autora tego projektu (<https://modelmianiak.pl>) znajdziesz m.in. także w archiwalnych wydaniach miesięcznika „Młody Technik” (zapytaj o nie w swojej szkolnej bibliotece)



2007/09 FEST – latawiec festynowy (prosty, łatwy w budowie latawiec o typowej wielkości do zabawy na zewnątrz) – jego plany znajdziesz na stronie: <https://bit.ly/3BBp2LI> (cały artykuł)



2011/08 DeLF II czyli Demontowalny Latawiec Foliowy wersja II (cały artykuł), <https://bit.ly/3CcxguXv>



2012/08 Składany jak parasolka latawiec skrzynkowy typu Pottera o długości 1 m, <https://bit.ly/3CFiQDH> (cały artykuł)



2015/09 Minilatawiec MSL (papierowy, wielkości ~A4) – jego plany znajdziesz na stronie autora: <https://modelmianiak.pl>



2017/09 Mikrokajtki (latawce kieszonkowe, ~A6), <https://bit.ly/3bxiOlR> (cały artykuł)



2018/09 Latawiec całopapierowy („samolotowy” wielkości ~A3), <https://bit.ly/3GS57Mu> (cały artykuł)

- <https://bit.ly/3jWwuLx>
- <https://bit.ly/3byuDaY> – Lu Ban
- <https://bit.ly/3pWErNC> – Kite History
- <https://bit.ly/3w9bru8>
- <https://bit.ly/3CE0yTt>
- <https://bit.ly/3q2Ckd5>



(np. w tym jak bezpiecznie trzymać i wypuszczać latawiec). Startujący operator powinien jednocześnie obserwować drogę przed sobą i zachowanie latawca, odpowiednio dostosowując swoją prędkość (lub oddawanie i wybieranie holu) do zachowania modelu w powietrzu.

Prawidłowo wykonany latawiec powinien po starcie wznosić się możliwie stabilnie w górę, aż do osiągnięcia ustalonego kąta na holu. Nieprawidłowości lotu oraz sposób przeciwdziałania tymże:

- ciągle skręcanie w jedną stronę; należy sprawdzić i w razie potrzeby skorygować: kierunek startu w stosunku do osi wiatru, długość przednich linek trzypunktowej uzdy lub równowagę obu skrzydeł po podwieszeniu na dwupunktowej uździe;
- szybkie wahania latawca wokół punktu zaczepienia ogona: prawdopodobnie zbyt duża prędkość względem strug powietrza: zwolnić bieg, jeśli był konieczny, ew. chwilowo oddawać hol, ew. zmienić uzdę z dwu- na trzypunktową, ew. wstrzymać loty do czasu zmniejszenia prędkości wiatru;
- latawiec wiruje w kręgach większych niż jego rozpiętość: za krótki ogon – przedłużyć;
- latawiec lata zbyt nisko (pod małym kątem linki holu w stosunku do ziemi) – dostosować długość tylnej linki uzdy: kąt ten można regulować w zależności od wiatru, skracając lub wydłużając

ogonową linkę uzdy (zwiększa lub zmniejsza się kąt natarcia płaszczyzny latawca w stosunku do strug powietrza). Przy mocniejszym wietrze kąt może być mniejszy (dłuższa linka ogonowa uzdy), przy słabszych nieco mniejszy (krótsza linka ogonowa uzdy). Należy dążyć do uzyskania możliwie najbardziej stabilnego lotu pod możliwie dużym kątem holu względem płaszczyzny lotniska.

Wersje rozwojowe

To drugi udany projekt latawca płaskiego (klasy FLP wg regulaminu Aeroklubu Polskiego) w naszych pracowniach, wykorzystujący przedłużane pręty szaszłykowe. Jeśli wykonawcy poradzą sobie ze zwijaniem rurek, dalsze prace są już raczej proste, a i walory eksploatacyjne latawca bardzo obiecujące. Nic więc nie stoi na przeszkodzie, żeby sprawdzone tu rozwiązania materiałowe i technologiczne wykorzystać w projektach latawców płaskich o innych kształtach oraz w projektach latawców skrzynkowych i specjalnych podobnej wielkości. Najprawdopodobniej powstaną takie i we wrocławskich pracowniach – może pojawią się również w przyszłych wydaniach „Na warsztacie”...

Dobrych wiatrów i tyłuż udanych startów co lądowań! ■

Paweł Dejnak



Szkoła Wynalazców

dozwolone do lat 15

Mielicie zadanie: zaproponować sposób podniesienia sprzedaży nożyków do warzyw tak, aby niezbędne zmiany w produkcji nic nie kosztowały.

Kiedy gospodyni kupuje kolejny nożyk? Zasadniczo wtedy, gdy aktualny nożyk staje się nieprzydatny: mocno stępiony, popękana rękojeść itp. I oczywiście wtedy, gdy go zgubi. Gospodyni kupi nowy nożyk, gdy będzie szczególnie atrakcyjny: ładny kształt, wygodna rękojeść i przyjemny dobór kolorów.

Co może tu zaproponować ekspert od marketingu i reklamy? W dzisiejszych czasach nożyki – jeśli idzie o ostrze – są na ogół bardzo dobre. Tępiemy je sobie sami, korzystając z podręcznych ostrzałek, w których nóż jest przeciągany przez szczelinę przyrządu, gdzie jest obustronnie „obstrugiwany” przez ostrza ze specjalnie twardej stali lub węglików spiekanych. Niestety tą metodą nie można noża ostrzyć „w nieskończoność”, bo traci on prawidłową geometrię kształtu ostrza. Wtedy trzeba naostrzyć nóż na szlifierce albo... kupić nowy. Pozostaje więc jedyna możliwość: zgubienie nożyka. Jak można gospodyniom „ułatwić” gubienie nożyka? Otóż można!

Wystarczy, gdy zmienimy kolory rękojeści (pamiętamy, że jest ona dwubarwna) na kolory identyczne z kolorem obierek ziemniaczanych, tzn. jasno kremowe, z beżowoszarzym. Tak jak wygląda łupina: jasna od wewnątrz i szarobeżowa na zewnątrz. Gospodynie pracują w pośpiechu, nożyk wpada pomiędzy łupiny i tak upodobniony, często jest wyrzucany do odpadów bio! I taka była rada eksperta od marketingu i reklamy. Inna rzecz, że taki zestaw kolorów rękojeści nie jest zbyt atrakcyjny. To nie szkodzi: po pewnym czasie pojawią się kolejne nożyki z ładnymi, żywymi barwami rękojeści i interes będzie się kręcił dalej!

Musimy sobie zdać sprawę, że w warunkach współczesnego przemysłu i handlu, my – klienci – jesteśmy substancją „obrabianą” przez liczne grono speców od marketingu. Przy zakupie czegokolwiek należy więc przede wszystkim myśleć! Zobaczmy, co wymyślili nasi czytelnicy:

Jacek Zieliński (2 pkt.) Najprościej byłoby wykonać ostrze w taki sposób, żeby po okresie bardzo dobrej ostrości – szybko się tępiło. Trzeba by tu zastosować ostrze z wąskim pasmem utwardzanej stali i z nieco bardziej miękkim pasmem tuż przy strefie przyostrzowej, tak żeby nóż się łatwo szczybił.

Byłoby to działanie zgodne z polityką niektórych firm, żeby nie produkować zbyt trwałych produktów, bo po nasyceniu rynku kto będzie kupował następne? Jest to problem, bo dziś wyprodukować produkt o bardzo wysokiej trwałości jest stosunkowo łatwo, ale faktycznie: co dalej? Dlatego mamy np. treпки base-nowe, wytrzymujące 2–3 miesiące eksploatacji i wiele podobnych wyrobów. Wbrew pozorom pomysł Jacka musiałby podnieść koszt produkcji, bo produkować ostrza na skalę masową i jednocześnie z ostrzem podobnym do japońskiej katany, wcale nie jest łatwo.

Mateusz Konarski (2 pkt.) Mateusz wie, że rękojeści dwubarwne produkuje się w formach wtryskowych, najczęściej tak, że w jednej formie wykonuje się wtrysk tworzywa jednej barwy i po przełożeniu nożyka do drugiej formy wtryskuje się tworzywo z drugą barwą. Należałoby zadbać o to, żeby te dwa tworzywa niedokładnie się sklejały i po pewnym czasie rękojeść po prostu rozlatywałaby się na kawałki.

To na pewno niemal bezkosztowy sposób. Korekcja procesu polegałaby jedynie na dodaniu do jednego z tworzyw jakiejś substancji pogarszającej sklejenie się obu barwnych wtrysków. Może odrobina talku?

Dziękuję za udział z konkursie i zapraszam do następnych.

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

Ranking Szkoły Wynalazców

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

Ranking Klubu Wynalazców

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

Nowe zadanie

Wszyscy wiedzą, że Robert Lewandowski – nasza gwiazda piłki nożnej – w każdym kolejnym kontrakcie zastrzega, że „musi” występować jako numer „9”. Jeden z wybitnych piłkarzy chilijskich – Iván Zamorano – miał też taką ambicję: być zawsze „dziewiątką”. Z jakichś tam powodów pewnego razu okazało się, że musi zagrać mecz jako numer 18. Zamorano pomyślał i umiejętnie obszedł zakaz i wystąpił jako numer 9, nie naruszając przy tym warunków trochę nieprecyzyjnie sformułowanej umowy.

Spróbujcie pomyśleć. Warunki umowy Zamorano musiał spełnić, ale w tym jednym meczu normalnie

było to niemożliwe, więc trzeba było sięgnąć do fortelu! Wasze zadanie:

Zaproponować sposób obejścia nakazu wystąpienia z numerem 18 zamiast „kultowego” numeru 9.

Spróbujcie wejść w sposób myślenia Zamorano, którego ambicja została mocno podrażniona, bo z jednej strony mógł zdecydować się nie zagrać w tym meczu, ale z drugiej – musiał wykonywać warunki kontraktu, a więc – grać!

Oczekujemy Waszych propozycji do końca grudnia 2021 r. Wszystkim życzymy błyskotliwego myślenia i dobrego pomysłu dla Zamorano.

Klub Wynalazców

bez ograniczeń wieku

Mieścieście zadanie: zaproponować jakiś nowy, w pełni bezpieczny system cumowania wielkich statków i okrętów do nadbrzeża, nienarażający obsługi i marynarzy na wypadki.

Zastanówmy się chwilę: na czym polega niebezpieczeństwo cumowania dużej jednostki pływającej. Może tu zaistnieć kilka przypadków:

- Lina, zarzucona już pętlą na poler, zrywa się i jej napięte końce uderzają w obsługę nadbrzeża. Mimo że liny cumownicze są obliczane wytrzymałościowo, to jednak energia dużego statku, odsuwającego się od nadbrzeża, jest ogromna, a w związku z tym lina może być naprężona do granic wytrzymałości i w skrajnych przypadkach może się zerwać. Napięta lina w chwili zerwania zachowuje się jak ciężka łuku: z dużą prędkością odskakuje w nieprzewidywalny sposób.
- Po narzuceniu pętli liny na poler obsługujący ją cumowniczy może zechcieć ją poprawić, a statek w tym momencie lekko się oddala. Jeżeli dłoń cumowniczego znajdzie się wewnątrz pętli, może zostać dociśnięta do polera i... nieszczęście gotowe!

Pomijamy tu dalsze niebezpieczeństwa, wynikające np. z uderzenia statku o nadbrzeże. Zajmujemy się wyłącznie cumą i ewentualnie jej odmianą – szpringiem. Nie zajmujemy się też współczesnym elektronicznym monitorowaniem procesu cumowania, z pomiarem odległości i prędkości statku od nadbrzeża. Skupiamy się na systemie: lina, poler, statek. Zasadniczym sposobem cumowania jest narzucanie liny z oczkiem zaplecionym na końcu na poler (1).



„Główki” polerów mają różne kształty i bywa, że lina w fazie wybierania luzu – żeby dociągnąć statek do nadbrzeża – za pomocą wciągarki okrętowej, zsuwa się i obsługa, chcąc ją złapać, ryzykuje całość dłoni. Obecnie niemal powszechną praktyką jest rzucanie ze statku tzw. rzutki – cienkiej, lekkiej linki, do której dowiązana jest lina cumownicza. Za pomocą rzutki przyciąga się oko cumy i nakłada na poler. W zasadzie nie powinno dochodzić do wypadków, które jednak się zdarzają. Jak wykluczyć w 100% możliwość zaistnienia wypadku? To było Waszym zadaniem, Zobaczmy, co koledzy wymyślili.

Zbigniew Przygodzki (4 pkt.) uważa, że bezpieczniejsze byłyby cumy, zamocowane jednym końcem na stałe na polerze statkowym lub na knadze, drugi



koniec założony już na bęben wciągarki. Zarzucałoby się całą, stosunkowo luźną linę, dzięki czemu odpadłoby „celowanie” okiem cumy na poler. Po zarzuceniu liny wciągarka wybrałaby luz i koniec. Obsługa brzegowa nie miałaby kontaktu z naprężoną liną. Wciągarka powinna mieć ogranicznik momentu, pozwalający kontrolować naciąg liny.

Rozwiązanie wydaje się dobre. Podwójna lina do cumowania nie stanowiłaby problemu; liny są stosunkowo tanie. Warunek bezpiecznego cumowania: „jak najdalej z rękami od cumy”, byłby tu spełniony.

Robert Pająk (5 pkt.) Proponuje zastosowanie amortyzatora, wpiętego w linę cumowniczą, który zdecydowanie zmniejszyłby energię szarpnięcia wywołanego przez przypadkową falę lub prąd morski.

Zamiast amortyzatora jako odrębnej części być może udałoby się wykonać linę mającą własności amortyzatora.

Bardzo dobra koncepcja; zabezpieczałaby przed niebezpiecznym zerwaniem się liny cumowniczej. Czy dałoby się wykonać linę z wewnętrznym tłumieniem, mającą cechy amortyzatora, to odrębne zagadnienie, bardzo ciekawe i godne zastanowienia.

Mirosław Lubicz (3 pkt.) pisze: lina cumownicza jest najbardziej niebezpieczna dla obsługi nabrzeżnej, znajdującej się w pobliżu polerów, na które trzeba zarzucić pętlę. Mirosław proponuje, żeby obsługę wyposażyć w specjalne kleszcze, które umożliwią pracę z cumą, ale dłonie będą z daleka od polera i jego główki.

Prawdopodobnie dałoby się takie kleszcze zaprojektować, ale czy pogorszona wygoda

sługiwanie się nimi okupiłaby zwiększone bezpieczeństwo? Ludzie na ogół myślą: wypadki są i będą, ale mnie się nic takiego nie zdarzy.

Kolegom gratuluję i zapraszam do następnych zadań.

Nowe zadanie

Zadanie nietypowe, bo „kto dziś czyta książki”? Jednakże na szczęście są tacy i im właśnie trzeba pomóc.

Osoby „starej daty” często posiadają księgozbiory liczące 2–3 tysiące woluminów. Zdarzają się jednak wyjątki także wśród młodych ludzi. Problemem przy tak bogatej bibliotece domowej jest szybkie odnalezienie potrzebnej książki. Ogólnie wiemy, gdzie ona powinna być, ale jak trzeba ją wyciągnąć, to robi się problem. Oczywiście są systemy katalogowania i profesjonalne biblioteki je mają, ale – jak wiadomo – liczba książek rośnie niemal niespostrzeżenie i jak już staje się problemem, wizja skatalogowania zbioru podnosi włosy na głowie! Spróbujcie więc:

Zaproponować system poszukiwania potrzebnej książki w zbiorze liczącym np. 2000 pozycji, tak jednak, żeby pominąć żmudny proces katalogowania zbioru.

Oczywiście pamiętamy wszyscy, że mamy już niemal w każdym domu komputer i że jest to na pewno zadanie dla niego. No, ale wprowadzić dane do pamięci komputera, to już zadanie katorżnicze i nikt tego nie lubi. Coś jednak trzeba zrobić. Co?

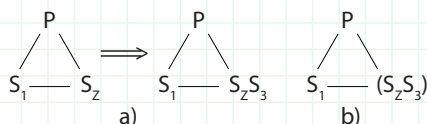
Wszystkim życzę dobrych pomysłów i zapraszam do zmierzania się z tym tematem. Termin nadsyłania propozycji – do końca grudnia br.

Vademecum Młodego Wynalazcy

Analiza wepolowa rozwijała się w miarę rosnących potrzeb przemysłu. Powstały zasady i prawa, ustalające warunki wykonania konkretnych analiz. Wiadomo, że rzeczywistość techniczna zmusza do korzystania z kompleksowych pól i substancji. Oznacza to, że w praktyce mamy do czynienia np. z polami elektromagnetycznymi, z substancjami złożonymi jak np. bimetale, a nawet stal damasceńska itp. Bardzo często substancją bywa uzupełniana przez inną substancję, łatwo poddającą się sterowaniu jakimś dogodnym polem. Oto prosty przykład:

Podczas modernizacji wydziału produkcyjnego fabryki zapalek zainstalowano nowe, wysokowydajne maszyny produkcyjne. Okazało się, że stworzyło to kolejne problemy: po pierwsze, zapalki w pudełkach

muszą być zorientowane łebkami w jedną stronę i w każdym pudełku powinna ich być jednakowa liczba. Na oko widać, że zapalki, małe i lekkie, są raczej niesterowalne. Chwytyki do takich maleńkich przedmiotów to spory problem. W zapisie wepolowym sytuacja wyjściowa będzie wyglądała następująco:



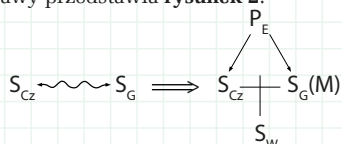
Technolog ma do dyspozycji tylko substancję S_1 – masę bezwładnie nasypanych do pojemnika zapalek. Z pierwszego prawa analizy wepolowej wynika, że system zawierający tylko jeden element

– substancję – nie może działać. Wiadomo więc, że aby powstało prawidłowe wepole, musi być dobudowana substancja i pole. Zapalka składa się z drewnianka i główki zapalającej. Obie te substancje nie są sterowalne – w sensie nadania im uporządkowanej orientacji – polem elektrostatycznym, mechanicznym ani magnetycznym. W takiej sytuacji musimy dodać do zapalki drugą substancję, dającą się łatwo sterować. Najprościej będzie do masy zapalającej, jeszcze na etapie jej przygotowywania, dodać niewielką ilość proszku ferromagnetycznego, czyli po prostu „pudru” żelaznego – S3 – **rysunek 1b**. Magnes ustawiony nad transporterem, podającym zapalki do pakowania, ustawia je wszystkie łebkami do góry!

W ogólnym sensie dodanie trzeciej substancji może się odbyć tak jak w przypadku zapalek, wewnątrz (ferroproszek w masie zapalającej) lub zewnątrz, np. blaszka miedziana przynitowana do blaszki stalowej dają element – „bimetal”. Jeżeli substancję wprowadzamy do wnętrza substancji podstawowej, to mamy do czynienia z „wewnątrznie kompleksowym wepolem”. Jeśli montujemy dwie oddzielne blaszki, żeby uzyskać bimetal – mówimy o „zewnątrznie kompleksowym wepolu”.

Uniwersalność symboliki wepoli i ich elastyczność jest ogromna. Potwierdzają to cztery przykłady wzięte z różnych dziedzin techniki:

Przykład 1. Zakład produkujący maszyny budowlane miał problem z koparkami. Na niektórych gruntach, szczególnie gliniasto-ilastych i nieco wilgotnych koparka nabierała kolejną porcję urobku, wykonywała obrót tak, aby czerpak znalazł się nad skrzynią ładunkową wywrotki, operator odwracał kosz i... nic. Głina z ilem „nie chciały” wypadać. Próbowano najrozmaitszych metod: zmiana kształtu czerpaka, pokrywanie czerpaka teflonem, wprowadzenia drgań przy rozładunku i wszystko daremne. Teflon był drogi i szybko się ścierał, drgania tak dużej masy wymagałyibratorów o dużej mocy. Zapis wepolowy tej sytuacji i jej naprawy przedstawia **rysunek 2**.



Sytuacja wyjściowa: dwie substancje: S_{Cz} – czerpak i S_G – glina oddziałują na siebie nawzajem niekorzystnie: glina przykleja się do czerpaka. Na schemacie wyjściowym widać, że potrzebne jest jakieś pole, żeby wepole było kompletne. Pole to powinno doprowadzić do rozdzielenia powierzchni czerpaka od gliny. Substancją, która już jest w systemie i jest „darmowa” – jest

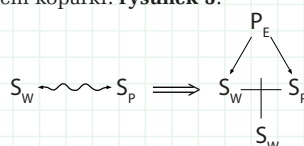
wilgotna glina z ilem, a konkretnie woda. Należy znaleźć takie pole, które przemieści wodę z urobku na powierzchnię czerpaka, a więc oddzieli blachę czerpaka od masy gliny. W banku informacyjnym TRIZ znajduje się m.in. wykaz efektów fizycznych, przydatnych w rozwiązywaniu sprzeczności w systemach. W tym wykazie znajdujemy grupę efektów: „przemieszczanie substancji” i elektroosmozę jako jedną z metod.

Należy więc zaopatrzyć czerpak koparki w elektrodę, odizolowaną od materiału kosza, i na kosz oraz tę elektrodę podać napięcie stałe z akumulatora koparki. Z wilgotnej gliny wydzieli się woda i przemieści się... no właśnie! gdzie się przemieści? Będzie to zależało od odczynu chemicznego gleby, tzn. czy jest to gleba kwaśna, czy raczej zasadowa i od biegunowości elektrod, podłączonych do kosza i listwy kontaktowej zamocowanej na powierzchni kosza. Zmiana biegunowości tego przyłącza może być realizowana jednym prostym przełącznikiem, co pozwala dostosować się do różnych rodzajów gleby.

Przykład 2. Zakład produkujący pompy wirnikowe do różnych zastosowań otrzymał zamówienie na serię pomp, do przetwarzania pulpy masy papierowej w fabryce papieru. Zadanie proste. Zakład miał w programie pompy do przetwarzania cieczy zanieczyszczonych, a więc pulpa nie powinna stanowić problemu. Okazało się jednak, że jeśli zawartość celulozy w pulpie nie przekraczała ok. 10%, to pompy działały prawidłowo. Jeśli jednak zawartość celulozy sięgała 12%, zaczynały się problemy. Pompy przestawały podawać pulpę. Rozebrano taką pompę i okazało się, że na powierzchni wirnika utworzyła się „klucha” celulozowa, czyli wirnik stał się okrągłą bryłą i oczywiście nie mógł pompować.

Po dwóch latach prób dobierania materiału na wirnik, zmianach kształtu łopatek, wyglądania ich powierzchni itp. doszło do przypadkowego spotkania ze specjalistą z TRIZ. Problem został rozwiązany w 2 godziny, wliczając w to czas wykonania uproszczonej próby (bez użycia pompy). Czy ten „trizowiec” był taki genialny? Oczywiście nie. Po prostu metodyka TRIZ, to bardzo silne narzędzie dla pokonywania „nierozwiązywalnych” problemów technicznych.

Ciekawe, że schemat wepolowy i metoda rozwiązania problemu pompy były niemal identyczne z problemem koparki: **rysunek 3**.



Odczytujemy to następująco: dwie substancje: wirnik i pulpa, wzajemnie niekorzystnie na siebie



działają. W wepolu brakuje pola. Oczywiście wchodzi w grę pole elektryczne i efekt elektroosmozy.

Wykonano próbę uproszczoną: dwie elektrody zanurzone w 12% pulpie podłączono do źródła prądu stałego. Po niespełna sekundzie jedna z elektrod była czysta, druga oblepiona pulpą. Oczywiście podczas technicznego zastosowania tego zjawiska w rzeczywistej pompie można było:

- dobrać materiał wirnika i korpusu tak, aby uzyskać efekt siły elektromotorycznej, wynikający z różnicy naturalnego potencjału obu materiałów,
- dobrać baterię termooogniwi tak, aby naturalne ciepło silnika pompy mogło dać napięcie do wywołania efektu elektroosmozy.

Jak widać idea to nie wszystko; konieczne są próby, doświadczenia, które niewątpliwie ujawnią nowe problemy...

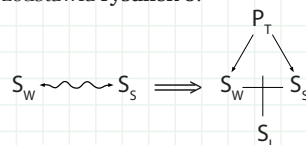
Przykład 3. W zakładzie ceramiki budowlanej pracował przenośnik ślimakowy, podający masę ceramiczną: wilgotną mieszkankę gliny kaolinowej z dodatkami do maszyny formującej kształtki. Jak już czytelnik się domyśla: podajnik nie działał prawidłowo, tzn. czasem, nawet dość długo, podawał glinę bez problemów, a czasem „nie chciał”.

Analiza, schemat wepolowy i metoda rozwiązania problemu identyczne jak w przypadku pompy.

Przykład 4. Jednym z ciekawszych obiektów pływających jest wodolot. Jego podwodne „skrzydła” stawiały mały opór czołowy, dzięki czemu może on poruszać się z prędkością rzędu 90 km/h. Nie ma róży bez kolców i tu też sprawdza się to stare powiedzenie. Problemem podwodnych skrzydeł wodolotu jest kawitacja. Pomijając szczegóły tego zjawiska, dość powiedzieć, że powoduje ono wżery na powierzchni skrzydeł, co prowadzi z czasem do ich dewastacji. Oczywiście podejmowano liczne próby: lepsze, bardziej wytrzymałe blachy lub odwrotnie: pokrycie skrzydeł warstwą gumy, jednakże nie dawało to zadowalającego rezultatu. Sytuacja wyjściowa to wepole: **rysunek 4.**

$$S_w \rightsquigarrow S_s$$

Jest to więc niepełne wepole, które odczytujemy następująco: substancja S_w – woda, działa niekorzystnie na substancję S_s – skrzydła, wywołując efekt kawitacji. Skrzydła dzięki zniszczonej powierzchni zwiększają opór wiskotyczny w wodzie. Wepole jest niepełne, należy dobudować pole, które wygeneruje substancję potrzebną do oddzielenia wody od skrzydeł. Substancja powinna być obecna w systemie, bardzo tania lub darmowa. Takiej substancji mamy tu pod dostatkiem: jest to oczywiście woda. Żeby woda mogła spełnić zadanie ochrony skrzydeł, konieczna jest zmiana jej stanu skupienia. Oczywiście się więc staje, że potrzebne jest pole termiczne, chłodzące, które może zamrozić cieniutką warstwę wody, mającą kontakt ze skrzydłami. W rezultacie kawitacyjnie będzie niszczonej lód, a ten nas już nie martwi. Schemat wepolowy nowej sytuacji przedstawia **rysunek 5.**



Pole termiczne – w tym przypadku generowane przez węzownice wbudowane w skrzydła i podłączone do agregatu chłodniczego, powoduje namarzanie cienkiej warstewki lodu na powierzchni blach skrzydeł. Jest to więc kolejny przykład zastosowania tej samej idei do różnych zagadnień.

Schematy wepolowe przypominają tzw. wzory strukturalne, stosowane w chemii. Słynny wzór strukturalny benzenu, na który Kekulé wpadł podczas poobiedniej drzemki, jest tego najbardziej wymownym i znanym przykładem. Symbolika wepolowa jest niezbędna do korzystania z tzw. standardów rozwiązywania zadań wynalazczych.

Prezes Klubu Wynalazców
instruktor TRIZ
Jan Boratyński

W cieniu terapeutki

Anna Krystaszek

Wydawnictwo MUZA S.A., cena: 39,99 zł

Dżgnięta nożem policjantka z Częstochowy ginie na służbie. Rozpoznała sprawcę, nie zdążyła jednak o tym nikomu powiedzieć. Od tamtej chwili minęło sześć lat, a zabójca wciąż nie został schwytany. Również przed sześciu laty doszło do wypadku samochodowego, w którym zginęli mąż i syn Magdy, psychoterapeutki. Długo nie mogła się pozbierać po tej tragedii. Któregoś dnia pod domem terapeutki zostaje znalezione ciało zamordowanej kobiety... Jak się okazuje, to jedna z jej pacjentek. Kim jest zabójca? Czy zbrodnia ma jakiś związek z wypadkami sprzed lat? Jaką rolę w tej sprawie odegrał Adam? Młody prokurator Jan Hejda próbuje dotrzeć do prawdy.



AR

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

Nieustannie czekamy na Wasze pomysły ulepszeń, innowacji, zmian.

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

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

Wystarczy podać numer wybranego pomysłu.

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

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

Pomysł miesiąca 11/2021

Obiecująca wydaje nam się koncepcja przeprojektowania patelni. Patrzymy na to nieco szerzej niż tylko jako na problem naleśników. Obrótowa i generalnie ruchoma część, nazwijmy to, robocza, patelnia ma w sobie potencjał, który pozwala myśleć o nowym podejściu do korzystania z tego sprzętu kuchennego.

Autorką pomysłu jest Monika Rachwalska.

„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:

1 Jan Wojton zauważył, że wysocy mężczyźni mają problem z zajęciem miejsca za kierownicą niskiego samochodu, takiego jak np. Toyota Corolla Liftback. Niemodny, ale bardzo dobry aerodynamicznie kształt nadwozia powoduje problem. Wsiadając, wielu uderza głową w górną krawędź otworu drzwiowego, a nawet w wewnętrzne lustro wsteczne. Znajomy Janka podpatrzył, jak wsiadają do takiego wozu panie, w wąskich minispódniczkach. Siadają one bokiem na fotelu i już siedząc i mając głowę w samochodzie, wciągają obie nogi do wnętrza. Sposób „niemęski” – odpada. Janek proponuje przeprojektować dachy samochodów tak, aby mogły unosić się nieco do góry na teleskopowych słupach i – po kłopotcie! **Projektanci karoserii samochodowych rysowali i widzieli już niejedno, Takiego pomysłu chyba nikt nie zrealizował. A że nowość i wygoda klientów przede wszystkim, to kto wie?**

2 Miłosz Kujawa – ma w domu ciągły kłopot z zawijaniem się krawędzi dywanu w wyniku wjeżdżania i wyjeżdżania kółek fotela komputerowego. Przykleić dywanu do parkietu nie można (władze domowe nie zezwolą). Miłosz kupił kółka o większej średnicy i utożyskowane na łożyskach kulkowych, co nieco złagodziło problem, ale go nie zlikwidowało. Miłosz proponuje rewolucję: zastąpienie kółek talerzykami o zaokrąglonych krawędziach i wypolerowanej powierzchni roboczej. Taki talerzyk dawałby o wiele mniejszy nacisk jednostkowy i po dywanie gładko by się przesuwiał. **Bardzo interesująca propozycja, ważna dla biur i takich mieszkań, w których podłoga jest pokryta wykładziną dywanową. Kółka istotnie maltretują dywan i wykładziny, a dzieje się to głównie podczas zwrotu „niby samonastawnych” kółek.**

3 Monika Rachwalska nadesłała genialny pomysł na pieczenie naleśników. Problemem jest uzyskanie naleśników o równej grubości i równo wypieczonych. Monika proponuje wykonanie patelni jako obrotowych, napędzanych małym silniczkiem elektrycznym. Porcję ciasta wlewałoby się na środek patelni, a wiry jej ruch rozprowadziłyby ciasto na całą powierzchnię. Pieczenie na obracającej się patelni zapewniłoby też równe upieczenie naleśnika.

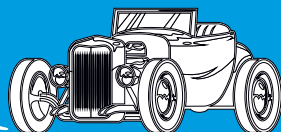
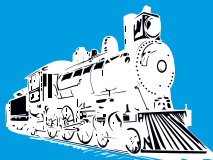
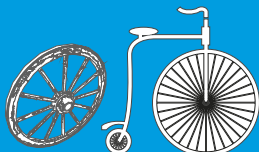
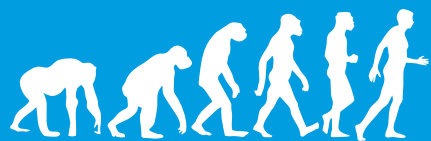
4 To wszystko prawda, ale Monika zapomniała o konieczności odwrócenia naleśnika, co zwykle robi się jednym zręcznym ruchem patelni; naleśnik wykonuje „salto mortale” i opada nieupieczoną stroną na patelnię. W pomysle Moniki też dałoby się to wykonać, pod warunkiem że patelnia byłaby łatwo zdejmowana, a jej rączka łatwo przypinana i odpinana.

5 Artur Gąsowski – ma dziadka, który zmuszony jest zażywać w różnych porach dnia różne pastylki. Dziadek ma już ponad 80 lat i oczywiście myli zarówno pory dnia, jak i rodzaj pastylek. Prosty sygnał ze smartfona to o wiele za mało. Artur uważa, że potrzebą chwili jest opracowanie specjalnego zegara – organizera, sprzężonego z kalendarzem, który wzywałby akustycznie do spojrzenia na jego ekran, na którym ukazałyby się nazwy leków, jakie w danej chwili dziadek powinien zażyć. Obsługa zwykłego smartfona dla dżentelmena w wieku powyżej 80 lat może być zbyt trudna. Zegar powinien mieć bardzo uproszczone programowanie i prostą obsługę.

Pewne działania w podobnym celu już uczyniono: pojawiły się elektroniczne organizery – pudełka z podziałem na godziny, w których można ułożyć tabletki i po prostu po kolei je tykać. Być może organizer, obejmujący cały tydzień, byłby właściwszym rozwiązaniem sprawy.

Stanisław Bizoń ostatnio próbuje jeździć konno. Zauważył, że jego koń okazał się inteligentnym i skłonym do żartów stworzeniem. Przy zakładaniu siodła i próbie dopięcia popręgów koń nadyma się i wydaje się, że popręgi będą „trzymać”. Tymczasem jest to sztuczka, bo gdy kandydat na jeźdźcę wkłada nogę w strzemień i próbuje mocno oprzeć się na tym strzemieniu, żeby przerzucić prawą nogę, koń wypuszcza powietrze i popręgi robią się zupełnie luźne, a jeździec ląduje na ziemi! Staszek proponuje opracować specjalne „rzepy” naszyte na dolnej powierzchni siodła, które z sierścią konia tworzyłyby pewne połączenie.

Może warto byłoby przedtem dokładnie obejrzeć konstrukcję rzepów. Sierść konia nie ma takich właściwości jak „przeciwrzep” i raczej wątpliwe, jest czy będzie trzymała dostatecznie mocno. Należałoby sprawę zbadać i być może coś by z tego wyszło.



INTERNET RZECZY

Nikola Tesla w wywiadzie dla magazynu „Colliers” mówi: „Kiedy łączność bezprzewodowa zostanie wdrożona, cała ziemia zostanie przekształcona w ogromny mózg (...) a instrumenty, dzięki którym będziemy mogli tego dokonać, będą zadziwiająco proste w porównaniu z naszym obecnym telefonem. Człowiek będzie mógł nosić jeden w kieszeni swojej kamizelki”.

Edward O. Thorp opracował koncepcję pierwszego komputera typu „wearable”. Analogowa maszyna o wielkości paczki papierosów, noszona w bucie (1), miała zastosowanie w hazardzie, służąc do przewidywania kolejnych trafień w ruletkę. Koncepcja była rozwijana przez kilka lat i testowana w Las Vegas w 1961 roku. Opinia publiczna nie była jednak o tym poinformowana. Świat dowiedział się o jej istnieniu dopiero w 1966 roku.

Internet, sam w sobie będący podstawą później rozwijanego IoT, rozpoczął działalność jako część DARPA (Defense Advanced Research Projects Agency) w 1962 r., a w 1969 r. przekształcił się w ARPANET. Później w latach 80. komercyjni dostawcy usług zaczęli wspierać publiczne wykorzystanie sieci ARPANET, co umożliwiło jej przekształcenie się we współczesny Internet.

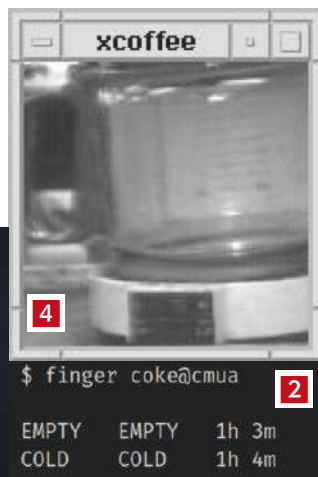
Hubert Upton, przez wielu nazywany ojcem wearables, w 1967 prezentuje montowany na okularach ekran, wspomagający czytanie z ruchu warg. Stało się to rok po ogłoszeniu znanej prognozy niemieckiego naukowca Karla Steinbucha, który ogłosił, że na przestrzeni kilku najbliższych dekad komputery będą implementowane w prawie każdym produkcie przemysłowym.

1926

1955

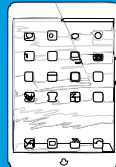
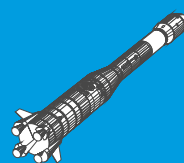
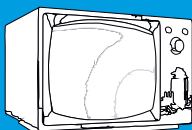
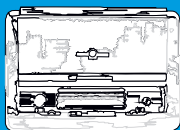
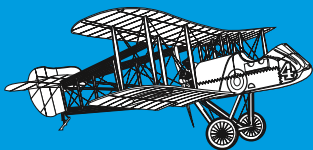
lata 60. XX wieku

1966–67



1. Noszony w bucie przenośny komputer Thorpa do przewidywania obrotów ruletki; **2.** Interfejs komunikacji z automatem do napojów na Uniwersytecie Carnegie Mellon; **3.** Prezentacja podłączonego do sieci tostera przez Johna Romkeya; **4.** Kamera monitorująca ilość kawy w dzbanku z 1993 roku; **5.** Steve Mann i jego noszone urządzenia

eprasa.pl 1078949591



lata 80. XX wieku

Na początku tej dekady członkowie zespołu Uniwersytetu Carnegie Mellon zainstalowali mikroprzetworniki w automacie Coca-Coli, łącząc je z administrującym komputerem, dzięki czemu widzieli na swoich terminalach, ile butelek aktualnie znajduje się w automacie oraz kiedy stają się odpowiednio schłodzone. Wystarczyło wpisać w terminalu sieci akademickiej „finger coke@cmua”, aby wiedzieć, czy przechadzka do automatu ma sens (2).

1990

John Romkey na targach branży urządzeń sieciowych prezentuje po raz pierwszy urządzenie (3), które może zostać włączone i wyłączone przez Internet (był to toster). Sprzęt był podłączony do komputera za pomocą sieci TCP/IP. Następnie za pomocą bazy informacji (SNMP MIB) włączał zasilanie.

1991

Członek ośrodka badawczego Xerox PARC Mark Weiser publikuje na łamach „Scientific American” artykuł „The Computer in 21st Century”, w którym opisuje wizję Internetu Reczy, w którym wyspecjalizowane elementy sprzętowe i oprogramowanie, połączone przez kable, fale radiowe i podczerwień, przenikną do świata codziennego do takiego stopnia, że nikt nawet nie zauważy ich obecności. Jest to prawdopodobnie pierwsza konceptualizacja idei „niewidzialnej sieci” i „zanurzenia” w cyberprzestrzeni tworzonej przez połączone urządzenia.

1993

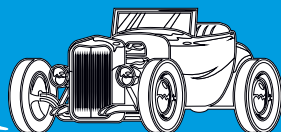
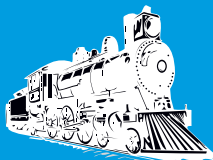
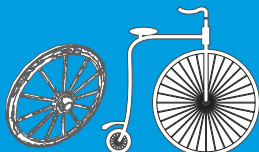
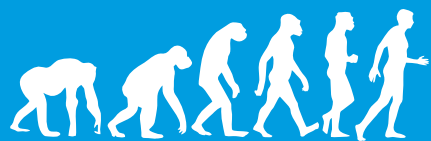
Grupa studentów z uniwersytetu w Cambridge wykorzystała kamerę internetową do informowania o dostępności kawy w ekspresie. Wpadli oni na pomysł wykorzystania pierwszego prototypu kamery internetowej do monitorowania ilości kawy dostępnej w dzbanku do kawy w ich laboratoriach komputerowych. Zrobili to poprzez zaprogramowanie kamery internetowej do robienia zdjęć dzbanka z kawą trzy razy na minutę. Zdjęcia te były następnie wysyłane do lokalnych komputerów, aby każdy mógł sprawdzić, czy kawa jest dostępna (4).

1993–94

Oficjalne uruchomienie ważnego dla rozwoju IoT systemu globalnego pozycjonowania satelitarnego (GPS). Do wiadomości publicznej podano, że GPS jest gotowy do eksploatacji (Initial Operational Capability). Ostatni satelita Block II został wprowadzony na orbitę w marcu 1994 roku, zamykając prace nad uruchomieniem systemu GPS.

1996–97

Nazywany niekiedy „ojcem elektroniki ubieralnej”, Steve Mann, tworzy serię urządzeń, które człowiek może założyć na siebie (5), m.in. WearComp („komputer ubieralny”) i WearCam („kamerę ubieralną”). Mann jako pierwszy zaproponował i wdrożył algorytm szacowania funkcji odpowiedzi aparatu na podstawie wielu różnie naświetlonych obrazów tego samego przedmiotu. Był również pierwszym, który zaproponował i wdrożył algorytm automatycznego rozszerzania zakresu dynamiki w obrazie poprzez łączenie wielu różnie naświetlonych zdjęć tego samego przedmiotu.



1999

Termin „Internet Rzeczy” zostaje spopularyzowany dzięki prezentacji Kevina Ashtona (6) dla koncernu Procter & Gamble. Ashton, dyrektor Auto-ID Center w Massachusetts Institute of Technology, wraz z Davidem Brockiem i Sanjayem Sarmą, zastosował technologię RFID (Radio – Frequency Identification) do identyfikacji poszczególnych urządzeń połączonych w ramach jednej sieci. Główną w jego rozumieniu istotą Internetu Rzeczy miało być wbudowanie mobilnych nadajników bliskiego zasięgu w różne gadzety i przedmioty codziennego użytku, aby umożliwić nowe formy komunikacji między ludźmi i rzeczami, a także między samymi rzeczami.

2000

Firma LG Electronics wprowadza na rynek pierwszą na świecie lodówkę podłączoną do Internetu. Umożliwiała ona konsumentom robienie zakupów spożywczych online i prowadzenie rozmów wideo. Jednak ta innowacja wówczas nie została przyjęta zbyt przychylnie przez potencjalnych klientów ze względu na cenę oraz, ich zdaniem, niepotrzebne funkcje.

2002–05

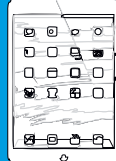
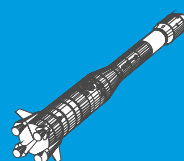
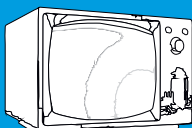
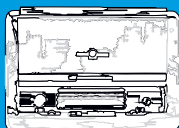
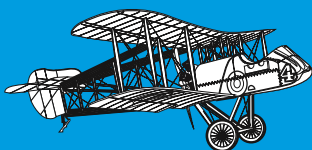
Kulisty gadżet Ambient Orb (7), stworzony przez Davida Rose i innych w ramach startupu wyrosłego z MIT Media Lab, pojawia się w sprzedaży. „New York Times Magazine” nazywa go jednym z Pomysłów Roku. To urządzenie, w swoim czasie niezwykle innowacyjne, monitoruje różne źródła danych, np. prognozy pogody lub raporty giełdowe, zmieniając swoją barwę na podstawie zmiennych parametrów i wizualnie sygnalizując np. zmianę pogody. Jest to typowy przykład komunikacji maszyna-maszyna. Podobnym co do charakteru i sposobu działania gadżetem był Nabaztag, stworzony w formie małego króliczka przez przez Rafiego Haladjiana i Oliviera Mévela. Łączył się z Internetem przez Wi-Fi, odbierał wiadomości i sygnalizował.

2005

We Włoszech w celu zbudowania urządzenia kontrolującego studenckie projekty interaktywne, jako tańsza alternatywa dla innych dostępnych wtedy systemów, rodzi się projekt Arduino (8). Przeobraża się w platformę programistyczną dla systemów wbudowanych, opartą na prostym projekcie open hardware przeznaczonym dla mikrokontrolerów montowanych w pojedynczym obwodzie drukowanym, z wbudowaną obsługą układów wejścia/wyjścia oraz standaryzowanym językiem programowania. Typowy zestaw Arduino zawiera kontroler, cyfrowe i analogowe linie wejścia/wyjścia oraz interfejs UART lub USB do połączeń z komputerem-hostem. Komputer jest wykorzystywany do programowania kontrolera oraz do interakcji w czasie działania z Arduino. Pomimo że płyty Arduino generalnie nie współpracują z siecią, częstym rozwiązaniem jest łączenie jednego lub kilku Arduino z hostem sieciowym, gdzie Arduino używa się w roli sprzętowych kontrolerów, a host przyjmuje rolę sieci lub interfejsu użytkownika. Możliwe jest programowanie interfejsu w kilkunastu językach programowania, m.in. w Javie, ActionScript, C/C++, C#, Perl, VBScript.

2008

Grupa przedsiębiorstw powołała do życia stowarzyszenie IPSO Alliance, którego celem jest wspieranie wykorzystania protokołu internetowego (IP) w sieciach „inteligentnych obiektów” a w konsekwencji umożliwienie powstania Internetu Rzeczy. Do sojuszu IPSO należy obecnie ponad pięćdziesiąt firm, w tym Bosch, Cisco, Ericsson, Intel, SAP, Sun, Google i Fujitsu.



2011

Publiczna premiera IPv6 – nowy protokół pozwala na ok. 340 sekstyliionów adresów. IPv6 został opracowany przez Internet Engineering Task Force (IETF) jako recepta na problem wyczerpania adresów IPv4. W grudniu 1998 roku IPv6 stał się projektem standardu dla IETF, który następnie został ratyfikowany jako standard internetowy 14 lipca 2017 roku. Urządzenia w Internecie mają przypisany unikalny adres IP w celu identyfikacji i określenia lokalizacji. IPv6 wykorzystuje 128-bitowy adres. Rzeczywista liczba dostępnych adresów jest nieco mniejsza, ponieważ wiele zakresów jest zarezerwowanych do specjalnego użytku lub całkowicie wyłączonych z użycia. Protokoły, starszy i nowszy, nie zostały zaprojektowane jako interoperacyjne, a zatem bezpośrednia komunikacja między nimi jest niemożliwa, co komplikowało przejście na IPv6. Opracowano jednak kilka mechanizmów przejściowych, aby temu zaradzić. Oprócz większej przestrzeni adresowej IPv6 oferuje inne korzyści techniczne. W szczególności umożliwia on hierarchiczne metody przydzielania adresów, które ułatwiają agregację tras w całym Internecie, a tym samym ograniczają rozbudowę tablic routingu. Rozszerzone i uproszczone jest stosowanie adresowania multicastowego, co zapewnia dodatkową optymalizację dostarczania usług.

2012

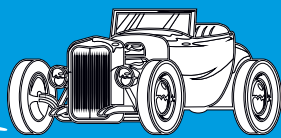
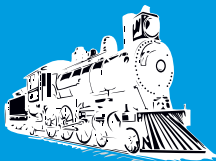
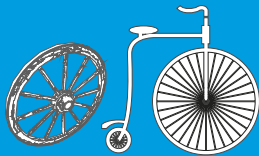
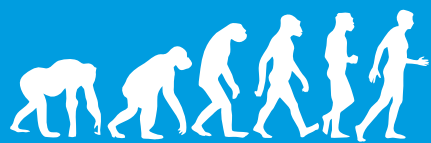
Premiera Raspberry Pi, platformy komputerowej stworzonej przez Raspberry Pi Foundation. Urządzenie składa się z pojedynczego obwodu drukowanego. Jego pierwotnym celem było wspieranie nauki podstaw informatyki. Urządzenie nie ma dysku twardego, ale w celu załadowania systemu operacyjnego i przechowywania danych oferuje złącze dla kart SD lub microSD albo pamięci USB. Raspberry Pi ma również złącze USB do podłączenia dowolnych urządzeń zewnętrznych. Działa pod kontrolą systemów operacyjnych opartych na Linuksie oraz RISC OS, a od modelu Raspberry Pi 2 B działa również pod kontrolą Windows 10 Internet of Things. Możliwe jest także uruchomienie Windows 10 dla procesorów ARM.

2017 do dziś

Początek i kolejne postępy wdrażania sieci 5G na świecie. Bezprzewodowa sieć piątej generacji uważana jest za klucz do przyspieszenia rozwoju Internetu Rzeczy, zwłaszcza w zastosowaniach przemysłowych. W 2017 roku przedsiębiorstwa Telia, Ericsson i Intel uruchomiły pierwszą sieć 5G działającą w czasie rzeczywistym w Europie. Testy dotyczyły Tallina i Sztokholmu. W tym samym roku zaczęto testować 5G w Londynie. W 2019 roku na rynku pojawiły się pierwsze komercyjne smartfony obsługujące standard 5G. Według Deutsche Telekom w 2020 roku standard opuści fazę prototypu i zacznie być udostępniany klientom biznesowym. Z kolei prezes Grupy Orange w grudniu 2018 zapowiedział, że na kolejny rok planowanych jest 17 komercyjnych wdrożeń sieci 5G w Orange. Unia Europejska chce wykorzystywać do tego pasmo 700 MHz, które, zgodnie z planem, jest przypisane do szerokopasmowych usług internetowych od czerwca 2020 roku.



6. Kevin Ashton; 7. Ambient Orb; 8. Arduino UNO; 9. Internet Rzeczy w 5G



WYNALEZKÓW HISTORIĘ ODKRYJ

Klasyfikacja Internetu Rzeczy według zastosowań

I. Zastosowania konsumenckie

1. Inteligentny dom

Urządzenia IoT są częścią większej koncepcji automatyki domowej, która może obejmować oświetlenie, ogrzewanie i klimatyzację, media i systemy bezpieczeństwa oraz systemy kamer. Celem jest przede wszystkim oszczędność energii dzięki automatycznemu wyłączeniu światła i elektroniki.

2. Opieka nad osobami starszymi i niepełnosprawnymi

Jednym z kluczowych zastosowań inteligentnego domu jest zapewnienie pomocy dla osób wymagających szczególnej troski. Mogą być również wyposażone w dodatkowe funkcje bezpieczeństwa. Funkcje te mogą obejmować czujniki, które monitorują nagłe sytuacje medyczne, takie jak upadki lub ataki.

II. Zastosowania zbiorowe i organizacyjne

1. Medycyna i opieka zdrowotna

Internet Rzeczy Medycznych (IoMT) bywa też określony jako „inteligentna ochrona zdrowia”. Urządzenia IoT mogą być wykorzystywane do umożliwienia zdalnego monitorowania zdrowia i powiadamiania o zagrożeniach. IoMT to zarówno monitory ciśnienia krwi i tętna, jak i zaawansowane urządzenia zdolne do monitorowania specjalistycznych implantów, takich jak rozruszniki serca, opaski elektroniczne Fitbit czy zaawansowane aparaty słuchowe.

2. Transport

Dynamiczna interakcja między elementami systemu transportowego umożliwia komunikację między i wewnątrz pojazdu, inteligentną kontrolę ruchu, inteligentne parkowanie, elektroniczne systemy poboru opłat, logistykę i zarządzanie flotą, kontrolę pojazdów, bezpieczeństwo i pomoc drogową.

3. Komunikacja V2X

Komunikacja pojazd-otoczenie (V2X) składa się z trzech głównych elementów: komunikacji pojazd-pojazd (V2V), komunikacji pojazd-infrastruktura (V2I) i pojazd-piesi (V2P). V2X jest pierwszym krokiem do pełnej autonomii i połączonej infrastruktury drogowej.

3. Automatyka w budynkach i domach

Urządzenia IoT mogą być wykorzystywane do monitorowania i sterowania systemami mechanicznymi, elektrycznymi i elektronicznymi stosowanymi w różnego rodzaju budynkach (np. publicznych i prywatnych, przemysłowych, instytucjach czy mieszkalnych), w systemach automatyki domowej i automatyki budynkowej.

III. Zastosowania przemysłowe

1. Produkcja

IoT może łączyć różne urządzenia produkcyjne wyposażone w funkcje wykrywania, identyfikacji, przetwarzania, komunikacji, uruchamiania i sieci. Głównie



elementy systemu to: kontrola sieciowa i zarządzanie urządzeniami produkcyjnymi, zarządzanie aktywami i sytuacją lub kontrola procesów.

2. Rolnictwo

Zbieranie danych o temperaturze, opadach, wilgotności, prędkości wiatru, inwazji szkodników i wartości gleby. Dane te mogą być wykorzystywane do automatyzacji technik rolniczych, podejmowania decyzji.

3. Branża morska

Urządzenia IoT są wykorzystywane m.in. do monitorowania środowiska a także ruchu łodzi i jachtów.

IV. Zastosowania infrastrukturalne

1. Wdrożenia na skalę metropolitalną

Duża część miasta ma być okablowana i zautomatyzowana i funkcjonować przy żadnej lub niewielkiej ingerencji człowieka. Na świecie testuje się obecnie wiele tego rodzaju systemów do realizacji różnych zadań.

2. Zarządzanie energią

Urządzenia IoT pozwalają na zdalne sterowanie przez użytkowników lub centralne zarządzanie za pomocą interfejsu opartego na chmurze, umożliwiając planowanie (np. zdalne włączanie lub wyłączenie systemów grzewczych, sterowanie piecami, zmiana warunków oświetlenia itp.).

3. Monitorowanie środowiska

Urządzenia IoT w tym zastosowaniu zazwyczaj pokrywają duży obszar geograficzny i mogą być również mobilne.

V. Zastosowania wojskowe

Znane na świecie projekty tego rodzaju to m.in. Internet of Battlefield Things, zainicjowany i realizowany przez Laboratorium Badawcze Armii Stanów Zjednoczonych oraz Ocean of Things, prowadzony przez DARPA program mający na celu stworzenie Internetu Rzeczy na dużych obszarach oceanów w celu gromadzenia, monitorowania i analizowania danych dotyczących środowiska i aktywności statków. ■

M.U.



Na trzech kołach

Trójkołowce to motoryzacyjna nisza, ale za to pełna niezwyklej konstrukcji o tradycji sięgającej samochodowej prehistorii. Dzisiaj takie pojazdy ponownie zyskują na znaczeniu w związku z rozwojem napędów elektrycznych.

Ponoć historia kołem się toczy, a już historia motoryzacji na pewno. Choć mówiąc potocznie o „czterech kółkach”, mamy na myśli samochód, to wśród pionierskich samochodów raczej dominowały pojazdy, w których zestaw kół był nieco uboższy. Pierwszy samochód z prawdziwego zdarzenia, za jaki uważa się Patent Motorwagen Nummer 1, zaprezentowany przez Karla Benz w 1886 r., miał przecież trzy koła. Podobnie jak jego przodkowie z bardziej odległej przeszłości – pojazdy parowe – na czele z najstawniejszym z nich, czyli „fardier à vapeur” z 1770 r. Nicolasa-Josepha Cugnota, inżyniera wojskowego w służbie króla Francji Ludwika XV. Fardier miał dwa koła z tyłu i jedno z przodu, kierowalne i jednocześnie podtrzymujące olbrzymi kocioł. Nietrudno sobie wyobrazić, że kierowanie taką machiną było wyzwaniem dla każdego, kto zasiadł za jej sterami (kierownicy wtedy jeszcze nie znano). Prymitywny mechanizm skrętu zawiódł już podczas jednej z jazd testowych, kiedy to pojazd Cugnota uderzył w ogrodowy mur, wybijając w nim wielką dziurę. Król w uznaniu wysiłków konstruktora wypłacił mu po każde uposażenie, ale dalszych eksperymentów zakazał, a fardiera zamknął w arsenale. Kiedy na początku XIX w. pojazdy parowe dopracowano i zdobyły większą popularność, nadal większość z nich jeździła na nieparzystym zestawie kół. Trójkołowe podwozia miały na przykład dyliżanse parowe w Anglii.

Choć historia uznana Karla Benz za ojca motoryzacji, to wielu Anglików kwestionuje ten tytuł, przyznając go swojemu rodakowi Edwardowi Butlerowi, który już w 1884 r. na Stanley Cycle Show w Londynie pokazał projekt

pojazdu spalinowego zwanego Butler Petrol Cycle. Problem w tym, że ze względu na skomplikowane prawo patentowe pierwszy jeżdżący prototyp powstał dopiero w 1888 r. Pojazd Butlera także był trójkołowcem, ale w odróżnieniu od motorwagena Benz oś z dwoma kołami znajdowała się z przodu, a pojedyncze koło z tyłu. We Francji trójkołowce konstruował w latach 90. XIX w. inżynier Léon Bollée, producent pierwszych samochodów w tym kraju napędzanych silnikami benzynowymi. Pionierski automobil Bollée był jednocześnie pierwszym seryjnie produkowanym autem wyposażonym w gumowe opony.

Oczywiście popularność rozwiązań trzykołowych nie wynikała z jego przewagi technologicznej nad bardziej stabilnym zestawem czterokołowym. Po prostu zarówno Cugnot, jak i jego „parowi” epigoni nie mogli sobie poradzić ze skonstruowaniem mechanizmu, który pozwoliłby zsynchronizować skręt każdego z napędzanych i sterowanych kół. Kąt skrętu musiał być przecież taki sam. Ten sam ból głowy mieli także pierwsi konstruktorzy samochodów. Przynajmniej do 1893 r., kiedy to Benz wpadł na pomysł zamontowania w swoich autach wynalazku znanego w zasadzie od starożytności, ale rozwiniętego i opatentowanego w 1876 r. przez producenta trójkołowych rowerów Jamesa Starleya. Od tej pory większość samochodów zostawiała za sobą cztery ślady.

Zresztą te ślady, a w zasadzie koleiny, jakie żłobiły na nieutwardzonych drogach powozy konne, także miały niebagatelną wpływ na popularyzację samochodów z dwoma kołami na każdej osi. Wyjeżdżone koleiny były przecież dwie, a między nimi zalegał często wał ziemi,

„Fardier à vapeur” Nicolasa-Josepha Cugnota z 1770 r.



Butler Petrol Cycle





co skutecznie uniemożliwiało jazdę po takich traktach pojazdom trójkołowym. Układ dwa koła z tyłu i jedno z przodu był także mało korzystny z punktu widzenia projektowania karoserii, zwłaszcza tych zamkniętych.

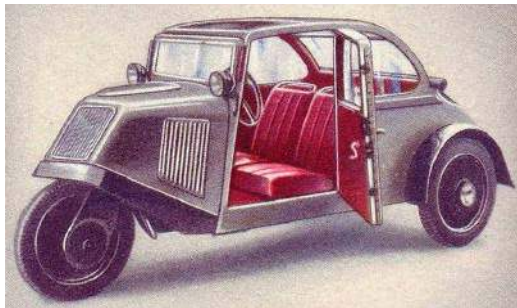
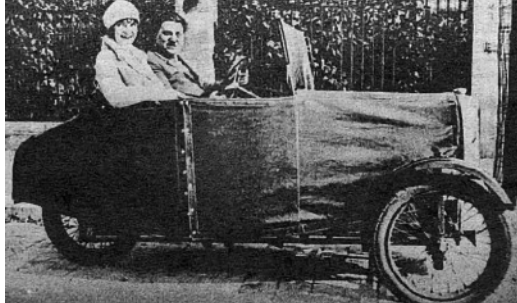
Wszystkie te przeciwności nie oznaczały jednak nagłej i całkowitej śmierci trójkołowców. Przetrwały i miały się dobrze, chociaż zawsze stanowiły margines rynku. Ich prosta konstrukcja przekładała się na niskie koszty produkcji i użytkowania. Złote lata przeżywały więc, nieco paradoksalnie, w czasach kryzysów ekonomicznych lub paliwowych, a także na rynkach krajów rozwijających się, ale ubogich, np. w niektórych krajach azjatyckich nadal skutecznie wypełniają lukę między motocyklami a samochodami.

Między motocyklem a samochodem

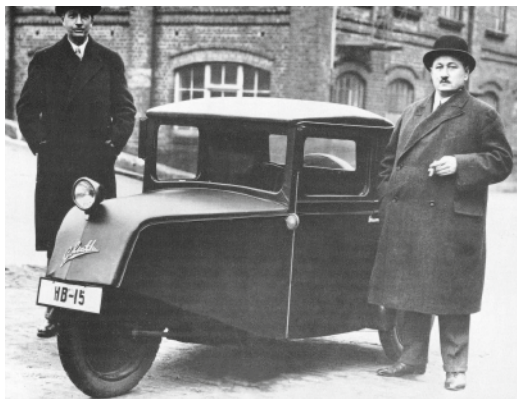
Podobna luka występowała także u zarania motoryzacji. Pierwsze automobile były towarem luksusowym, drogie w zakupie i utrzymaniu (paliwo, bardzo wysokie podatki), natomiast znacznie tańsze motocykle szybciej trafiały pod strzechy. W drugiej dekadzie XX w. zapanował boom na ogniwo pośrednie, którym stały się niezwykle popularne w Stanach Zjednoczonych i Europie Zachodniej tzw. cyclecary. Nie wszystkie one były trójkołowcami, np. popularna francuska Bédélia miała dwie pary kół, ale zdecydowana większość wyróżniała się uproszczoną budową, dzięki czemu były opodatkowane jak motocykle. Regułą było także sięganie przez konstruktorów po motocyklowe silniki, napęd łańcuchowy i dwubiegowe przekładnie. Chętnie stosowano również motory jednocyndrowe lub V2. Cyclecary zwykle były dwumiejscowe, siedzenie pasażera znajdowało się obok lub tuż za plecami kierowcy, chociaż zdarzały się także wersje trzyosobowe, jak np. produkowany w latach 1905–1927 niemiecki Phänomobil, w którym z tyłu montowana była niewielka kanapa dla dwójki pasażerów albo skrzynia ładunkowa.

Najpopularniejszym trójkołowym cyclecarem była z pewnością konstrukcja Morgana, ale o niej napiszemy nieco więcej w dalszej części tekstu, bowiem to chyba najstawniejszy i najlepszy trójkołowiec, a firma Morgan jako jedyna przetrwała wszelkie dziejowe zawieruchy i nadal istnieje. Chociaż warsztatów produkujących te niewielkie autka działało przed I wojną światową przeszło setka, większość z nich później zbankrutowała lub wróciła do składania zwykłych motocykli lub rowerów. Trudno im było rywalizować z coraz tańszymi (np. produkowany taśmowo Ford T) i coraz lepiej dopracowanymi czterokołowcami. Jednak w latach 20. nastąpił ponowny boom na trzy koła – wiele osób nauczyło się podczas wojny kierować pojazdami mechanicznymi, ale w trudniej powojennej sytuacji ekonomicznej nie było ich stać na auto z prawdziwego zdarzenia.

W Niemczech reaktywacja trójkołowej motoryzacji wiąże się z firmą Diabolo Kleinauto, która produkowała dwa modele oparte konstrukcyjnie na brytyjskim Morganie z wąską przednią osią i pojedynczym, tylnym kołem, napędzanym silnikiem V2 firmy Motosacoche o pojemności skokowej 1093 cm³ i mocy 16 KM. Powstała także



Tempo Front z przełomu lat 20. i 30. XX w.



Goliath Pioneer

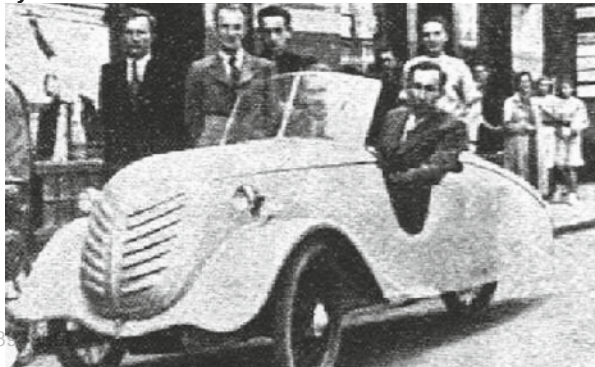
wersja sportowa, która miała o 10 KM więcej. Pod koniec lat 20. wiele szumu na rynku zrobił niezwykły trójkołowiec Zaschka. Jego twórca, Engelbert Zaschka, głowił się, jak rozwiązać rosnący problem braku miejsc parkingowych w dużych miastach i wpadł na pomysł zbudowania lekkiego autka, jeżdżącego na trzech kołach, opartego na konstrukcji rurowej oplecionej mocną tkaniną, które dwie osoby mogły zdemontować w pięć minut, a części składowe zabrać ze sobą na przykład do mieszkania.

Inni znani niemieccy producenci „trzechśladów” to firmy Tempo i Goliath. Ta pierwsza konstruowała

Phänomobil



Cyklonetka w Warszawie w 1939 r.





Czechosłowacki Velorex

m.in. w latach 30. niewielkie ciężarówki napędzane dwusuwowym silnikiem o pojemności zaledwie 200 cm³, jednak w historii motoryzacji zapisała się przede wszystkim osobowym modelem Hanseat, który po wojnie sprzedany na licencji do Indii był z różnymi modyfikacjami produkowany aż do początku XXI w. Na indyjskiej prowincji nadal bez problemu spotkamy jeżdżące egzemplarze. Goliata, będącego częścią koncernu Borgward, najbardziej kojarzono z modelem Pioneer produkowanym od 1931 r., wyróżniającym się zgrabną drewniano-materiałową zamkniętą karoserią, której znalazł aż 4 tys. nabywców. Po II wojnie światowej trójkołowy transportowiec Goliath Goli stał się z kolei istotnym elementem niemieckiego cudu gospodarczego lat 50.

We Francji w okresie międzywojennym sporym powodzeniem cieszyły się np. produkowane na licencji Morgana trójkołowce Darmont, które charakteryzowały się zacięciem sportowym (np. model Special z 1927 r. z silnikiem V2 o poj. 1100 cm³, rozpędzał się do 150 km/godz.!).

Trzy koła po polsku

Trójkołowa motoryzacja miała także swoich przedstawicieli nad Wisłą. Takie pojazdy od 1906 r. powstawały w warsztatach rodziny Kopeć, warszawskich producentów kas pancernych. Źródłem napędu był silnik o mocy 3,5 KM umieszczony nad przednim, kierowanym kołem. Podobno konstrukcja automobilu była tak prosta, „że w ciągu kwadransa każdy mógł obeznac się z mechanizmami i kierować samodzielnie”, wypuszczając się na wycieczki nawet w teren górzysy, pojazd bowiem „pod górę, na przykład na ulicy Bednarskiej, bieży tak samo jak po równym bruku”. Kwota 1000 rubli, za jaką oferowano te trójkołowce, była kilkakrotnie niższa od ceny przeciętnego zagranicznego samochodu czterokołowego.

Inny polski rozdział tej opowieści to Cyklonetka, stworzona w okresie międzywojennym przez braci Eugeniusza, Jana i Wacława Knaów z Kielc. Przyświecała im chwalebna idea – zbudowania polskiego auta dla mas, taniego, łatwego w obsłudze i naprawach, co najmniej dwuosobowego. Jako napędu użyli motorowerowego silnika o pojemności 98 cm³ wyprodukowanego w fabryce Steinhagen i Stransky. Pozostałe rozwiązania konstrukcyjne, łącznie z nowoczesnym opływowym nadwoziem, były ich pomysłem i powstały w rodzinnym garażu. Jeżdżący prototyp skonstruowali w ciągu siedmiu miesięcy i latem 1939 r. samochód sprawnie pokonał trasę z Kielc do Warszawy, gdzie bracia zamierzali szukać sponsorów

Messerschmitt KR 175



BMW Iso SpA Isetta

lub pozyskać dotacje państwowe. Co ciekawe, do prowadzenia takiego pojazdu nie było potrzebne prawo jazdy i wystarczyła rejestracja motocyklowa. Cyklonetka miała kosztować jedynie 1500 zł, ale rozwój projektu został przerwany przez wybuch II wojny światowej.

Lądowe messerschmitty

Krajobraz motoryzacyjny w powojennej Europie zmieniał się diametralnie. Zniszczony przemysł wymagał szybkiej odbudowy i przestawienia go z powrotem na tryb pokojowy. Samochody ponownie zastępowały na liniach produkcyjnych czołgi, ale prawie wszędzie potrzebą chwili były auta tanie, proste i praktyczne. Znowu nadeszły dobre czasy dla trójkołowców.

W Polsce i innych krajach Europy Środkowej i Wschodniej, które zostały podporządkowane ZSRR, nie było dłużej miejsca na tak powszechne przed wojną prywatne, garażowe inicjatywy, bo przemysłem samochodowym sterowano centralnie. Jedynie w Czechosłowacji sytuacja wyglądała nieco inaczej i właśnie tam pojawił się najbardziej znany powojenny trójkołowiec pochodzący zza żelaznej kurtyny – Velorex 16/350, początkowo produkowany pod nazwą Oskar. Był to pojazd przeznaczony dla inwalidów, ale okazało się, że w okresach socjalistycznych niedoborów i kryzysów może być atrakcyjny dla osób w pełni sprawnych. Velorex, jak wiele innych trójkołowców na przestrzeni dekad, inspirowany był konstrukcją Morgana. Nadwozie nieco przypominało wspomnianego już składaka Zaschki – na szkielet z rurek naciągnięta była powłoka z dermatoidu. Jednostka napędowa pochodziła z motocyklu Jawa i rozpędzała lekkie, ważące 310 kg pojazdy do 85 km/h. Autko produkowano w latach 1945–1973, a połowę z 12 tys. egzemplarzy wyeksportowano do innych krajów RWPG.

Znacznie więcej trójkołowców powstało w Europie Zachodniej. Rynek zdominowały wtedy pojazdy zwane skuterami kabinowymi. Potrafiły się mocno różnić, ale jedną cechą niewątpliwie miały wspólną – był nią koszmarny wygląd.

W Niemczech producenci samolotów tacy jak Messerschmitt czy Heinkel nie mogli kontynuować dotychczasowej militarnej produkcji, spróbowali więc swoich sił w motoryzacji. Nie zapomnieli jednak o swoich korzeniach – pokazany w 1953 r. na salonie w Genewie Messerschmitt KR 175 wyróżniał się samonośnym nadwoziem wzorowanym na kabinach lotniczych. Zamiast drzwi zamontowano odsuwającą owiewkę z pleksiglasu, a kierownicę zastępował drążek bezpośredniego połączenia z kołem. Motocyklowy silnik o mocy 9 KM zużywał tylko 3 litry mieszanki na 100 km. Dwa lata później Niemcy pokazali kolejnego trójkołowca – Messerschmitta KR200. W 1956 r. swój samochód zbudował także inny niemiecki potentat z dziedziny awiacji – Heinkel. W Niemczech znany był pod nazwą Kabine, a po sprzedaniu w latach 60. licencji do Irlandii i zamontowaniu nowego większego silnika, jako Trojan 200.

Messerschmitty i heinkle sprzedawały się nieźle, ale daleko było im do sukcesu, jaki odniosła włosko-niemiecka hybryda sygnowana literami BMW. Mowa o BMW Iso SpA Isetta. Pojazd ten, opracowany przez włoską firmę Iso SpA w 1954 r., początkowo nie wzbudził



zainteresowania. Sytuacja zmieniła się diametralnie rok później, kiedy BMW kupiło licencję i zamontowało w nim własne czterosuwowe silniki. Do 1962 r. wyprodukowano ponad 161 000 egzemplarzy, a dzisiaj to chyba najwyżej wyceniany oldtimer z gatunku skuterów kabinowych.

Samochód Bonda i „plastikowa świnia”

Samochodziki z trzema kołami były szczególnie popularne w Wielkiej Brytanii, a lista wyspiarskich trójkołowców jest bardzo długa, dlatego przypomnimy tylko najbardziej znane modele.

Tytuł tego rozdziału jest nieco przewrotny i żartobliwy, bo oczywiście Bond, James Bond, agent JKM z licencją na zabijanie, nigdy nie zamienił Aston Martina na dziwacznego trójkołowca. W tym wypadku chodzi o serię mikrosamochodów z dwoma kołami z tyłu i pojedynczym z przodu, produkowanych w latach 1949–1970 przez niewielką brytyjską manufakturę założoną przez Lawrence’a „Lawrie” Bonda, inżyniera z Preston. W okresie tuż po wojnie, kiedy benzyna była racjonowana, niewielkie auto z lekką aluminiową karoserią i silnikiem 125 cm³ okazało się ciekawą propozycją. W kolejnych latach Bond wypuszczał ulepszone wersje, a zwieńczeniem samodzielnej produkcji był model 875, napędzany motorem Hillman Imp o pojemności 875 cm³. W 1969 r. firmę Bonda wykupił Reliant Motor Company, który pod marką Bond Bug produkował do 1974 r. sportowego trójkołowca, który wyposażony w czterocylindrowy silnik o pojemności 700 cm³ potrafił rozpędzić się aż do 126 km/h. Oryginalnie pojazd był dostępny wyłącznie w kolorze jasnopomarańczowym, ale na potrzeby reklamy papierosów Rothmans wyprodukowano sześć białych egzemplarzy (dziś to unikat). Trójkołowy Bond nie towarzyszył co prawda na wielkim ekranie agentowi 007, ale pośrednio i tak zrobił karierę filmową, bowiem podwozie jednego z bugów posłużyło do budowy landspeedera, pojazdu znanego z sagi „Gwiezdne Wojny”.

Skoro już jesteśmy przy Bondzie Bugu, należy poświęcić kilka zdań firmie, która go wyprodukowała, bowiem Reliant to jeden z najbardziej cenionych producentów trójkołowców, znany przede wszystkim z modelu Robin, z którego uwielbiali naśmiewać się twórcy i prowadzący znanego brytyjski programu motoryzacyjnego „Top Gear”. Ze względu na brak urody i karoserię wykonaną z włókna szklanego Robina zwano na Wyspach „plastikową świnią”. Auto było dostępne w wersjach hatchback, sedan, a nawet... van. Produkowano go zaskakująco długo – od 1978 r. aż do 2001 r., a Reliant swego czasu był drugim co do wielkości brytyjskim producentem samochodów po British Leyland i światowym liderem w dziedzinie wytwarzania karoserii z włókna szklanego. Poza Robinem i Bugiem firma wypuściła na rynek także inne popularne trójkołowce – Regal i Rialto.

Wszystkie trójkołowce Bonda i Relianta, a także wiele innych tego typu pojazdów opierały się na dość ułomnej konstrukcji w układzie dwa koła z tyłu i jedno sterowane z przodu, co czyniło je pojazdami mało stabilnymi i niebezpiecznymi. Zdecydowanie praktyczniejszym rozwiązaniem jest układ odrotny, co najlepiej ilustruje przykład największego klasyka gatunku – Morgana.

Od Runabouta do 3Wheeler

Morgan to jedna z najstarszych firm motoryzacyjnych na świecie, niezwykle wierna swojej tradycji i mało podatna

Reliant Robin.
You can't beat it for all round economy.



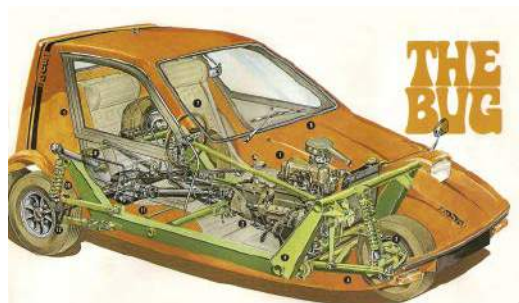
na wszelkie mody i zmiany. Nadal każdy Morgan jest wytwarzany ręcznie przy użyciu trzech podstawowych materiałów: drewna jesionu, aluminium i skóry. W fabryce w Malvern Link nieprzerwanie od 1910 r. działa stolarnia, w której powstaje m.in. drewniany szkielet, na którym osadzone są ręcznie formowane panele aluminiowe nadwozia, a od wewnątrz skórzane obszycia. Założył ją w 1910 r. Henry Frederick Stanley Morgan. Chciał stworzyć pojazd, który będzie bardziej stabilny i bezpieczny niż motocykl, a że automobile na czterech kołach były wtedy wysoko opodatkowane, wymyślił trójkołowca, napędzanego przez jednocylindrowy 4-kołowy silnik (potem pojawił się mocniejszy motor V2) z pierwszym na świecie niezależnym zawieszeniem. Pojazd, nazywany popularnie Runabout, początkowo mieścił jedynie kierowcę, potem miał dwa miejsca, a ostatecznie pojawiła się także wersja czteroosobowa.

Chociaż trójkołowce Morgana nie mają tak bogatej tradycji jak ich najbardziej znany model, czterokołowy Roadster 4/4 produkowany nieprzerwanie przez 84 lata, to jeszcze niedawno można było kupić współczesną wersję Runabouta, zwaną 3Wheeler. To całkowicie szalony, ekstremalny pojazd wywołujący dziką radość u każdego, kto spróbował zasiąść za kierownicą tego narowistego rumaka napędzanego benzynowym V2 amerykańskiej firmy S&S o pojemności 2 l. Dysponuje on mocą 83 KM i maksymalnym momentem obrotowym 140 Nm. Na papierze to niedużo, ale 3Wheeler waży niewiele ponad 525 kg, dzięki czemu przyspiesza do setki w 6 s i rozpędza się do prędkości maksymalnej 185 km/h. W otwartym roadsterowym nadwoziu 3Wheeler zapewnia to niezapomniane wrażenia. Niestety ta oldskulowa konstrukcja nie spełniała współczesnych wyśrubowanych norm emisji spalin i w zeszłym roku zniknęła z rynku. Brytyjczycy zapowiedzieli jednak wkrzeszenie swojego klasyka. Nadal będzie jeździł na trzech kołach, ale najprawdopodobniej przednią oś obciążą tym razem silnik elektryczny.

Na trzech kołach w przyszłość?

„Zielony” Morgan nie będzie wcale konstrukcją pionierską, bowiem kolejna, współczesna fala trójkołowców związana jest właśnie z elektromobilnością i wzbiera co najmniej od kilku lat. To także swoisty łącznik z motoryzacyjną prehistorią. 19 kwietnia 1881 r. na Rue de Valois w Paryżu pojawił się bcykl z trzema kołami napędzany silnikiem elektrycznym Siemens. Jego twórcą był paryski zegarmistrz i jednocześnie genialny wynalazca Gustave Trouvé. Był to pierwszy pojazd elektryczny w historii.

Henry Frederick Stanley Morgan w trójkołowcu własnej konstrukcji





Morgan 3Wheeler

Elektryczne trójkołowce znów budzą zainteresowanie, ponieważ lepiej niż tradycyjne samochody radzą sobie z ograniczeniami technologicznymi napędów EV. W małych i bardzo lekkich autach, które z założenia nie są przeznaczone do pokonywania znacznych odległości, nie ma konieczności instalowania dużych akumulatorów, a zasięg i tak pozostaje atrakcyjny dla użytkownika. Takie pojazdy są także przedmiotem dużego zainteresowania mikromobilności, której istotę stanowi użytkowanie w rozwiązaniach komunikacyjnych niewielkich, lekkich i bezemisyjnych, tzw. osobistych środków transportu (ang. Personal Mobility Device – PMD) umożliwiających pokonywanie krótkich dystansów – najczęściej pierwszego lub ostatniego odcinka zaplanowanej podróży.

Przykładem trójkołowego PMD jest projekt, który Grupa PSA (m.in. Citroën, Peugeot, Opel) wraz z partnerami stworzyła w ramach projektu Horyzont 2020 współfinansowanego przez Unię Europejską. Nosi on nazwę EU-LIVE (Efficient Urban Light Vehicle, czyli wydajny lekki pojazd miejski). W trybie bezemisyjnym pojazd EU-LIVE napędzany jest przez dwa silniki elektryczne o łącznej mocy około 8 kW zainstalowane w kołach. Dzięki nim jest w stanie poruszać się z szybkością do 70 km/h. Jest to jednak pojazd uniwersalny, bowiem „zielonemu” napędowi towarzyszy 1-cylindrowy motor spalinowy o mocy 31 kW. Tandem ten zapewnia ok. 300 km zasięgu, co pozwala pokonywać także dłuższe odcinki poza miastami. EU-LIVE jest niezwykle łatwy w prowadzeniu dzięki kontroli stabilności wykorzystującej komponenty hydrauliczne oraz hydropneumatyczne zawieszenie.

Własną trójkołową elektryczną hybrydę samochodu i motocykla od dawna promuje także Toyota: i-Road mieści dwie osoby, a tandemowe ułożenie siedzeń sprawia, że pasażer siedzi tuż za kierowcą. Toyota i-Road rozwija prędkość maksymalną 60 km/h, napęd przekazywany jest na tylne koło, a dwa przednie zmieniają kąt nachylenia i położenie za pomocą systemu Active Lean. Pojazd ma jedynie dwa metry długości, dlatego zaparkujemy go bez trudu – na miejscu parkingowym przeznaczonym dla tradycyjnego samochodu, zmieszczą się przynajmniej trzy i-Roady.

W Stanach Zjednoczonych ciekawie zapowiadają się m.in. projekty Arcimoto i Aptera. Ta pierwsza firma w Kalifornii, Oregonie, na Florydzie i w stanie Waszyngton sprzedaje małe elektryczne trójkołowce, z częściowo otwartą karoserią w kilku wersjach – osobowych FUV (Fun Utility Vehicle) oraz Roadster, transportowo-kurierskiej Deliverator i przeznaczonej dla służb ratowniczych Rapid Responder. Autka rozpędzają się do około 120 km/h, mają ok. 164 km zasięgu, a ceny zaczynają się od niecałych 18 tys. dolarów.

Z kolei kalifornijski startup Aptera Motors wyznaczył sobie jako cel stworzenie najbardziej energooszczędnego pojazdu elektrycznego na świecie. Kluczem do tego mają być przede wszystkim optywowa sylwetka z rekordowo niskim współczynnikiem oporu powietrza Cx wynoszącym 0,11 oraz superlekkie kompozytowe nadwozie pokryte panelami



EU-LIVE



i-Road

fotowoltaicznymi. Dzięki temu w słoneczne dni autko będzie się samo ładowało i zapewni zasięg do 65 km. Jeśli załączymy silnik elektryczny, to w zależności od konfiguracji przejedziemy bez konieczności ładowania od 400 do ponad 1600 km! Kompozytowa powłoka Aptery zawiera wiele mikrokanalików wypełnionych chłodziwem, które odprowadzają ciepło z akumulatorów, silników i paneli słonecznych. Dodatkowo Aptera ma mieć niezłe osiągi: od 0 do 97 km/h w 3,5 sekundy, a prędkość maksymalna to 175 km/h. Sprzedaż ma ruszyć w przyszłym roku, a ceny zaczynają się od 25 900 dolarów.

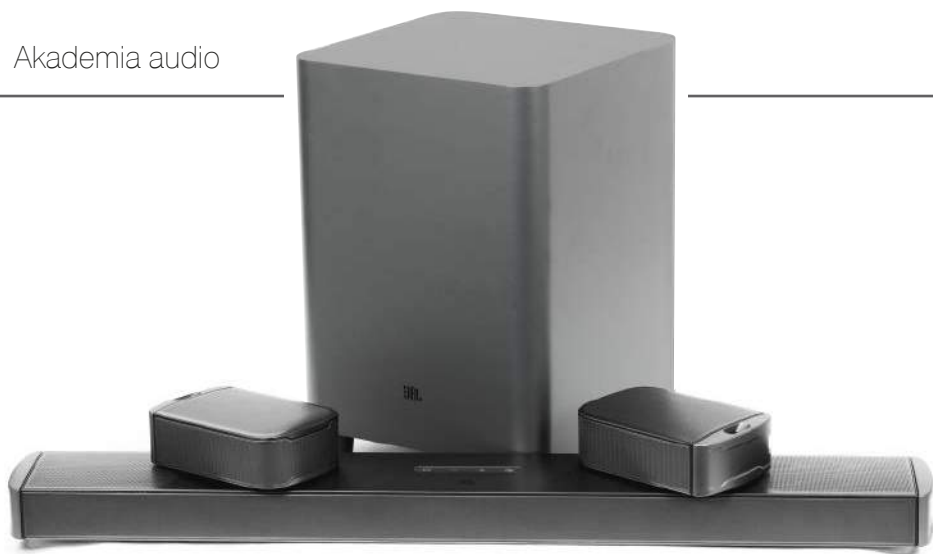
Autom na trzech kołach nie można odmówić oryginalności, a niektóre konstrukcje potrafią dać wiele frajdy każdemu fanowi motoryzacji. Trójkołowce nigdy jednak nie wyszły poza swoje ściśle specjalizowane nisze i postrzegane były i są jako przedziwna forma pośrednia pomiędzy zwinnymi i tańszymi motocyklami a bezpieczniejszymi i praktyczniejszymi samochodami czterokołowymi. Ich współczesne elektryczne inkarnacje ponownie mają swoje pięć minut, ale jest to raczej sytuacja tymczasowa, związana z przejściowymi ograniczeniami technologii EV. Niewykluczone jednak, że w obliczu gwałtownie rosnącego przeludnienia wielkich miast mogą się okazać ważnym elementem mobilności przyszłości, szczególnie tej w skali mikro. ■

Krzysztof Michał Józwiak



**Arcimoto w wersji FUV
Aptera**





Nowoczesne soundbary (1)

Soundbary to urządzenia bardzo popularne, a jednak dość tajemnicze. Enigmatyczne „belki” wieszane pod telewizorami pełnią coraz więcej funkcji, są coraz bardziej skomplikowane i różnicowane. Za miesiąc, na przykładzie trzech nowych modeli wysokiej klasy, przedstawimy najnowsze trendy i najambitniejsze rozwiązania, a teraz trochę historii i ogólna koncepcja soundbarów.

Na przełomie XX i XXI wieku bujnie rozwinęło się kino domowe – wykorzystując takie wynalazki, jak DVD i płaskie telewizory, a także nowoczesny zapis wielokanałowych, filmowych ścieżek dźwiękowych (Dolby Digital). Wszyscy zapragnęli dźwięku dookólnego, który miały zapewnić (i do dzisiaj zapewniają) rozbudowane systemy głośników rozstawionych dookoła słuchacza (widza). Popularyzacja tej koncepcji postępowała szybko, sklepy zapęłniły się niskobudżetowymi „zestawami z jednego pudełka”, zawierającymi komplet głośników. Często jednak, już po rozpakowaniu, szczęście trwało krótko, kiedy okazywało się, że kupujący źle oszacował swoje możliwości lokalowe i zapął do przeprowadzenia całej instalacji, więc sił starczało mu tylko na ustawienie głośników przednich, ewentualnie razem z centralnym, a tylne „surroundy”, wymagające pociągnięcia kabli i powieszenia na ścianie, czekały na lepsze czasy.

Problem nadmiernego skomplikowania dostrzegli producenci, którzy zaproponowali rozwiązanie znacznie prostsze (dla użytkownika), funkcjonalne, eleganckie i wciąż niedrogie – właśnie soundbary, czyli podtelewizorowe „grające belki”. Doskonale pasujące (wizualnie) do coraz bardziej płaskich i coraz większych telewizorów, które jednak emitowały dźwięk... gorszy niż dawne telewizory kineskopowe z dużymi pudłami.

Można więc było przekonywać również użytkowników wcześniej niezainteresowanych kinem domowym, że soundbarem warto poprawić dźwięk nowego telewizora. Oczywiście dla takiego klienta nie mógł to być soundbar drogi.

Nie mógł też zapewnić dźwięku dosłownie dookólnego, tak jak klasyczne systemy wielogłośnikowe, jednak dzięki sprytnym rozwiązaniom akustycznym rezultaty okazały się na tyle satysfakcjonujące, że bardzo duża część klientów „przeorientowała” się na soundbary. I od tego czasu zdobywają one coraz większe udziały, bowiem rozwijają się i zaspokajają różne potrzeby.

Wspólnym mianownikiem jest podłużna belka, zwykle jest ona cienka i dopasowana długością do telewizora – ale nie zawsze; pojawiają się krótkie „minisoundbary”, a także znacznie grubsze niż przeciętnie – gdy mają działać samodzielnie, bez wsparcia systemowym subwooferem. Koncepcja pojedynczego urządzenia szybko uległa modyfikacji i większość obecnie proponowanych soundbarów jest de facto dwuczęściowymi zestawami soundbara z subwooferem, bowiem bardzo cienka listwa nie zapewnia warunków odpowiednich do przetwarzania niskich częstotliwości (wymagających większych przetworników i większej objętości). Chociaż skrzynka subwoofera nigdy nie była ozdobą salonu,

to na taki „deal” większość się zgadza – na subwoofer znajdzie się gdzieś miejsce, byle soundbar był cieniutki. Większość, ale nie wszyscy – niektórzy wolą grubsze soundbary bez subwooferów. Są też takie, które bez subwoofera dają radę, ale osiągnięcie lepszego basu możliwe jest po dokupieniu i podłączeniu odpowiedniego subwoofera.

Często też specjalna sekcja przetwarzająca niskie częstotliwości, nawet gdy znajduje się w soundbarze (a nie w zewnętrznej obudowie), określana jest jako subwoofer (producenci nazywają to soundbarem ze zintegrowanym subwooferem).

Już tutaj zaznacza się różnicowanie konstrukcji, a w samym soundbarze różnych opcji akustycznych może być bardzo wiele, służących wykreowaniu mniej lub bardziej rozwiniętego dźwięku dookólnego. Przy czym nie należy bezkrytycznie wierzyć w obietnice, że systemy skomplikowane, zawierające więcej kanałów, zawsze są skuteczniejsze w tworzeniu dźwięku przestrzennego, a tym bardziej – w emitowaniu dźwięku ogólnie wysokiej jakości. Zależy to od wielu cech, które trudno ująć w podstawowej specyfikacji.

Najprostsze soundbary to systemy stereofoniczne – dwukanałowe, a więc z kanałami lewym i prawym (system 2.0), i ewentualnym dodatkiem subwoofera (zintegrowanego lub zewnętrznego – system 2.1). Dalej rozbudowa może iść w różnych kierunkach, w tym do standardowego systemu wielokanałowego 5.1, gdzie kolejne trzy sekcje przetwarzają sygnały kanałów centralnego (to akurat dość proste, skoro soundbar znajduje się naprzeciwko słuchacza) i kanałów surroundowych lewego i prawego; tutaj pojawia się wyzwanie – jak dźwięk emitowany z głośników przed nami ma być słyszany z tyłu czy choćby z boków?

Producenci obiecują cuda, sugestywnie przedstawiają je na grafikach, dowolnie zakrzywiając fale dźwiękowe, jednocześnie usilnie pracują nad jak najlepszymi efektami, wykorzystując do tego zarówno tradycyjną psychoakustykę, jak i nowoczesną technikę cyfrową. Zasadnicze sposoby są dwa. Pierwszy to operowanie przesunięciami fazowym w sygnałach poszczególnych kanałów, które potrafią nasz słuch „oszukać” co do pozycji źródła – sposób wykorzystywany już dawno temu, techniką analogową, do tworzenia efektu „superstereo”, rozszerzającego bazę, a zaawansowana technika cyfrowa pozwala na o wiele więcej (choć nie na wszystko). Drugi sposób to kierowanie promieniowania (fal dźwiękowych) kanałów surroundowych w stronę ścian bocznych, aby tam zostały odbite i nadbiegły z boku. Najwięksi optymiści zakładają, że mogą odbić się od ścian bocznych, potem od tylnych, i nadbiec stamtąd... ale to mało realistyczne, nawet w specjalnie pod

tym kątem przygotowanych warunkach akustycznych, jako że głośniki nie promieniają wiązek bardzo skupionych i przy takich odległościach, i kilku odbiciach, fala jest już bardzo rozproszona. Tym bardziej niezależne ukierunkowanie fal kanałów bocznych i tylnych (systemu 7.1) jest trudne w realizacji.

Soundbary – i to nie tylko te najdroższe – muszą jednak wychodzić naprzeciw modnym hasłom, a do nich obecnie należy Dolby Atmos, czyli system jeszcze bardziej rozbudowany, z dodatkową „warstwą” tworzącą wymiar pionowy przestrzeni dźwiękowej. Wzorcowe systemy Dolby Atmos mają więc kanały i głośniki „sufitowe” (aż cztery). Jak pogodzić soundbarową integrację z dalszym rozmnożeniem się kanałów, głośników i miejsc ich instalacji? Producenci soundbarów nauczyl się i nas, że dla chcącego nic trudnego... trzeba trochę pokombinować, zasugerować, aby użytkownik był pod wrażeniem dźwięku swobodnego, efektownego, który może odbiegać od wzorca, ale wystarczy, jeżeli będzie przyjemny. Namiastkę Atmos(fery) zapewnić mają głośniki promieniujące nie z góry, lecz... do góry – oczywiście po to, aby fala odbiła się od sufitu. Nie jest to zresztą kompromis przypisany tylko soundbarom; kłopot, jaki instalacja sufitowa sprawia większości użytkowników (również klasycznych systemów wielokanałowych, złożonych z niezależnych zespołów głośnikowych), musiał znaleźć rozwiązanie, aby nazbyt wysoko (dosłownie i w przenośni) zawieszona poprzeczka nie pogrzebała całej inicjatywy zmierzającej przecież do sprzedania wielu urządzeń nowej generacji – nie tylko głośników i soundbarów, ale też wzmacniaczy AV. Pojawiły się więc „atmosowe” głośniki-nadstawki, kładzione na kolumnach głównych (lewej i prawej), promieniujące do góry w znanym już celu. Do soundbarów tylko ten pomysł przeniesiono, zresztą nawet lepiej pasuje do ich kombinowanego sposobu działania.

Zdecydowana większość soundbarów to konstrukcje aktywne – analogicznie jak w przypadku zespołów głośnikowych oznacza to, że urządzenie zawiera wszystkie potrzebne wzmacniacze, więc wystarczy do niego podłączyć źródła dźwięku (a nawet obrazu, który prześle dalej do telewizora).

Ta sfera rozwinęła się i stała jednym z podstawowych zadań soundbarów. Nowoczesne soundbary są samowystarczającymi systemami audiowizualnymi z oddawczymi sieciowymi, nowoczesnym sterowaniem, komunikacją bezprzewodową – rozwiązaniem kompleksowym, wygodnym i relatywnie niedrogim. Wi-Fi, Bluetooth, Google Chromecast, Apple AirPlay 2, Spotify Connect, asystenci głosowi, aplikacje mobilne – takie wyposażenie nie jest już ekskluzywnie. ■

Andrzej Kisiel



Gdy weźmiemy pod uwagę stosunek ceny do jakości, możliwości, a także, może przede wszystkim, pojemności akumulatora, to Moto G9 Power wydaje się produktem dającym komuś, kto decyduje się na zakup tego urządzenia, całkowitą satysfakcję.

Duża rzecz za nieduże pieniądze

Smartfon Motorola Moto G9 Power

Obudowa telefonu z tworzywa sztucznego snobów może przyprawić o kręcenie nosem, jednak oprócz przypomnienia relatywnie niewysokiej ceny, warto zwrócić uwagę na praktyczne zalety takiego rozwiązania. Stosunkowo miękki i podatny plastik nie pęka przy upadkach. Poza tym jest lżejszy niż metalowe i szklane materiały, co przy sporych gabarytach urządzenia ma znaczenie niebagatelne.

A trzeba otwarcie przyznać, że smartfon należy do kategorii tych większych. Jego rozmiary to: 172,14 mm wysokości, 76,79 mm szerokości i 9,66 mm grubości. Z gabarytów tego sprzętu trzeba jasno zdać sprawę. Nie jest oczywiste to ani wada, ani zaleta. Po prostu jest to produkt dla ludzi, którzy takich właśnie rozmiarów poszukują.

Dzięki gabarytom G9 mieści duży 6,8-calowy wyświetlacz Max Vision IPS z otworem na kamerę do selfie umieszczonym w lewym rogu. Ekran ma proporcje 20,5:9. Obraz wyświetlany jest w rozdzielczości HD+, a dokładniej 720×1640 pikseli. Oprogramowanie ekranu oferuje filtr światła niebieskiego (można samodzielnie sterować nasileniem żółtej barwy, a także zdecydować, czy filtr ma aktywować się według zaplanowanego harmonogramu).

We współczesnych telefonach pewnej klasy standardem jest zestaw kilku obiektywów. Moto G9 Power oferuje kamerę główną o rozdzielczości 64 megapiksele z rozmiarem pojedynczego piksela 1,4 µm i przystoną obiektyw f/1.8. Rozdzielczość wideo to Full HD (1080p) przy 60 klatkach na sekundę. Druga kamero to dwumegapikselowy obiektyw makro z minimalną odległością ustawiania ostrości 2,5 cm. Jest jeszcze obiektyw odpowiedzialny za głębię ostrości z autofocusem (2 MP, f/2,4, 1,75 µm). Kamera po przedniej stronie telefonu ma rozdzielczość 16 megapikseli z otworem względnym obiektywu f/2,2 i rozmiarem pikseli 1 µm. Nagrywać filmy można za pomocą kamery przedniej w rozdzielczości 1080p przy 30 fps.

Przetwarzanie danych napędzane jest w tym urządzeniu układem Qualcomm Snapdragon 662, wy-



produkowanym w jedenastonanometrowym procesie litograficznym i składającym się z ośmiordzeniowego procesora (4× Kryo 260 Gold taktowany zegarem 2,0 GHz + 4× Kryo 260 Silver 1,8 GHz), wspartym przez układ graficzny Adreno 610 i połączonym z czterema gigabajtami pamięci RAM. Smartfon pracuje pod kontrolą systemu Google Android 10, prawie czystego, tzn. nie ma w nim specjalnych dodatków i przeróbek producenta, za co wielu użytkowników może okazywać Motoroli wdzięczność.

Oprócz budzącego szacunek akumulatora o pojemności 6000 mAh smartfon oferuje technikę szybkiego ładowania Turbo Power o mocy 20 W.

Telefon obsługuje protokół łączności NFC, czyli w praktyce – płatności zbliżeniowe, ma modem LTE, hybrydowy slot na karty nanoSIM (obie karty SIM ze wsparciem LTE), obsługuje Bluetooth 5.0 a także GPS (z AGPS, LTEPP, SUPL, Glonass, Galileo). Ma umieszczony na tylnym panelu czytnik linii papilarynych, za pomocą którego można szybko i bez trudności, przynajmniej my żadnych trudności podczas testów nie odnotowaliśmy, odblokować ekran. ■

M.U.

*** Pisownia oryginalna ***



PRZEGLĄD TECHNICZNY

Sukcesy Forda w kolejnictwie

Właściciel znanej fabryki samochodów, Ford, nabył w r. 1920-m zaniebaną powzglem technicznym i przynoszącą stale deficyt linję kolejową długości ok. 930 km, prowadzącą z miejscowości w południowej części stanu Ohio do miasteczka, położonego obok znacznego ośrodka przemysłowego, miasta Detroit w stanie Michigan. Finanse tej kolei były w tak złym stanie, że gdy jeden z banków amerykańskich ogłosił, że nabywa wszelkie ilości akcji tego towarzystwa za 60% ceny nominalnej, akcje uprzywilejowane za połowę ceny i wreszcie obligacje za 1/100 ceny nominalnej – 98% posiadaczy tych walorów pospieszyło skorzystać z tak korzystnej oferty. Nabywcami byli H. Ford i syn jego. Ku ogólnemu zdziwieniu fachowców pierwszym krokiem nowych właścicieli było podniesienie wszystkich płac o 20% przy jednoczesnej niższej taryfy kolejowej również o 20%. Rezultaty tego paradoksalnego rozporządzenia były zadziwiające: podczas gdy w czerwcu 1920 r. na każdy dolar dochodu przypadało 1,18 dol. wydatków, to pod nowymi rządami, w czerwcu 1921 r. liczba ta spadła do 53 ct.; innymi słowy kolej zaczęła przynosić dochody. Cały wyższy personel administracyjny został usunięty i zastąpiony przez osobistości z pośród pracowników fabryki samochodów Forda. Przykładano wszelkich starań, aby pośród pracowników kolejowych obudzić zainteresowanie do swej pracy. W tym celu, niezależnie od wspomniane-

go powyżej zwiększenia płac, wprowadzono 8-miogodzinny dzień pracy, starając się jednocześnie, w miarę możliwości, pozostawić również dzień niedzielnego do rozporządzenia pracowników. Mianowicie ruch pociągów w niedzielę został znacznie ograniczony. Z drugiej strony, nowy zarząd dbał o to, aby 8-miogodzinny dzień pracy był jej rzeczywiście poświęcony. Skasowano ścisły podział na specjalne rodzaje pracy i stopnie służbowe, dzięki czemu urzędnik, mający czas wolny (np. po odejściu pociągów), mógł pomagać innemu w jego pracy i, w następstwie, sprawy załatwiano szybko. W parze z tem szła redukcja ilości pracowników: z 2700 na 1650. Natomiast ilość towarów przewiezionych w ciągu 3 miesięcy nowej gospodarki dorównała ilości osiągniętej dawniej w ciągu 1/2 roku. Zmalała liczba reklamacji z powodu strat w towarach. Koszty biurowe spadły o 50%; koszty zużytego węgla, pomimo zwiększonego ruchu pociągów, spadł z 1/3 na 1/5 wydatków ogólnych. Wyraznym dowodem uzdrowienia stosunków na kolei tej jest wzrost dochodów z 319079 dol. (czerwiec 1920 r.) do 686355 dol. (czerwiec 1921 r.); odwrotnie wydatki obniżyły się z 476916 dol. na 376383 dol.

1 listopada 1921

Wystawa polskiej sztuki

drukarskiej w Warszawie Związek Polskich Grafików w Warszawie urządził w połowie grudnia r. b. wystawę polskiej sztuki drukarskiej w kamienicy Baryczków na Starem Mieście. Wystawa obejmuje trzy działy: I) retrospektywny, II) drukarstwo współczesne: wzory techniki drukarskiej, książki, ilustracje, ex-librisy, druki akcydensowe, papiery wartościowe, karty do gry, etykiety, marki handlowe, dewocjonalja drukarskie, afisze i III) Introligatorstwo retrospektywne i współczesne.

9 listopada 1921

PRZEGLĄD

PRZEMYSŁOWO-HANDLOWY „Elektryczność”

W dobie odbudowywania się Polski po tej długoletniej

wojnie, której ślady widzimy na każdym kroku, każde przedsiębiorstwo, współdziałające ze sprawą odbudowy staje się placówką doniosłego znaczenia. Musi ono jednak opierać się nie tylko na zdrowych podstawach materialnych, lecz przede wszystkim na sprzężym kierownictwie, świadomem swojej roli i pracy, popartej gruntowną znajomością zawodu oraz wszelkich nowoczesnych zdobyczy techniki. Do rzędu w ten sposób zorganizowanych i tak pojmujących swoją działalność przedsiębiorstw należy „Elektryczność”, Józef Nagórski i S-ka we Lwowie, przy ulicy 3 Maja pod № 15. Powyższe przedsiębiorstwo, założone w r. 1910, wykonuje wszelkie roboty i dostawy z dziedziny elektrotechniki, posiadając odpowiednio zorganizowany i liczny personel techniczny, skład wszelkiego rodzaju przyborów oraz własne warsztaty do wyrobu lamp. Posiadając znakomite, fachowe kierownictwo głównego firmanta, inżyniera Józefa Nagórskiego, „Elektryczność” daje zupełną gwarancję, iż każda z robót, powierzona temu przedsiębiorstwu, wykonana będzie solidnie, według ścisłej i uczciwej kalkulacji technicznej, z materiału pierwszorzędnej jakości i z zastosowaniem rozumnej oszczędności. Ten system pracy, stosowany w przedsiębiorstwie „Elektryczność”, Józef Nagórski i S-ka sprawił, iż obecnie wszystkie znaczniejsze instytucje, zarówno rządowe jak i prywatne, oraz liczne firmy przemysłowe i handlowe korzystają stale z usług „Elektryczności”, dającą tem zaufaniem, o którego pozyskanie najgłośniej zabiega, rozumiejąc, że tylko na gruncie zaufania może rozwijać się każda poważna działalność.

listopad 1921

Nasz tabor kolejowy szerokotorowy

Polska posiada na podstawie traktatu ryskiego z 18 marca 1920 r. znaczny tabor kolejowy szerokotorowy, a mianowicie: ośm tysięcy wagonów i dwieście kilkadziesiąt

parowozów. Tabor ten ułatwił nawiązanie stosunków ekonomicznych z Rosją i Ukrainą sowiecką. Wyzyskanie posiadanego taboru staje się aktualne w związku z ewentualnym zawarciem umowy handlowej z Rosją i Ukrainą. Sprawą tą zainteresowały się prywatne przedsiębiorstwa handlowe i transportowe. Następujące firmy wyraziły chęć eksploataowania taboru szerokotorowego: Polskie Domy Składowe we Lwowie, T-wo Akc. „Hartwig”, Sp. Akc. „Wawel”, Polsko-bałtyckie Towarzystwo Transportowe, Warszawskie Towarzystwo dla transportu i żeglugi, Sp. Akc. Zjednoczone Domy ekspedycyjne, Polskie T-wo Komunikacyjne, Przedsiębiorstwo szerokotorowej kolei Warszawa – granica rosyjska i Polski Glob. Również Szwajcarski Czerwony Krzyż zamierza wydzierżawić powyższy tabor w celu niesienia skutecznej pomocy głodnym w Rosji.

listopad 1921

Polski Komitet Chłódnotwa

Z inicjatywy Ministerstwa Przemysłu i Handlu oraz innych zainteresowanych Ministerstw powstał komitet chłódnotwa. Ma on siedzibę swoją w Warszawie i dąży do rozwoju chłódnotwa, które znajduje szerokie zastosowanie dla konserwowania produktów spożywczych. Sprawa ta ma szczególne znaczenie obecnie, kiedy ilość środków spożywczych jest b. ograniczona.

listopad 1921

Błoto defekacyjne

Pomorska Izba Rolnicza zwraca uwagę sfer rolniczych na możliwość wykorzystania dla celów nawozowych błota defekacyjnego, którego mogą dostarczyć cukrownie. Wobec dotkliwego braku nawozów sztucznych rolnicy powinni zainteresować się tą sprawą, tembardziej, że zagranica już zwróciła się do naszego Rządu z propozycją zakupienia tego błota. Błoto defekacyjne zawiera po wysuszeniu dosyć znaczny procent azotu, potasu, wapna i kwasu fosforowego.

listopad 1921

Nie przegap listopadowego wydania „Elektroniki dla Wszystkich”



W numerze między innymi:

Miernik wzmacniaczy operacyjnych

Znakomitym, praktycznym uzupełnieniem cyklu Droga do RRIO jest przedstawiony w tym numerze uniwersalny prosty miernik, który pozwala w bardzo łatwy sposób zmierzyć wszystkie kluczowe parametry większości wzmacniaczy operacyjnych.

Inteligentny dom także dla Ciebie! New Kid in Town, czyli odkrywanie „nowego ESP8266”

Mnóstwo urządzeń przeznaczonych do inteligentnego domu wykorzystuje znany ESP8266. Bardzo interesujące jest też inne rozwiązanie, prezentowane w artykule.

Współczesne akumulatory

Wskutek pytań Czytelników powstał obszerny materiał o współczesnych akumulatorach. Omówione będą głównie liczne tajemnice akumulatorów litowych, ale dla pełnego obrazu trzeba też wrócić do akumulatorów zasadowych i kwasowych.

Dokładność modułowych mierników prądu stałego

Niektóre tanie modułowe mierniki napięcia i prądu mają nieoczekiwanie dużą precyzję. Warto się dowiedzieć, z czego ona wynika oraz jak poznać takie najlepsze moduły.

Rezystancja kabli

Zasilacze, stabilizatory, mierniki oraz przystawki do pomiaru prądu mają coraz lepsze parametry, ale mogą one zostać popsute przez nieoczekiwanie dużą rezystancję przewodów i styków.

Ponadto w numerze:

- Filozofia sieci. Protokół TCP
- Czym się różni wróblek? Czyli o niedoskonałości elementów. Część 9 – Rezystory praktycznie
- Sieci energetyczne dla elektroników
- Szkoła Konstruktorów:
 - Zastosowanie elektroniki do kontroli procesu kompostowania
 - Zbadaj niedoskonałości elementów rezystancyjnych (na przykład różnego rodzaju rezystorów, przewodów, styków)

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