



styczeń 2026

www.mlodytechnik.pl



Tu przejrzysz
i kupisz ten numer

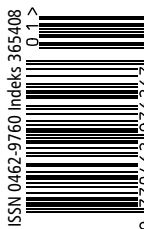
NEWS 24/7
przeglądaj codziennie
na swoim smartfonie

mlody **m.technik**

Ciekawi świata są zawsze młodzi

NAUKA 4.0

Najnowsza technika wspomaga i... zastępuje uczonych



ISSN 0462-9760 Indeks 365408

cena: **14,90 zł** (w tym 8% VAT)

Jak AI zamienia twarze w obrazach?

O sieciach adwersaryjnych

WYDANIA SPECJALNE

Młodego Technika

„**Na Warsztacie**” to połączenie praktycznej wiedzy z kreatywną zabawą. Niezależnie od tego, czy jesteś początkującym miłośnikiem techniki, czy masz już doświadczenie, znajdziesz tu projekty, które Cię zainteresują.



ZAMÓW PRZEZ QR KOD
LUB NA ULUBIONYKIOSK.PL



Temat okładkowy
Hi-tech znalazł ciekawe zastosowania w poszukiwaniach archeologicznych. Swój udział w tym ma też sztuczna inteligencja, choć jej rola we współczesnych badaniach naukowych nie jest jednoznacznie pozytywna...

Nauka niehumanitarna, ale coraz bardziej obiecująca

Z jednej strony czytamy i słyszymy, że coraz więcej treści w publikacjach mających status „naukowych” generowanych jest przez sztuczną inteligencję. Z drugiej – pojawiają się zagadkowe doniesienia o tym, że AI, którą angażuje się do badań, dochodzi do wyników lub proponuje rzeczy, które uczonym trudno pojąć.

W obu przypadkach nauka przestaje być czymś „ludzkim”, stając się obszarem... no właśnie, trudno określić – jakim. Odrzucanie wszystkiego, co pochodzi od sztucznej inteligencji jako „halucynacji”, błędów, pozbawionego sensu bełkotu, może być pochoptne. Nie można wykluczyć, że AI wskazuje nam takie aspekty i naukowe ścieżki, których przez swoje ludzkie uprzedzenia nigdy

*Nowa technika daje
zastanawiające
wyniki w badaniach
naukowych*

byśmy nie odkryli. Jest wielu widzących w tych narzędziach szansę na zupełnie nowe otwarcie w nauce.

Badamy „odciski palców AI” w nauce. Jest ich już

sporo. Ma to wydźwięk, jak sugerujemy już na wstępie, zarówno pozytywny, jak i negatywny. Niektórzy uważają, że nie powinniśmy oczekiwać od sztucznej inteligencji żadnych nowych odkryć. Jej zadaniem, wskazują, jest pomoc w radzeniu sobie z ogromem danych, które generują nasze instrumenty, laboratoria, ośrodki naukowe. Jak wskazuje przykład giga repozytorium „surówki” badawczej, wygenerowanej w akceleratorach CERN, ludzie, choćby byli najzdolniejszymi naukowcami, są w stanie przetworzyć, przeanalizować jedynie ich niewielki ułamek. AI być może da sobie z tym radę i dostarczy nam analiz i syntez, których potrzebuje nauka.

Najnowocześniejsze techniki, nie tylko AI, rewolucjonizują dziedzinę badań starożytności, archeologii, analizę artefaktów, które były wcześniej niedostępne lub niezrozumiałe. Być może uda nam się, dzięki nim, zgłębić tajemnice starych form „pisma”, które dręczą świat nauki od dekad. Już teraz, dzięki nowym technikom skanowania powierzchni Ziemi, odkrywamy artefakty starych kultur, w miejscach, które nas zaskakują. A nauka 4.0 dopiero się rozkręca.

Mirosław Usidus

PRENUMERATA

*Czytaj więcej,
płać mniej!*



Zyskaj
15%
rabatu

W prenumeracie tylko
178,80 zł

152,00 zł

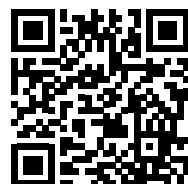
/roczna prenumerata drukowana

Dlaczego warto?

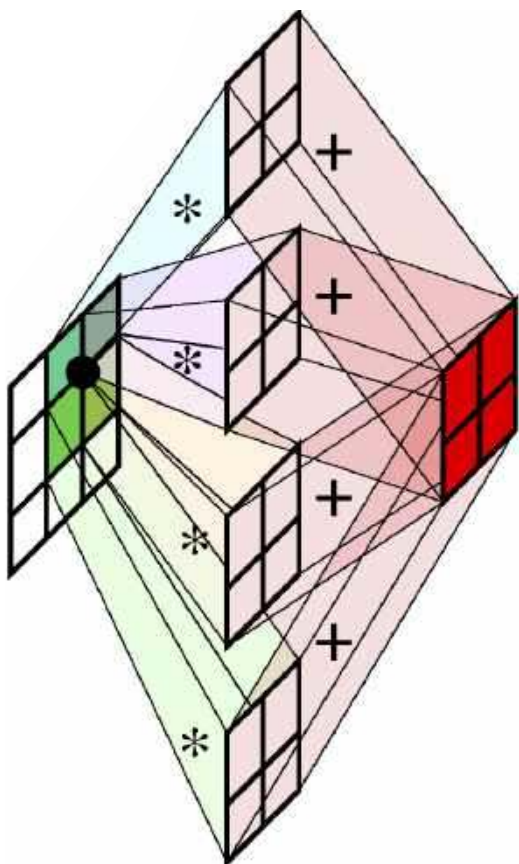
- ▶ Dostawa **gratis** prosto do Twojego domu
- ▶ Tylko dla prenumeratorów: **niższe ceny** przy zakupie czasopism na UlubionyKiosk.pl
- ▶ Pakiet 2w1 (papier + e-wydania):
–80% na równoległą e-prenumeratę PDF

Szczegóły na UlubionyKiosk.pl/promocje

Zamów prenumeratę na www.UlubionyKiosk.pl
lub zeskanuj kod QR i zaprenumeruj w 1 minutę



AVT-Korporacja sp. z o.o., ul. Leszczynowa 11, 03-197 Warszawa
prenumerata@avt.pl | 22 257 84 22 (godz. 10:00–14:00)
rachunek bankowy: ING Bank Śląski 18 1050 1012 1000 0024 3173 1013
eprasa.pl 236bf595de



W jaki sposób AI
generuje obrazy?

51

Neuronowe modele dyfuzyjne

W tym wydaniu MT m.in.:

- **Horyzonty mgłą spowite:**
Czyżby wreszcie przełom w bateriach?
Na horyzoncie pojawiają się dające nadzieję rozwiązania akumulatorów z elektrolitem stałym i półprzewodnikowych
- **Koniec i co dalej: Spotify?**
Spotify kończy się albo... wręcz przeciwnie

Temat numeru: Nauka 4.0 – najnowsza technika wspomaga i... zastępuje uczonych

- 24 • Czy pojmemy naukę, gdy zajmie się nią sztuczna inteligencja? Tureckie kazanie profesora AI
- 29 • Najnowsza technika w odkrywaniu przeszłości. Lidarowym okiem w głąb dziejów
- 37 • Maszyna w zespole, na dobre i na złe. Odciski palców AI w nauce
- 46 • Rewolucja w krainie królowej nauk? Przyszłość mAltematyki

Technika

- 8 Info Zoom
- 16 Dodaj do obserwowanych
Horyzonty mgłą spowite
- 17 • Czyżby wreszcie przełom w bateriach? Ognia, o jakie chodziło
- 20 • Projekt Starcloud. Serwerownia na orbicie
- 22 • Pokłosie wielkiej awarii AWS. Chmury nad chmurami
- 51 Sztuczna inteligencja: W jaki sposób AI generuje obrazy? Neuronowe modele dyfuzyjne

Fantastyka naukowa w „Młodym Techniku”

- 60 Wyniki konkursu literackiego PFFN
- 61 • Zlecenie
- 63 • Marsjanie czy jeszcze Ludzie?
- 64 • Ludzie... Ludzie!

Szkola

- 67 Koniec i co dalej: Spotify? Śmierć przez... dochodowość
- 70 MT studiuje: Inżynieria jakości
- 72 Fizyka bez granic: Zjawisko indukcji elektromagnetycznej
- 75 Chemia inna niż w szkole: Chemiczny Nobel 2025, czyli szukanie dziury w całym
- 80 Matematyka z ludzką twarzą: O dużych liczbach Klub i Szkoła Wynalazców
- 85 • Szkoła Wynalazców – dozwolone do lat 15
- 86 • Klub Wynalazców – bez ograniczeń wieku
- 87 • Vademecum Młodego Wynalazcy
Odkryj historię wynalazków
- 90 • Tkaniny
- 94 • Klasyfikacja tkanin
- 95 Pomysły genialne, zwariowane i takie sobie

Hobby

- 96 Akademia audio: W labiryntach, część 3

- 3 Od wydawcy
- 4 Prenumerata
- 6 Listy
- 99 Sędziwy Technik – 100 lat temu prasa pisała

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.

Konsultacja graficzna:
Małgorzata Jabłońska

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

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

Prenumerata:
www.ulubionykiosk.pl
tel. 22 257 84 22 (godz. 10.00–14.00)
e-mail: prenumerata@avt.pl

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

List miesiąca

Internet kwantowy

Piszę w odniesieniu do artykułu, który przeczytałem w Państwa magazynie, na temat przesyłu Internetu kwantowego przy użyciu kabli światłowodowych. Jest to temat szczególnie ciekawy dla mojej osoby i chciałbym wnieść do niego kilka wartościowych informacji.

Termin „Internet kwantowy” budzi zazwyczaj skojarzenia z czymś rewolucyjnym i całkowicie odmiennym od obecnej infrastruktury telekomunikacyjnej. Tymczasem rzeczywistość jest bardziej złożona i – co równie istotne – bardziej osiągalna, niż mogłoby się wydawać. Kluczową kwestią, którą warto przybliżyć szerszej publiczności, jest fakt, że rozwijające się dziś technologie komunikacji kwantowej mogą w dużym stopniu wykorzystywać istniejącą infrastrukturę światłowodową.

Zanim przejdę do sedna sprawy, pozwolę sobie na krótkie wyjaśnienie. Internet kwantowy to sieć komunikacyjna wykorzystująca zjawiska mechaniki kwantowej, przede wszystkim splątanie kwantowe i kwantową dystrybucję klucza (QKD – Quantum Key Distribution). Jego głównym celem nie jest przyspieszenie przesyłu danych w rozumieniu większej przepustowości, lecz zapewnienie absolutnie bezpiecznej komunikacji, odpornej nawet na przyszłe komputery kwantowe.

Tradycyjne metody szyfrowania, takie jak algorytm RSA, mogą zostać złamane przez wystarczająco potężne komputery kwantowe. QKD natomiast opiera się na fundamentalnych prawach fizyki – każda próba podsłuchiwania komunikacji kwantowej nieuchronnie wprowadza zakłócenia, które są wykrywane przez komunikujące się strony.

Tutaj dochodzimy do kluczowej kwestii: czy potrzebujemy całkowicie nowej infrastruktury dla Internetu kwantowego? Odpowiedź brzmi: niekoniecznie. Współczesne światłowody, które już teraz stanowią kręgosłup globalnej komunikacji internetowej, nadają się – przynajmniej w teorii i w coraz większym stopniu w praktyce – do przesyłu fotonów będących nośnikami informacji kwantowej.

Światłowod to nic innego jak bardzo cienkie włókno szklane lub plastikowe, które przewodzi światło dzięki zjawisku całkowitego wewnętrznego odbicia. Pojedyncze fotony, którymi operuje komunikacja kwantowa, mogą być przesyłane przez te same kanały, co strumienie fotonów używane w klasycznej komunikacji optycznej. To sprawia, że istniejąca infrastruktura może służyć jako fundament dla nowych zastosowań.

Oczywiście, nie jest to tak proste, jak mogłoby się wydawać. Przesył informacji kwantowej stawia przed światłowodami bardziej wymagające warunki niż klasyczna komunikacja. Po pierwsze, pojedyncze fotony są niezwykle wrażliwe na straty. W standardowym światłowodzie spora część fotonów zostaje zaabsorbowana lub rozproszona na odległości już kilkudziesięciu kilometrów. Dla klasycznej komunikacji nie stanowi to większego problemu – można wzmocnić sygnał za pomocą regeneratorów. Jednak w przypadku stanów kwantowych tradycyjne wzmacnianie prowadzi do ich zniszczenia zgodnie z zasadą niekopiowania informacji kwantowej.

Naukowcy pracują nad rozwiązaniami tego problemu. Jednym z nich są kwantowe repetytory – urządzenia pozwalające na przedłużenie zasięgu komunikacji kwantowej poprzez wykorzystanie splątania kwantowego i teleportacji kwantowej. Innym podejściem jest optymalizacja samych światłowodów – minimalizacja strat, redukcja szumów termicznych czy wykorzystanie specjalnych długości fal, dla których tłumienie jest najmniejsze.

Sz szczególnie interesującym zagadnieniem jest możliwość współistnienia w tym samym światłowodzie zarówno klasycznej komunikacji internetowej, jak i kwantowej dystrybucji klucza. Badania prowadzone przez zespoły naukowe na całym świecie pokazują, że jest to możliwe, choć wymaga starannego zarządzania widmem optycznym. Kanały kwantowe mogą wykorzystywać inne długości fal niż ruch klasyczny, co pozwala na jednoczesne funkcjonowanie obu systemów bez wzajemnych zakłóceń.

Takie rozwiązanie ma ogromne znaczenie praktyczne. Oznacza bowiem, że nie musimy budować całkowicie odrębnej sieci światłowodowej dla komunikacji kwantowej. Możemy stopniowo integrować elementy kwantowe z istniejącą infrastrukturą, co znacząco obniża koszty wdrożenia i przyspiesza proces implementacji.

Warto podkreślić, że nie mówimy tu o odległej przyszłości. Już dziś funkcjonują eksperymentalne i komercyjne sieci QKD wykorzystujące standardowe światłowody. Chiny uruchomiły sieć rozciągającą się na ponad dwa tysiące kilometrów, łącząc Pekin z Szanghajem. Podobne projekty realizowane są w Europie, między innymi w ramach inicjatywy European Quantum Communication Infrastructure (EuroQCI), która ma na celu stworzenie ogólnoeuropejskiej sieci komunikacji kwantowej opartej w dużej mierze na istniejących połączeniach światłowodowych.

Rozwój Internetu kwantowego przez zwykłe światłowody to przykład tego, jak nowe technologie mogą budować na fundamentach starych, nie wymagając każdorazowo rewolucyjnej przebudowy całej infrastruktury.

Z wyrazami szacunku,

Andrzej z Łęborka

Perowskity

Szanowna Redakcjo,

z wielkim zainteresowaniem śledzę rozwój technologii materiałowych i chciałbym zwrócić uwagę Czytelników na fascynujący temat, który został poruszony przez Państwa w ostatnim wydaniu „Młodego Technika” i może zrewolucjonizować medycynę w nadchodzących latach – chodzi oczywiście o zastosowanie perowskitów w diagnostyce i terapii medycznej.

Perowskity to materiały o strukturze krystalicznej nazywanej na cześć rosyjskiego mineraloga Lwa Perowskiego, przez długi czas kojarzone były głównie z zastosowaniami w elektronice i fotowoltaice. Te niezwykle materiały, charakteryzujące się ogólnym wzorem chemicznym ABX_3 , gdzie A i B to kationy o różnych rozmiarach, a X to anion (najczęściej tlen lub halogeny), tworzą unikalną strukturę krystaliczną przypominającą kostkę. To właśnie ta geometria nadaje perowskitom ich wyjątkowe właściwości fizyczne i chemiczne. Jednak ostatnie lata przyniosły przełom w badaniach nad ich potencjałem medycznym, otwierając przed nami zupełnie nowe możliwości diagnostyczne i terapeutyczne, których jeszcze dekadę temu nie mogliśmy sobie wyobrazić.

Jednym z najbardziej obiecujących kierunków jest wykorzystanie perowskitowych kropek kwantowych w obrazowaniu biomedycznym. Te nanostruktury wykazują wyjątkowe właściwości optyczne, w tym intensywną i regulowaną luminescencję, wysoką stabilność fotochemiczną oraz możliwość precyzyjnego dostrojenia długości fali emisji poprzez modyfikację składu chemicznego. W przeciwieństwie do tradycyjnych barwników organicznych, które szybko ulegają wyblaknięciu pod wpływem światła (zjawisko fotobieleńcia), perowskitowe kropki kwantowe zachowują stabilną emisję przez dłuższy czas, co umożliwia prowadzenie długotrwałych obserwacji procesów biologicznych w czasie rzeczywistym.

Co więcej, perowskity wykazują szerokie spektrum absorpcji przy jednoczesnej wąskiej emisji, co pozwala na jednoczesne wykorzystanie wielu różnych znaczników w tej samej próbce bez wzajemnych zakłóceń sygnałów. Ta właściwość, zwana multipleksowaniem, jest niezwykle cenna w zaawansowanych badaniach biologicznych, gdzie naukowcy muszą jednocześnie śledzić wiele różnych molekuł lub struktur komórkowych. Dzięki tym cechom perowskity mogą stać się alternatywą dla obecnie stosowanych barwników fluorescencyjnych i znaczników biologicznych, oferując znacznie lepszą rozdzielczość obrazowania komórkowego i tkankowego.

Szczególnie fascynujące są badania nad wykorzystaniem perowskitów w mikroskopii superrozdzielczej, technice obrazowania przekraczającej klasyczne ograniczenia dyfrakcyjne światła. Perowskitowe znaczniki umożliwiają wizualizację struktur subkomórkowych z rozdzielczością zbliżoną do nanometrowej, co otwiera zupełnie nowe możliwości w badaniach mechanizmów molekularnych chorób i działania leków na poziomie pojedynczych cząsteczek.

Co więcej, niektóre perowskity wykazują zdolność do konwersji światła podczerwonego na światło widzialne w procesie zwanym konwersją w górę. Ma to ogromne znaczenie praktyczne, ponieważ światło podczerwone

penetruje tkanki znacznie głębiej niż światło widzialne, co potencjalnie umożliwia leczenie guzów położonych głęboko w organizmie, niedostępnych dla konwencjonalnej terapii fotodynamicznej. Badania przedkliniczne prowadzone na modelach zwierzęcych wykazują obiecujące rezultaty, z istotnymi redukcjami objętości guzów i wydłużeniem czasu przeżycia przy minimalnych efektach ubocznych.

W neurologii badacze eksplorują możliwości wykorzystania perowskitów w interfejsach mózg-komputer. Ich właściwości elektrofizjologiczne mogą umożliwić bardziej precyzyjne rejestrowanie i stymulację aktywności neuronalnej, co ma ogromne znaczenie dla pacjentów z chorobami neurodegeneracyjnymi czy po urazach rdzenia kręgowego.

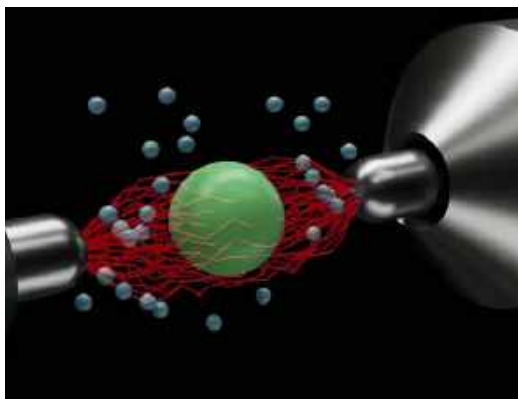
Oczywiście, droga od laboratoryjnych eksperymentów do praktycznych zastosowań klinicznych jest długa i wyboista. Głównym wyzwaniem pozostaje kwestia biotoksyczności niektórych perowskitów, szczególnie tych zawierających ołów. Ołów, będący jednym z najsilniejszych składników w najwydajniejszych perowskitach, jest znanym neurotoksyną i kumuluje się w organizmie, co stanowi poważne zagrożenie dla zdrowia. Naukowcy intensywnie pracują nad opracowaniem perowskitów bezołowiowych, w których ołów zastępowany jest innymi metalami, takimi jak cyna, bizmut, german czy srebro, lub nad skutecznymi metodami enkapsulacji, które zapobiegają uwolnieniu toksycznych jonów w organizmie.

Enkapsulacja perowskitów w powłokach z biozgodnych polimerów, tlenku krzemu czy innych materiałów barierowych, jest przedmiotem intensywnych badań. Idealna powłoka ochronna musi spełniać kilka pozornie sprzecznych wymagań – być całkowicie szczelna dla jonów metali ciężkich, jednocześnie pozostając przepuszczalna dla analitytów, które mają być wykrywane, oraz nie zakłócać właściwości optycznych czy elektrycznych perowskitu. To inżynierskie wyzwanie wymaga interdyscyplinarnego podejścia łączącego wiedzę z chemii, nauki o materiałach, biologii i medycyny. Kolejnym wyzwaniem jest stabilność perowskitów w środowisku biologicznym. Wiele z tych materiałów ulega degradacji pod wpływem wilgoci, co w środowisku wodnym organizmu ludzkiego stanowi poważny problem.

Perowskity w medycynie reprezentują jeden z najbardziej ekscytujących kierunków rozwoju współczesnej nauki biomedycznej. Od obrazowania komórkowego o bezprecedensowej rozdzielczości, poprzez ultraczułą diagnostykę nowotworów, innowacyjne terapie fotodynamiczne, biosensory nowej generacji, aż po rewolucyjne interfejsy mózg-komputer – potencjał tych materiałów wydaje się niemal nieograniczony.

Choć wyzwania związane z biotoksycznością, stabilnością i regulacjami pozostają znaczące, tempo postępu badawczego i rosnące zaangażowanie zarówno środowisk akademickich, jak i przemysłu, napawają optymizmem. Mamy uzasadnione podstawy, by sądzić, że już w najbliższej dekadzie pierwsze perowskitowe technologie medyczne znajdą zastosowanie w praktyce klinicznej, rozpoczynając nową erę w diagnostyce i leczeniu chorób.

Michał z Białegostoku



FIZYKA

Mikroskopijny silnik nagrany do milionów stopni

W artykule w czasopiśmie „Physical Review Letters” naukowcy opisują mikrosilnik, o rozmiarach mniejszych niż żywa komórka, wewnątrz drobiny uwięzionej w próżni w polu elektrycznym, który osiąga temperaturę sięgającą dziesięciu milionów kelwinów, przekraczającą znacznie temperaturę powierzchni Słońca.

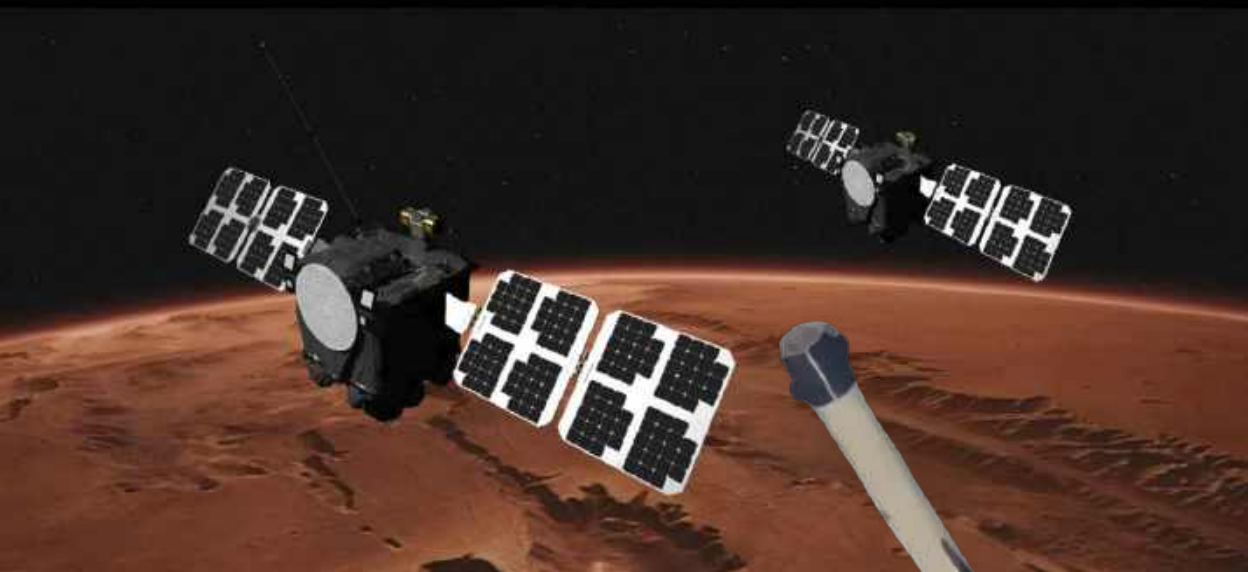
W silniku tym elektrody wychwytyują i utrzymują mikrocząsteczki w stanie zawieszenia w układzie zbliżonym do próżni, zwanym pułapką Paula. To rodzaj pułapki jonowej, która wykorzystuje dynamiczne pola elektryczne do wychwytywania naładowanych cząstek. Nazywana jest również pułapką częstotliwości radiowej. Kiedy naukowcy podłączyli do elektrod napięcie z zakłóceniami, cząsteczka zaczęła intensywnie drgać, powodując wykładniczy wzrost temperatury całego układu.

Wyniki były zastanawiające. Według naukowców silnik przechodził od wysokiej wydajności do całkowitego zaprzeczania podstawowym prawom termodynamiki. W niektórych cyklach moc wyjściowa silnika przekraczała energię, którą zużywał. W innych przypadkach silnik losowo się ochładzał w warunkach, które powinny powodować jego przegrzanie. Naukowcy zauważyli, że prawdopodobnie wynikało to z działania niewidocznych sił, biorąc pod uwagę niewielkie rozmiary systemu. Zastosowanie tego mikroukładu badacze widzą głównie w symulacji zjawisk z mikroświata, np. procesów kształtowania cząsteczek proteinowych. ■



Rakieta New Glenn firmy Blue Origin założyciela Amazona, Jeffa Bezosa, wystartowała ze stacji amerykańskich sił kosmicznych w Cape Canaveral, wynosząc międzyplanetarną misję i co może ważniejsze, gdyż nie był to pierwszy start tej rakiety, po raz pierwszy udało się przeprowadzić lądowanie głównym członem na morzu.

Rakieta Blue Origin wyniosła dwa statki kosmiczne w ramach misji ESCAPEDE Mars NASA, które podążają obecnie w kierunku Czerwonej Planety. ESCAPEDE to skrót od pełnej nazwy misji: Escape and Plasma Acceleration and Dynamics Explorers. Składają się na nią dwie komercyjne sondy, które będą badać, w jaki sposób Mars, na którego powierzchni niegdyś występowała woda w stanie ciekłym, z czasem utracił atmosferę, stając się suchą planetą, jaką dziś znamy. Zbudowane przez Rocket Lab dla NASA i Uniwersytet Kalifornijski w Berkeley kosztowały mniej niż 80 milionów dolarów, co w porównaniu z kosztami poprzednich misji marsjańskich NASA jest bardzo niewygórowana ceną. Misja ta ma przetestować sposób podróży na Marsa. Zamiast tradycyjnej, oszczędnej trasy stosowanej w poprzednich misjach, polegającej na korzystaniu z wąskiego okna startowego co dwa lata, wykorzystującego ustawienie planet, statek ESCAPEDE wybierze trasę, która najpierw zaprowadzi go do stabilnego punktu między Ziemią a Słońcem,



KOSMOS

Rakieta New Glenn startuje i ląduje tak samo jak rakiety Muska

tw. punktu Lagrange'a 2. Po roku krążenia wokół tego punktu statek zawróci w kierunku Ziemi, a następnie skieruje się w stronę Marsa. Start przebiegł bez żadnych zakłóceń, jeśli nie liczyć wcześniejszego jego przełożenia z powodu złej pogody słonecznej. Gdy drugi stopień rakiety kontynuował lot, aby wynieść ESCAPADE w przestrzeń kosmiczną, pierwszy stopień New Glenn rozpoczął serię manewrów hamujących w celu przeprowadzenia ladowania statku Blue Origin „Jacklyn”, który czekał 604 kilometry od miejsca startu na Oceanie Atlantyckim. Tym samym Blue Origin zrobiło wielki krok w kierunku dołączenia do SpaceX na lukratywnym rynku tanich lotów kosmicznych z odzyskiem rakiet.

Tymczasem SpaceX, która po przeprowadzeniu udanego testu rakiety Starship 11, szykuje się do kolejnych lotów swojej wielkiej rakiety w nowej

konfiguracji – misji nr 12, 13 i 14, opublikowała nowy, uproszczony plan budowy lądownika księżycowego, który mógłby przenieść astronautów na Księżyc w ramach programu Artemis. Konkurencja dotycząca budowy lądownika została ponownie otwarta przez NASA z powodu opóźnień innych projektów. Wszystkie te wydarzenia działy się równocześnie z dramatem trzech chińskich tajkonautów, Wanga Jie, Chen Zhongrui i Chen Donga, którzy z powodu uderzenia orbitalnego gruzu utknęli na pokładzie chińskiej stacji kosmicznej Tiangong. Udało się ich sprowadzić na Ziemię na pokładzie kapsuły, która wyniosła ich zmienników. ■



Krótki reportaż ukazujący start i udane lądowanie rakiety New Glenn: <https://youtu.be/pviGLY1PiHQ>



BIOMIMETYZM

Tajemnica kamuflażu ośmiornic wydarta przez uczonych

Ksantomatynę, pigment odpowiadający za niezwykle zdolności głowonogów do zmiany ubarwienia i wzorów na skórze, udało się wyprodukować w dużych ilościach zespołowi badawczemu kierowanemu przez naukowców z Uniwersytetu Kalifornijskiego w San Diego. Do tej pory pozyskiwanie tej substancji od zwierząt lub wytwarzanie jej w laboratorium było trudne, a wszelkie znane metody nie miały praktycznego potencjału.

Badacze sami nie wytworzyli tego pigmentu w sensie technicznym. Wykorzystali bioinżynierię, modyfikując bakterie, które nie tylko wytwarzają tę rzadką substancję, ale robią to z niespotykaną dotąd wydajnością, produkując nawet tysiąc razy więcej ksantomatyny, niż działa się to z wykorzystaniem poprzednich metod. Aby uzyskać tak dużą produkcję substancji z bakterii, zastosowali nową metodę, którą nazwali „biosyntezą sprzężoną ze wzrostem”, która zachęcała mikroorganizmy do wytwarzania dużych ilości ksantomatyny poprzez powiązanie mechanizmu ich przetrwania z produkcją pigmentu. Genetycznie zmodyfikowane komórki mogły rosnąć tylko



wtedy, gdy wytwarzały dwa związki – ksantomatynę i kwas mrówkowy. Ten ostatni służył jako paliwo, a ponieważ bakteria wytwarzała jedną cząsteczkę kwasu mrówkowego na każdą nową cząsteczkę pigmentu, miała wystarczającą ilość paliwa do wzrostu, o ile wytwarzała pigment.

Technika ta pozwoliła uzyskać do 3 gramów pigmentu na litr ośrodka mikrobiologicznego. Według uczonych, jest to znacznie więcej niż w starszych metodach, w których uzyskiwano najwyżej 5 miligramów na litr. Podkreślają w publikacji na łamach „Nature Biotechnology” potencjał swojej nowej metody, która może nie tylko pozwolić na praktyczne zastosowania w technikach kamuflażu opartych na biomimetyzmie, ale także stanowi postępowanie w dziedzinie produkcji mikrobiologicznej. ■



Reportaż na temat nowej metody produkcji ksantomatyny: <https://youtu.be/eMxtGLsdz0E>

GADŻETY

Smartfon z ekranem jak e-czytnik

Firma Bigme udostępniła dwa nowe typy smartfonów, pod nazwą HiBreak S, których wyświetlacze oparte są na znanej z e-czytników technice elektronicznego tuszu (E-Ink). Jeden z nowo oferowanych modeli ma ekran monochromatyczny, drugi – barwny. Poza wyświetlaczami nie różnią się znacząco rozmiarami ani parametrami technicznymi od innych smartfonów na rynku.

Urządzenie jest napędzane przez procesor z 6 GB pamięci RAM i 128 GB pamięci wewnętrznej, z możliwością rozszerzenia do 1 TB za pomocą karty microSD. Obsługuje sieć 4G LTE. Każdy model jest wyposażony w dwa aparaty, główny aparat o rozdzielczości 13 megapikseli z tyłu i aparat do selfie

o rozdzielczości 5 megapikseli z przodu oraz „technologię płynnego rozpoznawania znaków optycznych (OCR)” do skanowania dokumentów.

Obie odmiany smartfona Bigme mają ekran o przekątnej 5,84 cala. HiBreak S oferuje rozdzielczość 720×1440 pikseli (276 ppi) oraz technikę szybkiego odświeżania, obsługującą częstotliwość do 24 klatek na sekundę. Amatorzy obrazów w wersji kolorowej otrzymują rozdzielczość tylko 240×480 pikseli (92 ppi), co oznacza, że trudno nazwać urządzenie multimedialnym. Ale też nie taki raczej był cel projektantów. Producent twierdzi, że dzięki oszczędności w technice wyświetlania akumulator o relatywnie skromnej pojemności 3300 mAh będzie mógł zasiląć urządzenie bardzo długo, choć nie wiadomo dokładnie – jak długo. ■





INTERNET

ChatGPT Atlas, czyli wojna przeglądarek wkracza na poziom AI

Firma OpenAI, twórczyni ChatGPT, udostępniła swoją zapowiadaną i oczekiwaną od wielu miesięcy przeglądarkę AI, która nosi nazwę – ChatGPT Atlas. W pierwszej kolejności cieszyć się nią mogą jedynie użytkownicy systemu MacOS Apple'a. W reakcji na tę premierę silnie spadły notowania akcji spółki macierzystej Google'a – Alphabet, gdyż Atlas postrzegany jest jako konkurencja zarówno dla Chrome'a, jak i modelu Gemini, tworzono go przez wyszukiwarkowego potentata.

Podczas prezentacji nowego programu zespół OpenAI opisywał Atlasa jako przeglądarkę AI opartą na ChatGPT, umożliwiającą użytkownikom komunikację z przeglądarką w sposób, którego nie oferują tradycyjne przeglądarki. Program integruje ChatGPT z każdym wyszukiwaniem i kartą, w czym podobny jest do udostępnionej latem przez Perplexity przeglądarki Comet. Reaguje na załadowywane do kart treści w module dialogowym – można chatbotowi zlecić np. tłumaczenie czy streszczenie wyświetlanego artykułu.

Atlas oferuje również eksperymentalny „tryb agenta”, w którym ChatGPT może przejąć nawigację i interakcję ze stroną za użytkownika. Może wykonać za niego zadania takie jak np. opracowanie planu

ćwiczeń lub posiłków, sporządzenie listy składników i dodanie produktów spożywczych do koszyka zakupowego gotowego do dostawy. Agent ten ma jednak ograniczenia. Nie można uruchamiać kodu w przeglądarce, pobierać plików ani instalować rozszerzeń. Nie można uzyskać dostępu do innych aplikacji na komputerze lub w systemie plików, odczytywać ani zapisywać pamięci ChatGPT, uzyskiwać dostępu do zapisanych haseł ani korzystać z danych autouzupełniania. Strony odwiedzane w trybie agenta nie są dodawane do historii przeglądania. Można również zdecydować się na uruchomienie agenta w trybie wylogowania, wówczas ChatGPT nie będzie korzystać z żadnych istniejących plików cookie i nie zaloguje się na żadne z kont internetowych użytkownika bez jego wyraźnej zgody. ■





NEUROTECHNIKA

Mózg przetwarza zaledwie 10 bitów na sekundę

Naukowcy z Kalifornijskiego Instytutu Technologicznego podali w publikacji na łamach czasopisma „Neuron” informację, która brzmi o tyle zaskakująco, o ile dotyczy rzeczy, która wydawała się nam już znana, a okazuje się, że jednak nie. Według nich, nasze mózgi przetwarzają informacje z prędkością zaledwie dziesięciu bitów na sekundę. W porównaniu z miliardami bitów, które nasze zmysły zbierają z otaczającego nas świata w tym samym czasie, wydaje się to bardzo wolne tempo.

Badacze, pod kierownictwem Markusa Meistera i Jieyu Zheng, wykorzystali teorię informacji do analizy czynności mózgu podczas czytania, pisania i rozwiązywania zagadek. Pomimo że jest on wyposażony w ponad 85 miliardów neuronów, ogólny proces myślowy naszego umysłu wydaje się ograniczony przez jego architekturę. Wyniki badań wykazały, że szybkość ludzkiego myślenia jest znacznie wolniejsza niż szybkość działania np. Wi-Fi, które przetwarza około 50 milionów bitów na sekundę.

Zdaniem naukowców rezultaty te podkreślają znany fakt, iż w przeciwieństwie do naszych układów sensorycznych, które działają równolegle, w mózgu możemy przetwarzać tylko jedną myśl naraz. Sugerują, że ta ograniczona szybkość ludzkiego myślenia może mieć ewolucyjne korzenie. Wczesne stworzenia z układem nerwowym wykorzystywały swoje mózgi przede wszystkim do nawigacji, poszukiwania pożywienia i unikania drapieżników. Z biegiem czasu, wraz z ewolucją mózgu, ten mechanizm „jednej ścieżki naraz” prawdopodobnie przeniósł się do bardziej złożonych funkcji poznawczych. Badacze twierdzą, że ludzkie myślenie można porównać do nawigacji w przestrzeni abstrakcyjnych pojęć, do wytyczania ścieżki na mentalnej mapie. ■



ROBOTYKA

Zamiast wózka inwalidzkiego

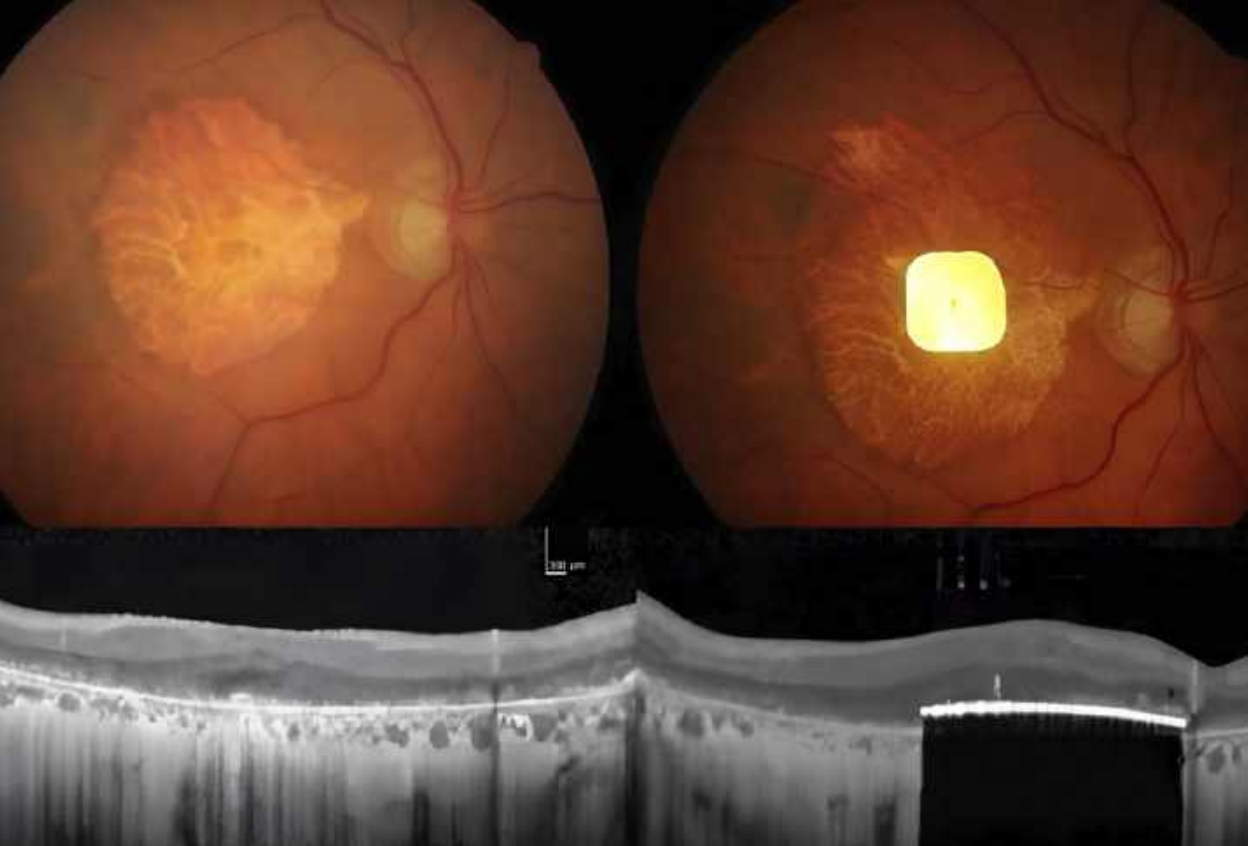
Walk Me – krzesło-robot zaprezentowane podczas targów Japan Mobility Show 2025 może poruszać się w złożonym środowisku dzięki czterem przegubowym nogom. Ma umożliwić użytkownikom o ograniczonej sprawności ruchowej pokonywanie schodów lub innych przeszkód, które byłyby nie do pokonania dla tradycyjnych wózków inwalidzkich. Potrafi również podnosić użytkownika tak, by mógł on dostać się do wnętrza samochodu lub do innych miejsc na podwyższeniu.

Prototyp Toyoty ma cztery składane nogi i siedzisko zaprojektowane tak, by pomagać utrzymywać prawidłową postawę. Całkowicie niezależne nogi są owinięte miękkim, kolorowym materiałem, który pełni podwójną funkcję: chroni wrażliwe elementy wewnętrzne konstrukcji (czujniki i siłowniki) przed uszkodzeniami zewnętrznymi, a jednocześnie nadaje przyjazny wygląd. Każda z nich może się zgiąć, podnosić lub składać. Gdy nogi nie są używane, składają się, dzięki czemu łatwiej umieścić urządzenie w bagażu. System może się również rozkładać i stabilizować bez pomocy użytkownika.

Robot jest wyposażony w funkcje, które pozwalają mu poruszać się po trudnym terenie, m.in. LiDAR, który wykorzystuje światło laserowe do pomiaru odległości i tworzenia bardzo dokładnych, trójwymiarowych obrazów otoczenia, które robot wykorzystuje do omijania przeszkód. Podczas wchodzenia po schodach urządzenie najpierw sprawdza wysokość przednimi nogami, a następnie przesuwa się w górę tylnymi kończynami. Wyposażony jest również w radary kolizyjne, które pozwalają uniknąć kontaktu z ludźmi lub przedmiotami. Zasilane jest akumulatorem ukrytym za siedziskiem, który ma wystarczać na cały dzień pracy. Ładuje się z domowego gniazdka. ■



Prezentacja robota-krzesła Walk Me: <https://youtu.be/VzD9bhUlaK4>



BIONIKA

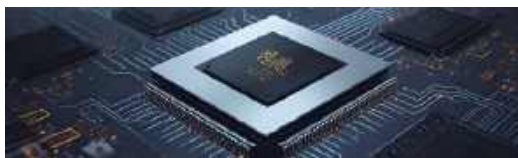
Chip wszczepiony do oka przywraca widzenie

Według publikacji, która ukazała się w periodyku „New England Journal of Medicine”, niewielki chip wszczepiony do oczu osób cierpiących na utratę wzroku spowodowaną nieodwracalnym zwyrodnieniem plamki żółtej, związanym z zaawansowanym wiekiem, przywrócił im widzenie centralne, czyli najbardziej precyzyjną część wzroku, odpowiadającą za szczegółowe rozpoznawanie obiektów, kolorów i czytanie.

System ten nosi nazwę PRIMA, a opracowany został przez międzynarodowy zespół lekarzy i naukowców. Składa się z dwu części. Implant to niewielki bezprzewodowy czujnik krzemowy o wymiarach 2×2 mm i grubości mniejszej niż włos, zawierający 378 pikseli fotowoltaicznych. Umieszcza się go za siatkówką, w miejscu, gdzie zanik komórek jest największy. Drugą częścią systemu PRIMA jest połączona z procesorem para okularów. Okulary rejestrują obrazy i przekształcają je w światło bliskiej podczerwieni o długości fali około 880 nanometrów,

a następnie przesyłają je do implantu. Długość fali ma znaczenie – światło bliskiej podczerwieni jest niewidoczne dla ludzkiego oka, więc nie jest odbierane przez zdrowe fotoreceptory siatkówki ani nie zakłóca pozostałego widzenia peryferyjnego pacjenta. Implant z kolei przekształca sygnały podczerwone na sygnały elektryczne i wysyła je do mózgu, by zostały odebrane, podobnie jak naturalne sygnały z oka. Ponieważ implant jest zasilany światłem, nie potrzebuje zewnętrznego źródła energii.

Testy systemu przeprowadzono w siedemnastu europejskich szpitalach. Rezultat był taki, że przywrócił widzenie centralne u 26 z 32 pacjentów, którzy korzystali z niego przez rok. Wielu z nich mogło nawet ponownie czytać. Średni wiek pacjentów wynosił 79 lat i wszyscy cierpieli na utratę wzroku spowodowaną degeneracją plamki żółtej. Obecnie PRIMA działa tylko w trybie czarno-białym. Naukowcy pracują nad stworzeniem wersji w skali szarości oraz zwiększeniem rozdzielczości systemu. ■



ELEKTRONIKA

Analogowe układy z Chin zamiast topowych chipów Nvidia

Naukowcy z Uniwersytetu Pekńskiego opracowali nowy typ rezystancyjnego układu pamięci RAM, który jest określany jako „analogowy” w odróżnieniu od typowych krzemowych cyfrowych komponentów. Jak piszą w „Nature Electronics”, może on osiągać prędkość nawet tysiąc razy większą niż procesory graficzne Nvidia H100 i AMD Vega 20.

Analogowość nowego chińskiego chipa polega na tym, że obliczenia są w nim wykonywane na własnych obwodach fizycznych, a nie za pomocą binarnych jedynek i zer, jak w standardowych procesorach cyfrowych. Choć układy analogowe zostały dekady temu uznane za mniej praktyczne niż cyfrowe rozwiązania binarne, mają swoje zalety, głównie szybkość działania i wydajność. Ponieważ nie muszą rozbić obliczeń na długie ciągi kodu binarnego, układy analogowe mogą przetwarzać jednocześnie duże ilości informacji, zużywając znacznie mniej energii. Ma to szczególne znaczenie w zastosowaniach wymagających dużej ilości danych i energii, takich jak sztuczna inteligencja, a także w projektowanej na przyszłość sieci 6G, gdzie systemy będą musiały przetwarzać ogromne ilości nakładających się sygnałów bezprzewodowych w czasie rzeczywistym.

Chiński układ zbudowany jest z układów komórek rezystancyjnej pamięci o dostępie swobodnym (RRAM), które przechowują i przetwarzają dane poprzez regulację przepływu prądu elektrycznego przez każdą komórkę. Zastosowany do rozwiązywania złożonych problemów, w tym problemów związanych z odwracaniem macierzy, których rozwiązywanie jest typowym zadaniem w ogromnych systemach wielowęściowych wielowyjściowych (MIMO) i w technice bezprzewodowej, układ scalony dorównał, według opublikowanej pracy, dokładnością standardowym procesorom cyfrowym, zużywając jednocześnie około stu razy mniej energii. ■



AERONAUTYKA

Formuła 1 latających maszyn Jetson ONE

Podczas UP.Summit, imprezy gromadzącej co roku setki światowych innowatorów w dziedzinie transportu, która w tym roku odbyła się w Bentonville w stanie Arkansas, zorganizowano pierwszy w historii wyścig elektrycznych osobowych pojazdów latających pionowego startu i lądowania (eVTOL) firmy Jetson ONE.

Wyścig, w którym brały udział cztery maszyny tego typu, miał na celu zademonstrowanie ich dojrzałości technologicznej. Wśród pilotów biorących udział w wyścigu nazywanym przez firmę Formułą 1 pojazdów latających, był Tomasz Patan, dyrektor techniczny i współzałożyciel firmy, z którym „Młody Technik” rozmawiał latem 2024 roku.

Pojazd Jetson ONE waży ok. 55 kg bez baterii i 115 kg z bateriami. Akumulatory zasilają osiem silników i śmigieł, a przy odpowiedniej wadze pilota na pokładzie (maksymalnie 95 kilogramów) samolot osiąga ograniczoną programowo prędkość maksymalną 100 km/h. Maszyna została zaprojektowana tak, aby cały lot można go było kontrolować jedną ręką za pomocą dżojstika. Firma twierdzi, że dzięki komputerowi pokładowemu i intuicyjnemu systemowi sterowania lotem zasady sterowania można opanować w mniej niż pięć minut. Według reklam, czas lotu wynosi około 20 minut. Dostawy dla klientów (firma ma podobno kilkaset zamówień) nie będą miały miejsca wcześniej niż w 2028 r. ■



Relacja z wyścigu
pojazdów Jetson ONE:
<https://youtu.be/ysE8FMhAPH0>



SAMOLOTY NADDŹWIĘKOWE

X-59 w swoim pierwszym locie

Budowany od lat przez NASA samolot naddźwiękowy X-59, o którego projekcie pisaliśmy w MT wielokrotnie, wzbił się w końcu po raz pierwszy w powietrze. X-59 skonstruowano tak, by przekraczał barierę dźwięku bez generowania charakterystycznego dla lotów naddźwiękowych huku. Wystartował z lotniska Sił Powietrznych Stanów Zjednoczonych (USAF) w Palmdale w Kalifornii do dziewięcioletniego lotu testowego.

NASA nie ogłaszała publicznie tego lotu ani nie wydała oświadczenia po jego zakończeniu. Jednak miłośnicy lotnictwa i fotografowie amatorzy opublikowali w mediach społecznościowych filmy i zdjęcia, na których widać wydłużoną konstrukcję samolotu X-59 startującą z Palmdale i lecącą na północ. Na podstawie śladu pozostawionego przez maszynę

w powietrzu ustalono, że przebywała w powietrzu nieco ponad godzinę, latając po owalnym torze nad wojskową bazą lotniczą Edwards, po czym wylądowała na terenie tego obiektu.

X-59 został zaprojektowany przez NASA i zbudowany przez firmę Lockheed Martin w jej słynnym zakładzie Skunk Works w Palmdale. Założeniem było opracowanie konstrukcji, która podczas przekraczania bariery dźwięku nie wywołuje bardzo głośnego huku sonicznego. Z powodu takich huków od 1973 r. loty naddźwiękowe nad lądem w Stanach Zjednoczonych są zabronione. Jeśli X-59 udowodni, że „cichy” lot naddźwiękowy jest możliwy, ograniczenia dotyczące przekraczania bariery dźwięku mogą pewnego dnia zostać zniesione, otwierając drogę do nowej ery komercyjnych lotów naddźwiękowych. ■



TECHNIKA MEDYCZNA

♦ Wykorzystujący skupione impulsy ultradźwiękowe do mechanicznego niszczenia tkanki nowotworowej bez konieczności wykonywania operacji, radioterapii lub chemioterapii, system o nazwie Edison Histotripsy, opracowany przez amerykańską firmę HistoSonics, zajmującą się produkcją urządzeń medycznych, został zastosowany w testach medycznych na stu pacjentach w szpitalu w Hongkongu i, według oficjalnych komunikatów, wyniki są bardzo dobre, rokując nadzieje na skuteczną walkę z rakiem. ♦ Firma Snke OS GmbH ogłosiła wprowadzenie do sprzedaży okularów SnkeXR, pierwszych okularów rozszerzonej rzeczywistości klasy medycznej, zaprojektowanych do pracy w różnych obszarach praktyki klinicznej, w tym m.in. do ortopedii, neurochirurgii, elektrofizjologii, radiologii interwencyjnej i ginekologii. ♦

ROBOTYKA

♦ Według publikacji w „Advanced Functional Materials”, południowokoreańscy naukowcy opracowali innowacyjną, polimerową strukturę chemiczną sztucznego mięśnia, który może unieść ciężar nawet cztery tysiące razy przekraczający wartość jego masy własnej – według twórców może on znaleźć zastosowanie w przyszłych robotach humanoidalnych. ♦ Z wewnętrznych dokumentów strategicznych, do których dostarli dziennikarze gazety „The New York Times”, wynika, że kierownictwo firmy Amazon planuje zastąpienie robotami ponad pół miliona stanowisk pracy. ♦

NAPĘDY

♦ Jau Tang z uniwersytetu w chińskim Wuhan twierdzi, że udało im

się skonstruować silnik, który wykorzystuje technikę mikrofalową i plazmową, bez potrzeby stosowania jakiegokolwiek paliwa i baterii zasilającej, zaś zbudowany przezeń prototyp jest w stanie za pomocą siły ciągu poruszać stalową kulą o masie ponad kilograma. ♦ GE Aerospace testuje nowy silnik strumieniowy Solid-Fuel Ramjet (SFRJ), w którym eliminuje się paliwo płynne, zwykle w takich konstrukcjach wtryskiwane do komory spalania, i zamiast tego pokrywa się wnętrze silnika paliwem węglowodorowym o konsystencji stałej, podobnej nieco do gumy. ♦



TECHNIKA WOJSKOWA

♦ Firma Shield AI zaprezentowała koncept o nazwie X-BAT, autonomiczny myśliwiec pionowego startu i lądowania (VTOL) o zasięgu dwóch tysięcy mil morskich i pułapie do ok. 15 km, który, według niej, może startować z odległych wysp i statków, bez potrzeby wykorzystania lotnisk i pasów startowych. ♦ Jin Zhang z Uniwersytetu Pekńskiego i jego zespół opracowali nową metodę wkomponowania nanorurek węglowych w łańcuchy polimeru aramidowego, która to metoda pozwoliła na stworzenie kompozytu o grubości 1,8 milimetra wytrzymalszego niż kevlar i, jak podają źródła, być może najwytrzymalszego materiału ze znanych nam na świecie, co pozwoliłoby stworzyć kamizelki kuloodporne nowej generacji. ♦ System broni laserowej australijskiej firmy zbrojeniowej Electro Optic Systems (EOS) znany jako Apollo High Energy Laser Weapon (HELW), o mocy do 150 kW, podczas prezentacji udowodnił, że potrafi zniszczyć dwieście średniej wielkości dronów w jednym cyklu uruchomienia, korzystając wyłącznie z własnego wewnętrznego źródła zasilania, czyli nie wymagając zasilania (kabla) z sieci energetycznej. ■

M.U.



1. BMW i7 wyposażone w półprzewodnikowe ogniwa akumulatorowe EV firmy Solid Power © BMW Group

Czyżby wreszcie przełom w bateriach?

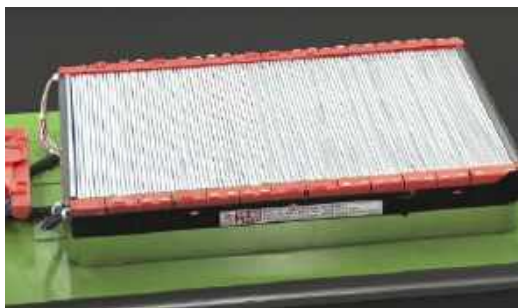
Ogniwa, o jakie chodziło

Kto obserwuje historię postępów w dziedzinie techniki akumulatorowej, tak jak my w MT robimy od lat, ten wie, że jest pełna niespełnionych (lub spełnionych w taki sposób, że rozczarowują) obietnic i prognoz. Niedawno na horyzoncie pojawiły się rozwiązania akumulatorów z elektrolitem stałym i półprzewodnikowych, dające nadzieję, że to może wreszcie jest nadzieja na prawdziwy przełom.

Pierwsze pojazdy elektryczne BMW zasilane akumulatorami półprzewodnikowymi z elektrolitem całkowicie stałym są obecnie testowane na drogach. BMW korzystało ze swojego modelu i7 (1) do przetestowania od lat obiecywanego świętego Graala technologii akumulatorów do pojazdów elektrycznych, mającego oznaczać większy zasięg przy niższych kosztach. Sprawdzane w praktyce są ogniwa akumulatorowe firmy Solid Power, określane jako ASSB (skrót z ang. „all-solid-state batteries”). BMW podąża śladami firmy Mercedes-Benz, która w lutym 2025 r. ogłosiła, że dzięki współpracy z amerykańską firmą Factorial Energy wprowadziła na rynek „pierwszy samochód napędzany litowo-metalowym akumulatorem półprzewodnikowym”. Mercedes wykorzystał zmodyfikowany model EQS, wyposażony w akumulatory półprzewodnikowe. Dzięki oczekiwanej redukcji masy o 40 proc. w porównaniu z obecnymi

akumulatorami litowo-jonowymi, Factorial zamierza osiągnąć zasięg ponad 1000 km.

Niemieccy producenci samochodów nie są jedynymi, którzy zaangażowali się w testy nowej techniki akumulatorów. Oczekuje się, że w ciągu najbliższych kilku lat światowi liderzy w dziedzinie akumulatorów, chińskie firmy CATL i BYD, wprowadzą na rynek pojazdy elektryczne z akumulatorami półprzewodnikowymi. Sun Huajun z BYD powiedział na początku 2025 r., że firma spodziewa się wprowadzenia pierwszych pojazdów elektrycznych z akumulatorami półprzewodnikowymi w 2027 r., a do 2030 r. akumulatory ASSB firmy BYD mają wejść na rynek masowy. W początkowej fazie BYD będzie stosować rozwiązanie oparte na siarczkuach w niektórych swoich modelach z wyższej półki. Kilka innych firm motoryzacyjnych, w tym Hyundai, Nissan, Stellantis, Toyota (2) i Honda, przyspiesza prace nad bateriami nowej generacji. Toyota wprowadzić ma

**2. Prototyp akumulatora ASRB firmy Toyota**

na rynek pierwszy na świecie samochód elektryczny z akumulatorem półprzewodnikowym do 2027 r.

ASRB różni się od litowo-jonowych konsystencją elektrolitów. Tradycyjne Li-Ion mają elektrolit w postaci płynnej lub żelowej, zaś ASRB wykorzystują elektrolity stałe. Obecnie badany jest cały wachlarz elektrolitów stałych, w tym ceramiczne, polimerowe, oparte na żywicach i kompozytach szklanych. Jak wiadomo, tradycyjne akumulatory litowo-jonowe budzą obawy dotyczące bezpieczeństwa ze względu na ich łatwopalność. To ryzyko w ASRB jest mniejsze. Ponadto ograniczone jest w nich niebezpieczeństwo wycieku, przegrzania i tworzenia się dendrytów w elektrolicie. Ogniwa z elektrolitem stałym oferują również wyższą gęstość energii w porównaniu z tradycyjnymi litowo-jonowymi. ASRB mogą również potencjalnie znacznie przyspieszyć ładowanie.

Chińczycy już postawili na krzem w bateriach, ale...

Przy okazji warto zwrócić uwagę, że niektóre nowe rozwiązania, którymi kuszą producenci nie są jeszcze w pełni dopracowane. W ciągu ostatnich dwóch lat wiele marek smartfonów zaczęło stosować baterie krzemowo-węglowe, dzięki czemu niektóre flagowe modele telefonów z systemem Android mają teraz bardzo pojemne baterie. Dotyczy to głównie producentów chińskich. Dostępny wyłącznie w Chinach Xiaomi 15 Pro ma baterię o pojemności 6100 mAh, vivo X200 Pro ma baterię o pojemności 6000 mAh, a OPPO Find X8 Pro ma baterię o pojemności 5910 mAh. Przecieki z platformy X zapowiadały modele OPPO Find X9 Pro, wyposażone w baterię o pojemności sięgającej nawet 7550 mAh.

Jednak wielu ekspertów jest sceptycznych, zwracając uwagę na wady tych reklamowanych jako „przełom technologiczny” akumulatorów. Być może dlatego Apple, Google lub Samsung, producenci spoza Chin, się wstrzymują, a baterie tego typu ujrzymy w ich modelach dopiero w 2026 roku albo później.

Stosowane w chińskich smartfonach baterie krzemowo-węglowe (Si/C) to ewolucja tradycyjnej technologii litowo-jonowej, a nie całkowicie nowa koncepcja. Modyfikują konwencjonalną anodę grafitową, wzbogacając ją o krzem, który może magazynować znacznie więcej energii – teoretycznie nawet dziesięć razy więcej niż grafit (372 mAh/g w porównaniu z ~4200 mAh/g dla czystego krzemu). Jednak anody z krzemu oznaczają poważne wyzwania. Najbardziej problematyczne jest ich silne „puchnięcie”, w wyniku której struktura pęcznieje nawet o 300 proc. po pełnym naładowaniu. Powoduje to poważne obciążenie mechaniczne baterii, skracając jej żywotność i powodując uszkodzenia strukturalne. Ponadto krzem reaguje silnie z elektrolitem, rozbijając i ponownie tworząc warstwę międzypfazową elektrolitu stałego (SEI) przy każdym cyklu, co prowadzi do stopniowej utraty litu i zmniejszenia pojemności. Krzem ma również niższą przewodność elektryczną niż grafit, co może spowolnić proces ładowania i rozładowania oraz zwiększyć straty spowodowane oporem wewnętrznym, a to może skutkować większym wydzielaniem ciepła, co jest kolejnym czynnikiem negatywnie wpływającym na żywotność baterii.

Aby rozwiązać te problemy, zamiast czystego krzemu stosuje się kompozyt krzemowo-węglowy (Si/C). Węgiel zapewnia wsparcie strukturalne, pomagając złagodzić wzrost objętości i ustabilizować warstwę SEI. Węgiel poprawia również przewodność elektryczną, zapewniając lepszy przepływ jonów litowych i zwiększając wydajność. Kompromis polega na tym, że anody Si/C nie osiągają pełnego dziesięciokrotnego wzrostu pojemności czystego krzemu, oferując bardziej umiarkowany wzrost gęstości energii o 10–20 proc. Nie ma nic za darmo – baterie Si/C poprawiają wydajność, ale mają ograniczenia i wady. Biorąc pod uwagę powyższe kompromisy, jasne jest, że Si/C nie będą działać tak długo, jak najtrwalsze obecnie dostępne na rynku ogniwa litowo-jonowe na bazie grafitu.

Poprawianie elektrolitu i kwanty

W ośrodkach badawczych wre praca nad nowymi ogniwami i doskonaleniem tych znanych od lat. Elektrolit stały to niejedyny nurt poszukiwań. Niektórzy szukają sposobu na ograniczenie niekorzystnych efektów w starszych typach ogniw. Naukowcy z Korei zastosowali niedawno np. kreatywną chemię, otwiera to drogę do akumulatorów, które zapewniłyby zasięg 800 kilometrów po 12-minutowym ładowaniu, przez walkę z dendrytami w elektrolicie. Ich baterie litowo-metalowe różnią się od standardowych

baterii litowo-jonowych tym, że anoda grafitowa została zastąpiona litem metalicznym. Prawdziwy jednak sekret tkwi w nowym rodzaju ciekłego elektrolitu, „hamującego kohezję”, czyli wzrost dendrytów (kryształów fraktalnych, które w elektrolicie są niekorzystne, powodując straty pojemności i zwarcia). Zespół badawczy odkrył, że podstawową przyczyną powstawania dendrytów była „nierównomierna spójność międzyfazowa na powierzchni metalu litowego”, czyli w praktyce jony litowe nie osadzają się równomiernie na anodzie podczas ładowania, tworząc słabe punkty, w których mogą zacząć tworzyć się dendryty. Opracowali więc płynny elektrolit o strukturze chemicznej, która pomaga zapewnić bardziej równomierne osadzanie się jonów na powierzchni anody, zapobiegając ich gromadzeniu się w dendryty. W testach laboratoryjnych bateria ładowała się od 5 do 70 proc. w ciągu 12 minut i utrzymywała tę prędkość przez 350 cykli. Opis badań został opublikowany w czasopiśmie „Nature Energy”.

Bada się też nieustannie alternatywy dla litu, na przykład sód (3). Naukowcy odkryli niedawno, jak ustabilizować wysokowydajny związek sodu, by nadać bateriom półprzewodnikowym na bazie sodu moc i stabilność, których im brakuje. Nowy materiał, opisany w czasopiśmie „Joule”, znacznie skuteczniej przewodzi jony i pozwala na stosowanie grubszych katod o dużej gęstości energii. Sód stanowi znacznie tańszą i łatwiej dostępną alternatywę, a jego wydobycie jest znacznie mniej szkodliwe dla środowiska niż górnictwo litu. Jak wyjaśnia autor, Sam Oh z A*STAR Institute of Materials Research and Engineering w Singapurze, „przełom polega na tym, że faktycznie stabilizujemy strukturę metastabilną hydrydoboranu sodu, która ma bardzo wysoką przewodność jonową”. Aby stworzyć tę strukturę, naukowcy podgrzali metastabilną formę hydrydoboranu sodu do punktu, w którym zaczęła krystalizować, a następnie szybko ją przechłodziли, aby utrwalić strukturę. Połączenie fazy metastabilnej z katodą typu O_3 pokrytą elektrolitem stałym na bazie chlorku pozwala uzyskać grube katody o dużej powierzchni.

Baterie kwantowe to alternatywa brzmiąca nieco bardziej futurystycznie, ale faktem jest, że prace nad różnego rodzaju rozwiązaniami i prototypami takich urządzeń także są prowadzone. W 2022 r. badacze z australijskiego Royal Melbourne Institute of Technology w Australii opracowali prototyp baterii kwantowej, która może utrzymać energię tysiąc razy dłużej niż poprzednie urządzenia. Francesco Campaioli wraz ze swoimi współpracownikami podjął próbę wydłużenia żywotności



3. Wizualizacja baterii opartych na sodzie

baterii kwantowych przez eksperymenty ze stanami trypletowymi, czyli stanami elektronów w cząsteczce po absorpcji światła. Jedna warstwa, zawierająca barwnik rodaminę 6G, bardzo dobrze pochłania światło, przekazując energię do warstwy wykonanej ze związku zwanego tetrafenyloporfiryń palladu, która magazynuje ją w postaci ciemnych stanów trypletowych. Nowy prototyp magazynował energię przez mikrosekundy zamiast nanosekund. Dzięki takim efektom jak splątanie kwantowe i superabsorpcja, gdzie szybkość absorpcji światła wzrasta wraz z liczbą cząsteczek, baterie kwantowe mogą ładować się znacznie szybciej niż nawet najlepsze baterie klasyczne. „To tak, jakby mieć telefon, który ładuje się w 30 minut i wyczerpuje się po około 20 dniach, jeśli nie jest używany. Nie jest to takie złe”, wyjaśnia Campaioli w czasopiśmie „PRX Energy”. W innym artykule, opublikowanym w czasopiśmie „Physical Review Letters”, uczeni z Uniwersytetu Badawczego PSL w Paryżu i uniwersytetu w Pizie opisują model baterii kwantowej w skali mikroskopowej. „Nasz model składa się z dwóch sprzężonych oscylatorów harmonicznych: jeden działa jako ‘ładowarka’, a drugi jako ‘bateria’”, mówią autorzy serwisowi „Phys.org”. Wcześniej grupa naukowców opracowała model baterii kwantowej wielkości atomu, która wykorzystywała pośredniczące wnęki w celu uniknięcia „dekoherencji”, czyli procesu, w wyniku którego system traci swoje właściwości kwantowe. Jednak oba ostatnie przykłady są jedynie modelami – skonstruowanie urządzenia w praktyce to zupełnie inna sprawa. ■

Mirosław Usidus



1. Wizja Starcloud na orbicie okołoziemskiej

Projekt Starcloud

Serwerownia na orbicie

„Za 10 lat prawie wszystkie nowe centra danych będą budowane w przestrzeni kosmicznej”, twierdzi firma Starcloud. Jej nazwę nosi też mało znany dotychczas projekt stworzenia wielkiego centrum przetwarzania danych na orbicie okołoziemskiej, które zasilane ma być przez gigantyczną kosmiczną farmę słoneczną rozciągającą się na kilometry (1).

Pierwszym krokiem do realizacji projektu jest wysłanie w kosmos procesora graficznego NVIDIA H100, sto razy wydajniejszego niż jakikolwiek procesor, który dotychczas poleciał na orbitę, na pokładzie satelity Starcloud-1, rakietą SpaceX Falcon 9. Starcloud z siedzibą w Redmond, w stanie Waszyngton, wykorzysta tę pierwszą misję do przetestowania, jak przetwarzanie danych działa w kosmosie.

Zdaniem zwolenników projektu przeniesienie przetwarzania danych w kosmos zmniejszyłoby obciążenie środowiska na Ziemi. Centra danych zużywają ogromne ilości energii elektrycznej i wody. „W kosmosie można uzyskać niemal nieograniczoną, tanią energię odnawialną”, powiedział w komunikacie dla mediów Philip Johnston, współzałożyciel i dyrektor generalny Starcloud. „Jedynym kosztem naszego projektu dla środowiska będzie start. Następnie w całym okresie eksploatacji centrum danych

nastąpi dziesięciokrotna redukcja emisji dwutlenku węgla w porównaniu z zasilaniem centrum danych na Ziemi”. Przeniesienie centrów danych w kosmos będzie jednak wymagało znacznego obniżenia kosztów startu. Firma Starcloud oczekuje, że kalkulacja kosztów będzie odpowiednia, gdy wielka rakietą Starship firmy SpaceX osiągnie pełną operacyjność.

Firma Starcloud została założona w 2024 roku przez Philipa Johnstona, Ezrę Feildena i Adiego Olteana. Zebrała 21 milionów dolarów między grudniem 2024 r. a lutym 2025 r. Została wsparta m.in. przez NFX, In-Q-Tel i program NVIDIA Inception, 468 Capital.

Jednym z oczywistych pomysłów założycieli była energia słoneczna pozyskiwana w kosmosie. Jednak zamiast skupiać się na przesyłaniu energii w celu zasilania naziemnych centrów danych, ludzie ze Starcloud zadali sobie inne pytanie – co by

było, gdybyśmy przenieśli centrum danych w kosmos? Przeprowadzili obliczenia kosztów uruchomienia i sporządzili dokument, który stał się załącznikiem projektu Starcloud.

Jest sporo założeń technicznym, które muszą być spełnione, by projekt zadziałał. Po pierwsze, sprzęt musi przetrwać start, co wymaga, aby cały system był odporny na znaczne wibracje i obciążenia mechaniczne. Po drugie, na orbicie chipy muszą funkcjonować w środowisku o wysokim poziomie promieniowania, co stanowi wyzwanie dla niezawodności systemu. Po trzecie, zarządzanie temperaturą staje się bardziej złożone, ponieważ brak konwekcji w powietrzu czy cieczach, i wymaga nowych podejść do rozpraszania ciepła, czyli chłodzenia w przestrzeni kosmicznej.

Jeśli chodzi o układ odprowadzania ciepła, to wybrano pomysł, by wykorzystywać chłodzenie radiacyjne, czyli proces rozpraszania ciepła przez emisję promieniowania podczerwonego. Starcloud jeszcze na Ziemi ma zanurzać tranzystory, rezystory i inne elementy tworzące płytę w cieczy przewodzącej ciepło, takiej jak olej lub wosk, tak aby każdy element był równomiernie chłodzony. W przestrzeni kosmicznej nadmiar energii odprowadzamy ma być w postaci promieniowania podczerwonego za pomocą dużych, lekkich, rozkładanych radiatorów. Opracowanie rozkładanych radiatorów, które są wystarczająco duże, aby poradzić sobie z tymi obciążeniami, a jednocześnie wystarczająco lekkie, aby można je było wystrzelić w opłacalny sposób, jest złożonym problemem inżynierskim, nad którym firma wciąż pracuje.

Przetwarzanie orbitalne

Wysokiej rozdzielczości obrazy optyczne i radarowe dużo „wagę”, co wymaga od satelitów przesyłania na Ziemię ogromnych zbiorów danych. Stacje naziemne nie zawsze są dostępne, a przepustowość jest ograniczona. Przetwarzanie danych na orbicie wyeliminowałoby część utrudnień, ponieważ najlepsze obrazy mogłyby być identyfikowane bezpośrednio na orbicie i wysyłane na Ziemię podczas przelotu stacji naziemnej. Satelita Starcloud-1 będzie również obsługiwał otwarty model językowy kategorii Gemma firmy Google, co stanowi kolejną ważną nowość.

Jeśli wszystko pójdzie zgodnie z planem, firma wprowadzi w najbliższych latach na rynek bardziej wydajne układy, testując jeszcze mocniejsze procesory graficzne NVIDIA, w tym platformę Blackwell, która zapewnić ma nawet dziesięciokrotny wzrost wydajności. Pierwsze komercyjne satelity, Starcloud-2, mają zostać wystrzelone w 2026 roku. Wyposażone mają być one w klaster GPU, pamięć trwała,

ciągły dostęp oraz systemy termiczne i zasilające, a wszystko to w obudowie SmallSat. Satelity te będą działać na orbicie heliosynchronicznej i będą służyć zarówno użytkownikom kosmicznym, jak i naziemnym. Trwała pamięć masowa ma służyć do przechowywania danych na orbicie, które byłyby dostępne nawet wtedy, gdy satelita nie transmituje aktywnie, podobnie jak w przypadku przechowywania i pobierania danych w naziemnych centrach danych. Ciągły dostęp ma być możliwy dzięki łączom o dużej przepustowości do konstelacji satelitów, takich jak Starlink, a nie tylko w krótkich oknach komunikacyjnych, gdy satelita znajduje się bezpośrednio nad punktem obioru. Satelity Starcloud będą wyposażone w terminale laserowe, który zapewniają przepustowość, sięgającą setek gigabitów na sekundę, oraz opóźnienie wynoszące powyżej 50 milisekund.

Klientom naziemnym Starcloud-2 zaoferować chce w przyszłości bezpieczne globalne przechowywanie danych i suwerenne przetwarzanie w chmurze, które będzie działać całkowicie niezależnie od infrastruktury naziemnej. W najbliższej przyszłości skupi się na komunikacji i obróbce danych z innych satelitów. Są to zazwyczaj zadania związane z wnioskowaniem na podstawie obrazów, danych z czujników, hiperspektralnych lub radarowych.

Obecnie Starcloud zakłada, że będzie w stanie umieścić na orbicie centrum danych o mocy 40 megawatów. Powyżej tej skali platforma ta będzie musiała ewoluować w kierunku architektury modułowej (2). Modułowość w centrum danych odnosi się do budowy z oddzielnych jednostek, takich jak obliczenia, pamięć masowa lub chłodzenie, które można dodawać lub wymieniać niezależnie. Dzięki temu skalowanie i dostosowywanie byłoby łatwiejsze niż w przypadku pojedynczego monolitycznego systemu.

Gdyby przetwarzanie kosmiczne stało się znaczącą częścią globalnego rynku, na przykład 25 proc.

2. Wizualizacja montażu modułów Starcloud na orbicie





całkowitego popytu, firmy takie jak NVIDIA mogłyby mieć motywację do projektowania układów scalonych zoptymalizowanych specjalnie pod kątem środowiska orbitalnego.

Koncept „serwerowni” lub centrum obliczeniowego w kosmosie wzbudził pewien ferment wśród innych firm technologicznych. Google zapowiedziało na początku listopada 2025 r. swój projekt o nazwie Project Suncatcher. Firma przewiduje, że jego jednostki przetwarzające Tensor (TPU) będą krążyć wokół Ziemi na satelitach wyposażonych w panele słoneczne generujących energię elektryczną, czyli dość podobnie do projektu Starcloud. Konkurencja z centrami danych na lądzie „wymaga połączeń między satelitami, które obsługują dziesiątki terabitów na sekundę”, pisze Google. Pomóc to osiągnąć ma manewrowanie konstelacjami satelitów w ciasnych formacjach, w odległości

„kilometrów lub mniej” od siebie. To znacznie bliżej niż obecnie działają satelity, a już teraz rosnące ryzyko stanowią kosmiczne śmieci powstałe w wyniku kolizji. Google przetestowało procesory Trillium TPU pod kątem odporności na promieniowanie i twierdzi, że „wytrzymują one całkowitą dawkę promieniowania jonizującego równoważną 5-letniej misji bez trwałych awarii”. Analiza kosztów przeprowadzona przez firmę sugeruje, że uruchomienie i prowadzenie centrum danych w przestrzeni kosmicznej może stać się „mniej więcej porównywalne” z kosztami energii równoważnego centrum danych na Ziemi w przeliczeniu na kilowat/rok do połowy lat 30. XXI wieku. Google podaje, że planuje wspólną próbną misję z firmą Planet, chcąc do 2027 r. wystrzelić kilka prototypowych satelitów w celu przetestowania swojego sprzętu na orbicie. ■

Mirosław Usidus

Pokłosie wielkiej awarii AWS

Chmury nad chmurami

A może za bardzo ufamy chmurom – zaczęli się zastanawiać ludzie z branży po serii awarii, które wstrząsnęły siecią. Kulminacją katastrof była październikowa zapaść usług serwerowych AWS Amazona. Niektórzy postanowili wyciągnąć wnioski.

Awaria największej na świecie platformy przetwarzania w chmurze, Amazon Web Services (AWS), dotknęła tysiące organizacji, w tym banki, platformy oprogramowania finansowego oraz platformy mediów społecznościowych, np. Snapchata. Była spowodowana błędną konfiguracją jednego z centrów danych AWS zlokalizowanego w północnej Wirginii w Stanach Zjednoczonych. Amazon twierdzi, że naprawiono problem, ale niektórzy użytkownicy Internetu mimo to zgłaszali zakłócenia w działaniu usług. Nie brakuje opinii, że takie, a nawet większe, awarie będą się powtarzać. Zdaniem wielu ekspertów incydent ten to dzwonek ostrzegawczy dla wszystkich, którzy obdarzyli powszechne dziś w technice sieciowej i komputerowej chmury obliczeniowe, nadmiernym zaufaniem.

Przetwarzanie w chmurze to dostarczanie na żądanie różnych zasobów, takich jak moc obliczeniowa, pamięć baz danych i aplikacje przez Internet. Mówiąc prościej, jest to wynajem (a nie posiadanie – przynajmniej w najczęstszym przypadku) własnej

infrastruktury IT. Przetwarzanie w chmurze pojawiło się i stopniowo zyskiwało popularność wraz z rozwojem sieci, gdy firmy zajmujące się technologią cyfrową zaczęły dostarczać oprogramowanie przez Internet. W miarę jak firmy takie jak Amazon doskonaliły swoje umiejętności w dziedzinie nazywanej z czasem „oprogramowaniem jako usługa” sieciowa, zaczęły one oferować innym możliwość wynajmowania swoich wirtualnych serwerów za opłatą. Jednym z ważnych aspektów wpływających na wzrost popularności przetwarzania w chmurze był model płatności zgodnie z rzeczywistym zużyciem, podobny do rachunku za energię czy wodę w domu, zamiast ogromnych inwestycji początkowych wymaganych do zakupu, obsługi i zarządzania własnym centrum danych. W rezultacie według statystyk ponad 94 proc. wszystkich przedsiębiorstw w krajach rozwiniętych korzysta z usług opartych na chmurze w takiej czy innej formie.

Globalny rynek usług w chmurze jest zdominowany przez trzy firmy. Największy udział ma AWS



1. Awaria usług AWS Amazona © AI

W planowanej nowej elektrowni Amazon będzie miał reaktor modułowy (SMR) Xe-100 firmy X-Energy, który zostanie zainstalowany w pobliżu istniejącej elektrowni Columbia Generating Station firmy Energy Northwest w stanie Waszyngton. Xe-100 to reaktor wysokotemperaturowy chłodzony gazem (HTGR), zasilany paliwem TRISO-X, składającym się z ziaren uranu otoczonych warstwami węgla i ceramiki, połączonych w kulki. Są one podawane przez lejek do zbiornika reaktora w celu wywołania samoregulującej się reakcji rozszczepienia i chłodzone helem, który odprowadza ciepło do wymiennika ciepła w celu wytworzenia pary dla generatorów elektrycznych. Zużyte paliwo gromadzi się na dnie zbiornika. SMR-y Xe-100 wytwarzają tylko 80 MWe mocy.

Według komunikatu firmy Amazon, budowa obiektu Cascade Advanced Energy Facility rozpocznie się pod koniec dekady, a obiekt zostanie oddany do użytku w latach 30. XXI wieku. Obiekt będzie składał się z czterech reaktorów o łącznej mocy 320 MWe i będzie miał potencjał rozbudowy do 12 reaktorów o mocy 960 MWe. Jeśli inicjatywa zakończy się sukcesem, Amazon planuje wprowadzić do 2039 r. w Stanach Zjednoczonych łącznie 5 GW energii jądrowej. ■

Mirosław Usidus

(ok. 30 proc.). Kolejne miejsca zajmują Microsoft Azure (20 proc.) i Google Cloud Platform (13 proc.). Wszyscy trzech dostawcy usług doświadczyli w ostatnim czasie dużych awarii, które miały znaczący wpływ na działanie platform usług cyfrowych. Na przykład w 2024 r. problem z oprogramowaniem strony trzeciej poważnie wpłynął na działanie Microsoft Azure, powodując rozległe awarie operacyjne w przedsiębiorstwach na całym świecie. Google Cloud Platform doświadczył w 2025 roku poważnej awarii spowodowanej błędną konfiguracją usługi. A jesienią uderzyła potężna awaria AWS Amazona.

Zdaniem specjalistów, tak duża zależność globalnego Internetu od zaledwie kilku wielkich dostawców stwarza poważne zagrożenia zarówno dla przedsiębiorstw, jak i zwykłych użytkowników. Jak widać na przykładzie ostatniego zdarzenia w AWS, prosty błąd konfiguracyjny w jednym centralnym systemie może wywołać efekt domina, który natychmiast paraliżuje ogromne segmenty Internetu. Po drugie, dostawcy ci często narzucają uzależnienie od jednego dostawcy. Zmiana platformy jest niezwykle trudna i kosztowna ze względu na złożoną architekturę danych i nadmiernie wysokie opłaty za przenoszenie dużych ilości danych z chmury. Skutkuje to skutecznym uwięzieniem klientów, którzy stają się zakładnikami warunków jednego dostawcy. Wreszcie dominacja amerykańskich dostawców usług w chmurze wiąże się z innymi problemami. Dane przechowywane w tych ogromnych systemach podlegają prawu amerykańskiemu i wymogom rządu. Firmy te mają prawo cenzurować lub ograniczać dostęp do usług.

Czy da się to zmienić? Obecnie najlepszym sposobem na ograniczenie tego ryzyka jest przyjęcie podejścia wielochmurowego, które umożliwia decentralizację. Oznacza to uruchamianie krytycznych aplikacji u wielu dostawców w celu wyeliminowania punktowego ryzyka. Podejście to można uzupełnić o „przetwarzanie brzegowe” (tzw. egde computing), w ramach którego przechowywanie i przetwarzanie danych jest przenoszone z dużych, centralnych centrów danych do mniejszych, rozproszonych węzłów (serwery lokalne), które firmy mogą kontrolować bezpośrednio.

Na problemy – SMR-y?

Mniej więcej w tym samym czasie, gdy wydarzyła się wielka awaria AWS, pojawiła się wiadomość, że Amazon buduje pierwszą z serii modułowych elektrowni jądrowych m.in. po to, by chronić swoje usługi przed awariami.

Łatwo zbyć niezrozumiałe dla uczonych pomysły i twory AI jako pozbawiony sensu bełkot. Jednak podejście takie niekoniecznie ma coś wspólnego z nauką. Poważny naukowiec mimo wszystko spróbuje zrozumieć to, co proponuje sztuczna inteligencja. W niektórych przypadkach spokoju nie daje wrażenie, że to, czego ludzie nie pojmują, może być czymś nowym i świeżym.

Czy pojmiemy naukę, gdy zajmie się nią sztuczna inteligencja?

TURECKIE KAZANIE PROFESORA AI

Urania, model AI opracowany przez zespół kierowany przez Mario Krenna z Laboratorium Sztucznej Inteligencji w Instytucie Maxa Plancka, zaproponowała np. nowe projekty detektorów fal grawitacyjnych, które, według naukowców, „mają potencjał”, przewyższający pod wieloma względami możliwości projektów stworzonych przez człowieka. Zespół Krenna dążył do stworzenia, nowatorskich projektów detektorów wykorzystujących interferometrię, która mierzy bardzo subtelne zmiany w interakcjach własnych fal. Wykorzystując uczenie maszynowe, Urania pomagała zespołowi w opracowaniu szeregu nowatorskich projektów eksperymentalnych detektorów fal grawitacyjnych. Według publikacji w „Physical Review X” w kwietniu 2025 r., Urania nie ograniczyła się jedynie do powielenia istniejących koncepcji. Według komunikatu laboratorium Krenna, które współpracuje z zespołem LIGO, zajmującym się detekcją fal grawitacyjnych od lat, części nowych projektów autorstwa sztucznej inteligencji „naukowcy nie rozumieją jeszcze w pełni”. Niektóre z nich wydają się przewyższać nawet najlepsze istniejące koncepcje czysto ludzkie, potencjalnie znacznie poprawiając ich czułość i poszerzając ogólny zakres wykrywalności docelowych sygnałów fal grawitacyjnych (1). Inne były bardzo niekonwencjonalne i trudne do pojęcia dla badaczy. Krenn i jego zespół podają, że zebrali pięćdziesiąt najlepiej działających projektów i udostępnili je publicznie w ramach „Detector Zoo”, zapraszając innych naukowców do dalszych badań.



1. Wizualizacja fal grawitacyjnych i ich detekcji

Gdyby jeszcze było wiadomo, jak AI dochodzi do swoich wyników

Kwestia niezrozumiałości wyników czy produktów naukowej AI dołącza do znanego od lat problemu wyjaśnialności, czyli braku wglądu w to, jak dokładnie modele sztucznej inteligencji rozumują i jak dochodzą do rezultatów, które widzimy „na wyjściu”. To zagadnienie starsze niż nowa fala generatywnej sztucznej inteligencji, która zaczęła się od premiery ChatGPT, o którym pisaliśmy w MT już lata temu. Powodem do zaniepokojenia od dawna jest fakt, iż nawet specjaliści i firmy, które tworzą modele i algorytmy AI, nie wiedzą dokładnie, dlaczego i jak modele działają. Nazywane jest to czasem problemem „czarnej skrzynki”, choć to mylące, bo czarna skrzynka w samolotach służy właśnie do tego, by dowiedzieć się dokładnie, co się wydarzyło. Niedawno tworząca modele z rodziny Claude firma Anthropic uruchomiła nawet program badawczy, którego celem jest zrozumienie, co tam tak naprawdę w środku modeli zachodzi, gdy

rozumują, dają odpowiedzi i wnioskuje. Jeśli AI miałyby zająć się nauką, to oczywiście potęguje problem.

Już w 2015 r. grupa badawcza ze szpitala Mount Sinai w Nowym Jorku została poproszona o zastosowanie techniki deep learning do analizy rozległej bazy danych tamtejszych pacjentów. Powstały w toku prac program analityczny naukowcy nazwali Deep Patient. Został on przeszkolony z użyciem danych ok. 700 tys. osób, a po przetestowaniu na nowych rejestrach okazał się niezwykle skuteczny w prognozowaniu chorób. Bez asysty ludzkich specjalistów Deep Patient odkrył w danych szpitalnych wzorce wskazujące, który pacjent znajduje się na ścieżce do choroby, np. raka wątroby. W ocenie specjalistów skuteczność prognostyczno-diagnostyczna systemu była znacznie wyższa niż jakiegokolwiek inne znane metody. Jednocześnie badacze odnotowali, że Deep Patient działa w sposób zagadkowy. Okazało się np., że doskonale sprawdza się w rozpoznawaniu zaburzeń psychicznych, takich jak schizofrenia, co jest niezwykle trudne dla lekarzy. Wzbudziło to zdziwienie, zwłaszcza że nikt nie miał pojęcia, jak to się dzieje, iż system AI tak dobrze widzi chorobę psychiczną na podstawie samej tylko dokumentacji medycznej pacjenta. Owszem, specjaliści byli bardzo zadowoleni z pomocy tak skutecznego maszynowego diagnosty, byłiby jednak znacznie bardziej usatysfakcjonowani, gdyby rozumieli, w jaki sposób AI dochodzi do swoich wniosków.

Głośny był przed laty artykuł redaktora naczelnego serwisu Axios, Scotta Rosenberga, pt. „Najstraszniejsza tajemnica sztucznej inteligencji”, w którym opisuje powszechną wśród twórców sztucznej inteligencji świadomość faktu, że nie zawsze potrafią oni wyjaśnić lub przewidzieć zachowania swoich systemów. A skoro oni nie rozumieją, zwracał uwagę autor, to tym bardziej nie rozumieją wszyscy dookoła, od polityków, którym na rozwoju AI tak bardzo zależy, bo ma dać przewagę strategiczną nad rywalami, po zwykłych ludzi, którym przecież wyjaśnienia się należą.

Przypomnijmy, że duże modele językowe (LLM), ChatGPT firmy Open AI, Claude firmy Anthropic i Gemini firmy Google, nie są tradycyjnymi systemami oprogramowania działającymi zgodnie z jasnymi instrukcjami napisanymi przez ludzi, takimi jak Microsoft Word. Program Word robi dokładnie to, do czego został zaprojektowany. Natomiast ogromne sieci neuronowe, w zamierzeniu podobne do mózgu, które przetwarzają ogromne ilości informacji, by nauczyć się generować odpowiedzi, tak nie działają. Wprawdzie inżynierowie wiedzą, co uruchamiają i z jakich źródeł danych czerpie to, co uruchamiają, jednak rozmiar modelu LLM, ogromna,

niehumaniczna liczba zmiennych w każdym wyborze „najlepszego następnego słowa”, w praktyce oznacza, że nawet eksperci nie potrafią dokładnie wyjaśnić, dlaczego model wybiera konkretną odpowiedź. Możemy obserwować pracę LLM, to wszystko, co generuje, ale proces, w którym decyduje o odpowiedzi, jest w dużej mierze niejasny. Jak otwarcie stwierdzili badacze z OpenAI, „nie opracowaliśmy jeszcze zrozumiałych dla człowieka wyjaśnień, dlaczego model generuje określone wyniki”.

Tak samo jak nie wiadomo dokładnie, jak działa, nie wiadomo też, dlaczego nie działa, gdy nie działa tak, jak ma działać. Gdy OpenAI zmodyfikowało architekturę modelu w GPT-4, w odpowiedziach pojawiło się więcej halucynacji niż we wcześniejszych wersjach. Firma nie tylko przyznała, że tak się stało, ale przyznała, że nie wie, dlaczego tak się dzieje. Anthropic z kolei przyznał, że nie wie, dlaczego Claude 4 posunął się do szantażu w jednym z przypadków interakcji. „Z pewnością nie rozwiązyaliśmy kwestii interpretowalności”, powiedział Altman na konferencji w Genewie kilka lat temu. Dario Amodei, dyrektor generalny firmy Anthropic, w eseju zatytułowanym „The Urgency of Interpretability” (z ang. „Pilna potrzeba interpretowalności”) ostrzegał: „Osoby spoza branży często są zaskoczone i zaniepokojone, gdy dowiadują się, że nie rozumiemy, jak działają nasze własne rozwiązania AI. Ich obawy są uzasadnione. Ten brak zrozumienia jest zasadniczo bezprecedensowy w historii technologii”. Amodei nazwał to poważnym zagrożeniem dla ludzkości, co nie przeszkadza jego firmie pracować nad coraz bardziej zaawansowaną AI, choć z drugiej strony Amodei zapewnia, że jego firma pracuje także nad problemem



2. Zrzut ekranowy artykułu Dario Amodei – „The Urgency of Interpretability”

„interpretowalności”. „Mamy zespół badawczy, który koncentruje się na rozwiązaniu tej kwestii i poczynił znaczące postępy w poszerzaniu wiedzy branży na temat wewnętrznego funkcjonowania sztucznej inteligencji”, pisał (2).

Aby nieco uspokoić zaniepokojonych, warto dodać, że firma Apple opublikowała niedawno pracę pt. „The Illusion of Thinking” (ang. „Iluzja myślenia”), w którym podkreśla się, że nawet najbardziej zaawansowane modele rozumowania sztucznej inteligencji tak naprawdę nie „myślą” i mogą zawieść w testach warunków skrajnych. Badanie wykazało, że najnowocześniejsze modele (o3-min firmy OpenAI, DeepSeek R1 i Claude-3.7-Sonnet firmy Anthropic) nadal nie są w stanie rozwiązać ogólnych umiejętności rozwiązywania problemów, a ich dokładność ostatecznie spada do zera „po przekroczeniu pewnego poziomu złożoności”.

Od początku, czyli od momentu, gdy pojęcie sztucznej inteligencji stało się znane, istniały dwie szkoły myślenia o AI. Pierwsza zakładała, że najbardziej sensowne jest budowanie maszyn, które rozumują zgodnie ze znanymi nam zasadami i ludzką logiką, czyniąc ich wewnętrzne funkcjonowanie przejrzystym dla każdego. Inni uważali, że inteligencja wyłoni się łatwiej, jeśli maszyny będą uczyć się przez obserwacje i kolejno powtarzane doświadczenia. To drugie oznacza odwrócenie typowego programowania komputerowego. Zamiast programisty piszącego polecenia mające rozwiązać problem, program generuje własny algorytm na podstawie przykładowych danych i pożądanego wyjścia. Techniki uczenia maszynowego, które później wyewoluowały w najpotężniejsze znane dziś systemy AI, podążyły właśnie drogą, która polega na tym, że zasadniczo maszyna sama się programuje. Przy tym wszystkim jednak nie można po prostu zajrzeć do głębokiej sieci neuronowej, aby zobaczyć, jak działa „wewnątrz”. Procesy rozumowania sieci są osadzone w zachowaniach tysięcy symulowanych neuronów, ułożonych w dziesiątki, a nawet setki misternie połączonych ze sobą warstw. Każdy z neuronów w pierwszej warstwie odbiera sygnał wejściowy, np. intensywność piksela w obrazie, a następnie wykonuje obliczenia, zanim wyprowadzi sygnał wyjściowy. Te są podawane, w złożonej sieci, do neuronów w następnej warstwie – i tak dalej, aż do ostatecznego sygnału wyjściowego. Ponadto istnieje proces znany jako back-propagation, dostosowujący obliczenia dokonywane w poszczególnych neuronach w sposób, który pozwala sieci uczyć produkować pożądaną sygnał wyjściowy.

Już przed laty próbowano wyjaśnić, co dokładnie dzieje się w tego rodzaju systemach. W 2015 r.



3. Jeden z nieco przerażających obrazów w stylu Deep Dream

naukowcy z Google’a tak zmodyfikowali algorytm rozpoznawania obrazu oparty na głębokim uczeniu, aby zamiast dostrzegać obiekty na zdjęciach, generował on je lub modyfikował. Uruchamiając algorytm w odwrotnym kierunku, chcieli odkryć cechy, które program wykorzystuje do rozpoznawania, powiedzmy, ptaka lub budynku. Eksperymenty te, znane publicznie pod nazwą Deep Dream, doprowadziły do wygenerowania niesamowitych, czasem cokolwiek przerażających (3), obrazów, przedstawiających groteskowe, dziwaczne zwierzęta, krajobrazy i postacie. Ujawniając pewne tajniki postrzegania maszynowego, np. fakt wielokrotnego powracania i powtarzania określonych wzorców, pokazały zarazem, jak bardzo głęboka nauka maszynowa różni się od ludzkiej percepcji – np. w tym sensie, że rozbudowuje i powiela artefakty, które my podczas naszego procesu percepcyjnego ignorujemy bez zastanowienia.

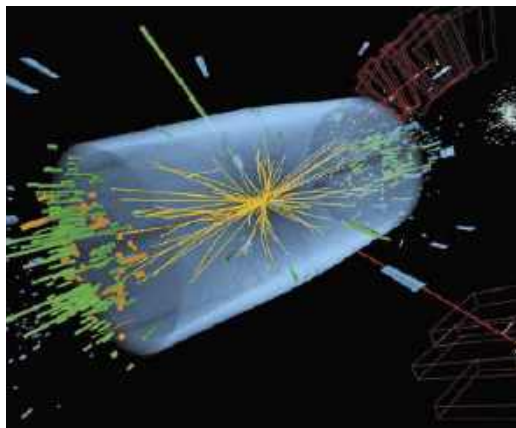
Dla wielu naukowców takie badania są jednak nieporozumieniem, gdyż do zrozumienia systemu, do rozpoznawania wzorców wyższego rzędu i mechanizmów podejmowania skomplikowanych decyzji trzeba, ich zdaniem, poznać całość interakcji obliczeniowych wewnątrz głębokiej sieci neuronowej. To gigantyczny labirynt funkcji matematycznych i zmiennych. Nie do ogarnięcia.

Scenariusz zautomatyzowanej nauki

Brandon Boesch (4), filozof z Uniwersytetu Morningside w amerykańskim stanie Iowa, jest jednym z pierwszych naukowców, którzy postanowili zmierzyć się z problemem dzieł naukowych autorstwa sztucznej inteligencji i ich potencjalnej niezrozumiałości dla ludzi. W swoich wypowiedziach przytacza on innego filozofa, Paula Humphreysa, który opracował koncepcję „hybrydowego scenariusza”

nauki, w którym część procesu naukowego jest powierzana komputerom. Chociaż zaczął pisać o tych ideach na długo przed pojawieniem się generatywnej sztucznej inteligencji (AI), Humphreys dalekowzrocznie, przewidywał, że dni, w których ludzie przewodzą procesowi naukowemu, mogą być policzone. Zidentyfikował także następną, późniejszą, fazę nauki, nazwaną przez niego „scenariuszem zautomatyzowanym”, w którym komputery całkowicie przejmą naukę. W tej wersji przyszłości moce obliczeniowe do rozumowania naukowego, przetwarzania danych, tworzenia modeli i teoretyzowania znacznie przewyższałyby nasze własne możliwości, do tego stopnia, że my, ludzie, nie będziemy już potrzebni. Maszyny będą kontynuować rozpoczętą przez nas pracę naukową, wznosząc nasze teorie na nowe wyżyny, których nie mogliśmy nawet przewidywać.

„Według niektórych źródeł, koniec ludzkiej dominacji nad nauką jest już bliski”, pisze Boesch w jednym z esejów. „Niedawne badanie przeprowadzone wśród badaczy sztucznej inteligencji wykazało 50-procentową szansę, że w ciągu stulecia AI mogłaby realnie zastąpić nas w każdym zawodzie. (...) To byłby naprawdę dziwny świat. Po pierwsze, AI



5. Wizualizacja zderzenia protonów w detektorze CMS z 2012 r. powszechnie używana jako ilustracja odkrycia bozonu Higgsa

mogłaby zdecydować się na zgłębianie problemów naukowych, którymi naukowcy nie chcą się zajmować, tworząc zupełnie nowe ścieżki odkryć. Mogłaby nawet zdobyć wiedzę o świecie wykraczającym poza to, co nasze mózgi są w stanie pojąć”.

Jak wskazuje Boesch, zautomatyzowany scenariusz Humphreysa dla nauki byłby w gruncie rzeczy kolejnym krokiem w trwającym od wieków procesie. Ludzie nigdy tak naprawdę nie uprawiali nauki sami – podkreśla. Od dawna polegaliśmy na narzędziach, które wzbogacały naszą obserwację świata, mikroskopach, teleskopach, standardowych liniijkach i zlewkach itd. Istnieje wiele zjawisk fizycznych, których nie możemy bezpośrednio precyzyjnie obserwować bez instrumentów, takich jak termometry, liczniki Geigera, oscyloskopy, kalorymetry i tym podobne. Wprowadzenie komputerów stanowiło kolejny krok w kierunku zmniejszenia roli człowieka w nauce. Pierwsze amerykańskie loty kosmiczne wymagały obliczeń wykonywanych przez matematyków. Do czasu misji księżycowych, niecałą dekadę później, większość tych obliczeń została już przekazana komputerom. Kolejne dekady były świadkami ciągłego wzrostu mocy obliczeniowej a także skorelowanego z tym spadku kosztów obliczeń. Symulacje obliczeniowe odegrały kluczową rolę w odkryciu bozonu Higgsa (5). „Obecnie znajdujemy się na etapie, który moglibyśmy nazwać zaawansowanym, hybrydowym stadium nauki, z jeszcze większym wykorzystaniem systemów obliczeniowych”, pisze Boesch.

I dodaje, że obecnie coraz większy wpływ na naukę zaczęła wywierać sztuczna inteligencja. Na przykład model AlphaFold zaprojektowany do przewidywania geometrii skręcania cząsteczek białkowych



4. Brandon Boesch © Uniwersytet Morningside

na podstawie ich składu chemicznego. Chociaż ludzie mogliby wykonywać tę pracę niezależnie od komputerów, jest to czasochłonne, pracochłonne i kosztowne. Twórcy AlphaFold z Google DeepMind twierdzą, że pozwoliło to zaoszczędzić „setki milionów lat czasu badawczego”. Podobne korzyści można zaobserwować w różnych dziedzinach nauki, w analizie niezwykle dużych zbiorów danych w astronomii i genomice, opracowywaniu nowatorskich dowodów w matematyce, przewidywaniu pogody, opracowywaniu nowych leków i wielu innych.

Kiedy wkład sztucznej inteligencji obliczeniowej zaczyna być mierzony w „setkach milionów lat”, zaczynamy odczuwać, pisze Brandon Boesch, że my, ludzie, jesteśmy mało efektywnymi elementami w tych wszystkich projektach. „Możemy się więc zastanawiać, czy jesteśmy już w scenariuszu zautomatyzowanym, ale raczej tak nie jest... jeszcze”. uważa. „Nasz wkład w naukę pozostaje kluczowy. to my, ludzie, wciąż podejmujemy decyzje, identyfikujemy pytania naukowe, interpretujemy wyniki i ostatecznie decydujemy o jej postępie. Jeśli podążymy ścieżką Humphreysa, całkowita abdykacja człowieka z tronu nauki nastąpi dopiero na późniejszym etapie. Będzie tak wtedy, gdy sztuczne superinteligencje byłyby nie tylko zdolne do wykonywania zadań, które im wyznaczyliśmy (kontynuacja scenariusza hybrydowego), ale także do wyznaczania własnych zadań, własnego programu badawczego, gromadzenia danych, modelowania i teorii, zgodnie z własnym, niezależnie zidentyfikowanym zestawem teoretycznych cnót i wartości”.

Logiczną konsekwencją w ocenie filozofa jest sytuacja, w której, gdyby superinteligencje realizowały własne programy badawcze, ich praca stałaby się dla nas niezrozumiała. Chociaż sztuczne superinteligencje mogłyby nadal działać w ramach paradygmatów naszych obecnych teorii, nie ma powodu, dla którego musiałyby to robić – mogłyby zdecydować się na nowy początek z nową teorią świata. Podobnie, chociaż mogą używać matematyki i symboli znanych ludziom, nie byłyby ograniczone tymi konwencjami – mogłyby szybko opracować nową matematykę i systemy wyrażania. „Z naszego punktu widzenia ich badania byłyby nauką stworzoną do celów teoretycznych, których nie znamy, z celami, których nie znamy, interpretowanymi w sposób, którego nie znamy”, zwraca uwagę Boesch. „Jest również taki problem, że rozumowanie sztucznej superinteligencji może w praktyce przekraczać nasze możliwości intelektualne. Zrozumienie nauki w scenariuszu zautomatyzowanym

może wymagać na przykład jednoczesnego rozważenia setek złożonych modeli, każdy z setkami parametrów, z których żaden nie jest powiązany ze znanym ludzkim pojęciem. Choć możliwe, że moglibyśmy zrozumieć parametry (a może nawet modele) indywidualnie, brakowałoby nam zdolności przetwarzania w mózgu, by łączyć je wszystkie naraz”. Czyli wywoły AI byłyby czymś w rodzaju przysłowiowego tureckiego kazania dla najmądrzejszych nawet ludzi.

Jeśli rezultaty całkowicie zautomatyzowanego scenariusza wykraczają poza nasze zrozumienie, to dlaczego mielibyśmy chcieć przeznaczać zasoby ekonomiczne i talent intelektualny na jego rozwój? Jednym z takich powodów może być, jak przypuszcza badacz, to, że wierzymy w pozytywny postęp. Być może superinteligencje będą tworzyć nowe elementy technologii, zasoby lub nowe sposoby rozwiązywania problemów. Inżynierowie ludzcy mogliby następnie wykorzystać te nowe technologie i znaleźć dla nich zastosowania, nawet jeśli nie rozumieją dokładnie, jak działają. Ludzkość wielokrotnie korzystała z pomysłów i odkryć, których (jeszcze) nie rozumiała w sensie naukowym, co nie przeszkadzało jej korzystać. Są jednak argumenty przemawiające za rezygnacją z realizacji scenariusza automatyzacji. „Być może odkrycia dokonane i przekazane nam przez sztuczną superinteligencję przyniosą nową i straszną broń”, ostrzega uczony. „Może zwiększyłyby ryzyko scenariusza katastroficznego, takiego jak zniewolenie lub zagłada ludzkości. Można też wyobrazić sobie obawę, że niektóre superinteligencje zaczną działać z nienaturalnie ludzką pychą, eksperymentując w sposób niebezpieczny, nieetyczny lub sprzeczny z podzielanymi przez ludzkość wartościami”.

Jednak, podsumowując, Boesch sądzi, że pomimo tych obaw wydaje się mało prawdopodobne, abyśmy byli w stanie powstrzymać rozwój scenariusza zautomatyzowanego. Nie możemy uciec przed siłami kapitału i konkurencji. Jednak nie zgadza się z twierdzeniami Humphreysa, że scenariusz zautomatyzowany zastąpi ludzką naukę. „Nasze pragnienia zrozumienia, wyjaśnienia, wiedzy i kontroli będą trwać”, wierzy. „Nie możemy powstrzymać się od podjęcia działań, by realizować te pragnienia i kontynuować uprawianie nauki. Będziemy tkwić w naszej ciekawości, aby zrozumieć i wyjaśnić otaczający nas świat przyrody”. Zaś scenariusz zautomatyzowany byłby w tym alternatywną ścieżką, która nie zastępuje, lecz uzupełnia naukę ludzi. ■

Mirosław Usidus



1. Wizualizacja pustej przestrzeni w Wielkiej Piramidzie na podstawie skanowania mionowego

Znany egipski archeolog Zahi Hawass niedawno zasugerował, że w 2026 roku zostanie ujawnione odkrycie dotyczące wnętrza Wielkiej Piramidy w Gizie, która ma „przepisać historię” na nowo. Nowe odkrycia we wnętrzu piramid wykorzystują zaawansowaną technikę skanowania przekrojów kamiennych obiektów, która dekady temu nie była dostępna.

Najnowsza technika w odkrywaniu przeszłości

LIDAROWYM OKIEM W GŁĘB DZIEJÓW

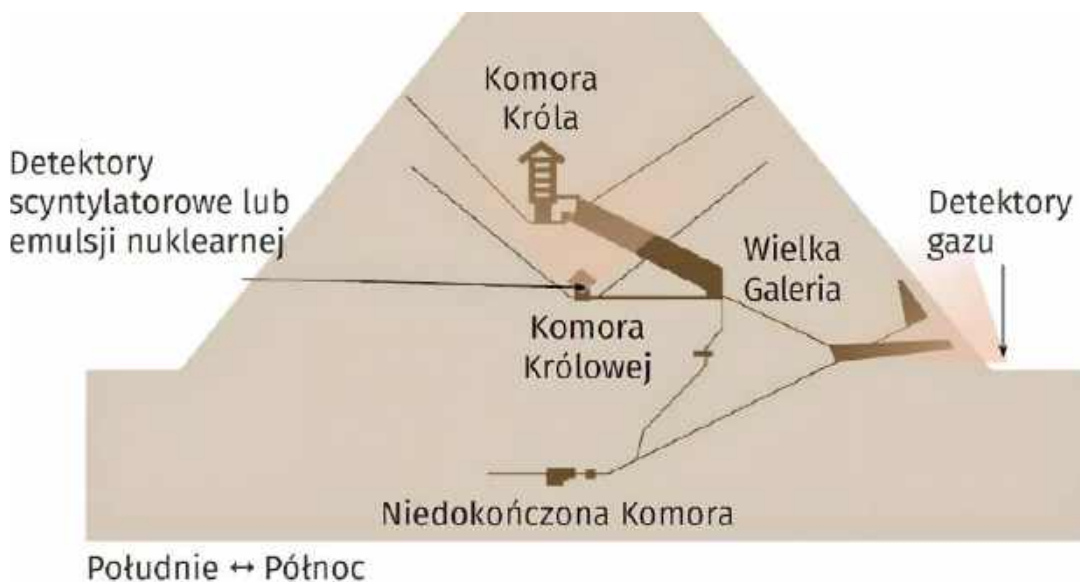
Uchylając dalej rąbka tajemnicy, Hawass dodał, że „to wielkie odkrycie to nowy korytarz o długości 30 metrów” (1), który, jak powiedział, został „wykryty przy użyciu zaawansowanego sprzętu” i wydaje się prowadzić do ukrytych drzwi wewnątrz Wielkiej Piramidy.

Nie jest to pierwszy i jedyny raz, kiedy ogłasza się odkrycia dotyczące tajemniczych egipskich piramid dokonywane za pomocą sprzętu i technik nowej

generacji. W 2023 r. w ramach projektu ScanPyramids ogłoszono odkrycie w Wielkiej Piramidzie korytarza o długości dziewięciu metrów, opierając się na wcześniejszych badaniach, w których naukowcy zidentyfikowali „pustkę” jeszcze w 2017 r. Zespół ScanPyramids wykorzystywał termografię w podczerwieni, tomografię mionową, symulacje 3D oraz techniki rekonstrukcji (2).

Wypowiedzi Hawassa wywołały spekulacje, że odkrycia mogą być związane z odkryciem grobowca legendarnego Imhotepa, uchodzącego tradycyjnie za głównego budowniczego piramid, chociaż słynny starożytny egipski architekt żył w III dynastii, prawie sto lat przed budową Wielkiej Piramidy. Imhotepowi przypisuje się budowę słynnej piramidy schodkowej w Sakkarze dla faraona Dżesera z III dynastii.

Hawass powstrzymał się od komentarza na temat tego, co znajduje się za przejściem w Wielkiej Piramidzie, twierdząc, że tegoroczna publikacja międzynarodowego zespołu naukowców dostarczy pełniejszych informacji dopiero po dokładnej analizie uzyskanych



2. Naukowcy umieścili trzy różne typy detektorów mionowych w Wielkiej Piramidzie i wokół niej, aby sporządzić mapę gęstości struktury i poszukać ukrytych komór

danych, które zostały zebrane przy użyciu technologii mapowania 3D i radiografii mionowej.

Tymczasem także we wnętrzu innej z trzech piramid w Gizie, najmniejszej, zwanej Piramidą Mykerinosa, odkryto znów dwie kolejne „pustki”. Zostały ujawnione w skanach wykonanych przez naukowców z uniwersytetu w Kairze w Egipcie i politechniki w Monachium (TUM) w Niemczech. Badania, które to odkryły, są częścią projektu ScanPyramids. Naukowcy zastosowali szereg nieniszczących metod obrazowania, w tym georadar, ultradźwięki i tomografię oporu elektrycznego. Naukowcy planują również badania mionograficzne innych obiektów i struktur archeologicznych na świecie, np. piramidy Majów w Chichén Itzá w Meksyku.

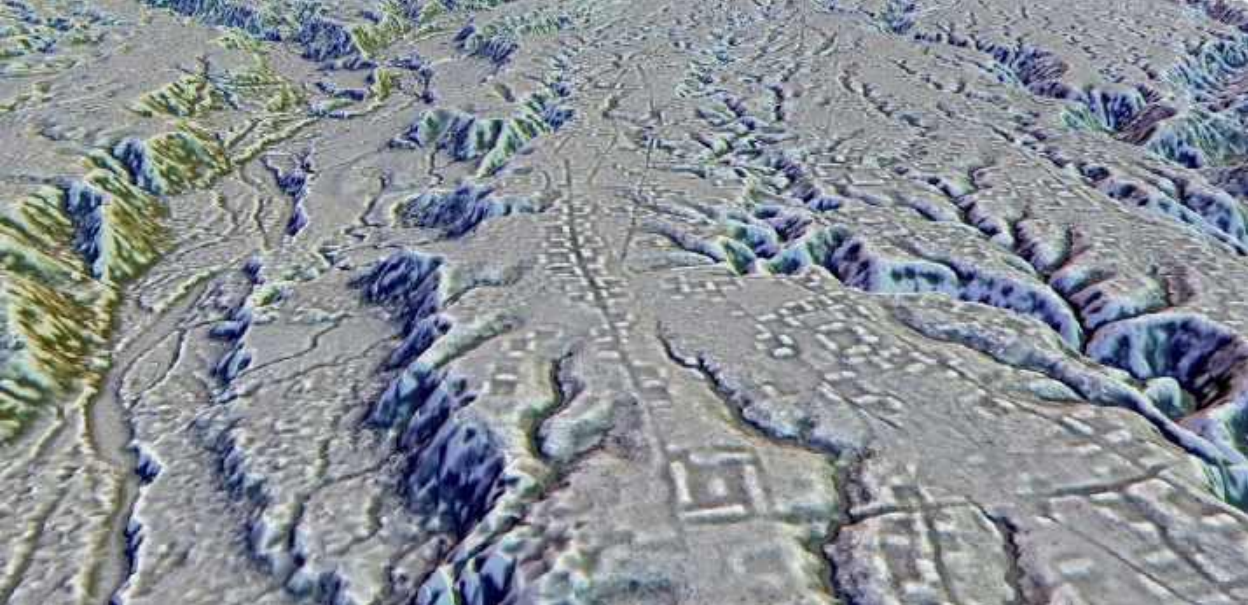
Sieć ulic skryta pod dżunglą

Archeolodzy wierzyli kiedyś, że porośnięte dżunglą tereny Amazonii były miejscem niegościnnym, słabo zaludnionym najwyżej przez rozproszone przez grupy łowców zbieraczy. Jednak pozostałości ogromnych fortyfikacji ziemnych, piramid i dróg od Boliwii po Brazylię, odkryte w ciągu ostatnich dwóch dekad, dowiodły, że Amazonia była miejscem rozwoju dużych, złożonych społeczeństw na długo przed przybyciem europejskich kolonizatorów. Gęsta sieć połączonych ze sobą miast, obecnie ukryta pod lasem w dolinie Upano w Ekwadorze, została odkryta dzięki technologii mapowania laserowego o utrwalonej już polskiej nazwie lidar (od angielskiego akronimu

LIDAR, utworzonego od wyrażenia: Light Detection and Ranging).

Odkryte w Amazonii osady, opisane w czasopiśmie „Science”, mają co najmniej 2500 lat, czyli są o ponad 1000 lat starsze niż jakiegokolwiek inne znane złożone społeczeństwo amazońskie. Mapa lidarowa miasta Kunguints w ekwadorskiej Amazonii (3) ujawniła sieć ulic pomiędzy domami. Lidar pozwala badaczom zajrzeć pod pokrywę lasu i odtworzyć starożytne stanowiska archeologiczne znajdujące się pod nią.

Stéphane Rostain, archeolog z CNRS, francuskiej agencji badawczej, rozpoczął wykopaliska w dolinie Upano prawie 30 lat temu. Jego zespół skupił się na dwóch dużych osadach, zwanych Sangay i Kilamope, i znalazł kopce zorganizowane wokół centralnych placów, ceramikę zdobioną farbą i wrytymi liniami oraz duże dzbany zawierające pozostałości tradycyjnego piwa kukurydzianego chicha. Datowanie radiowęglowe wykazało, że stanowiska w Upano były zamieszkane od około 500 r. p.n.e. do okresu między 300 a 600 r. n.e. Badacz miał przeczucie, że to, co widać z powierzchni ziemi gołym okiem, to nie wszystko, że jego zespół nie ma pełnego obrazu regionu. Sytuacja zmieniła się, gdy Ekwadorski Narodowy Instytut Dziedzictwa Kulturowego sfinansował w 2015 r. badania doliny za pomocą lidar. Specjalnie wyposażone samoloty wysyłały impulsy laserowe do lasu i mierzyły ich powrót, ujawniając cechy topograficzne, które w innym przypadku byłyby niewidoczne pod drzewami. Dane z lidar pozwoliły



3. Mapa lidarowa miasta Kunguints w ekwadorskiej Amazonii

Rostainowi i jego współpracownikom dostrzec powiązania między osadami, a także odkryć wiele nowych. Zespół zidentyfikował pięć dużych osad i 10 mniejszych na obszarze 300 kilometrów kwadratowych w dolinie Upano, z których każda była gęsto zabudowana budynkami mieszkalnymi i ceremonialnymi. Miasta są poprzecinane prostokątnymi polami uprawnymi i otoczone tarasami na zboczach wzgórz, gdzie ludzie uprawiali rośliny, w tym kukurydzę, maniok

i słodkie ziemniaki, znalezione podczas wcześniejszych wykopalsk. Szerokie, proste drogi łączyły miasta między sobą, a ulice biegły między domami i dzielnicami w obrębie każdej osady.

Autorzy twierdzą, że zakres zmian krajobrazu w Upano dorównuje „miastom-ogrodom” klasycznej cywilizacji Majów. A to, co odkryto do tej pory, to „tylko wierzchołek góry lodowej” tego, co można znaleźć w ekwadorskiej Amazonii, uważa wielu badaczy. Sić

4. Kompozytowy widok lidarowy Tugunbulak © SAIElab



dróg łączących stanowiska archeologiczne w Upano sugeruje, że wszystkie one istniały w tym samym czasie. Są one o tysiąc lat starsze od innych złożonych społeczeństw amazońskich, w tym Llanos de Mojos, niedawno odkrytego starożytnego systemu miejskiego w Boliwii.

Nowe ślady cywilizacji od gór Uzbekistanu po lasy w Michigan

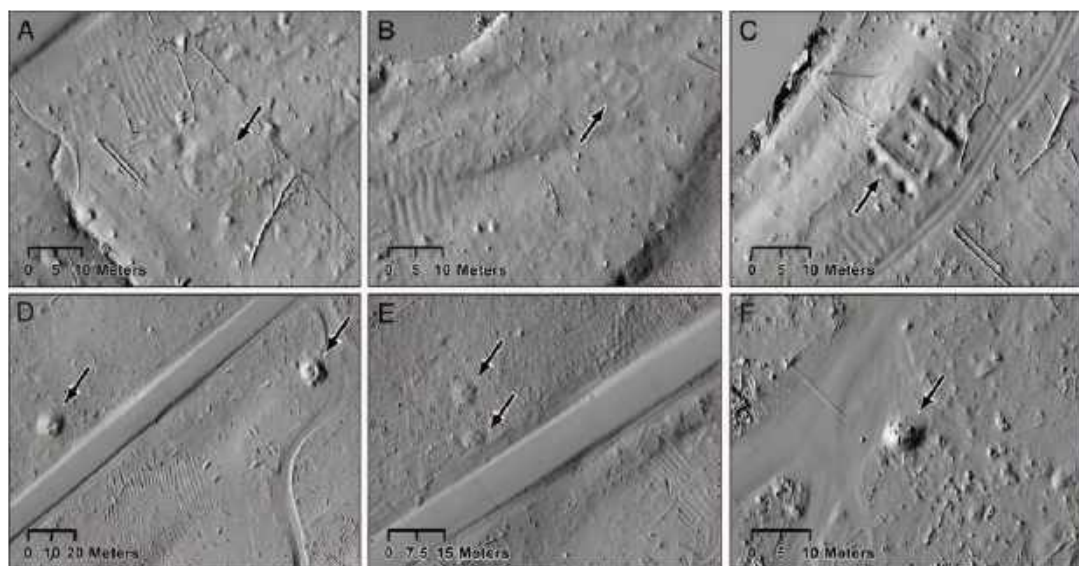
Pierwsze w historii zastosowanie najnowocześniejszego lidar, opartego na dronach, w Azji Środkowej pozwoliło archeologom uchwycić niesamowite szczegóły dwóch nowo udokumentowanych miast handlowych położonych wysoko w górach Uzbekistanu. Zespół naukowców pod kierownictwem Michaela Frachetti z Uniwersytetu Waszyngtana w St. Louis, oraz Farhoda Maksudova, dyrektora Narodowego Centrum Archeologii w Uzbekistanie, wykorzystał lidar oparty na dronach do sporządzenia mapy archeologicznej i planu dwóch niedawno odkrytych wysoko położonych stanowisk archeologicznych w Uzbekistanie. Średniowieczne miasta należą do największych, jakie kiedykolwiek udokumentowano w górzystych częściach Jedwabnego Szlaku, rozległej sieci starożytnych szlaków handlowych łączących Europę i Azję Wschodnią.

Dron uchwycił obrazy Tugunbulak (4) w 2018 roku. Skanowanie za pomocą drona z lidarem zapewniło niezwykle szczegółowy obraz placów, fortyfikacji, dróg i osad, które kształtowały życie i gospodarkę społeczności górskich, handlarzy i podróżników od VI do XI wieku w Azji Środkowej. Oba miasta położone są na trudnym terenie, na wysokości od 2000 do 2200 metrów nad

poziomem morza. Mniejsze miasto, obecnie nazywane Tashbulak, zajmowało powierzchnię około 12 hektarów, podczas gdy większe miasto Tugunbulak osiągało powierzchnię 120 hektarów, „co czyniło je jednym z największych miast regionalnych tamtych czasów”, wyjaśnia w komunikacie Frachetti. „Były to ważne ośrodki miejskie w Azji Środkowej, zwłaszcza gdy opuszczało się nizinne oazy i przenosiło się do trudniejszych warunków wysokogórskich”.

Technologia lidar jest powszechnie stosowana do mapowania krajobrazów archeologicznych zasłoniętych gęstą roślinnością, ale ma dodatkową wartość w miejscach, gdzie roślinność jest rzadka, takich jak góry Uzbekistanu. Skanowanie z dokładnością do centymetra umożliwiło zaawansowaną analizę komputerową starożytnych powierzchni archeologicznych. „Są to jedne z najwyższej rozdzielczości obrazów lidarowych stanowisk archeologicznych, jakie kiedykolwiek opublikowano”, mówi Frachetti. „Ostateczne mapy w wysokiej rozdzielczości były kompozycją ponad siedemnastu lotów dronów w ciągu trzech tygodni”.

Naukowcy z Dartmouth wykorzystujący technologię lidarową zamontowaną na dronach do badania lasów Półwyspu Górnego (Upper Michigan) w USA donieśli niedawno o odkryciu tam śladów bytowania, liczącej 1000 lat społeczności rolniczej, co podważa dotychczasowe założenia dotyczące przedkolonialnej Ameryki. Badania zespołu z Dartmouth skupiały się na stanowisku archeologicznym na Półwyspie Górnym Michigan, zwanym Anaem Omot, co w języku Menominee oznacza „psi brzuch”. Ponieważ w wielu



5. Skany terenów w Michigan za pomocą lidara

wcześniejszych badaniach tego miejsca, sięgających lat pięćdziesiątych XX wieku, nie była dostępna technika lidar, władze plemienne Menominee poprosiły naukowców o zbadanie i udokumentowanie tego obszaru przy użyciu tej właśnie nowej technologii. Pierwsza wyprawa rozpoczęła się w maju 2023 r., po stopieniu się śniegu, ale przed pojawieniem się wiosennych liści na drzewach.

Profesor powiedział, że gęste lasy są szczególnie „mylące” dla archeologów, co powoduje, że polegają oni na publicznie dostępnych badaniach lidarowych. Jednak większość tych skanów została wykonana przez samoloty lecące na dużej wysokości, co często ograniczało dokładność danych. Zespół uznał, że drony zdolne do latania znacznie niżej niż samoloty mogą znacznie zwiększyć rozdzielczość skanów i ogólną jakość danych. Drony wyposażone w lidar zbadały obszar o powierzchni 330 akrów. Szacują oni, że obszar ten stanowi około 40 proc. większego stanowiska, a pozostała część nadal nie została zmapowana. Analiza skanów ujawniła kilka interesujących cech, w tym zestawy równoległych „grzbietów”, które tworzyły wzory przypominające „kołdrę”, rozciągające się na starożytnym krajobrazie. Jednym z najbardziej zaskakujących odkryć była ogólna skala 1000-letnich terenów rolniczych. Według zespołu badawczego, żadna z dotychczasowych publikacji na temat przedkolonialnych populacji tubylczych nie sugerowała, że uprawiali oni ziemię na taką skalę i przy takim poziomie organizacji wspólnotowej. „Jeśli spojrzeć na skalę upraw, wymagałoby to takiego rodzaju organizacji pracy, która zazwyczaj kojarzy się z dużo większym, hierarchicznym społeczeństwem na poziomie państwowym”, mówi McLeester.

Miasto Majów, które zniknęło pod dżunglą w Meksyku, zostało odkryte przypadkowo. Archeolodzy twierdzą, że jest ono wielkości Edynburga, stolicy Szkocji, i znajduje się w południowo-wschodnim stanie Campeche, gdzie otaczają je piramidy, boiska sportowe i groble łączące dzielnice i amfiteatry. Aby znaleźć ukryty kompleks o nazwie Valeriana, archeolodzy wykorzystali lidar. Kompleks jest ogromny i ustępuje jedynie Calakmul, uważanemu za największe stanowisko archeologiczne Majów w starożytnej Ameryce Łacińskiej. Miasto zostało nazwane Valeriana na cześć pobliskiej laguny.

AI czyta zwoje, które się spaliły i analizuje skany terenu

Wiele fizycznych dowodów na istnienie ludzkości zostało wymazanych, poddanych erozji, zakopanych, wysadzonych w powietrze i na inne sposoby

zatarzonych. Od prawie dwóch stuleci poszukiwaniami tego, co da się odzyskać, zajmują się archeolodzy. Kiedyś musieli oni w dużym stopniu polegać na łopatach, pędzlach, szklach powiększających, żmudnej pracy w terenie, jednej kości, fragmencie ceramiki lub kawałku węgla drzewnego naraz. Ten rodzaj pracy jest kontynuowany, ale dziś archeolodzy wykorzystują również zdumiewający zestaw technologii XXI wieku do odkrywania przeszłości.

Te nowe narzędzia to nie tylko opisywany wyżej lidar, ale również sztuczna inteligencja, sekwencjonowanie DNA, zdjęcia satelitarne, drony przenoszące termiczne kamery na podczerwień i małe roboty, które mogą czołgać się w dół szybu grobowca. Technologie teledetekcji, a także wykorzystanie sztucznej inteligencji i uczenia maszynowego, które przeszukują duże zbiory danych, znacznie zwiększają prawdopodobieństwo, że miejsce docelowe archeologów przyniesie ważne artefakty i nie będzie stratą wysiłku.

Pierwszy technologiczny przełom w badaniach przeszłości nastąpił w połowie ubiegłego wieku, kiedy fizycy opracowali datowanie radiowęglowe, wykorzystując przewidywalny rozpad izotopu węgla-14, który jest absorbowany przez żywe organizmy, do stabilnych form azotu-14 po śmierci organizmów. Naukowcy mogą obliczyć liczbę pozostałych atomów węgla-14 i określić, jak dawno temu były one częścią czegoś żywego. Daje to archeologom możliwość datowania rzeczy, które powstały w ciągu ostatnich 50 tys. lat.

21-letni student informatyki wygrał światowy konkurs Vesuvius Challenge na odczytanie tekstu zapisanego w zwęglonym zwoju papirusowym ze starożytnego rzymskiego miasta Herkulanum. Napisał był dotychczas nieczytelny, gdyż rolka została spalona wskutek wybuchu wulkanu w 79 r. n.e., tego samego, który pogrzebał pobliskie Pompeje. Oparta na rentgenowskiej tomografii komputerowej i algorytmach AI metoda może pozwolić odczytać setki tekstów pochodzących z jedynej nienaruszonej biblioteki, która przetrwała z kultury grecko-rzymskiej.

Zwęglonych z powodu pożaru starożytnych dokumentów nie można było bez zniszczenia rozwinąć. A nawet gdyby się to udało, to odczytanie ze spalonego papirusu byłoby wyzwaniem. Skanowanie tomografią pozwoliło „wirtualnie rozwinąć” zwoje, wykrywając warstwy i zapisane powierzchnie. Metoda ta była już wcześniej wykorzystywana do odczytywania spalonych starożytnych dokumentów, ale w Herkulanum w odróżnieniu od wcześniejszych przypadków tusz nie zawierał metalu, który w tomografii pomaga wykrywać napisy.



6. Zwęglone zwoje spod Wezuwiusza i Luke Farritor, zwycięzca Vesuvius Challenge

To było pierwsze wyzwanie. Drugim było czytanie znaków i łączenie ich w odpowiadające rzeczywistości słowa i zdania. W konkursie na najlepszą i najtrafniejszą metodę, w którym brało udział ok. 1,5 tysiąca osób na całym świecie, wygrał Luke Farritor z Uniwersytetu Nebraska-Lincoln, który opracował algorytm uczenia maszynowego, wykrywający greckie litery na kilku liniach zwiniętego papirusu, w tym *πορφύρας* (*porphyras*), co oznacza „purpurowy”. Farritor wykorzystał subtelne, niewielkie różnice w fakturze powierzchni, by wyszkolić sieć neuronową (6).

Jak wiadomo, AI już w tej chwili dokonuje zdumiewających rzeczy w dziedzinie językowej, zwłaszcza tłumaczeń. Zapewne wkrótce dokona jeszcze więcej. Google buduje uniwersalny translator mowy. Szef Google Brain, Zoubin Ghahramani, zapowiadał model tłumaczący wzajemnie pomiędzy tysiącem najważniejszych języków świata, jednak współczesne języki, choćby nawet najbardziej egzotyczne, to dopiero początek wyzwania. Algorytmy analizujące sygnały z elektroencefalografów i innych technik odczytywania aktywności mózgu, nad którymi pracuje wiele firm, w tym wspomniana Meta, mogłyby pomóc porozumiewać się z otoczeniem ludziom mającym problemy z mową lub np. sparaliżowanym. Niektórzy chcą też wykorzystać algorytmy AI do rozszyfrowania języków starożytnych, zapomnianych, nierozszyfrowanych artefaktów nieznanymi form pisma. Gdyby to się udało, to wiele zagadek historii zostałoby wreszcie wyjaśnionych.

W 2019 roku Jiaming Luo wraz z zespołem kolegów naukowców z MIT opracował algorytm wyszkolony na wzorcach zmian języków i pisma w czasie. Przetestowali swój model na dwóch starożytnych skryptach, które zostały już rozszyfrowane, piśmie ugaryckim oraz piśmie linearnym B, które zostało odkryte na greckiej wyspie Kreta. Rozszyfrowanie tego ostatniego zajęło badaczom stosującym tradycyjne metody prawie sześć dekad. Algorytm już po kilku godzinach umiał poprawnie przetłumaczyć 67,3 proc. słów z linearnego B na ich współczesne greckie odpowiedniki. Jeśli chodzi o pismo ugaryckie, które było wcześniej już analizowane przez algorytmy, narzędzie z MIT robiło to szybko.

W środowisku naukowym pojawiło się pytanie, czy uczenie maszynowe mogłoby pomóc w odczytaniu próbek pisma, które do tej pory opierały się wszelkim próbom tłumaczenia? Na przykład znaków znajdujących na tabliczkach-pieczęciach pochodzących z cywilizacji Doliny Indusu. Wyzwanie jest trudniejsze, gdyż nie mamy zbyt wielu informacji o tej starożytnej kulturze. Nie było też całkowitej pewności, czy te znaki to rzeczywiście pismo, a nie coś innego. W 2004 roku naukowcy z Harvardu, Steve Farmer, Richard Sproat i Michael Witzel, opublikowali pracę, w której twierdzili, że pieczęcie te to nic innego jak zbiór symboli religijnych lub politycznych, podobnych do, powiedzmy, znaków drogowych, a wszelkie próby rozszyfrowania ich jako języka były stratą czasu. Badacze starożytności odrzucali te hipotezy, ale

potrzeba było twardszych dowodów, że chodzi o pismo. Zespół matematyków pod wodzą Rajesha Rao opracował oparty na AI program sprawdzający, czy pismo Doliny Indusu to zapis języka. Badacze nakarmili swój program czterema tysiącami próbek znaków, które tworzą całość pisma Doliny Indusu. Doszli do wniosku, że pieczęcie ukazują pismo odpowiadające językowi mówionemu. Nie oznacza to, że pismo zostało odszyfrowane, ale stanowi punkt wyjścia do dalszych badań.

Archeolodzy wykorzystali też łączone techniki oparte na nowoczesnych radarach i sztucznej inteligencji i zdołali odkryć na pustyni niedaleko Dubaju w Zjednoczonych Emiratach Arabskich zaginioną cywilizację sprzed pięciu tysięcy lat, ukrytą i pogrzebaną pod największą pustynią świata, wraz z szlakami handlowymi i starożytnymi miastami. Dzięki połączeniu sztucznej inteligencji i tak zwanego radaru z syntetyczną aperturą (SAR), projekt ten, kierowany przez Marię Gonzalez, pozwolił na badania pustyni w Dubaju i walcie przyczynił się do odkryć śladów działalności ludzkiej w pozostałościach osad i długich szlakach handlowych.

Technologia SAR potrafi „zeskanować” tysiące kilometrów kwadratowych w ciągu kilku godzin, analizując dane w odniesieniu do istniejących stanowisk archeologicznych w celu wykrycia subtelnych różnic w ukształtowaniu terenu. Pozwala to zaoszczędzić nie tylko czas, ale także pieniądze, ponieważ wysyłanie specjalistów w teren jest niezwykle kosztowne, podczas gdy te nowe innowacje pozwalają przewidzieć wcześniej ukryte lokalizacje, które byłyby prawie niemożliwe do wykrycia przez samych ludzi. Co więcej, połączenie sztucznej inteligencji i technologii SAR pozwala nie tylko odkryć miejsca, w których mogły istnieć starożytne cywilizacje i ruchy ludności, ale także wyobrazić sobie, jak mogły wyglądać osady i życie sprzed wielu tysięcy lat.

Przez długi czas region pustyni Rub’ al-Khali, położony w pobliżu Arabii Saudyjskiej i Dubaju, był uważany za martwą otchłań. Do tego stopnia, że obszar ten nazywano Pustą Ćwiartką. Jednak nie wszystko było takie, jak się wydawało, ponieważ sekret ukryty pod piaskiem został w końcu odkryty. Wszystko zaczęło się w 2002 roku, kiedy szejkh Mohammed bin Rashid Al Maktoum, obecny władca Dubaju, podczas lotu helikopterem nad tą częścią pustyni zauważył dziwne formacje na wydmach. To przeczcucie skłoniło zespół naukowców do dalszych badań, w wyniku których przeprowadzono wykopaliska. Dzięki temu zespół odkrył Saruq Al-Hadid. Na głębokości 3 metrów pod piaskiem naukowcy odkryli artefakty,

w tym ceramikę i kości zwierzęce. Zespół badawczy z Uniwersytetu Khalifa w Abu Zabi wykorzystał radar z syntetyczną aperturą (SAR), aby zrozumieć, co kryje się pod powierzchnią pustyni, bez konieczności kopania. Wyniki badań sugerują, że w miejscu, gdzie obecnie znajduje się pustynia, niegdyś istniały osady i drogi. Wyniki badań opublikowano w czasopiśmie, w którym napisano: „Studium przypadku stanowiska Saruq Al-Hadid ilustruje potencjał tych technologii w zakresie usprawnienia badań archeologicznych i przyczynienia się do ochrony dziedzictwa kulturowego”.

Archeolodzy stworzyli modele głębokie nauki maszynowej, w szczególności z wykorzystaniem DeepLab V3+, do segmentacji semantycznej w celu identyfikacji nieznanymi wcześniej zbiorników wodnych z okresu Angkor. Angkor, położony w prowincji Siem Reap w Kambodży, był stolicą imperium Khmerów między IX a XV wiekiem naszej ery. Po upadku imperium w 1431 r. miasto zostało w większości opuszczone.

W poprzednich badaniach międzynarodowy zespół naukowców wykorzystał zdjęcia satelitarne i technologię lidarową do zlokalizowania pozostałości miejskiej zabudowy, odkrywając ponad 25 tysięcy obiektów pod koronami drzew. Jednak nakład pracy wymagany do sporządzenia mapy terenu poza centralną częścią Angkor okazał się niezwykle czasochłonny, co zostało rozwiązane w nowym badaniu opublikowanym w czasopiśmie „PLOS one” dzięki wykorzystaniu AI. Szkolenie sztucznej inteligencji wymagało oczyszczenia i opracowania dużego zbioru danych dotyczących znanych kształtów zbiorników wodnych. Spośród ponad 11 tys. przykładów, po dodatkowej ręcznej weryfikacji wykorzystano około 3600 wysokiej jakości próbek kompatybilnych ze sztuczną inteligencją. Według autorów badania sztuczna inteligencja potrafiła wykrywać złoże z akceptowalną dokładnością. Metoda ta, w połączeniu z weryfikacją przez ekspertów, oznacza, że sztuczna inteligencja może skrócić czas wykonywania zadań związanych z mapowaniem nawet o 90 proc.

Prześwietlanie starożytnego betonu i tatuaży

Kolejne zdobycze techniki umożliwiają nowe odkrycia, wcześniej nieosiągalne. Zaginione miasta są raz po raz odkrywane w gęstych lasach deszczowych albo, jak całkiem niedawno, bo w listopadzie 2025 r., na pustyniach Kazachstanu; student wykorzystuje sztuczną inteligencję do odczytania starożytnych zwojów, archeolodzy odkrywają nieznanymi wcześniej

starożytny rzymski obóz wojskowy na szczycie góry, naukowcy wykorzystują technikę akceleratorów cząstek w Lawrence Berkeley National Laboratory do prześwietlania za pomocą promieni rentgenowskich fragmentów rzymskiego betonu, co pomogło ujawnić jego okrytą mgłą tajemniczą recepturę, pozwalającą, jak wiadomo, Rzymianom budować rzeczy, które nie zawałyby się przez dwa tysiące lat, choćby Panteon.

Jedno z najnowszych odkryć dotyczy DNA kości słoniowej morsa ze średniowiecznej Europy. Naukowcy odkryli, że część kości słoniowej można przypisać zwierzętom, które prawdopodobnie żyły w wodach arktycznych na północny zachód od Grenlandii, a Grenlandzcy Nordycy albo zabili zwierzęta w tych odległych, mroźnych wodach, albo zaangażowali się w handel kością słoniową z rdzenną ludnością, Eskimosami z Thule. Postęp w zakresie możliwości badania starożytnego DNA odegrał również rolę w niedawnym (choć makabrycznym i niepokojącym) odkryciu na słynnym stanowisku Majów w Chichén Itzá. Rodrigo Barquera i jego koledzy z Max-Planck Institute for Evolutionary Anthropology przeanalizowali DNA wyodrębnione z kości znalezionych w 1967 roku w cysternie podczas wykopalisk. Wcześniej naukowcy wiedzieli, że kości należały do dzieci złożonych w rytualnej ofierze tysiąc lat temu. Nowa technologia pomogła im zidentyfikować, że wszystkie dzieci były płci męskiej, w tym dwa zestawy bliźniąt.

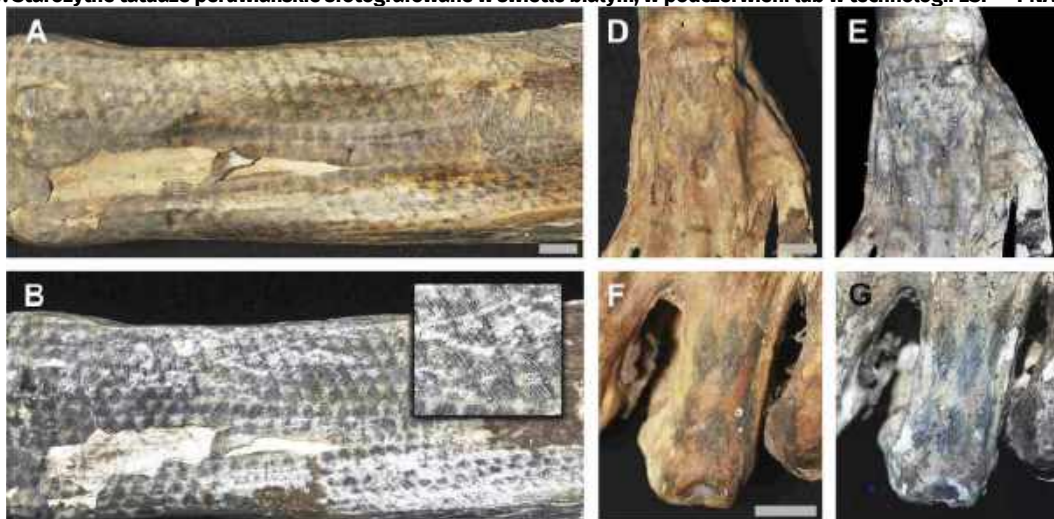
Międzynarodowy zespół naukowców wykorzystał z kolei lasery do ujawnienia tatuaży na mumiach z Peru (7). Dzięki tej technice, opisaniej szczegółowo w badaniu opublikowanym 13 stycznia 2025 r. w czasopiśmie

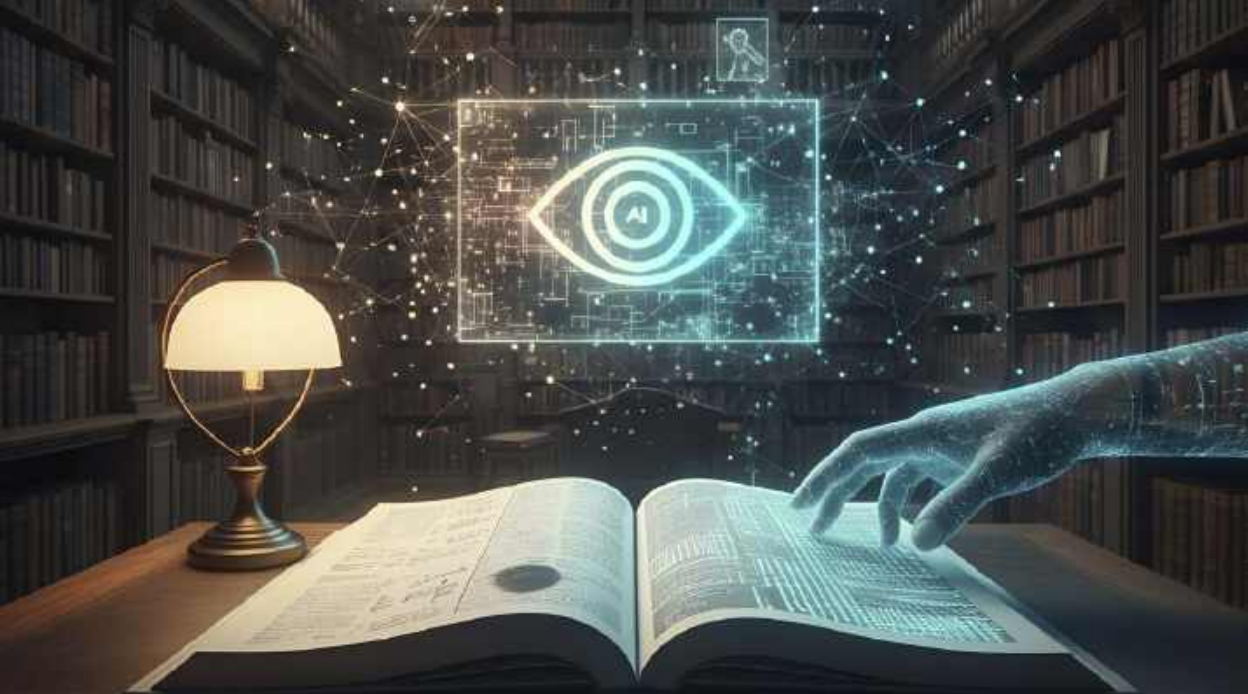
PNAS, naukowcy odkryli skomplikowane wzory i rzucili światło na zawiłane praktyki tatuażu starożytnej kultury Chanca. Ich odkrycia ujawniają wyższy poziom umiejętności artystycznych w prekolumbijskim Peru, niż wcześniej sądzono. Zespół naukowców wykorzystał technikę zwaną fluorescencją stymulowaną laserowo (LSF), która wykorzystuje lasery do ujawnienia szczegółów w tkankach miękkich, aby zbadać tatuaże na peruwiańskich mumiach. Według Science Alert paleontolog od lat wykorzystują LSF do badania szczątków dinozaurów, ale po raz pierwszy technika ta została zastosowana do analizy starożytnych tatuaży na zmumifikowanych szczątkach ludzkich – a wyniki były fantastyczne. Naukowcy wykorzystali lasery, aby skóra bez tatuaży świeciła w ostrym kontraście do skóry z tatuażami, ujawniając nawet delikatne wzory niewidoczne gołym okiem. Niektóre wzory miały linie o grubości od 0,0039 do 0,0079 cala (0,1 do 0,2 milimetra), napisali naukowcy w badaniu. Szczegóły te „odzwierciedlają fakt, że każda kropka tuszu została umieszczona celowo ręcznie z wielką wprawą, tworząc różnorodne, wykwitzne wzory geometryczne i zoomorficzne”.

Ogólny wniosek, który płynie z przedstawionej powyżej, niepełnej listy odkryć archeologicznych i historycznych, które zawdzięczamy najnowszej technice, jest taki, że świat, w którym żyjemy, kryje w sobie znacznie więcej śladów bujnej i ciekawej przeszłości, niż nam się wydawało, gdy nie byliśmy jeszcze wyposażeni w arsenał środków od lidaru po sztuczną inteligencję. Najciekawsze oczywiście jest pytanie – ile rzeczy zostało jeszcze nieodkrytych. ■

Mirosław Usidus

7. Starożytne tatuaże peruwiańskie sfotografowane w świetle białym, w podczerwieni lub w technologii LSF © PNAS





1. Generowanie treści naukowych i AI © AI

W lipcu 2025 r. pojawiła się informacja, że w ramach szeroko zakrojonych badań wykryto ślady sztucznej inteligencji w milionach artykułów naukowych. W miarę jak ChatGPT, Google Gemini i inne narzędzia AI stają się coraz bardziej biegłe w generowaniu tekstów o jakości zbliżonej do ludzkiej, coraz trudniej jest odróżnić takie treści (1). Wzbudziło to obawy społeczności akademickiej.

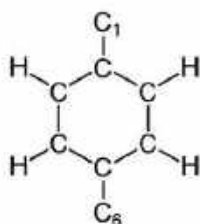
Maszyna w zespole, na dobre i na złe

ODCISKI PALCÓW AI W NAUCE

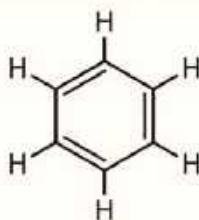
Artykuł na łamach „Science Advances” relacjonuje, jak zespół badaczy z USA i Niemiec przeanalizował ok. piętnastu milionów streszczeń biomedycznych w serwisie PubMed, by ustalić, czy modele typu LLM (ang. skrót od „large language models” – „duże modele językowe”) miały wykrywalny wpływ na wybór konkretnych słów w artykułach naukowych. Ich badania wykazały, że od momentu ich pojawienia się i udostępnienia nastąpił wzrost częstotliwości występowania niektórych stylistycznych form słownych

w literaturze naukowej. Dane te sugerują, że co najmniej 13,5 proc. artykułów opublikowanych w 2024 r. zostało napisanych przy użyciu przetwarzania przez AI w mniejszej lub większej proporcji. Naukowcy wzorowali swoje badania na wcześniejszych badaniach dotyczących zdrowia publicznego w kontekście COVID-19, które pozwoliły wywnioskować wpływ pandemii koronawirusa na śmiertelność przez porównanie nadmiarowej liczby zgonów przed pandemią i po niej. Stosując to samo podejście w nowym badaniu, przeanalizowano wzorce doboru słownictwa przed pojawieniem się modeli LLM i po nim. Naukowcy odkryli, że po wprowadzeniu modeli LLM nastąpiło znaczące przesunięcie w kierunku nadmiernego użycia „stylistycznych i kwiecistych” wyrazów, takich jak „prezentowanie”, „kluczowy” i „zmaganie się”. Autorzy ustalili, że przed 2024 r. 79,2 proc. nadmiernie używanych słów to były rzeczowniki. W 2024 r. nastąpiła wyraźna zmiana. 66 proc. nadmiernie używanych słów stanowiły czasowniki, a 14 proc. – przymiotniki.

ChatGPT



Microsoft Copilot



Google Gemini



2. Próby narysowania cząsteczki benzenu przez różne modele AI

Nie tylko samo pisanie prac naukowych przez AI jest problemem, ale także błędy i halucynacje, których się dopuszcza, co oczywiście w publikacjach naukowych raz jeszcze silniej niż w zwykłych tekstach; np. naukowcy zwrócili niedawno uwagę, że społeczność naukowa chemików powinna zakazać rysowania struktur chemicznych za pomocą generatywnej sztucznej inteligencji. Ostrzegają, że wykorzystywanie AI do tworzenia obrazów ilustrujących treści związane z chemią, takich jak chociażby schematy struktur molekularnych, prowadzi do poważnych błędów i może zaszkodzić nie tylko wizerunkowi autorów prac, ale także edukacji adeptów tej dziedziny. Szereg analizowanych przez badaczy modeli sztucznej inteligencji miało trudności z czymś tak, wydawałoby się, prostym dla chemika, jak narysowanie struktury chemicznej benzenu (2). Co więcej, powtarzające się polecenia mogą generować za każdym razem inne struktury, z nowymi i nieznanymi wcześniej błędami.

Audrey Moores, chemiczka z Uniwersytetu McGill w Kanadzie i Vânia Zuin Zeidler u Uniwersytetu Leuphana w Niemczech, w artykule na łamach „Nature Reviews Chemistry” wezwały do „zakazu stosowania generatywnej AI do przedstawiania cząsteczek”. Moores twierdzi, że ostatecznym bodźcem do napisania artykułu była okładka czasopisma poświęconego zielonej chemii, na której wykorzystano sztuczną inteligencję do wygenerowania struktur chemicznych, które, jak twierdzi, były ewidentnie błędne. „Poprosiłam Copilota o przedstawienie wizualne z zakresu chemii nieorganicznej i o układ okresowy pierwiastków, ale to, co stworzył, to nie był wcale układ okresowy. Próbowalam wiele razy poprawić polecenie, aby uzyskać lepszy wynik, ale program najwyraźniej nie był w stanie tego zrobić”, pisze autorka artykułu. Podkreśla jednak, że problem nie leży w samej sztucznej inteligencji i że AI ma wiele przydatnych zastosowań dla chemików. Problem, jak twierdzi, dotyczy przede wszystkim

i konkretnie wizualizacji struktur chemicznych oraz faktu, że w wielu przypadkach nie stosuje się należytej staranności podczas korzystania ze sztucznej inteligencji w przygotowywaniu publikacji. Chociaż niektóre narzędzia AI w testach były w stanie wygenerować prawidłowe struktury bardzo prostych związków chemicznych, to już nie poradziły sobie chociażby z kofeiną. Żaden z trzech testowanych modeli nie narysował jej prawidłowej struktury. W swoim artykule Moores i Zeidler ostrzegają, że niezdolność sztucznej inteligencji generatywnej w dostępnych obecnie formach do tworzenia dokładnych i użytecznych reprezentacji chemii może mieć nie tylko istotne konsekwencje w dziedzinie szkolnictwa i w nauce chemii.

Magazyn „Nature” zdecydował, że w dającej się przewidzieć przyszłości, z wyjątkiem artykułów dotyczących samej sztucznej inteligencji, „nie będzie publikować żadnych treści, w których zdjęcia, filmy lub ilustracje zostały stworzone w całości lub w części przy użyciu generatywnej sztucznej inteligencji”. Inni wydawcy ogłosili własne wytyczne dotyczące wykorzystania AI w treściach wizualnych. Opublikowane reguły brytyjskiego Królewskiego Towarzystwa Chemicznego dotyczące treści wizualnych dla poszczególnych czasopism stanowią: „W przypadku korzystania z narzędzi AI do tworzenia (grafiki/okładki) autorzy muszą potwierdzić, że narzędzie AI zostało przeszkolone przy użyciu w pełni licencjonowanych zbiorów danych, a warunki licencji na korzystanie z wyników AI pozwalają na komercyjne ponowne wykorzystanie”. Deklaracje i reguły dotyczące użycia AI w publikacjach ogłosiły inne wydawnictwa i ośrodki naukowe.

Dotrzeć do danych i je udostępnić – wszystko dzięki AI

Mówimy o odciskach palców AI i jej wpływie na naukę w kontrowersyjnym lub całkiem negatywnym sensie. Jednak oczywiście sztuczna inteligencja ma wielki potencjał pozytywny i może badaczom



3. Znak graficzny projektu FAIR²

ogromnie pomóc, np. w dziedzinach, które wymagają przetwarzania wielkich zasobów danych, z którymi ludzie ze względu na swoje ograniczenia sobie nie radzą. Ogromne ilości cennych danych badawczych pozostają niewykorzystane, uwięzione w laboratoriach lub utracone w czasie. Wydawnictwo Frontiers zamierza to zmienić dzięki FAIR² Data Management, systemowi opartemu na sztucznej inteligencji, którego zadaniem jest ponowne dotarcie do zbiorów danych w celu weryfikacji i umożliwienia cytowania (3). Łącząc redakcję, legalne wykorzystanie, recenzje i interaktywne wizualizacje w jednej platformie, FAIR² umożliwia naukowcom dzielenie się swoją, zapomnianą lub niedostrzeżoną nalezycie, pracą i zdobycie uznania. Usługa oparta na sztucznej inteligencji została zaprojektowana tak, by dane były zarówno ponownie wykorzystywane, jak i odpowiednio przypisywane, przez połączenie wszystkich niezbędnych etapów, selekcji, kontroli zgodności, formatowania dostosowanego do sztucznej inteligencji, recenzji, interaktywnego portalu, certyfikacji i stałego hostingu, w jeden płynny proces.

FAIR² opiera się na zasadach FAIR (z ang. „Findable, Accessible, Interoperable and Reusable” – „łatwe do znalezienia, dostępne, interoperacyjne i możliwe do ponownego wykorzystania”) z rozszerzoną otwartą strukturą, która gwarantuje, że każdy zbiór danych jest kompatybilny z AI i może być ponownie wykorzystany w sposób etyczny zarówno przez ludzi, jak i maszyny. Praca, która kiedyś wymagała miesięcy wysiłku ludzkiego, od organizowania i weryfikowania zbiorów danych po generowanie metadanych i wyników nadających się do publikacji, jest teraz wykonywana w ciągu kilku minut przez AI Data Steward, mechanizm oparty na technologii Senscience. Naukowcy, którzy przesyłają dane, otrzymują wyniki, certyfikowany pakiet danych, recenzowany i cytowalny artykuł dotyczący danych, interaktywny portal danych z wizualizacjami i czatem

AI oraz certyfikat FAIR². Każdy element zawiera kontrole jakości i podsumowania, które sprawiają, że dane są łatwiejsze do zrozumienia dla zwykłych użytkowników i bardziej kompatybilne w różnych dyscyplinach badawczych.

Pilotażowe zbiory danych projektu FAIR² to m.in. opracowania na temat właściwości wariantów SARS-CoV-2, dane z przedklinicznych badań MRI urazów mózgu, prace na temat różnorodności biologicznej atoli Indo-Pacyfiku.

Medycyna sięga odważnie

Przy wszystkich podanych wyżej wątpliwościach i zastrzeżeniach jednak świat badań naukowych wydaje się bardziej zdecydowany we wdrażaniu rozwiązań opartych na AI niż biznes; np. od razu wzbudził duże zainteresowanie model sztucznej inteligencji (AI) opracowany w Pacific Northwest National Laboratory (PNNL), który może identyfikować wzorce na obrazach materiałów z mikroskopu elektronowego bez interwencji człowieka, umożliwiając dokładniejsze i spójniejsze badania w materiałoznawstwie. Zazwyczaj, by wyszkolić model sztucznej inteligencji w celu badania zjawiska takiego jak uszkodzenie radiacyjne, naukowcy skrupulatnie musieliby tworzyć ręcznie oznakowany zbiór danych, ręcznie śledząc obszary uszkodzone przez promieniowanie na obrazach z mikroskopu elektronowego. Ten ręcznie oznakowany zbiór danych byłby następnie wykorzystywany do szkolenia modelu, który identyfikowałby wspólne cechy regionów zidentyfikowanych przez człowieka i starał się zidentyfikować podobne regiony na nieoznakowanych obrazach. Zamiast tego PNNL zastosowało nienadzorowany model, który potrafi analizować dane, obrazy z mikroskopu elektronowego, bez angażowania ludzi. Naukowcy zastosowali nowy model do badania uszkodzeń spowodowanych promieniowaniem w strukturach materiałów stosowanych w reaktorach jądrowych. Potrafi dokładnie wyłapać zdegradowane obszary i posortować obraz na grupy reprezentujące różne poziomy uszkodzeń radiacyjnych.

Najbardziej chyba odważne w sięganiu po narzędzia sztucznej inteligencji są dziedziny biologii i medycyny. Uczeń wykorzystują sztuczną inteligencję np. do rozpoznawania różnic na poziomie komórkowym w mózgach mężczyzn i kobiet. Według publikacji, która ukazała się w „Scientific Reports” w maju 2024 r., modele AI zidentyfikowały płeć w skanach MRI z dokładnością 92...98 proc. Różnice stwierdzono w istocie białej mózgu, kluczowej dla komunikacji między obszarami mózgu. Uważa się, że poznanie

tych różnic może ulepszyć narzędzia diagnostyczne i metody leczenia zaburzeń mózgu, stwardnienia rozsianego i autyzmu. Wcześniejsze badania mikrostruktury mózgu w dużej mierze opierały się na modelach zwierzęcych i próbkach tkanek ludzkich. Wzbudzały wątpliwości ze względu na oparcie się na analizach statystycznych „ręcznie rysowanych” regionów zainteresowania, co oznacza, że badacze musieli podejmować subiektywne decyzje co do kształtu, rozmiaru i lokalizacji wybranych regionów.

Model sztucznej inteligencji opracowany wspólnie przez Google i Uniwersytet Yale doprowadził do powstania przełomowej hipotezy dotyczącej interakcji komórek nowotworowych z ludzkim układem odpornościowym. Hipoteza została sformułowana na podstawie modelu podstawowego o 27 miliardach parametrów, zwanego Cell2Sentence-Scale 27B (C2S-Scale), opracowanego przez naukowców z Google DeepMind i Uniwersytetu Yale. Model C2S-Scale, oparty na otwartym modelu sztucznej inteligencji Gemma firmy Google, jest przeznaczony do analizy pojedynczych komórek, umożliwiając naukowcom przewidywanie zachowania komórek nowotworowych w organizmach żywych. Odkrycie opiera się na wcześniejszych badaniach Google, które wykazały, że biologiczne modele sztucznej inteligencji podlegają jasnym prawom skalowania – większe modele wykazują wyższy poziom rozumowania warunkowego, podobny do zachowania systemów sztucznej inteligencji opartych na języku naturalnym. Według Google model C2S-Scale 27B może interpretować „język” poszczególnych żywych komórek, co pozwala mu przekształcić trudne do wykrycia „zimne” guzy w „gorące”. Proces ten sprawia, że komórki nowotworowe są bardziej widoczne dla układu odpornościowego i lepiej reagują na terapię. Wyjaśniając, w jaki sposób C2S-Scale został przeszkolony do rozumowania w złożonych warunkach biologicznych, Google stwierdził, że jego naukowcy zaprojektowali tak zwany „wirtualny ekran podwójnego kontekstu”, symulujący działanie ponad czterech tysięcy leków na rzeczywistych próbkach guzów pacjentów i izolowanych danych linii komórkowych bez kontekstu immunologicznego. Poproszony o zidentyfikowanie leków, które mogłyby selektywnie wzmocnić prezentację antygenów w pierwszym kontekście, model wskazał kilku kandydatów, tylko 10...30 proc. z nich było wcześniej znanych jako skuteczne w leczeniu raka. Pozostałe przewidywania nie miały wcześniej znanego związku z ekranem ani immunoterapią raka. Przewidywania te zostały następnie zweryfikowane w zastosowaniach klinicznych.

Zarówno Gemma, jak i C2S-Scale 27B są publicznie dostępne na Hugging Face i GitHub. Firma Google opublikowała również pracę naukową na bioRxiv. Naukowcy ostrzegają jednak, że wszystkie prognozy będą wymagały recenzji i walidacji klinicznej przed przyjęciem do użytku terapeutycznego.

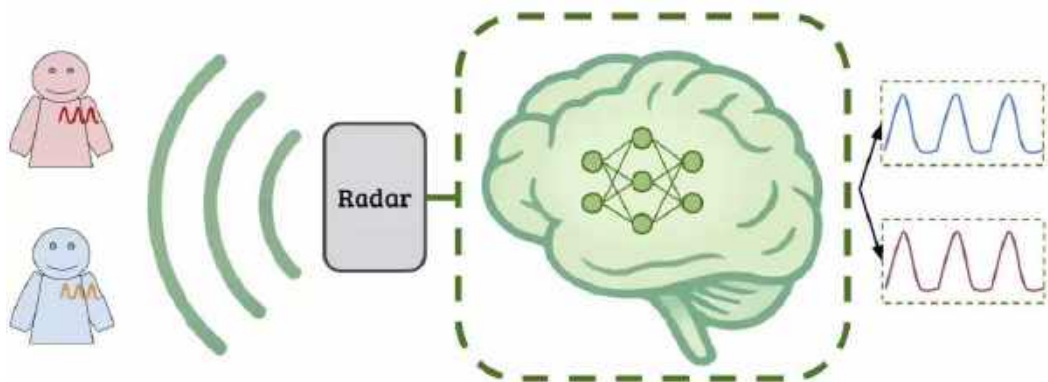
Tekst, na którym pracuje generatywna AI, to sekwencja znaków, które mogą być łączone w złożony sposób w celu uzyskania niezliczonej liczby złożonych znaczeń. Podobnie podstawą życia jest zaledwie kilku podstawowych znaków (tylko cztery dla DNA i dwadzieścia dla białek), a ich niezliczone kombinacje pozwalają uzyskać całą różnorodność biologiczną, jaką znamy. Jeśli jesteśmy zbudowani z sekwencji, a modele językowe mogą być zdolne do analizowania sekwencji, dlaczego nie wykorzystywać modeli językowych z sekwencjami DNA i białek? Zatem AlphaFold2, dzięki wykorzystaniu modelu językowego wyszkolonego na sekwencjach białkowych, rekonstruuje struktury białek na podstawie jedynie sekwencji znaków. Model nauczył się reprezentacji białek i wzorców obecnych w ich sekwencji (sekwencje te, podobnie jak sekwencje tekstowe, nie są przypadkowe, ale mają znaczenie funkcjonalne i własną semantykę). Reprezentacja ta pozwala nam następnie przewidzieć strukturę i funkcję białka lub inne parametry. Duże modele języka białek uczą się wystarczającej ilości informacji, aby umożliwić dokładne przewidywanie struktury białek na poziomie atomowym. Jeśli model rozumie, które części sekwencji białka grają określoną rolę funkcjonalną lub są odpowiedzialne za określone zachowanie, może następnie wykorzystać je we wnioskowaniu. Na przykład możemy poprosić model o wygenerowanie białka zdolnego do cięcia pierścieni aromatycznych. Mogłoby to być enzym, który mógłby być sztucznie produkowany i wykorzystywany do oczyszczania wody zanieczyszczonej ropą naftową. Wykorzystując duży model językowy, naukowcy stworzyli funkcjonalne białka o sekwencjach, które nie istnieją w naturze z pożądanymi właściwościami biofizycznymi. Połączenie sztucznej inteligencji i techniki „edycji” genów CRISPR-Cas pozwala nam wyobrazić sobie przyszłość, w której edycja genów może być wykorzystywana do leczenia niemal każdej choroby. Zakłada się, że w przyszłości lekarz podczas diagnozy będzie sekwencjonował genom, zauważy, jakie mutacje leżą u podstaw choroby, a edycja genów służyć ma leczeniu pacjenta. LLM posłuży do identyfikacji nowych mutacji i nowych środków edycji genów. W przyszłości LLM mają oferować także arsenał technik (szybkie projektowanie,

dostrajanie itp.) do generowania białek o pożądanych funkcjach, które nie istnieją w naturze.

W Cold Spring Harbor Laboratory (CSHL) opracowano m.in. model AI mózgu muszki owocowej w celu badania, w jaki sposób widzenie świata kieruje zachowaniem. Wyciszając genetycznie określone neurony wzrokowe i obserwując zmiany w zachowaniu, wyszkolili sztuczną inteligencję do dokładnego przewidywania aktywności neuronalnej i zachowania. Benjamin Cowley i jego zespół udoskonalili swój model sztucznej inteligencji za pomocą opracowanej przez siebie techniki zwanej „treningiem nokautującym”, który polega na zaburzaniu sieci podczas szkolenia, by dopasować ją do zaburzeń rzeczywistych neuronów podczas eksperymentów. Najpierw nagrali zaloty samca muszki owocowej. Następnie genetycznie wyciszili określone typy neuronów wzrokowych u samca muszki i wyszkolili swoją AI, aby wykrywała wszelkie zmiany w zachowaniu, prowadząc w efekcie do dokładnego przewidywania, jak zachowa się prawdziwy samiec muszki na widok samicy. „Możemy obliczeniowo przewidzieć aktywność neuronową i zapytać, w jaki sposób określone neurony przyczyniają się do zachowania”, wyjaśnia Cowley w komunikacie badawczym. Wyobraźnia podpowiada, że konsekwencją dalszych badań może być zdolność sztucznej inteligencji do przewidywania ludzkich zachowań. Nie tak szybko. Mózg muszki owocowej zawiera około 100 tys. neuronów. Ludzki mózg ma ich prawie 100 miliardów.

Naukownicy z uniwersytetów w Kalifornii i firm, takich jak Precision Neuroscience z siedzibą w Nowym Jorku, poczynili postępy w generowaniu naturalnej mowy poprzez połączenie implantów mózgowych i sztucznej inteligencji. Edward Chang, neurochirurg z Uniwersytetu Kalifornijskiego w San Francisco, i jego współpracownicy,

w tym z Uniwersytetu Kalifornijskiego w Berkeley, opublikowali w czasopiśmie „Nature Neuroscience” artykuł opisujący ich pracę z kobietą cierpiącą na paraliż kończyn i tułowia, która nie była w stanie mówić od 18 lat po udarze mózgu. Kobieta ta szkoliliła sieć neuronową głębokiej nauki, próbując w milczeniu wypowiedzieć zdania złożone z 1024 różnych słów. Dźwięk jej głosu został stworzony poprzez przesyłanie jej danych neuronowych do wspólnego modelu syntezy mowy i dekodowania tekstu. Technika ta zmniejszyła opóźnienie między sygnałami mózgowymi pacjentki a wynikowym dźwiękiem z ośmiu sekund do jednej sekundy. Jest to znacznie bliższe 100...200-milisekundowej przerwie w normalnej mowie. Średnia prędkość dekodowania systemu wynosiła 47,5 słowa na minutę, czyli około jednej trzeciej tempa normalnej rozmowy. Jeśli technika ta okaże się skuteczna, naukowcy mają nadzieję, że będzie można ją rozszerzyć, aby pomóc osobom mającym trudności z wysławianiem się z powodu takich schorzeń. Potencjał neuroprotezy głosowej wzbudził też zainteresowanie biznesu. Firma Precision Neuroscience twierdzi, że rejestruje sygnały mózgowie o wyższej rozdzielczości niż naukowcy akademicy, ponieważ elektrody jej implantów są gęściej rozmieszczone. Precision otrzymała zgodę organów regulacyjnych na pozostawienie swoich czujników wszczepionych na okres do 30 dni. Umożliwi to naukowcom szkolenie systemu przy użyciu tego, co w ciągu roku może stać się „największym repozytorium danych neuronowych o wysokiej rozdzielczości istniejącym na Ziemi”, co podkreśla dyrektor generalny Michael Mager. Kolejnym krokiem będzie „zminiaturyzowanie komponentów i umieszczenie ich w hermetycznie zamkniętych opakowaniach, które są biokompatybilne, dzięki czemu można je wszczepić w ciało na zawsze”.



4. Sztuczna inteligencja przekształca surowe sygnały radarowe w wykresy funkcji życiowych

Naukowcy opracowują techniki wykorzystujące sztuczną inteligencję, które umożliwiają monitorowanie parametrów życiowych człowieka na odległość. Jedną z tych technologii jest radar. Bezkontaktowe monitorowanie stanu zdrowia może być wygodniejsze i łatwiejsze w użyciu niż tradycyjne metody, szczególnie dla osób, które chcą monitorować swoje parametry życiowe w domu.

Radar, jak wiadomo, działa przez wysyłanie fal elektromagnetycznych, oczekiwanie na ich odbicie od obiektów znajdujących się na ich drodze oraz wykrywanie ich, gdy powracają do urządzenia. Radar może być wykorzystywany do monitorowania funkcji życiowych, takich jak oddychanie i tętno (4). Każdy oddech lub uderzenie serca powoduje nieznaczne poruszenie klatki piersiowej. Dzisiejsze radary są wystarczająco czułe, aby wykrywać te niewielkie ruchy, nawet z drugiego końca pokoju. Istnieją też inne technologie, które można wykorzystać do zdalnego pomiaru stanu zdrowia. Techniki oparte na kamerach mogą wykorzystywać światło podczerwone do monitorowania zmian na powierzchni skóry w taki sam sposób, jak pulsoksymetry, ujawniając informacje o aktywności serca. Systemy widzenia komputerowego mogą również monitorować oddychanie i inne czynności, takie jak sen, a także wykrywać upadki. Takie rozwiązania mają swoje wady – nie są zbyt

przydatne w przypadku zasłonięcia lub zniekształceń barwnych. Obrazy radarowe mają z kolei niską rozdzielczość. Wszystko więc ma swoje plusy i minusy. Stworzono więc rodzaj „mózgu”, który sprawia, że radar jest inteligentniejszy. Ten mózg, nazwany mm-MuRe, to sieć neuronowa, która uczy się bezpośrednio z surowych sygnałów radarowych i szacuje ruchy klatki piersiowej. Podejście to nazywa się uczeniem end-to-end. Oznacza to, że w przeciwieństwie do innych technik radarowych w połączeniu ze sztuczną inteligencją, sieć sama ustala, jak ignorować szumy i skupiać się tylko na ważnych sygnałach. Wykorzystano sztuczną inteligencję do przekształcenia surowych, nieprzetworzonych sygnałów radarowych w wykresy parametrów życiowych jednej lub dwóch osób. Ulepszenie AI nie tylko zapewnia dokładniejsze wyniki, ale także działa szybciej niż tradycyjne metody. Obsługuje wiele osób jednocześnie i dostosowuje się do nowych sytuacji, nawet tych, których nie widziało podczas szkolenia, na przykład, gdy ludzie siedzą na różnych wysokościach, jadą samochodem lub stoją blisko siebie.

Maszynowy psycholog

Także psychologowie sięgają po sztuczną inteligencję, aby odkrywać ukryte sygnały psychologiczne w mowie, od doboru słów po ton i tempo.

5. Wizualizacja narzędzia AI analizującego mowę ludzką





6. Will Downs prezentuje swój model w amerykańskiej telewizji

Sygnały te mogą ujawniać cechy osobowości, a nawet wczesne oznaki zaburzeń zdrowia psychicznego, ale lekarze mogą je przeoczyć. Sztuczna inteligencja może zapewnić szybszą i dokładniejszą analizę, chociaż naukowcy ostrzegają przed stronniczością, jeśli modele nie są szkolone na podstawie zróżnicowanych danych. Dzięki starannemu opracowaniu sztuczna inteligencja może zmienić ocenę psychologiczną, oferując potężne nowe narzędzia wspierające lekarzy.

Narzędzia sztucznej inteligencji (AI) wyszkolone do wykrywania charakterystycznych oznak w mowie (5) mogą zrewolucjonizować dziedzinę diagnoz psychologicznych, sądzi psycholog z WashU Josh Oltmanns. „Psychologowie są ludźmi, a ludzie są omylni, więc nawet dobry klinicysta nie zawsze może wychwycić ważne sygnały werbalne”, powiedział Oltmanns w Wydziale Sztuki i Nauk Ścisłych Uniwersytetu Waszyngtona w St. Louis. „Ale odpowiednio wyszkolony model komputerowy wychwyci te sygnały”. Teoretycznie psycholog mógłby poprosić pacjenta o opisanie swojego życia i problemów, co jest standardową częścią wstępnej oceny. Oprócz wykorzystania własnej wiedzy klinicznej, psycholog mógłby wprowadzić tę rozmowę do programu zaprojektowanego do wykrywania cech osobowości i oznak problemów ze zdrowiem psychicznym. „Program komputerowy mógłby pomóc w potwierdzeniu ich obserwacji lub ostrzec ich o czymś, co mogło umknąć ich uwadze”. Oltmanns współpracuje

ze swoimi współpracownikami nad opracowaniem narzędzi AI, które pomogą psychologom odkrywać te ukryte sygnały w języku. Niedawno opisał potencjał takich narzędzi w czasopiśmie „Advances in Methods and Practices in Psychological Science”.

Ponieważ każda rozmowa zawiera tak wiele potencjalnych informacji, psychologowie od dawna pragnęli pomocy komputerów w analizowaniu mowy. Ponad dwadzieścia lat temu naukowcy opracowali program Linguistic Inquiry and Word Count, który pozwala oceniać różne aspekty psychologiczne na podstawie tekstu pisanego. Narzędzia te były z czasem udoskonalane, ale pojawienie się sztucznej inteligencji otwiera nowe możliwości. Programy AI mogą być znacznie szybsze, dokładniejsze i bardziej precyzyjne niż poprzednie modele komputerowe. Oltmanns ostrzega jednak, że AI wiąże się również z ryzykiem. „Często jest ona szkolona na podstawie informacji znalezionych w Internecie, co oznacza, że może być stronnicza”, podkreśla.

Nadzieje i sceptycyzm

Medycyna, jak napisaliśmy, jest stosunkowo odważna w sięganiu po AI. Inne dziedziny też się przełamują. Na przykład naukowcy politechniki ETH w Zürichu z powodzeniem zautomatyzowali proces modelowania turbulencji w płynach przez połączenie mechaniki płynów ze sztuczną inteligencją. Ich podejście opiera się na połączeniu algorytmów wzmocnionego uczenia maszynowego z symulacjami

przepływów turbulentnych przeprowadzonymi w superkomputerze Piz Daint ze Szwajcarskiego Narodowego Centrum Superkomputerowego. Według publikacji na ten temat, która ukazała się w periodyku „Nature Machine Intelligence”, opracowana w tym projekcie symulacja przepływów turbulentnych ma kluczowe znaczenie w wielu dziedzinach nauki i techniki, od projektowania samochodów i prognozowania pogody po implanty zastawek sercowych i wyjaśnianie procesów narodzin galaktyk.

Model AI okazał się niezwykle skuteczny w identyfikowaniu i śledzeniu tropikalnych zjawisk atmosferycznych, skupisk chmur i wiatrów, które mogą przekształcić się w potężne huragany. W ramach pierwszej tego rodzaju techniki Will Downs z uniwersytetu w Miami opracował model sztucznej inteligencji, który śledzi tropikalne fale, skupiska chmur i wiatru, przemieszczające się przez Ocean Atlantycki (6). Niektóre z tych fal przekształcają się w burze tropikalne i huragany. Istniejące algorytmy miały trudności z identyfikacją i śledzeniem tych fal, model oparty na sztucznej inteligencji Downsa okazał się bardziej skuteczny. Aby go stworzyć, wykorzystał on historyczne dane pogodowe zebrane przez meteorologów z Oddziału Analiz i Prognoz Tropikalnych (TAFB) Narodowego Centrum Huraganów (NHC). Dane te, które meteorolodzy zapisali w szczegółowych raportach, obejmowały lokalizacje tropikalnych fal wschodnich na Morzu Karaibskim w ciągu ostatnich kilku dekad. Następnie połączył te historyczne obserwacje z danymi z ponownej analizy przeszłych warunków pogodowych i klimatycznych, szkoląc swój model sztucznej inteligencji, aby dokładnie wykrywał nie tylko tropikalne fale wschodnie, ale także ważne zjawiska pogodowe.

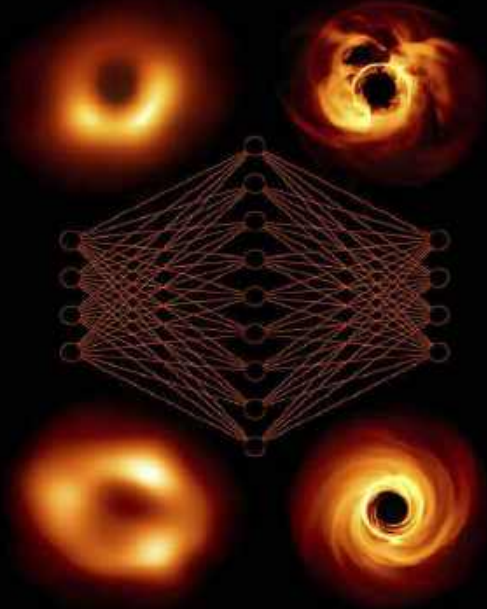
Z kolei to, jak daleko zaszła AI w inżynierii, pokazuje przykład opracowanego przy jej wykorzystaniu nowego silnika raketowego typu aerospike, spalającego tlen i naftę, zdolnego do wytworzenia

ciągu 5000 N (7). Został zaprojektowany od początku do końca przy użyciu zaawansowanego modelu inżynierii obliczeniowej Large Computational Engineering Model przez firmę o nazwie LEAP 71 z siedzibą w Dubaju.

W rakietach konwencjonalnych stosuje się znany dzwon, który kieruje i rozpręga gorące gazy wychodzące z silnika po przejściu przez dyszę Venturiego. Konstrukcja ta działa bardzo dobrze, ale ma jedną poważną wadę. Krzywizna dzwonu musi być specjalnie zaprojektowana, aby działała wydajnie na określonej wysokości, więc rakieta, która działa bardzo dobrze podczas startu, będzie działać gorzej w miarę wznoszenia się w atmosferze i spadku ciśnienia powietrza. Dlatego silniki raketowe drugiego i trzeciego stopnia różnią się od silników pierwszego stopnia. Inżynierowie chcieliby mieć silnik, który samoczynnie dostosowuje się do zmian ciśnienia powietrza. Silnik aerospike osiąga to poprzez nadanie mu kształtu kolca lub wtyczki z krzywizną podobną do wewnętrznej strony dzwonu raketowego. Gdy spaliny przepływają z silnika, krzywizna działa jak jedna strona dzwonu, a otaczające powietrze jak zewnętrzna krzywizna. Wraz ze zmianą ciśnienia powietrza zmienia się również kształt wirtualnego dzwonu. Od lat 50. XX wieku opracowano wiele silników aerospike, a jeden z nich wzbili się w powietrze, ale przekształcenie obiecującego pomysłu w praktyczny silnik kosmiczny to długa droga. Wykorzystany przez LEAP 71 model jest zaprogramowany i wyszkolony przez ekspertów z dziedziny lotnictwa i kosmonautyki, na podstawie określonego zestawu parametrów wejściowych tworzy projekt spełniający te parametry, wnioskując o fizycznych interakcjach różnych czynników, w tym zachowań termicznych i przewidywanej wydajności. Wyniki tych obliczeń są następnie wprowadzane z powrotem do modelu sztucznej inteligencji w celu jego dostrojenia, ponieważ przedstawia on obliczone parametry wydajności,

7. Aerospike LEAP 71 w fazach odpalania





8. Model sztucznej inteligencji pomógł dopracować szczegóły na pierwszych w historii zdjęciach czarnych dziur

geometrię silnika, parametry procesu produkcyjnego i inne szczegóły. Według firmy Noyron potrafi samodzielnie zaprojektować nowy aerospike w ciągu około trzech tygodni. Następnie powstał prototyp przy użyciu przemysłowej technologii druku 3D Selective Laser Melting. Pod koniec 2024 r. pomyślnie przeszedł pierwsze testy, podczas których został poddany działaniu temperatury gazu wynoszącej 3500°C. Wyniki tego testu są, jak uważa firma, bardzo zachęcające.

Międzynarodowy zespół uczonych podjął próbę wykorzystania potencjału sztucznej inteligencji w celu uzyskania dodatkowych informacji na temat Sagittarius A* na podstawie danych zebranych przez teleskop Event Horizon (EHT). Składa się on z szeregu połączonych instrumentów rozszaniach po całym świecie, które działają razem. EHT wykorzystuje długie fale elektromagnetyczne, o długości do milimetra, do pomiaru promienia fotonów otaczających czarną dziurę. Jednak technika ta jest bardzo podatna na zakłócenia spowodowane parą wodną w atmosferze ziemskiej. Oznacza to, że naukowcom może być trudno zrozumieć informacje zebrane przez instrumenty. „Bardzo trudno jest pracować z danymi z teleskopu Event Horizon Telescope”, wyjaśniał „Live Science” Michael Janssen, astrofizyk z Uniwersytetu Radboud w Holandii, współautor badania. „Sieć neuronowa idealnie nadaje się do rozwiązania tego problemu”. Janssen i jego zespół przeszkolili model sztucznej inteligencji na danych EHT, które wcześniej zostały odrzucone jako zbyt zakłócone. Innymi słowy, zakłócenia atmosferyczne były zbyt duże, aby rozszyfrować informacje przy użyciu klasycznych technik. Dzięki

technice sztucznej inteligencji wygenerowali nowy obraz struktury Sagittarius A*, a ich zdjęcie ujawniło kilka nowych cech. Na przykład czarna dziura wydaje się obracać z „prawie maksymalną prędkością”, jak stwierdzili naukowcy w oświadczeniu, a jej oś obrotu również wydaje się skierowana w stronę Ziemi. Wyniki ich badań zostały opublikowane w „Astrophysics”.

Jednak niektórzy eksperci podchodzą do wyników sceptycznie. Model, częściowo oparty na złożonych danych teleskopowych, które wcześniej uważano za zbyt zakłócone, by były użyteczne, ma na celu stworzenie najbardziej szczegółowych obrazów czarnych dziur w historii (8). „Bardzo sympatyzuję z tym, co robią i jestem tym zainteresowany”, powiedział serwisowi „Live Science” Reinhard Genzel, astrofizyk z Instytutu Fizyki Pozaziemskiej im. Maxa Plancka w Niemczech i jeden z laureatów Nagrody Nobla w dziedzinie fizyki w 2020 roku. „Ale sztuczna inteligencja nie jest cudownym lekarstwem”. Według Genzela stosunkowo niska jakość danych wprowadzonych do modelu mogła spowodować nieoczekiwane odchylenia. W rezultacie nowy obraz może być zniekształcony i nie należy go traktować jako ostateczny – uważa.

Ta ostatnia wypowiedź jest charakterystyczna, bo ilustruje stanowisko wielu ludzi ze świata nauki wobec AI. To połączenie wątpliwości i sceptycyzmu z zainteresowaniem zabarwionym nadzieją, że sztuczna inteligencja, pomimo swoich obecnych wad i niedoskonałości, ostatecznie zrewolucjonizuje wiele dziedzin badań. ■

Mirosław Usidus



1. AlphaProof Google DeepMind

Podczas Międzynarodowej Olimpiady Matematycznej (IMO) w 2024 r. jeden z uczestników osiągnął tak dobry wynik, że otrzymałby srebrny medal, gdyby... nie był to system sztucznej inteligencji. Po raz pierwszy w historii imprezy AI osiągnęła wynik na poziomie medalowym.

Rewolucja w krainie królowej nauk?

PRZYSZŁOŚĆ MATEMATYKI

Sztuczna inteligencja to AlphaProof, program opracowany przez Google DeepMind (1), który uczy się rozwiązywać złożone problemy matematyczne. Osiągnięcie podczas IMO było imponujące, ale tym, co naprawdę wyróżnia AlphaProof, jest, według publikacji na łamach „Nature”, jego zdolność do wykrywania i poprawiania błędów. Chociaż duże modele językowe (LLM) potrafią rozwiązywać problemy matematyczne, często nie są w stanie zagwarantować dokładności swoich rozwiązań. W ich rozumowaniu mogą czaić się błędy. Odpowiedzi AlphaProof mają natomiast

być w 100 proc. poprawne. Gwarantować ma to specjalistyczne oprogramowanie o nazwie Lean (pierwotnie opracowane przez Microsoft Research), które działa jak drobiazgowy nauczyciel weryfikujący każdą logiczną operację algorytmu. „Jeśli porozmawiasz z matematykami i zapytasz ich o wyniki [AlphaProof] w IMO, spotkasz się z różnymi reakcjami. Myślę, że większość z nich powie, że są to dość trudne zadania na poziomie szkoły średniej, ale być może niektórzy inni matematycy uznają je za trywialne”, mówi Thomas Hubert, inżynier badawczy w DeepMind.

Szkolenie tego potężnego systemu podzielono na trzy etapy. Najpierw naukowcy zaaplikowali AlphaProof około 300 miliardów tokenów ogólnego kodu i tekstu matematycznego, aby zapewnić mu szerokie rozumienie pojęć takich jak logika, język matematyczny i struktura programowania. Następnie szkolono go na 300 tysiącach dowodów matematycznych napisanych przez ekspertów, dostępnych w środowisku Lean. Ostatni etap polegał na tym, że system nauczył się samodzielnie rozwiązywać problemy. Otrzymał ogromne zadanie domowe z 80 milionami

formalnych problemów matematycznych. Wykorzystując naukę maszynową ze wzmocnieniem, opartą na metodzie prób i błędów, AlphaProof był nagradzany za każde udane dowodzenie. Rozwiązując zadania matematyczne na tak ogromną skalę, system nauczył się nowych i złożonych strategii rozumowania, wykraczających poza kopiowanie ludzkich przykładów. W przypadku najtrudniejszych zadań AlphaProof wykorzystał opracowaną przez naukowców technikę o nazwie Test-Time RL (TTRL), która tworzy i rozwiązuje miliony uproszczonych wersji danego zadania, aż znajdzie rozwiązanie.

Wprawki w rozwiązywaniu największych problemów

Czy rewolucja w dziedzinie sztucznej inteligencji zmienia oblicze matematyki? Niektórzy wybitni matematycy uważają, że tak, wskazując imponujący wzrost możliwości i potencjał nowych modeli. W czerwcu 2025 r. około stu czołowych matematyków z całego świata zebrało się na Uniwersytecie Cambridge na konferencji skupiającej się na odpowiedzi na pytanie, czy komputery mogą pomóc matematikom w sprawdzaniu poprawności ich dowodów. Proces ten, znany jako formalizacja, niekoniecznie musi obejmować sztuczną inteligencję i rzeczywiście podczas podobnego wcześniejszego spotkania, które odbyło się w Cambridge w 2017 r., nie było jeszcze mowy o sztucznej inteligencji. Osiem lat później sztuczna inteligencja poczyniła ogromne postępy, co wywołało nową falę zainteresowania kwestią wpływu sztucznej inteligencji na matematykę. Naukowcy analizują różne kwestie, od automatycznego tłumaczenia dowodów napisanych przez ludzi na formalny język sprawdzalny komputerowo, po samodzielne konstruowanie dowodów.

Dwa z wykładów zostały zorganizowane przez Google DeepMind, który mówił oczywiście o AlphaProof, Morph Labs, amerykański start-up zajmujący się sztuczną inteligencją, zaprezentował narzędzie AI o nazwie Trinity, które zostało zaprojektowane do całkowitego automatycznego tłumaczenia dowodów, zaczynając od odręcznych zapisów matematycznych i generując w pełni sformalizowane i sprawdzone dowody w języku Lean. Bhavik Mehta z Imperial College London, który współpracował z Morph Labs, demonstrował przykład Trinity dowodzącej twierdzenia związanego ze znaną w świecie matematyki hipotezą ABC, zaliczaną do największych problemów matematycznych. Chociaż dowód ten stanowił jedynie niewielki element całego dowodu potrzebnego do potwierdzenia hipotezy ABC,

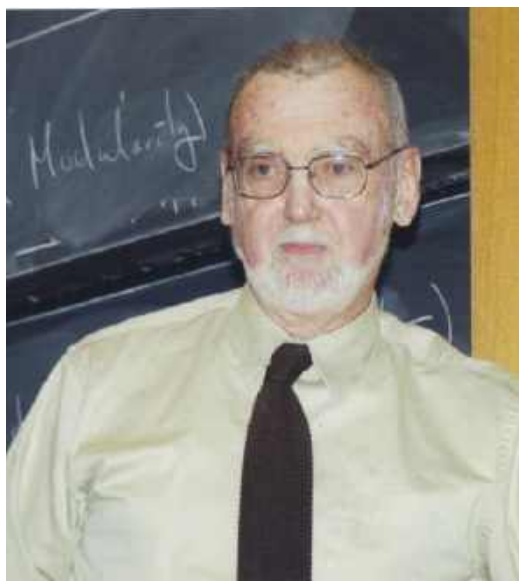
a Trinity wymagało nieco bardziej szczegółowej wersji ręcznie sporządzonego dowodu niż oryginalny artykuł, wiele osób było zaskoczonych tym, jak wiele poprawnych kodów matematycznych zostało wygenerowanych przez to narzędzie. Christian Szegedy z Morph Labs twierdzi, że gdy narzędzie będzie w pełni gotowe do działania, szybko osiągnie postępy. „Gdy powstanie pętla sprzężenia zwrotnego i nie będziemy potrzebowali takiego wsparcia, jakiego wymagało to twierdzenie... wtedy stanie się to w zasadzie reakcją łańcuchową i będziemy mogli zająć się całą matematyką naraz”, mówił.

Jednak nie wszyscy matematycy zgodzili się, że wyniki Morph Labs były aż tak imponujące. Rodrigo Ochigame z uniwersytetu w Lejdzie w Holandii ostrzegał, że nie znamy wszystkich szczegółów tej pracy. „Opublikowali tylko jeden, prawdopodobnie starannie wybrany wynik działania swojego systemu, nie pisząc artykułu ani nie dzieląc się podstawowymi informacjami na temat swoich metod. Nie powiedzieli nawet, czy przetestowali swój system na innych twierdzeniach”, mówił. „Kiedy publiczność zapytała o ilość obliczeń wykorzystanych przez model, wielokrotnie odmawiali odpowiedzi, co utrudnia ocenę znaczenia ich wyników”. Nadal istnieje więc sceptycyzm co do przydatności narzędzi AI.

Czy AI pomoże w dotarciu do teorii wszystkiego z pogranicza matematyki i fizyki?

W ciągu ostatniego roku naukowcy spełnili swoje wieloletnie marzenie, przedstawiając dowód geometrycznej hipotezy Langlandsa, kluczowego elementu grupy powiązanych ze sobą problemów zwanych programem Langlandsa. Dowód ten, ich zdaniem, potwierdza złożony program Langlandsa, który został przez niektórych okrzyknięty wielką zunifikowaną teorią matematyki. Pozostaje on jednak wciąż w dużej mierze niesprawdzony. W środowisku entuzjastów narzędzi AI w matematyce rośnie nadzieja, że dzięki nim w końcu się to uda.

Program Langlandsa sięga 1967 roku, do pracy młodego wówczas kanadyjskiego uczonego Roberta Langlandsa (2), który przedstawił swoją wizję w ręcznie napisanym liście do sławnego matematyka André Weil'a. Przez dziesięciolecia program przyciągał coraz większą uwagę matematyków, którzy podziwiali jego wszechstronność. To właśnie ta cecha skłoniła Edwarda Frenkela z Uniwersytetu Kalifornijskiego w Berkeley, który wniósł kluczowy wkład w część geometryczną, do nazwania jej wielką



2. Robert Langlands

zunifikowaną teorią matematyki. Celem Langlandsa było połączenie dwóch odrębnych głównych gałęzi matematyki, teorii liczb i analizy harmonicznej (nauki o tym, jak skomplikowane sygnały lub funkcje rozkładają się na proste fale). Szczególnym przypadkiem programu Langlandsa jest opublikowany w 1995 roku przez Andrew Wilesa dowód ostatniego twierdzenia Fermata – że żadne trzy dodatnie liczby całkowite a , b i c nie spełniają równania $a^n + b^n = c^n$, jeśli n jest liczbą całkowitą większą niż 2. Geometryczna hipoteza Langlandsa została po raz pierwszy sformułowana w latach 80. przez Vladimira Drinfelda. Podobnie jak pierwotna lub arytmetyczna forma hipotezy Langlandsa, hipoteza geometryczna również tworzy pewien rodzaj powiązania – sugeruje zgodność między dwoma różnymi zbiorami obiektów matematycznych. Chociaż dziedziny powiązane arytmetyczną formą Langlandsa są odrębnymi „światami” matematycznymi, różnice między dwiema stronami hipotezy geometrycznej nie są tak wyraźne. Obie dotyczą właściwości powierzchni Riemanna, które są „złożonymi rozmaitościami”, strukturami o współrzędnych będących liczbami zespolonymi (z częścią rzeczywistą i urojoną). Wielu matematyków podejrzewa, że „bliskość” tych dwóch stron oznacza, że dowód geometrycznej hipotezy Langlandsa może ostatecznie pomóc w rozwoju wersji arytmetycznej. „Aby naprawdę zrozumieć korespondencję Langlandsa, musimy zdać sobie sprawę, że ‘dwa światy’ w niej zawarte nie są tak bardzo różne – są raczej dwoma aspektami tego samego świata”, zauważa Edward Frenkel, jeden z matematyków, którzy

badali implikacje tych hipotez, zwłaszcza w odniesieniu do geometrii i teorii strun.

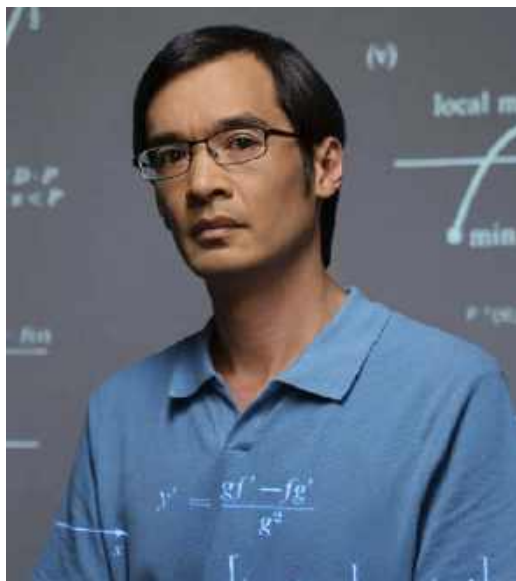
Udowodnienie geometrycznej hipotezy Langlandsa od dawna uważane jest za jedno z najgłębszych i najbardziej enigmatycznych przedsięwzięć współczesnej matematyki. Rozwiązanie tego problemu wymagało pracy zespołu dziewięciu matematyków, którzy opublikowali serię pięciu artykułów liczących prawie tysiąc stron. Grupą kierowali Dennis Gaitsgory z Instytutu Matematyki im. Maxa Plancka w Bonn w Niemczech i Sam Raskin z Uniwersytetu Yale w New Haven w stanie Connecticut. Wielkość ich osiągnięcia została doceniona przez społeczność matematyczną. Gaitsgory otrzymał nagrodę Breakthrough Prize in Mathematics, a Raskin został uhonorowany nagrodą New Horizons dla obiecujących matematyków. Podobnie jak wiele przełomowych wyników w matematyce, dowód ten zapowiada stworzenie mostów między różnymi dziedzinami, umożliwiając wykorzystanie narzędzi jednej dziedziny do rozwiązywania trudnych problemów w innej.

Jedna strona geometrycznej hipotezy Langlandsa dotyczy cechy zwanej grupą fundamentalną. W uproszczeniu grupa fundamentalna powierzchni Riemanna opisuje wszystkie sposoby, w jakie można zawiązać pętle wokół niej. Na przykład w przypadku pączka pętla może przebiegać poziomo wokół zewnętrznej krawędzi lub pionowo przez otwór i wokół zewnętrznej krawędzi. Geometryczna hipoteza Langlandsa dotyczy „reprezentacji” grupy fundamentalnej powierzchni, która wyraża właściwości grupy jako macierze (siatki liczb). Druga strona programu geometrycznego Langlandsa dotyczy specjalnych rodzajów „wiązek”. Te narzędzia geometrii algebraicznej to reguły, które przypisują „przestrzenie wektorowe” (gdzie wektory – strzałki – można dodawać i mnożyć) do punktów na rozmaitości w podobny sposób, jak funkcja opisująca pole grawitacyjne może przypisywać liczby odpowiadające sile pola do punktów w standardowej przestrzeni 3D.

Według niektórych badaczy, jednym z najbardziej zaskakujących pomostów, jakie zbudował geometryczny program Langlandsa, jest połączenie z fizyką teoretyczną. Od lat 70. fizycy badają kwantowy odpowiednik klasycznej symetrii – zamiany pól elektrycznych i magnetycznych w równaniach Maxwella, które opisują interakcje między tymi dwoma polami, nie zmieniając tych równań. Ta elegancka symetria stanowi podstawę szerszej koncepcji w kwantowej teorii pola, znanej jako dualność S. W 2007 roku Edward Witten z Instytutu Studiów Zaawansowanych (IAS) w Princeton w stanie New Jersey oraz Anton Kapustin

z Kalifornijskiego Instytutu Technologicznego w Pasadena wykazali, że dualność S w niektórych czterowymiarowych teoriach grawitacyjnych ma tę samą symetrię, która pojawia się w geometrycznej korespondencji Langlandsa. „Pozornie ezoteryczne pojęcia geometrycznego programu Langlandsa”, napisali naukowcy z jego zespołu, „wynikają naturalnie z fizyki”. Chociaż teorie obejmują hipotetyczne cząstki, zwane superpartnerami, które nigdy nie zostały zaobserwowane, naukowcy sugerują, że geometryczny program Langlandsa nie jest tylko abstrakcyjną ideą z zakresu czystej matematyki, ale może być postrzegany jako odzwierciedlenie głębokiej symetrii w fizyce kwantowej. A zatem to, co zaczęło się jako zbiór głębokich i trudnych do zrozumienia dla laików hipotez łączących abstrakcyjne gałęzie matematyki, przekształciło się w prężną, multidyscyplinarną inicjatywę, która rozciąga się od podstaw teorii liczb do granic fizyki kwantowej. Korespondencja Langlandsa może nie jest jeszcze wielką zunifikowaną teorią matematyki, ale dowód jej geometrycznego ramienia jest zbiorem idei, które prawdopodobnie będą kształtować tę dziedzinę przez wiele lat. „Korespondencja Langlandsa wskazuje na znacznie głębsze struktury matematyki, które dopiero zaczynamy odkrywać”, uważa Frenkel. „Nie rozumiemy jeszcze, czym one są. Nadal pozostają one za kurtyną”.

W książce pt. „The AI Langlands Program” H. Peter Alesso opisuje, jak wiosną 2022 roku student Alexey Pozdnyakov odkrył coś, co nie powinno istnieć. Uruchamiając algorytmy uczenia maszynowego na krzywych eliptycznych, znalazł „płynne wzory przypominające struktury lotu stada szpaków”. Książka pokazuje, jak sztuczna inteligencja rewolucjonizuje czystą matematykę przez odkrywanie wzorców, automatyczne generowanie dowodów i zarządzanie złożonością. Uczenie maszynowe analizuje jednocześnie miliony obiektów matematycznych, odkrywając ukryte struktury. Program Langlandsa, łączący teorię liczb, geometrię i algebrę, stanowi narracyjny kręgosłup książki, która opowiada m.in. o Andrew Sutherlandzie analizującym miliard krzywych eliptycznych za pomocą sztucznej inteligencji, Ninie Zubrilinie wykorzystującej spostrzeżenia obliczeniowe do wyprowadzania wzorów matematycznych oraz zespołach rozwiązujących czterdziestoletnie problemy dzięki uczeniu maszynowemu. Czytamy w niej o tym, jak sztuczna inteligencja rozwija formę intuicji matematycznej, rozpoznając obiecujące strategie i sugerując kierunki badań. I o tym, jak techniki sztucznej inteligencji mogą być stosowane do odkrywania wzorców i powiązań w duchu programu Langlandsa.



3. Terence Tao

„Przeciętny, ale nie do końca niekompetentny” student

Terence Tao (3), profesor matematyki na Uniwersytecie Kalifornijskim w Los Angeles (UCLA), nazywany „Mozartem matematyki” i powszechnie uważany za najwybitniejszego żyjącego matematyka na świecie, opublikował jakiś czas temu w Internecie swoje wrażenia na temat modelu o1 firmy OpenAI. Porównał go do „przeciętnego, ale nie do końca niekompetentnego” studenta. Jednak dostrzega też swoisty rodzaj „matematyki na skalę przemysłową” opartej na sztucznej inteligencji, która nigdy wcześniej nie była możliwa. W jego ocenie AI w niedalekiej przyszłości nie będzie kreatywnym współpracownikiem sama w sobie, a raczej środkiem smarującym hipotezy i podejścia matematyków. Ten nowy rodzaj matematyki, który mógłby otworzyć nieznane obszary wiedzy, pozostanie, jak przypuszcza, w swej istocie ludzkiej, uwzględniając to, że ludzie i maszyny mają różne mocne strony, które należy postrzegać jako uzupełniające się, a nie konkurujące.

Opowiada o swoich doświadczeniach z ChatGPT: „Zadałem kilka trudnych zadań matematycznych i otrzymałem dość absurdalne wyniki. To był spójny angielski, zawierał właściwe słowa, ale brakowało w nim głębi. Wczesne GPT były mało imponujące. Nadawały się do zabawnych rzeczy, na przykład do wyjaśnienia jakiegoś zagadnienia matematycznego w formie wiersza lub opowiadania dla dzieci”. Późniejsze wersje były, jak ocenia, lepsze, jednak tylko w niektórych zastosowaniach, np. analizie wariantów.

„To już się dzieje w niektórych innych obszarach”, mówi. Sztuczna inteligencja podbiła szachy lata temu (...). Teraz całkiem dobry szachista może teraz spekulować, które ruchy są dobre w jakich sytuacjach i może użyć silników szachowych, aby sprawdzić dwadzieścia ruchów do przodu. Widzę, że coś takiego może się kiedyś wydarzyć w matematyce. Masz projekt i pytasz: A co, jeśli wypróbuję to podejście? Zamiast spędzać godziny na próbach, aby to zadziałało, instruujesz GPT, by zrobił to za ciebie. Z o1 można to w pewnym sensie zrobić. Dałem mu problem, który wiedziałem, jak rozwiązać, i starałem się pokierować modelem. Najpierw dałem mu wskazówkę, a on ją zignorował i zrobił coś innego, co nie zadziałało. Kiedy to wyjaśniłem, przeprosił i powiedział: Dobrze, zrobię to po twojemu. Potem wykonał moje instrukcje całkiem dobrze, a potem znowu się zawiesił i musiałem go ponownie poprawić. Model nigdy nie wymyślił najsprytniejszych kroków. Potrafił wykonywać wszystkie rutynowe czynności, ale był bardzo pozbawiony wyobraźni. Jedną z kluczowych różnic między studentami a sztuczną inteligencją jest to, że studenci się uczą. Mówisz sztucznej inteligencji, że jej podejście nie działa, ona przeprosza, może chwilowo skoryguje swój kurs, ale czasami po prostu wraca do tego, co próbowała wcześniej. A jeśli zaczynasz nową sesję z AI, wracasz do punktu wyjścia”.

Tao zauważ jednak, że, jak ich określa, „asystenci dowodów”, to przydatne narzędzia komputerowe, które sprawdzają, czy dowód matematyczny jest poprawny, czy nie. Przy stuprocentowo poprawnie działającym asystencie „możesz wtedy wykonywać obliczenia matematyczne na skalę przemysłową, w stylu produkcji fabrycznej”. „Klasyczna idea matematyki polega na tym, że wybierasz bardzo trudny problem, a potem jedna lub dwie osoby są zamknięte na strychu na siedem lat i po prostu nad nim pracują”, mówi matematyk. „Problemy, z którymi chcesz walczyć za pomocą sztucznej inteligencji, są odwrotne. Naiwnym sposobem wykorzystania sztucznej inteligencji jest dostarczenie jej najtrudniejszego problemu, jaki mamy w matematyce. Nie sądzę, żeby to się sprawdziło i mamy już ludzi, którzy nad tymi problemami pracują”. Jak dodaje, pracuje nad projektem, który dotyczy dziedziny matematyki zwanej algebrą uniwersalną, która bada, czy pewne stwierdzenia matematyczne lub równania implikują prawdziwość innych stwierdzeń. „W przeszłości ludzie badali to w ten sposób, że wybierali jedno lub dwa równania i studiowali je do upadłego. Teraz mamy fabryki. (...) Zamiast wąskiej, głębokiej matematyki, gdzie ekspert ciężko pracuje nad wąskim



4. Wizualizacja matematycznej fabryki w której AI pracuje nad żmudnymi obliczeniami

zakresem problemów, można mieć szerokie, zleczone przez społeczność problemy z dużym wsparciem sztucznej inteligencji, które są może płytsze, ale na znacznie większą skalę. I może to być komplementarny sposób zdobywania wiedzy matematycznej (...) Sto pięćdziesiąt lat temu matematyka była przede wszystkim użyteczna w rozwiązywaniu równań różniczkowych cząstkowych. Istnieją już pakiety komputerowe, które robią to automatycznie. Sześćset lat temu matematycy tworzyli tablice sinusów i cosinusów, niezbędne do nawigacji, ale teraz komputery mogą je generować w kilka sekund. Nie chodzi o powielanie rzeczy, w których ludzie są już dobrzy. Wydaje się to nieefektywne. Myślę, że będziemy potrzebować ludzi jak też sztucznej inteligencji. Mają one wzajemnie uzupełniające się mocne strony. Sztuczna inteligencja jest bardzo dobra w przekształcaniu miliardów danych w jedną dobrą odpowiedź. Ludzie są dobrzy w zbieraniu dziesięciu obserwacji i formułowaniu trafnych hipotez”.

Jak widać, znani matematycy nie są skłonni pochopnie powierzać AI najważniejszych i najtrudniejszych problemów matematycznych. Widzą w niej jednak dość obiecujące narzędzie, które może ulżyć w uciążliwych i trudnych dla pojedynczych ludzi zadaniach (4). ■

Mirosław Usidus

W jaki sposób AI generuje obrazy?

Neuronowe modele dyfuzyjne

Sztuczna inteligencja to w ostatnich czasach nie tylko modele językowe, umożliwiające konstrukcję rozwiązań takich, jak ChatGPT, gdzie możliwa staje się „rozmowa” z maszyną. W ostatniej dekadzie rozwinęły się również znacznie zastosowania sztucznej inteligencji w zakresie przetwarzania obrazów. Ma to oczywiście związek z rosnącą mocą obliczeniową sprzętu komputerowego – zagadnienia, których sformułowanie dawniej było zbyt złożone obliczeniowo do rozwiązania, obecnie są w zasięgu techniki. Sztuczna inteligencja w przetwarzaniu obrazów przestaje być stosowana tylko do prostych zadań, jak wykrywanie uśmiechu na twarzy czy zmian nowotworowych na zdjęciu preparatu biopsyjnego. Ostatnio otwierają się nowe możliwości, związane z generowaniem całkowicie sztucznych obrazów, które na pierwszy rzut oka wydają się fotografiami. Jak to działa?

W branży tej istnieje wiele rozwiązań, ale do najpopularniejszych należą obecnie modele dyfuzyjne, którymi zajmę się w tym artykule. W modelach tych sieć neuronowa uczy się oczyszczać obraz z zakłóceń (mniejszych lub większych, a nawet drastycznych), co ostatecznie skutkuje możliwością **wyłonięcia obrazu z czystego szumu**.

Drugą klasą rozwiązań, która święciła sukcesy do około roku 2020, są sieci adwersaryjne (GAN – Generative Adversary Networks). Sieci GAN w ciekawy sposób formułują funkcję celu i mogą być w związku z tym wykorzystywane również w zagadnieniach odbiegających od generowania obrazów fotograficznych. Ale zostawimy je sobie na inną okazję.

Sposób zaszumiania obrazu na potrzeby uczenia

Generujące modele dyfuzyjne podczas treningu sieci opierają się na idei rekonstrukcji obrazu źródłowego, który zostaje stopniowo coraz mocniej zaszumiany. Stąd właśnie nazwa modeli dyfuzyjnych – podczas treningu rozpoczynamy od ostrej fotografii świata rzeczywistego, a następnie dodajemy stopniowo do obrazu szum i obniżamy średnią wartość piksela do zera (przymiemy się ułomkowy zakres wartości pikseli od -1 do 1 , a nie całkowity od 0 do 255). Ogólnie, dla każdego piksela można napisać:

$$x_{n+1} = \sqrt{\alpha_n} x_n + \sqrt{1 - \alpha_n} \xi_n$$

gdzie x_{n+1} jest wartością piksela w kolejnym etapie zaszumiania, x_n jest wartością piksela w chwili bieżącej, α_n jest współczynnikiem regulującym siłę zaszumiania (w oryginalnym rozwiązaniu czynnik $1 - \alpha_n = 10^{-4} + 0.0199n/n_{\max}$), a ξ_n jest liczbą losową o rozkładzie Gaussa $N(0,1)$. Równanie to obrazuje stopniowe „zapominanie” wartości piksela ($\sqrt{\alpha_n} < 1$) oraz losowe zastępowanie zapominanych wartości.

Równanie zaszumiania w sensie matematyczno-fizycznym jest przykładem równania Langevina, opisującym proces dyfuzji, tj. ruchy Browna w jakimś specyficznym potencjale (przyciągającym cząstkę Browna w okolice zera). W takim procesie położenie cząstki Browna (np. pyłku na wodzie czy cząstki gazu w powietrzu) ulega zmianom opisywanym analogicznym równaniem. W zastosowaniu do przetwarzania obrazów – zamiast zmiany położenia w przestrzeni – jesteśmy zainteresowani zmianą wartości piksela (czy precyzyjniej: pikseli obrazu).



1. Proces zaszumiania obrazu oryginalnego (czerwone strzałki) oraz jego odszumiania (niebieskie). Warto zwrócić uwagę na nazewnictwo rozkładów q oraz p

W takim opisie przejście od wartości x_n do x_{n+1} odpowiada pojawieniu się pewnej losowości. Można ją opisać prawdopodobieństwem warunkowym

$$q(x_{n+1}|x_n) = \frac{1}{\sqrt{2\pi(1-\alpha_n)}} \exp \left[-\frac{(x_{n+1} - \sqrt{\alpha_n}x_n)^2}{2(1-\alpha_n)} \right]$$

Co jest zwykłym rozkładem Gaussa, wynikającym z dodania do wartości $\sqrt{\alpha_n}x_n$ (stanowiącej obecnie **średnią** tego rozkładu) – losowej liczby ξ_n o wariancji $(1-\alpha_n)$. W ogólności, połączenie n kroków zaszumiania obrazu – począwszy od ostrego obrazu o indeksie x_0 , skutkuje rozkładem o coraz niższej średniej, z coraz większą amplitudą szumu, który przyjmuje postać (można ją dowiedzieć metodą indukcji matematycznej – Dodatek):

$$q(x_n|x_0) = \frac{1}{\sqrt{2\pi(1-\alpha_n)}} \exp \left[-\frac{(x_n - \sqrt{\alpha_n}x_0)^2}{2(1-\alpha_n)} \right] \quad \alpha_n = \alpha_1\alpha_2 \dots \alpha_{n-1}\alpha_n$$

Co bardzo ważne – do reprezentacji na n -tym kroku przetwarzania wg powyższej relacji możemy dojść w **jednym** kroku! Wystarczy wygenerować liczbę losową z powyższego rozkładu, nie musimy znać rozkładów dla niższych wartości n ! To ułatwia i przyspiesza uczenie.

Probabilistyczny model rekonstrukcji obrazu zaszumionego

Jak na razie idea jest prosta i nic się nie dzieje w związku z **generowaniem** obrazów. Ale oto podczas treningu, po pełnym zaszumieniu obrazu, pojawia się zagadnienie **rekonstrukcji** obrazu źródłowego z szumu (niebieskie strzałki na rysunku 1). Wydaje się to zupełnie bez sensu – w końcu w pełni zaszumiony obraz nie zachowuje nawet śladu informacji o obrazie źródłowym. I o to chodzi. Chcemy, w kontekście obejmowanym przez treningowy zbiór obrazów, nauczyć sieć neuronową generować coś, co byłoby **prawdopodobnym** źródłem szumu. Innymi słowy – **prawdopodobnym** obrazem ze zbioru uczącego. Chcemy maksymalizować prawdopodobieństwo łączne:

$$q(x_0)p_0(x_0) \rightarrow \max$$

lub równoważne, ponieważ $q(x_0)$ jest ostrym obrazem i $q(x_0)=1$, gdy x_0 ma wartość piksela obrazu oraz $q(x_0)=0$, w przeciwnym wypadku,

$$\int dx_0 q(x_0)p_0(x_0) = E_{x_0 \sim q} p_0(x_0) \rightarrow \max$$

Czytelniku! Modelowanie probabilistyczne dyfuzji w obrazach jest **wymagające pod względem matematycznym**. Jeśli czujesz, że nie masz na to siły – przyjmij po prostu, że moduł odszumiania musi nauczyć się (metodą prób i błędów przy zmienianiu współczynników sieci neuronowej, akceptując tylko korzystne zmiany) oddzielać szum ϵ od obrazu x_i z i -tego poziomu zaszumiania. Stopniowe odejmowanie tego szumu umożliwia wyłonienie coraz lepszego obrazu. Z tą wiedzą możesz przeskoczyć do paragrafu o uczeniu dyfuzyjnej sieci neuronowej, gdzie odszumiony obraz oznaczamy literą μ , a oszacowanie szumu – ϵ_θ . θ oznacza zbiór wag sieci neuronowej, która identyfikuje ϵ w obrazie x_i .

(całka biegnie po **wartościach** wybranego piksela – kryterium musimy zastosować do wszystkich pikseli obrazu). Lub jeszcze inaczej, ponieważ wygodniej minimalizować logarytm prawdopodobieństwa niż samo prawdopodobieństwo,

$$\int dx_0 q(x_0) \log p_\theta(x_0) = E_{x_0 \sim q} \log p_\theta(x_0) \rightarrow \max$$

Pojawiła nam się tutaj bariera dla czytelników bez przygotowania z zakresu matematyki wyższej – **całka**. Ale nie trzeba się przejmować. Całka zapisana we wszystkich relacjach na prawdopodobieństwa, jakich tu użyjemy, jest równoważna operacji sumowania po wszystkich wartościach zmiennej całkowania, np. $\int dx_i p(x_i) \approx \sum_{x_i \in R} P(x_i)$, przy czym $P(x_i) = p(x_i) dx$, a dx to różnica między sąsiednimi wartościami x_i . Małe $p(x)$ jest zatem **gęstością prawdopodobieństwa**, a duże $P(x)$ jest **prawdopodobieństwem**. Taka suma z nieskończoną liczbą nieskończenie małych różniących się sąsiednich składników jest dość dziwna (a częściej łatwiejsza do wysumowania niż sumy dyskretne!) i dlatego Newton wynalazł całkowanie.

Drugim pojęciem, które warto wyjaśnić, jest pojęcie wartości oczekiwanej. Załóżmy, że chcemy policzyć średnią wartość funkcji $f(x)$ dla liczby X wyrzuconej w rzucie kostką do gry (może być asymetryczna) po wykonaniu n rzutów. Policzylibyśmy to za pomocą średniej arytmetycznej $f(x) = \frac{1}{n} \sum_{i=1}^n f(x_i)$. Ale jeśli jest to kostka, to X może przyjąć np. tylko 6 możliwych wartości i zamiast sumować po wszystkich uzyskanych wynikach, możemy wykonać sumę po możliwych wartościach: $f(x) = \sum_{i=1}^6 \frac{n_i}{n} f(x_i) \approx \sum_{i=1}^6 P(x_i) f(x_i)$. Tak liczymy wartość oczekiwaną (czyli swego rodzaju „średnią”) dla rozkładu o dyskretnej, skończonej liczbie wartości. Dla rozkładu ciągłego, zgodnie z poprzednią ramką, wprowadzamy całkę $f(x) = \int dx p(x) f(x)$ – i to właśnie z tej postaci korzystamy w tym artykule.

Wyobraźmy sobie teraz stopniowy proces „odszumiania” zaszumionego obrazu:

$$p_\theta(x_0) = \int dx_n dx_{n-1} dx_{n-2} \dots dx_2 dx_1 p_\theta(x_0 | x_1) p_\theta(x_1 | x_2) \dots p_\theta(x_{n-1} | x_n) p_\theta(x_n) = \\ = \int dx_n dx_{n-1} dx_{n-2} \dots dx_2 dx_1 p_\theta(x_0, x_1, \dots, x_n)$$

Czyli: dysponując w pełni zaszumionym obrazem x_n , możemy obliczyć prawdopodobieństwa przejść do wszystkich możliwych wartości pikseli obrazów: x_{n-1} pod warunkiem, że x_n miało określoną wartość, następnie x_{n-2} pod warunkiem, że x_{n-1} miało określoną wartość, itd., aż do x_1 , by na koniec uzyskać zbiór prawdopodobieństw dla pikseli obrazu x_0 . Z tego zbioru wybierzemy piksele najbardziej prawdopodobne, np. maksymalizując proponowaną już wartość oczekiwaną $E_{x_0 \sim q}[\log p(x_0)]$. **Proszę zwrócić uwagę na zastosowanie symbolu p_θ i q – to bardzo ważne, by rozróżniać rozkłady konstruowane od rekonstruowanych.**

Warto zauważyć, że wymnożenie prawdopodobieństw warunkowych całej tej sekwencji (jeszcze przed wykonaniem całkowania), z definicji **prawdopodobieństwa łącznego zdarzeń** przedstawia prawdopodobieństwo łącznego wystąpienia sekwencji $x_0, x_1, \dots, x_n$. Zgodnie z treścią równania powyżej, wycalkowanie (wysumowanie) tego rozkładu po wszystkich możliwych wartościach $x_n, \dots, x_1$, przy ustalonym x_0 – daje $p_\theta(x_0)$. Pojęcie prawdopodobieństwa łącznego umożliwi nam zaraz odwrócenie sekwencji mnożeń prawdopodobieństw warunkowych.

Pomnożymy teraz i podzielimy $p_\theta(x_0, \dots, x_n)$ przez szereg iloczynów prawdopodobieństw $q(x_i | x_{i-1})$, a prawdopodobieństwo łączne $p_\theta(x_0, \dots, x_n)$ wyrazimy we wspomnianej „odwróconej kolejności” iloczynów prawdopodobieństw warunkowych. W ten sposób zaczynamy kierować naszą uwagę do kluczowego dla modeli dyfuzyjnych pojęcia **dywergencji Kullbacka-Leiblera**, o której więcej powiemy za chwilę:

$$p_\theta(x_0) = \int dx_{n-1} dx_{n-2} \dots dx_2 dx_1 \overbrace{p_\theta(x_0, \dots, x_n)}^{\text{rozwinac} \rightarrow} \prod_{i=1}^n \overbrace{\frac{q(x_i | x_{i-1})}{q(x_i | x_{i-1})}}^{\leftarrow \text{zwinac}} = \\ = \int dx_{n-1} dx_{n-2} \dots dx_2 dx_1 q(x_n, x_{n-1}, \dots, x_1 | x_0) p_\theta(x_0) \prod_{i=1}^n \frac{p_\theta(x_i | x_{i-1})}{q(x_i | x_{i-1})} = \\ = E_{x_1, \dots, x_n \sim q(x_1, \dots, x_n | x_0)} \left[p_\theta(x_0) \prod_{i=1}^n \frac{p_\theta(x_i | x_{i-1})}{q(x_i | x_{i-1})} \right]$$

Wylączylismy tu prawdopodobieństwo łączne $q(x_1, x_2, \dots, x_{n-1}, x_n | x_0) = q(x_1 | x_0) q(x_2 | x_1) \dots q(x_n | x_{n-1})$, a następnie skorzystaliśmy z definicji wartości oczekiwanej względem tego rozkładu (notabene, jest to rozkład, jakim dysponujemy z danych).

Przypomnijmy sobie teraz, że chcemy maksymalizować logarytm z prawdopodobieństwa. Umożliwia to nam wykonanie pewnej sztuczki, z wykorzystaniem tzw. nierówności Jensena (że logarytm z wartości oczekiwanej jest większy lub równy wartości oczekiwanej logarytmu):

$$\begin{aligned}\log p(x_0) &= E_{x_1 \sim q} \log E_{x_1, \dots, x_n \sim q} \left[p_\theta(x_0) \prod_{i=1}^n \frac{p_\theta(x_i | x_{i-1})}{q(x_i | x_{i-1})} \right] \\ &\geq E_{x_1, \dots, x_n \sim q} \log \left[p_\theta(x_0) \prod_{i=1}^n \frac{p_\theta(x_i | x_{i-1})}{q(x_i | x_{i-1})} \right] \\ &= E_{x_1, \dots, x_n \sim q} \log \left[\frac{p_\theta(x_0, x_1, \dots, x_{n-1}, x_n)}{q(x_n, \dots, x_2, x_1 | x_0)} \right]\end{aligned}$$

W toku uczenia sieci neuronowej będziemy maksymalizować właśnie to oszacowanie, które za chwilę drastycznie uprościmy. Aby określić błąd tej nierówności, możemy odjąć od $\log p_\theta(x_0)$ jego przybliżenie:

$$\begin{aligned}\log p(x_0 | x_n) - E_{x_1, \dots, x_n \sim q} \log \left[\frac{p(x_0, x_1, \dots, x_{n-1}, x_n)}{q(x_n, \dots, x_2, x_1 | x_0)} \right] &= E_{x_1, \dots, x_n \sim q} \log \left[\frac{q(x_n, \dots, x_2, x_1 | x_0)}{p(x_n, \dots, x_2, x_1 | x_0)} \right] \\ &= D_{KL}[p(x_n, \dots, x_2, x_1 | x_0) \parallel q(x_n, \dots, x_2, x_1 | x_0)]\end{aligned}$$

Wynik – wartość oczekiwana z logarytmu ilorazu rozkładów – okazuje się znanym obiektem – **dywergencją Kullbacka–Leiblera**, miarą rozbieżności rozkładów prawdopodobieństwa. Jest to centralny obiekt w teorii modeli dyfuzyjnych (czytelnik natknie się na niego w większości związanych z nimi publikacji) i dlatego warto zajrzeć do poniższej ramki, aby zrozumieć, jak ta miara działa.

Dywergencja KL: Wiąże się ona ze znanym z teorii informacji pojęciem entropii źródła (generującego n symboli o prawdopodobieństwach p_i – np. litery w dokumentach tekstowych):

$$h = - \sum_{i=1}^n p_i \log_2 p_i$$

gdzie w optymalnym kodowaniu przypisuje się $-\log_2 p_i$ bitów symbolowi, występującemu z prawdopodobieństwem p_i . Innymi słowy, najczęściej występujące symbole dostają najkrótsze kody, a najrzadziej występujące symbole koduje się długimi kodami. W ten sposób uzyskujemy najlepsze upakowanie, czy „kompresję” danych. W tym sensie entropia wyraża średnią liczbę bitów, potrzebną na kodowanie jednego symbolu źródła. Aby przekonać się, że to działa – niech czytelnik policzy, ile bitów potrzeba na zakodowanie 256 równoprawdopodobnych symboli z $p_i = 1/256$. Wyjdzie 8.

Jeśli teraz wykorzystalibyśmy kodowanie optymalne względem symboli o rozkładzie q_i , do kodowania symboli o rozkładzie p_i , to na ogół będzie to skutkowało pogorszeniem „kompresji” danych – symbole o krótkich kodach mogą wystąpić rzadziej i mniejszy będzie pożytek z tych krótkich kodów, a symbole o kodach długich mogą wystąpić częściej, dając znaczny dodatkowy wkład do h . Dywergencja KL jest właśnie różnicą tak liczonych entropii:

$$D_{KL} = - \sum_i p_i \log_2 q_i - p_i \log_2 p_i$$

i wyraża liczbę bitów „niespodzianki” związanej z niezgodnością prawdopodobieństw.

Wróćmy teraz do oszacowania, jakie uzyskaliśmy z nierówności Jensena dla $\log p_\theta(x_0)$ i wstawmy to do wzoru na wartość oczekiwaną $E_{x_0 \sim q} \log p_\theta(x_0)$:

$$\begin{aligned}E_{x_0 \sim q} E_{x_1, \dots, x_n \sim q} \log \left[\frac{p_\theta(x_0, x_1, \dots, x_{n-1}, x_n)}{q(x_n, \dots, x_2, x_1 | x_0)} \right] &= \\ E_{x_0, \dots, x_n \sim q} \underbrace{\log p_\theta(x_n)}_{\text{pominac}} + \log \left[\frac{p_\theta(x_0, x_1, \dots, x_{n-1}, x_n)}{q(x_n, \dots, x_2, x_1 | x_0)} \right] &= \\ \rightarrow E_{x_0, \dots, x_n \sim q} \log \left[\prod_{i=1}^n \frac{p_\theta(x_{i-1} | x_i)}{q(x_i | x_{i-1})} \right] &= \\ \sum_{i=1}^n E_{x_0, \dots, x_n \sim q} \log \left[\frac{p_\theta(x_0, x_1, \dots, x_{n-1}, x_n)}{q(x_i | x_{i-1})} \right]\end{aligned}$$

W wyrażeniu po prawej stronie u góry $p_\theta(x_n)$ jest rozkładem normalnym $N(0,1)$, niezależnym od parametrów uczących θ , dlatego nie bierzemy pod uwagę tego członu w procesie optymalizacji (tj. modyfikacji θ w celu

uzyskania największego prawdopodobieństwa $p_\theta(x_0)$). Ostatecznie, po rozpisaniu prawdopodobieństw łącznych i zamianie logarytmu iloczynu na sumę logarytmów, dostaliśmy wynik widoczny w prawym dolnym rogu.

W tym ostatnim członie chcielibyśmy wprowadzić dywergencję KL, ale aby tego dokonać, musimy działać na tym samym argumentcie, a tutaj licznik dotyczy x_{i-1} , a mianownik x_i , w dodatku warunkowane są różnymi zdarzeniami i niepodobna traktować je jako odpowiadające sobie rozkłady. Aby z tego wybrnąć, prawdopodobieństwo p można wyznaczyć z relacji na prawdopodobieństwo warunkowe:

$$P(B|A) = \frac{P(A,B)}{P(A)}$$

czyli w naszych realiach:

$$q(x_{i-1}|x_i, x_0) = \frac{q(x_i|x_{i-1}, x_0)q(x_{i-1}|x_0)}{q(x_i|x_0)}$$

Procesy dyfuzyjne są **procesami Markowa**, w których kolejny krok czasowy zależy **wyłącznie** od kroku poprzedniego, wobec czego można uprościć pierwszy wyraz w liczniku powyższego wyrażenia po prawej (zapomnieć o zależności od x_0):

$$q(x_{i-1}|x_i, x_0) = \frac{q(x_i|x_{i-1}, x_0)q(x_{i-1}|x_0)}{q(x_i|x_0)}$$

Co ciekawe, wszystkie obiekty po prawej stronie są krzywymi Gaussa, o zdefiniowanych już wcześniej średnich i wariancjach. Mnożąc i dzieląc przez siebie te krzywe gaussowskie, uzyskamy wypadkowy rozkład o wartości oczekiwanej i wariancji (wyprowadzenie w Dodatku):

$$q(x_{i-1}|x_i, x_0) = N \left(\frac{\sqrt{\alpha_{i-1}}}{1 - \alpha_i} x_0 + \frac{\sqrt{\alpha_i} (1 - \alpha_{i-1})}{1 - \alpha_i} x_i, \frac{1 - \alpha_{i-1}}{1 - \alpha_i} \beta_i \right)$$

Wyliczając (ze wzoru o jeden wcześniejszego) $q(x_i, x_{i-1})$, uzyskamy:

$$q(x_i|x_{i-1}) = \frac{q(x_{i-1}|x_0)}{q(x_i|x_0)q(x_{i-1}|x_i, x_0)}$$

Podstawiamy do wzoru na logarytm prawdopodobieństwa rekonstrukcji:

$$\begin{aligned} E_{x_0 \sim q} \log p(x_0) &= E_{x_0, \dots, x_n \sim q} \sum_{i=1}^n \log \left[\frac{p_\theta(x_{i-1}|x_i)}{q(x_i|x_{i-1})} \right] = \\ &= E_{x_0, \dots, x_n \sim q} \log \left[\frac{p_\theta(x_0|x_1)}{q(x_1|x_0)} \right] + \sum_{i=2}^n \log \left[\frac{p_\theta(x_{i-1}|x_i)}{q(x_i|x_{i-1})} \right] = \\ &= E_{x_0, \dots, x_n \sim q} \log \left[\frac{p_\theta(x_0|x_1)}{q(x_1|x_0)} \right] + \sum_{i=2}^n \log \left[\frac{p_\theta(x_{i-1}|x_i)}{q(x_{i-1}|x_i)} \right] + \log \left[\frac{q(x_{i-1}|x_0)}{q(x_i|x_0)} \right] = \\ &\quad \{ \text{po wysumowaniu ostatniego wyrazu} \} \\ &= E_{x_0, \dots, x_n \sim q} \log \left[\frac{p_\theta(x_0|x_1)}{q(x_1|x_0)} \right] + \log \left[\frac{q(x_1|x_0)}{q(x_n|x_0)} \right] + \sum_{i=2}^n \log \left[\frac{p_\theta(x_{i-1}|x_i)}{q(x_{i-1}|x_i)} \right] = \\ &= E_{x_0, \dots, x_n \sim q} \log p_\theta(x_0|x_1) - \underbrace{\log q(x_n|x_0)}_{N(0,1)} + \sum_{i=2}^n \log \left[\frac{p_\theta(x_{i-1}|x_i)}{q(x_{i-1}|x_i)} \right] \end{aligned}$$

W związku z tym, że środkowy człon zależy od unormowanego rozkładu normalnego, nie podlega optymalizacji. Pozostają człony pierwszy i ostatni, które można zapisać, wzięwszy je na funkcję celu z minusem (zwykle w procesie uczenia dokonujemy minimalizacji straty), jako:

$$\begin{aligned} L &= -E_{x_0, \dots, x_n \sim q} \log p_\theta(x_0|x_1) + \sum_{i=1}^n \log \left[\frac{p_\theta(x_{i-1}|x_i)}{q(x_{i-1}|x_i)} \right] = \\ &= \sum_{i=1}^n D_{KL} [p_\theta(x_{i-1}|x_i) \parallel q(x_{i-1}|x_i)] - E_{x_0, x_1 \sim q} \log p_\theta(x_0|x_1) \end{aligned}$$

Uproszczona funkcja celu

Spróbujmy w powyższym wyrażeniu na dywergencję KL wyrazić prawdopodobieństwa za pomocą funkcji wykładniczych o określonej wariancji i średniej. Wówczas logarytm obecny w dywergencji skasuje funkcję

wykładniczą rozkładu Gaussa i będziemy mogli policzyć po prostu wartość oczekiwaną z argumentów eksponenty Gaussa. Czynniki przedwykładnicze, biorąc pod uwagę zbliżone wariancje obu rozkładów dla x_{i-1} , pomijamy.

$$\begin{aligned} D_{KL}[p_\theta(x_{i-1}|x_i) \parallel q(x_{i-1}|x_i)] &= \\ E_{x_{i-1}, x_i} \left[\frac{(x_{i-1} - \mu(x_i))^2 - (x_{i-1} - \mu_\theta(x_i))^2}{2\sigma_i^2} \right] &= E_{x_{i-1}, x_i} \left[\frac{\mu^2 - 2x_{i-1}(\mu - \mu_\theta) - \mu_\theta^2}{2\sigma_i^2} \right] = \\ E_{x_i} \left[\frac{\mu^2 - 2\mu_\theta(\mu - \mu_\theta) - \mu_\theta^2}{2\sigma_i^2} \right] &= E_{x_i} \left[\frac{(\mu - \mu_\theta)^2}{2\sigma_i^2} \right] \end{aligned}$$

Teraz musimy oszacować wartości oczekiwane: dla procesu q (wartość μ) oraz dla p (wartość μ_θ). Zależą one ogólnie od x_i oraz x_0 . Aby znaleźć x_0 , wykorzystamy równania na ewolucję x_i :

$$x_i = \sqrt{\alpha_i} x_0 + \sqrt{1 - \alpha_i} \epsilon \rightarrow x_0 = \frac{1}{\sqrt{\alpha_i}} \left(x_i - \sqrt{1 - \alpha_i} \epsilon \right)$$

W rozkładzie $q(x_{i-1}|x_i)$ mieliśmy:

$$\begin{aligned} \mu &= \frac{\sqrt{\alpha_{i-1}} \beta_i}{1 - \alpha_i} x_0 + \frac{\sqrt{\alpha_i} (1 - \alpha_{i-1})}{1 - \alpha_i} x_i = \\ \frac{\sqrt{\alpha_{i-1}} \beta_i}{1 - \alpha_i} \frac{1}{\sqrt{\alpha_i}} \left(x_i - \sqrt{1 - \alpha_i} \epsilon \right) &- \frac{\sqrt{\alpha_{i-1}} \beta_i}{1 - \alpha_i} \sqrt{\frac{1 - \alpha_i}{\alpha_i}} \epsilon + \frac{\sqrt{\alpha_i} (1 - \alpha_{i-1})}{1 - \alpha_i} x_i = \\ \frac{\beta_i}{1 - \alpha_i} \frac{x_i}{\sqrt{\alpha_i}} - \frac{\beta_i}{\sqrt{1 - \alpha_i}} \frac{1}{\sqrt{\alpha_i}} \epsilon &+ \frac{\alpha_i (1 - \alpha_{i-1})}{\sqrt{\alpha_i} (1 - \alpha_i)} x_i = \\ \frac{x_i}{\sqrt{\alpha_i}} - \frac{\beta_i}{\sqrt{1 - \alpha_i}} \frac{1}{\sqrt{\alpha_i}} \epsilon &= \frac{1}{\sqrt{\alpha_i}} \left(x_i - \frac{\beta_i}{\sqrt{1 - \alpha_i}} \epsilon \right) \end{aligned}$$

Ponieważ rozkład $p(x_{i-1}|x_i)$ ma odwzorowywać q , przyjmujemy podobną parametryzację dla drugiej średniej, aczkolwiek zamiast losowego ϵ użyjemy tutaj wyuczonego ϵ_θ , służącego do odwracania zaszumienia:

$$\mu_\theta = \frac{1}{\sqrt{\alpha_i}} \left(x_i - \frac{\beta_i}{\sqrt{1 - \alpha_i}} \epsilon_\theta \right)$$

Podstawiając obie parametryzacje, dostajemy dla dywergencji KL:

$$\begin{aligned} D_{KL}[p_\theta(x_{i-1}|x_i) \parallel q(x_{i-1}|x_i)] &= \\ E_{x_i} \left[\frac{(\mu - \mu_\theta)^2}{2\sigma_i^2} \right] &= E_{x_0, \epsilon} \left[\frac{\beta_i^2 (\epsilon_\theta - \epsilon)^2}{\sigma_i^2 \alpha_i (1 - \alpha_i)} \right] \end{aligned}$$

Przy czym czynnik x_i idealnie się anuluje w obydwu składowych wyrażenia.

Uczenie dyfuzyjnej sieci neuronowej

Uczenie sieci neuronowej polega na uzyskaniu informacji o współczynnikach ϵ_θ , koniecznych do odtworzenia wartości oczekiwanej obrazu z $i-1$ -szego etapu zaszumienia, gdy dysponuje się obrazem z etapu i . Matematycznie możemy to wyrazić zgodnie z przedstawioną wcześniej relacją na oczekiwaną wartość kroku $i-1$:

$$\mu_{\theta} = \frac{1}{\sqrt{\alpha_i}} \left(x_i - \frac{\beta_i}{\sqrt{1 - \alpha_i}} \epsilon_{\theta} \right)$$

Cały algorytm zapiszemy następująco:

1. Wybierz obraz x_0 ze zbioru uczącego
2. Wybierz etap zaszumiania i (rozkład równomierny)
3. Wylosuj $\epsilon \in N(0,1)$
4. Popraw wagi θ sieci neuronowej wyliczającej predykcję $\epsilon_{\theta}(x_i, i)$, aby zminimalizować funkcję celu – czynnik $(\epsilon - \epsilon_{\theta})^2$
5. Jeśli nie osiągnięto zbieżności, wykonaj od punktu 1.

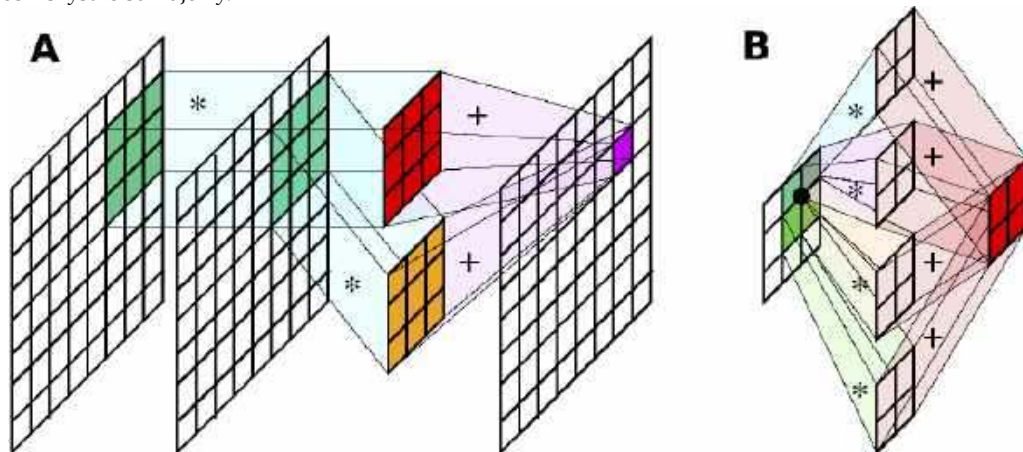
Funkcja obliczająca $\epsilon_{\theta}(x_i)$

Dużo już wiemy, ale nie powiedzieliśmy, jak wyznacza się przewidywanie szumu ϵ_{θ} . Robi to sieć neuronowa, która na wejściu otrzymuje obraz na poziomie zaszumienia „ i ”, i wyznacza wartości szumu, jakimi ten obraz jest obciążony. Część tej predykcji – proporcjonalna do przyczynku z i -tego kroku zaszumiania – jest odejmowana od x_i , aby uzyskać x_{i-1} .

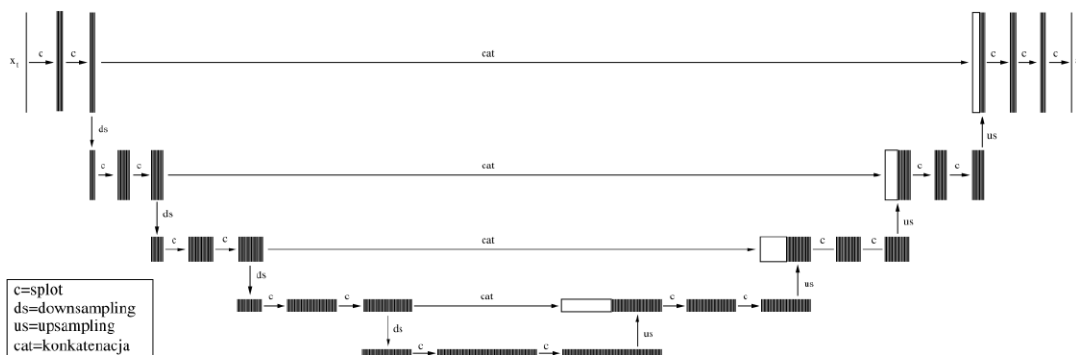
Zagadnienie to właściwie sprowadza się do znanego problemu **segmentacji** obrazu. np. w zagadnieniach biomedycznych problem ten mógł dotyczyć izolowania na preparacie mikroskopowym pojedynczych komórek. Natomiast w zagadnieniach dyfuzyjnego generowania obrazów chcemy wykrywać miejsca, obciążone szumem.

Do zagadnienia segmentacji najlepiej nadają się sieci neuronowe typu **U-net**: są to sieci, które obliczają z obrazu różne charakterystyki za pomocą **splotu** (CNN – ang. *convolutional neural networks* – sieci konwolucyjne), zapisując je w postaci (wielu) nowych warstw, a następnie redukują wymiar analizowanego obrazu, co umożliwia odróżnienie ważnych właściwości od tych, które są zbyt lokalne i zabezpieczają przed przypadkowym dopasowaniem predyktora do tych drugich.

Splot, wspomniany powyżej, to charakterystyka punktu obrazu obliczana na podstawie pikseli w jego otoczeniu. Przykładowymi operatorami splotowymi są np. filtr uśredniający lub krawędziowy, ale samo pojęcie splotu jest znacznie szersze i umożliwia obliczanie charakterystyk, dla których brak bezpośredniej intuicyjnej interpretacji. W takiej operacji neuron odpowiadający za piksel obrazu w kolejnej warstwie jest połączony z sąsiednimi pikselami warstwy poprzedniej (np. z lewym, prawym, górnym, dolnym sąsiadem). Splot często wykonuje się na wielu warstwach: wówczas każda warstwa ma swój własny filtr, mnożący piksele, a na koniec wszystko sumujemy.



2. A. Ilustracja działania splotu: w celu wyznaczenia piksela wynikowego każdy z pikseli pola sąsiedztwa (każdej warstwy wejściowej) jest mnożony przez odpowiadającą wagę w polu filtra i sumowany z pozostałymi. B. Splot transponowany: piksele sąsiedztwa (czarny punkt) mnożą – każdy z osobną maskę (środkowe tablice) i wynikowe tablice są sumowane. Charakterystyka jednego punktu jest zbiorem 2×2 pikseli – stąd jest to metoda zwiększania rozdzielczości (*upsampling*)



3. Architektura sieci U-net. Pionowe linie to warstwy obrazu o rozdzielczości proporcjonalnej do ich wysokości. Poszerzenie bloku takich linii oznacza pojawienie się nowych warstw (w wyniku zastosowania do wejściowych warstw obrazu wielu różnych operacji splotu, tzw. różnych filtrów)

Po uzyskaniu niskowymiarowej reprezentacji obrazu w postaci różnych charakterystyk – sieć typu U-net stopniowo odtwarza początkową rozdzielczość obrazu (wykorzystując tzw. **splot transponowany**), dokonując przy tym klasyfikacji jego obszarów w oparciu o uzyskane charakterystyki. W miarę wzrostu wymiaru rekonstruowanych warstw schemat sieci przyjmuje symetryczny kształt litery U (3).

Oczywiście – obliczanie miejsc obrazu, które są obciążone szumem, najlepiej się sprawdza na etapach niskiego zaszumienia. Podobnie jak my sami, spoglądając gołym okiem na obraz, widzimy wtedy obce wkłady szumu, niepasujące do obrazu – tak wagi predyktora można wyuczyć dość precyzyjnie. W przypadku dużego zaszumienia rekonstrukcja będzie wprowadzała znaczną przypadkowość. Ta przypadkowość jednak będzie potrzebna w zagadnieniu **generowania** całkiem nowych obrazów. Generator zaczyna od przypadkowego kształtu, a następnie formuje go w coś coraz bardziej rzeczywistego, do czego ten kształt był najbardziej – w rozumieniu generatora – podobny.

Szum wyliczany jest dla pełnego zakłócenia

Warto w tym miejscu zwrócić uwagę, że charakterystyka $\epsilon_\theta(x_i)$ oznacza wypadkowy szum, jaki działał między obrazem x_0 oraz x_i , a nie między x_{i-1} oraz x_i . Teoretycznie moglibyśmy od razu odjąć ten przyczynik od naszego obrazu x_i , aby uzyskać oszacowanie x_0 . Byłoby to jednak oszacowanie bardzo słabej jakości, gdyż siłą modeli dyfuzyjnych jest właśnie etapowa rekonstrukcja, wykonywana małymi krokami.

Generowanie obrazu

Wyobraźmy sobie, że nauczyliśmy nasz model dyfuzyjny możliwie najlepiej zdejmować zaszumienie z obrazu x_i , aby uzyskać x_{i-1} . A jak wygenerować zupełnie nowy obraz, dysponując taką siecią neuronową?

1. Wylosuj realizację w pełni zaszumionego $x_n \in N(0,1)$
2. Wylosuj $z \in N(0,1)$

$$3. \text{ Oblicz } x_{n-1} = \frac{1}{\sqrt{\alpha_n}} \left(x_n - \frac{1-\alpha_n}{1-\alpha_n} \epsilon_\theta(x_n, n) + \sigma_n z \right)$$

4. Jeśli $n>1$, powtórz od punktu 2.

Po wykonaniu tego algorytmu uzyskujemy obraz x_0 , który jest już w pełni odzsumionym obrazem. Obraz taki będzie losowo odpowiadał obrazom ze zbioru uczącego (np. jeśli były w nim zwierzęta, może to być obraz psa, kota czy konia).

Warunkowanie obrazu tekstem

Teoria, którą omówiliśmy do tej pory, umożliwia generowanie obrazu o podobnej treści jak obrazy zbioru uczącego. Ale jak przejść od tego etapu do etapu, na którym AI generuje nam obraz na podstawie opisu tekstowego?

Do opisu danych sformułowanych w **języku naturalnym** najlepiej nadają się sieci neuronowe typu **transformer** (opisywałem je w MT w serii 3 artykułów wiosną 2025). Na wyjściu enkodera transformera, po przejściu

przez bloki uwagi, uzyskujemy bardzo dobrą reprezentację zdania w maszynowej przestrzeni zanurzeń. Nie tylko koduje ona poszczególne słowa, ale też grupuje je w przestrzeni (512-wymiarowej!) w skupiska o podobnym znaczeniu. Ta przestrzeń zanurzenia umożliwia też grupowanie w odrębnych miejscach poszczególnych relacji między przedmiotami oraz przedmiotów obdarzonych pewnymi cechami. Po szczegóły odsyłam do artykułów o transformerach. Dość powiedzieć, że zanurzenia transformera są idealnym warunkiem dla sieci dyfuzyjnej, aby umożliwić jej **odmienne** trenowanie współczynników zależnie od wprowadzonego kontekstu!

Np. jeśli oprócz danych pikselowych obrazu dodamy informację, że obraz przedstawia „drzewo” lub „samochód”, to sieć neuronowa nauczy się w wyspecjalizowany sposób rekonstruować drzewa i samochody. Jeśli na etapie generowania losowo zainicjalizujemy piksele obrazu, ale przedstawiając etykietkę „drzewo”, to generator użyje wag, jakich nauczył się na zbiorze próbnym obrazów przedstawiających drzewa – a nie samochody! I wygeneruje drzewo.

Jest to krok we właściwym kierunku, ale takie proste warunkowanie to za mało, aby – na przykład – wygenerować obraz „samochodu pod drzewem”. Potrzeba jakoś takie relacje – i obiekty w nich występujące – zakodować w sposób czytelny dla maszyny cyfrowej. I okazuje się, że najlepszą do tego celu jest reprezentacja wychodząca z mechanizmów uwagi sieci typu transformer.

Co więcej – zamiast opisanego wyżej warunkowania – bardziej elegancko problem można rozwiązać za pomocą **mechanizmu uwagi krzyżowej** – również znanego z transformerów. Wówczas, jeśli przypomnimy sobie proces tłumaczenia tekstu – nie wchodząc w niepotrzebne szczegóły – reprezentację zakodowanego tekstu możemy potraktować jako „zdanie w języku źródłowym”, a generowany krok po kroku (!) obraz jako „zdanie w języku docelowym”. Jest to dokładnie zagadnienie translacji, dla którego oryginalnie stworzono transformery.

W tym ostatnim podejściu jednakże konieczne jest jeszcze przejście z obrazem od reprezentacji pikselowej do reprezentacji wewnętrznej (ang. *latent space*), gdzie zamiast zbioru pikseli dysponujemy zbiorem właściwości obrazów (np. opisując obraz człowieka, moglibyśmy na nasze potrzeby zaproponować wektor określający kąty ugięcia stawów, kolor włosów, rodzaj ubrania, itd.).

Taka zmiana reprezentacji to temat na osobny artykuł, lecz w praktyce przejście do reprezentacji wewnętrznej opiera się na sieci neuronowej typu **autoenkoder** – jest to sieć, która w kolejnych warstwach zmniejsza liczbę neuronów (od liczby równej wymiarowości obrazu do liczby odpowiadającej wymiarowi zredukowanej przestrzeni), a sieć taką trenuje się tak, aby umożliwiała rekonstrukcję obrazu źródłowego. Dopóki takie wytrenowanie jest możliwe, dopóty warstwa zawężona zawiera wszystkie istotne cechy obrazu.

Dodatek

Indukcyjny dowód relacji dla postaci $q(x_n|x_0)$. Sprawdzenie dla $n=1$ z definicji $q(x_n|x_{n-1})$; dowód do pobrania ze strony: <https://tiny.pl/cctpbsz9z> lub po zeskanowaniu kodu QR. ■

Przemysław Borys

Wydział Chemiczny Politechniki Śląskiej w Gliwicach



Literatura

1. J. Ho, A. Jain, P. Abbeel, Denoising Diffusion Probabilistic Models, 34th Conference on Neural Information Processing Systems (NeurIPS 2020), Vancouver, Canada [sformułowanie modeli dyfuzyjnych w generowaniu obrazów]
2. J. Sohl-Dickstein, E. A. Weiss, N. Maheswaranathan, S. Ganguli, Deep Unsupervised Learning using Nonequilibrium Thermodynamics, Proceedings of the 32th International Conference on Machine Learning, Lille, France, 2015. JMLR: W&CP volume 37 [preliminaria termodynamiczne modeli dyfuzyjnych]
3. R. Rombach, A. Blattmann, D. Lorenz, P. Esser, B. Ommer, High-Resolution Image Synthesis with Latent Diffusion Models, 2022 IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR) [sformułowanie popularnej implementacji modelu dyfuzyjnego – Stable Diffusion]
4. O. Ronneberger, P. Fischer, T. Brox, U-Net: Convolutional Networks for Biomedical Image Segmentation, arXiv:1505.04597, 2015 [architektura U-Net]
5. Y. Song, J. Sohl-Dickstein, D. P. Kingma, A. Kumar, S. Ermon, B. Poole, Score-Based Generative Modeling Through Stochastic Differential Equations, ICLR 2021.
6. K. Erdem, Step by Step Visual Introduction to Diffusion Models, Medium 2023
7. D. P. Kingma, R. Gao, Understanding Diffusion Objectives as the ELBO with Simple Data Augmentation, NeurIPS 2023
8. C. Zhang, Cha. Zhang, M. Zhang, I. S. Kweon, J. Kim, Text-to-image Diffusion Models in Generative AI: A Survey, arXiv:2303.07909v3 [cs.CV] 08 Nov 2024
9. P. Cao, F. Zhou, Q. Song, L. Yang, Controllable Generation with Text-to-Image Diffusion Models: A Survey, IEEE Transactions on Pattern Analysis and Machine Intelligence, 2024

FANTASTYKA NAUKOWA W „MŁODYM TECHNIKU”

POLSKA FUNDACJA
FANTASTYKI NAUKOWEJ

m.technik

W tym numerze „Młodego Technika” publikujemy opowiadania nagrodzone w konkursie literackim Polskiej Fundacji Fantastyki Naukowej pod patronatem „Młodego Technika” dla uczniów szkół średnich. Konkursowe opowiadania nadsyłane były na adres PFFN od 1 marca do 31 maja 2025 r. Latem 2025 r. jury wybrało zwycięzców.

Celem konkursu było promowanie wśród młodzieży zainteresowania nauką i literaturą w konwencji fantastyki naukowej, a także popularyzacja Dnia Polskiej Fantastyki Naukowej, ustanowionego przez Polską Fundację Fantastyki Naukowej w dniu 14 lipca 2024 r. w 150. rocznicę urodzin Jerzego Żuławskiego.

Konkurs przebiegał w dwóch etapach. W pierwszym etapie szkoła ustalała własny tryb wyłaniania pracy konkursowej i przysyłała jedną najlepszą pracę o maksymalnej objętości ok. siedmiu tysięcy znaków ze spacjami do organizatora. W drugim etapie jurorzy oceniali nadesłane prace i wybierali trzy najlepsze utwory w skali ogólnopolskiej. Nadesłane prace oceniał zespół jurorów w składzie: dr inż. Adam Froń, dr Marcin Kowalczyk i Marta Magdalena Lasik. Jurorzy przyznawali każdemu opowiadaniu indywidualną ocenę w zakresie od 0 do 10 punktów. Ocena końcowa stanowiła sumę punktów.

Zespół jurorów przyznał najwyższą liczbę punktów następującym opowiadaniom:

- „Ludzie... Ludzie!”, Bartłomiej Oreńczak, 24 pkt.
III Liceum Ogólnokształcące im. Mikołaja Kopernika w Wałbrzychu
- „Marsjanie czy jeszcze ludzie”, Michał Miasek, 20 pkt.
Liceum Ogólnokształcące Niepubliczne nr 40 w Warszawie
- „Zlecenie”, Hanna Bułat, 17 pkt.
X Liceum Ogólnokształcące im. KEN w Krakowie

Autorzy otrzymali nagrody rzeczowe od Wydawnictwa IX w postaci płyty muzycznej „Commentarii Lunares: In memoriam Jerzy Żuławski” oraz książki „Na srebrnym globie” Jerzego Żuławskiego z dedykacją od redaktora naczelnego patronującego konkursowi miesięcznika „Młody Technik”, w którym publikowane są, w numerze styczniowym 2026 roku, zwycięskie teksty.



Zlecenie

Była ciemna, deszczowa noc. Kojot spał niespokojnym snem w swoim mieszkaniu. Śniło mu się poprzednie zlecenie. Widział twarz swojej ofiary, kiedy ją zabijał, a był to straszny obraz. Zerwał się z łóżka zły po tym, z przyspieszonym oddechem. To nie był zły człowiek, pomyślał. Przesunął dłonią po twarzy i wstał. Był to postawny mężczyzna około czterdziestu lat z licznymi bliznami. Podeszedł do szafy i, wybierając ubranie, powiedział do białego pudełka na stole:

- Komputer, czy są jakieś wiadomości?
- Trzy: powiadomienie o przysłaniu zapłaty za ostatni kontrakt, reklama nowego syntetyzatora jajek i zaproszenie na rozmowę z prezesem firmy Humanux o pierwszej.
- Będę – mruknął. Zarzucił plecak na ramię i podeszedł do drzwi. Gdy wychodził, komputer zawołał za nim:
- A syntetyzator mam kupić, bo coś marnie ostatnio wyglądasz?
- Komputer, spadaj – odparł i wyszedł.

Kiedy znalazł się na zewnątrz, wsiadł do swojego antygrawu, zaczekał chwilę, aż pojazd dostosuje się do jego budowy ciała i podryfował do firmy Humanux. Kiedy przybył, miał jeszcze dużo czasu do spotkania. Postanowił go wykorzystać na dokładne zbadanie tego miejsca. Systematycznie penetrował każdy możliwy poziom, ale gdy chciał zobaczyć, co jest w podziemiach, ochroniarze nie dali mu przejść. Pewnie spierałby się z nimi, ale za parę minut miał spotkanie, więc udał się w kierunku biura prezesa.

Kiedy stał już pod drzwiami, podeszedł do niego jakiś mężczyzna w garniturze i poprosił, aby poszedł z nim. Po drodze ów człowiek wyjaśnił mu, że prezes miał bardzo ważną rzecz do zrobienia i nie mógł przyjść.

Na to Kojot odparł:

- Skoro pan prezes taki zajęty, to niech nie zawraca głowy innym zajętym.
- Ależ szanowny panie, właśnie dlatego wyznaczył mnie, abym przekazał panu wszystko.
- No to streszczaj się pan, bo nie mam całego dnia.
- Oczywiście, pójdźmy do mnie, tam jest znacznie spokojniej – powiedział przewodnik Kojota i zaprowadził go do swojego gabinetu. Tam, zamknąwszy drzwi, usiedli, a człowiek w garniturze przeszedł do konkretów. – Jestem zastępcą prezesa tej firmy, pana Henryka Vitrowskiego. Istnieje pewna sprawa, dość delikatna. Chodzi o człowieka, który posiada zagrażające naszej firmie informacje. Pana zadaniem byłoby usunięcie tej przeszkody. Co pan na to?

– O jakie informacje chodzi?

– Takie, które stanowią zagrożenie dla wizerunku firmy, więcej nie mogę powiedzieć.

– Potrzebuję jego imienia, nazwiska, cokolwiek, co pomoże mi go wytropić.

– Wiemy tylko, że nazywa się Adam Liberat.

– Pan żartuje? Jak pana zdaniem mam to zrobić?

– Wiemy, że wczoraj wieczorem widziano go w barze Aleksandria. Podobno jest pan najlepszy.

– Jestem.

– To dobrze, wykonaną pracę proszę udokumentować zdjęciem. Jaka stawka?

– Podwójna. Zlecenie jest bardzo trudne.

– Dobrze. Potrzeba panu czegoś jeszcze?

– Wstępu do podziemia.

– Da się zrobić.

Kojot wstał i wyszedł. Spróbował ponownie wejść na poziom minus jeden i tym razem ochrona go nie zatrzymała. Kiedy znalazł się w środku, zrozumiał, czemu nie chcieli go wpuścić. To tu mieściła się fabryka ludzi. Procedura polegała na tym, że każdy, z odpowiednim budżetem, mógł sobie zamówić dziecko, które można było odebrać już kilka dni później. Kojot nie pochwalał projektowania potomstwa – uważał to za coś nienaturalnego.

Rozmyślając, uważnie studiował otoczenie. Miał w plecaku kilka ładunków wybuchowych i zastanawiał się, jak je rozmieścić, aby w razie problemów z wypłatą wynagrodzenia wyrządzić maksymalne szkody. W tym pomieszczeniu wystarczyło podłożyć je w kilku miejscach, aby zburzyć cały budynek. Kojot dyskretnie rozmieścił ładunki, które od razu zakamuflowały się hologramem.

Kiedy skończył, podryfował do Aleksandrii. Tam zapytał o Liberata, ale nawet za dużą sumę barman nie chciał sobie przypomnieć. Większość ludzi w takiej sytuacji poddałaby się, ale nie Kojot. Dziwnym zbiegiem

okoliczności cały personel dostał zbiorowej amnezji, jakby ów człowiek wcześniej im zapłacił. Trafiłem na godnego przeciwnika, pomyślał.

Szczęście jednak się do niego uśmiechnęło, bo idąc korytarzem, usłyszał fragment rozmowy między pracownikami, z którego jasno wynikało, że Liberat zatrzymał się w hotelu Wawel. Kojot przybył tam godzinę później, ale okazało się, że już od pół dnia jego ofiary tam nie ma. Sytuacja z personelem powtórzyła się. Nie podoba mi się to, pomyślał.

Wiedząc, że niczego już się tu nie dowie, wyszedł. Wsiadając do swojego antygrawu, zobaczył, że na fotelu obok leży jakaś koperta. Zmarszczył brwi. Obecnie były inne metody komunikacji, a papier należał do przeżytku. Otworzył ją. W środku znajdował się krótki list, a kiedy na niego spojrzał, doznał czegoś, czego nie czuł już od dawna – strach. Nie chodziło o wiadomość, tylko o charakter pisma – był niezwykle podobny do jego własnego.

Zmuszając się, przeczytał wiadomość od Liberata. Napisał, że chce porozmawiać, a jako miejsce wybrał opuszczony magazyn dwie przecznice dalej.

To spotkanie mogło być jego jedyną szansą na wypełnienie zadania, ale jednak cała sprawa robiła się coraz dziwniejsza. I nie chodziło tylko o niezwykłą zbieżność charakterów pisma. Kojot cały czas nie mógł się pozbyć wrażenia, że jego ofiara ciągle jest krok przed nim.

Po chwili namysłu zdecydował się stawić na spotkanie. Kiedy znalazł się przed odpowiednim magazynem, wyciągnął broń laserową i ostrożnie wszedł. Po kilku krokach zamajaczyła przed nim niewyraźna postać.

– Witaj, już dawno chciałem cię spotkać.

– Kim jesteś?

– Imię nie jest istotne – powiedziała postać i weszła w krąg światła. Kojot poczuł, jak kolana się pod nim uginają. Stał przed nim on sam, tylko bez blizn i w innym ubraniu.

– Jesteś mną? – zapytał.

– Niezupełnie. Nasi, twoi rodzice skorzystali z usług firmy Humanux. Przez pewne komplikacje jest nas dwóch, nie jeden. Skończyłbym na śmietniku, gdyby pewnej pracownicy nie poruszyło sumienie i się mną nie zaopiekowała.

– Który z nas jest prawdziwy?

– A jak ci się wydaje?

– Chyba... ja.

– Może ty, może ja, może obaj jesteśmy, a może żaden z nas.

Po tych słowach nastąpiło dłuższe milczenie.

– Wiesz, że mam na ciebie zlecenie?

– Wiem. Zabijesz mnie?

– Nie wiem, nic już nie wiem! – wykrzyknął Kojot. Spróbował się uspokoić. – Ale wiem, że chcę odpowiedzieć, a udzieli ich pan prezes.

Poszli z Liberatem do jego antygrawu i podryfowali do siedziby Humanuxa. Kiedy dotarli na miejsce, nie pukając, weszli do gabinetu prezesa. Tym razem obaj stanęli jak wryci. Za biurkiem siedział mężczyzna identyczny z nimi, tylko otyły. Pierwszy otrząsnął się Kojot.

– Mam tego serdecznie dość! Albo mi to ktoś natychmiast wytłumaczy, albo nie rękę za siebie! – wykrzyknął i wyciągnął z plecaka małe, czarne pudełko.

– Jesteście klonami z mojego DNA, taka specjalna usługa naszej firmy.

– Ale czemu jesteś w naszym wieku? – zapytał Liberat.

– To mój kolejny klon z moją świadomością. Rozczarowałeś mnie, Kojocie, myślałem, że twoja reputacja nie kłamie.

– Ja zawsze wywiązuję się z kontraktów. Tylko widzisz, jest pewien problem: miałem pozbyć się Liberata, ale to nie jest jedna osoba.

– Co przez to rozumiesz? – zapytał ostrożnie prezes.

– Adam Liberat to trzy osoby. Są w tym pokoju, a ja mam możliwość skończenia zadania. – To powiedział, spojrzał w oczy Adamowi i rzekł: – Wybacz mi, ale nie mogę żyć, wiedząc, że nie jestem prawdziwy.

Kojot zamknął oczy, pojedyncza łza spłynęła mu po policzku. Wcisnął guzik detonatora. Wcześniej rozmieszczone ładunki eksplodowały, a cały budynek w kilka sekund runął. ■

Hanna Bułat

Marsjanie czy jeszcze Ludzie?

Czerwony pył unosił się, przysłaniając widok na Olympus. Kapitan Eliaz Mazur oparł się o transporter, trzymając w rękach karabin. Ludzie od paru lat przylatywali na Marsa jak pasożyty. Najpierw badacze, potem żołnierze, aż w końcu całe bataliony. Kolonizacja? Nie, to było coś więcej. Oni chcieli go zdobyć. Ale Mars nie był pusty. Marsjanie, z którymi kolejne starcia okazywały się trudniejsze od poprzednich, nie byli prymitywnymi zwierzętami, jak głosiła propaganda na Ziemi.

Mazur sądził, że znajdują tu tylko czerwone skały i może jakieś mikroorganizmy. Zamiast tego natrafili na nich – wysokich, smukłych, o ciemnej skórze i wielkich oczach, wydających niezrozumiałe, basowe dźwięki, często przy nich gestykulując. Kapitan spojrzał na niebo. Bez trudu odnalazł Ziemię. Była wystarczająco daleko, by nikt tam nie przejmował się losem żołnierzy.

– Widzę ich! Kierują się w waszą stronę. – Głos sierżanta trzeszczał w komunikatorze.

Eliaz wziął głęboki oddech i wstał. Dookoła rozciągała się jedynie martwa, czerwona pustka.

– Przecież tu nikogo nie ma! – wrzasnęła, zirytowana.

– Są tuż przy... – Głosnik znów zaskrzeczał, ale dźwięk urwał się w połowie.

Uwagę Mazura odwróciło pikanie – spojrzał na skanery, które wykazywały ruch pod jego pozycją. Nagle ziemia obsunęła się, a on z drużyną spadł w dół. Gdy uderzył o twardą, chłodną powierzchnię, echo jęków i przekleństw odbiło się od ścian. Mazur włączył latarkę i próbował wstać. Oparł się o ścianę – była gładka, sztucznie obrobiona.

– To nie jest naturalna jaskinia... – wyszeptał.

Na ścianach tunelu widniały symbole. Nie były przypadkowe. Część linii, choć wyblakłych, układała się w słowo NASA. Studiowanie napisów przerwały odległe kroki. Wszyscy skierowali broń w głąb tunelu. Z mroku wyłoniły się wysokie, smukłe sylwetki. Marsjanie długo wpatrywali się w żołnierzy, aż jeden z nich uniósł dłoń w literę V. Nowak ścisnął mocniej broń, ale coś sprawiło, że się zawahał.

– Ne pafu – wybełkotał, ale było już za późno. Ktoś spanikował i pociągnął za spust. Salwa wystrzałów rozzerwała ciszę tunelu. Światła latarek zatańczyły po ścianach, a Marsjanie szybko rozplynęli się w ciemności.

Ich głos... nie brzmiał jak przypadkowa zbitka dźwięków. To coś próbowało mówić... po ludzku? Czy oni uczyli się ziemskiego języka? Ale od kogo? Ponownie rozbrzmiały kroki – tym razem była to cała armia Marsjan. Mazur spojrzał na swój oddział. Nie był gotowy na to, co się może wydarzyć.

Korytarz zapełniła setka wysokich postaci. Większość z nich trzymała zardzewiałą broń, podobną do ziemskiej.

– Kapitulacu vin! – odezwał się Marsjanin w zakurzonej czapce oficerskiej.

Nowak wziął głęboki wdech, próbując opanować narastający strach. Wiedział, że nie ma szans w walce, po woli opuścił broń. Marsjanin zdziwił się, jakby nigdy nie widział poddających się ludzi.

– Sojourner, Curiosity, ligas ilin – dowódca Marsjan wywołał dwóch swoich kompanów.

Wyszli z szeregu, w rękach trzymając liny. Eliaz nie próbował się opierać. Gdy wszyscy ludzie byli już splątani, Marsjanie ruszyli tunelami, ciesząc się ze zwycięstwa. Kapitan szedł, rozglądając się w zamyśleniu. Curiosity? Sojourner? Nazwy wydawały się znajome. Tylko czemu te istoty używały ich jako imion? Czyżby ich język aż tak nie różnił się od ludzkiego?

Po długiej wędrówce stanęli przed wrotami, ozdobionymi różnymi znakami i kolejnymi ziemskimi literami ESA. Brama otworzyła się, a ich oczom ukazała się cała podziemna metropolia. Niemożliwe... jak oni to zrobili? Przecież Marsjanie mieli być niższą formą życia!

Zaprowadzono więźniów do obszernej sali na środku miasta, w której siedziało dziesięciu Marsjan. Byli inni – mieli jaśniejszą karnację oraz wyglądali na niższych i starszych. Istota w czapce oficerskiej ukloniła się i zaczęła coś tłumaczyć. Starszyzna siedziała i spoglądała co jakiś czas na ludzi.

W końcu jeden z nich, najstarszy, wstał i podszedł do Mazura. Kapitan cofnął się lekko. Był przerażony, a z drugiej strony zaciekawiony. Marsjanin zdjął z szyi mały srebrny medalik, podając go człowiekowi. Widniała na nim postać kobiety oraz jakieś słowa.

– Maryjo... módl się... – zaczął czytać wytartą inskrypcję. Przerwał, a jego oczy rozszerzyły się. To była inskrypcja z Ziemi. Przypominała słowa zapomnianej modlitwy, której się kiedyś nauczył.

Marsjanin wskazał na sufit. Żołnierze podnieśli głowy, a ich oczom ukazał się malunek z ich rodzimą planetą. Tylko inną – Afryka była połączona z Azją, a zamiast wielkiego Oceanu Europejskiego widniał jakiś kontynent.

Tero... du unu du kvin... Dwa. Jeden. Dwa. Pięć. O co chodzi? Czy te liczby oznaczają rok? Niemożliwe. Oznaczałoby to, że wiedzą o Ziemi od ponad trzech tysięcy lat.

Marsjanin spostrzegł zdziwienie i zakłopotanie ludzi. Odwrócił się i zaczął dyskutować z resztą starszych. Po długiej wymianie zdań dziesiątka starszyny zebrała się na środku sali. Spod pyłu odgrzebali starą klapę. Żołnierze weszli do katakumb, rozglądając się i oświetlając pomieszczenie latarkami. Pokój był zawałony starymi pudłami i papierami, gdzieś wieszali różne plany i zdjęcia.

Mazur przetarł ramkę starego obrazu, który tam leżał. Na zdjęciu widniała grupa ludzi w bazie, a za nimi rozciągał się widok na Marsa. Na odwrocie kartki było napisane: „Założa badawcza. Dnia 27 czerwca 2178 roku”. Jak to możliwe? Przecież na Ziemi mówili, że nikt nie był na Marsie przed nimi.

Kapitan zaczął czuć się przygnieciony tymi wszystkimi pytaniami, na które brakło odpowiedzi. Stary Marsjanin podszedł do panelu kontrolnego i sprawnym ruchem uruchomił go. Ekrany ożyły, migocząc niebieskim blaskiem. Pojawiły się na nich archiwalne materiały przedstawiające Marsa. Głosy w tle były radosne z postępów i odkryć, jakich dokonano.

Potem coś się zmieniło – barwy głosów sugerowały zdenerwowanie. Mówiły o wojnie, która pochłonęła Ziemię, o tym, że komunikacja zamilkła, a statki przestały przylatywać. Nagrania stawały się coraz bardziej chaotyczne. Błagania i przekleństwa rzucane w eter: „Wróćcie po nas!”, „To nie miało tak być!”, „Zdrajcy! Zostawili nas!”. Aż wreszcie został tylko szum zakłóceń.

Na ostatnim filmie pojawił się stary człowiek. Jego oczy były zmęczone, a głos dziwnie spokojny.

– Ludzie zaczęli się przyzwyczajać do Marsa. Niestety, deszcz asteroid zmusił nas do zejścia pod ziemię – tłumaczył powoli. – Zaczęliśmy nawet... to znaczy dzieci... te urodzone na Marsie zmieniły się... Ewolowały! – Zawahał się, jakby czuł ciężar własnych słów. – Nieprawdopodobne, nieprawdaz? Początek nowego gatunku człowieka... Homo martianus! – Po tych słowach ekrany zgasły.

Myśli Mazura wirowały jak szalone. Jak mogliśmy tego nie dostrzec? Te korytarze, te symbole, ten język... To nigdy nie była obca cywilizacja. To byli cały czas... ludzie. Poczuł ciężar tego odkrycia. Ten konflikt to nie walka z obcą rasą, tylko wojna bratobójcza. Nieporozumienie, zapomnienie własnej historii.

Ale czy ludzie na Ziemi kiedykolwiek uznaliby ich za swoich? Nie. Oni nie zobaczą w nich ludzi – jedynie podludzi, mutanty, które trzeba zgładzić. Co, u diabła, mają teraz zrobić? Czy którakolwiek droga, którą podejmą, nie będzie prowadzić do kolejnej tragedii? Ktokolwiek jest w stanie przewidzieć konsekwencje tych decyzji? ■

Michał Miąsek

Ludzie... Ludzie!

Na początku jest ciemno. Zupełnie tak, jakby wszystkie fotony dawno uciekły. Sytuacja poprawia się jednak z czasem. Przynajmniej na chwilę.

Oczy dostrzegają w tej ciemności delikatne zarysy komory. Ciało tak strasznie pragnie zerwać się do biegu. Oddech znów staje się naturalny, choć pierwsza myśl zapiera mu dech w piersiach.

Gdzie jestem? – pyta siebie. To pytanie rodzi w nim kolejne, dość podobne. Kim jestem?

– Nazywasz się Frank Sobun – jak na zawołanie odpowiada mu zimny, cyfrowy głos.

Frank Sobun, powtarza w myślach. Tak, już pamięta. Niestety pamięta.

Nagle przez matową szybę wpadają zaginione fotony. Frank mruży oczy. Głośnie trzaski, mija chwila i klapy nie ma już nad jego twarzą.

Frank podnosi ociężałe ciało i wygląda przez palce na otaczający go świat. Z pięciu niemal identycznych komór wychylają kolejne ciała. Rozglądają się. Mrużą oczy. Świat otaczający Franka i resztę sześcioosobowej załogi kończy się na metalowej puszcze, którą przemierzają próżnię.

Trzy kobiety, trzech mężczyzn. Wszyscy milczą. Obudzili się, a to znaczy, że prawie dotarli do celu. Na siłę próbują zachować obojętność, ale melancholia bije z ich oczu. Wiedzą, że nigdy nie wrócą do tego, co zostawili, wsiadając do statku.

* * *

Komputer pokładowy piszczy. Angelo, ciemnowłosy okularnik, próbuje wyczytać z niego coś konkretnego. Co jakiś czas gładzi nerwowo swoje krótkie, acz bujne loki. Cała załoga – Frank, Isa, Ash, Kai i Maya – pochyłają się nad Angelem.

– Co mówi? – pyta Frank.

– Że na Ziemi II rozwinęła się cywilizacja.

Niewiele to zmienia w ich planach. Nie są w stanie nawiązać komunikacji, więc ktoś musi tam wylądować. Zgodnie z wytycznymi na planetę pierwsi mają zejść Frank i Angelo. Wsiadają do kapsuły, zamykają drzwi, zapinają pasy. Reszta dołączy do nich, gdy będzie pewność, że Ziemia II jest bezpieczna. Bezpieczna dla nas, dla cywilizacji, dla rozwoju, które jest ostatnią nadzieją na ocalenie gatunku ludzkiego.

Angelo ciągnie wajchę. Kapsuła odrywa się od statku i za chwilę rusza z pełną mocą silników ku Ziemi II.

* * *

Frank się budzi. Jego myśli zaćmiewa ból. Czerwone światło rozświetla kapsułę i razi jego oczy. Ma pewność, że razem z Angelem dotarli na Ziemię II, ale nie pamięta lądowania. Angelo dochodzi do siebie dość szybko. Razem zabierają się za otwarcie wrót kapsuły. Te z początku nie dają za wygraną, ale w końcu mechaniczne zawiasy puszczają i do środka wpada blask.

Razem z nim docierają również odgłosy: szoku, przerażenia, szczęścia i wrogości. O dziwo, wydają się znajome. Zrozumiałe.

Angelo cofa się, gdy słyszy te dźwięki. Brzmia zbyt ludzko. Podchodzi do gablotki, rozbija łokciem szkło i wyjmując karabin plazmowy. Frank sam musi unieść drzwi do końca. Robi to, choć niechętnie.

Blask Heliosa oślepia go, zasłania oczy rękoma. Przez palce dostrzega ruchome kształty, źródło głosów. Pierwszy wdech nowego powietrza jest niewyobrażalnie rześki, błogi. Nie chce w to uwierzyć, nie może w to uwierzyć, ale jego głowa cały czas krzyczy jedno:

Ludzie! Ludzie!

Jedni wiwatują, inni unoszą ku nim jakieś prostokąty. Za nimi majaczą wysokie wieże. Frank zaniemówił, Angelo również. Nadejżdża jakaś czterokołowa maszyna. Ostre światło zmusza Franka do zamknięcia oczu. Coś nagle wbija się w jego kark.

* * *

Wpierw Frank panikuje. Rzuca się, wije. Jest w obcym miejscu, na obcej planecie, a jednak wszystko wygląda znajomo. Ludzie! Tu są ludzie! Później dwójka lekarzy podaje mu mocne środki i mężczyzna się uspokaja.

Mierzą go wzdłuż i wszerz, ważą, prześwietlają. Montują nieznaną urzędniczą sztycę. Wszystko to jakimś prehistorycznym sprzętem. Zachowują się normalnie. Tak jakby codziennością było, że ktoś spada z nieba.

Lekarz mówi nagle coś niezrozumiałego w obcej mowie.

– Ciesz się, że spadliście do nas – tłumaczy jakiś inny człowiek. – Są tacy, co was rozcinają. To właśnie oni zestrzelili wasz statek. Reszta załogi martwa... Naczynie rozwojowe zniszczone...

Nikt nie tłumaczy Frankowi, gdzie jest albo co się dzieje. Nikt nie zwraca na niego szczególnej uwagi.

Gdy badania dobiegają końca, lekarze i wojskowi pakują Franka do pojazdu. Jest w nim więcej ludzi, tak samo zmęczonych jak on. Jest tam również Angelo. Po wielu godzinach drogi docierają do wielkiego kompleksu. Jakiś mężczyzna przy kości mówi bardzo dużo niezrozumiałych rzeczy. Tłumacz przekłada te słowa, jednak niewiele dociera do Franka.

Frank i inni zostają wrzuceni pod prysznic. Nie protestują. Później wychodzą do wielkiego pomieszczenia. Jest tam więcej ludzi, takich jak Frank. Takich jak Angelo. Jedni po prostu siedzą. Inni jedzą, rysują, gapią się w wielkie świecące pudła, biegają albo śpią. Wszędzie wokół są lustra, ale Frank czuje, że ktoś ich obserwuje.

Leki powoli puszczają, cała grupa nowych lokatorów doznaje załamania. Wtedy urządzenie na ich szyjach dostarcza im kolejne środki. Ten cykl powtarza się wiele razy, ale w końcu ustaje. Stają się bierni.

Frank szybko zapoznał się z tym miejscem. Lustra co jakiś czas gasną i po drugiej stronie pojawiają się wtedy setki ciekawskich oczu. Mieszkańcy Ziemi II obserwują, robią zdjęcia, pukają w szybę, mimo wyraźnego zakazu.

Kanciasta puszka to telewizor i stąd biorą się wiadomości. Jest wiele telewizorów, bo lokatorzy mówią w wielu językach. Frank siada przy tym, który wyświetla napisy w jego ojczystym. Siedzi z nim Angelo. Obaj słuchają tego, co mówi kobieta w garniturze.

Ludzie... Ludzie! – ma w głowie Frank, gdy słucha tego wszystkiego:

Tego, że lodowce prawie stopniały.

Tego, że jakiś człowiek chce znieść ministerstwo klimatu.

Tego, że gigant atakuje kolejne państwo.

Tego, że ktoś chce wszczepiać ludziom czipy do mózgów.

Tego, że ktoś ośmiesza się w Internecie, by zarobić fortunę.

– Chcesz? – pyta jakiś facet siedzący obok Franka. Podaje mu miskę z czipsami. Facet śmieje się co chwila z tego, co słyszy.

– Nie – odpowiada Frank.

– Jak wolisz – bierze chipsa do ust. – Jakie to pyszne!

* * *

Frank zadomowił się w zoo. Jest tu już chyba rok. Jego misja ocalenia ludzkości legła w gruzach. Skończyła się tak samo, jak każdego innego lokatora. Wybudowali im nową część zoo. Bardziej rozrywkową. Frank udaje się tam z Anielem dwa razy w tygodniu. Alkohol, piękne kobiety...

Tam też są lustra, a że w telewizji dla lokatorów co chwila pojawiają się reklamy tej części Zoo, Frank doszedł do wniosku, że popyt musi być. Ludzie płacą, by oglądać.

Pewnego dnia przybywa dostawa nowych lokatorów. Jest wśród nich David. Środki uspokajające nie działają na niego dobrze. Już pierwszego dnia biega po ośrodku i krzyczy, nawołuje, ale nikt już nie pamięta.

Frank ignoruje go większość czasu, oglądając telewizję, jedząc czipsy. Przytył przez ostatni rok jakieś 60 kilo.

– Naprawdę cię to śmieszy? – pyta go kiedyś David, spoglądając w telewizor.

– Spokojnie, ciebie też zaczniesz – mówi Frank, śmiejąc się.

– Ale oni skończą tak jak my.

Franka ciekawi to, co słyszy.

– Co masz na myśli?

– Kiedyś ich Ziemia będzie tak samo niezdatna do życia, jak Ziemia każdego z nas.

Frank myśli chwilę.

– Twoja teoria może mieć sens. Taka nieskończona pętla – śmieje się. – Ale co z tego? Nasze domy już nie istnieją. A im można mówić, tak jak można było mówić naszym przodkom. A może... za sto lat wyślą ekspedycję na nową Ziemię?

Frank śmieje się znowu i obraca w stronę telewizora. Po raz kolejny nadawany jest na żywo strajk tych, którzy mówią, że zoo są nieetyczne, że ci ludzie zasługują na wolność i lepsze życie. Wszyscy przy telewizorze, poza Davidem, wybuchają śmiechem.

Chwilę później głowę Franka zajmuje już tylko: Ludzie! Ludzie! Jakie te czipsy są zaje... ■

Bartłomiej Oreńczak

Ocaleni 2.0

Tomasz Wilk, Joanna Zając, Szymon Szymański, Ignacy Stankiewicz, Tymoteusz Adam Leonard Sokołowski, Jakub Sokołowski, Małgorzata Sanchez, Andrzej Piętka, Michał Paszkowski, Mikołaj Marks, Sebastian Kmiecik, Maciej Ciembroniewicz, Grzegorz Blinowski

Wydawnictwo Wilki Trzy i Kot, liczba stron: 392, cena sugerowana: 49,90 zł

Poszukiwacz ocalonych historii, wracasz do świata, w którym przetrwanie wciąż wymaga odwagi, a człowieczeństwo nabiera nowych znaczeń. W drugiej antologii *Ocaleni* spotkasz zarówno autorów znanych z ubiegłorocznego tomu – naszą starą literacką watahę – jak i nową krew: tegorocznych finalistów konkursu Polskiej Fundacji Fantastyki Naukowej, wyłonionych i rekomendowanych przez Łukasza Marka Fiemę. To właśnie z połączenia doświadczenia i świeżego spojrzenia powstały opowieści o granicach poznania, odwadze i decyzjach, które kształtują to, kim jesteśmy – niezależnie od czasu, miejsca i świata, w którym przychodzi nam żyć. Zanurz się w te opowieści. Daj się ocalić – ponownie.

<https://www.wilkitrzyikot.pl/produkt/ocaleni-2/>



Spotify?

Śmierć przez...

dochodowość

Spotify kończy się (1) albo... wręcz przeciwnie. Zależnie od tego, gdzie przyłożymy ucho i jakie wiadomości wybierzemy, możemy zetknąć się z przeciwstawnymi opiniami. Dla jednych to pewno koniec ich przygody z muzycznym serwisem streamingowym. Dla innych – nowa era Spotify.

Co się stało? Od takich pytań zaroilo się na forach i blogach internetowych w 2024 r. Na hasło „koniec Spotify” z łatwością znajdziemy w Google mnóstwo wpisów krytykujących, gorzko podsumowujących przygody użytkowników serwisu i w końcu żegnających się „na dobre” z platforma streamingową. Jednym z najbardziej znanych i cytowanych głosów w tej fali był artykuł Kyle’a Chayki „Why I Finally Quit Spotify” („Dlaczego w końcu zrezygnowałem ze Spotify”) opublikowany w magazynie „New Yorker”.

Serwis ma się dobrze, ale...

Krytyka zawarta w tym tekście a także w wielu innych dotyczyła zmian w interfejsie, architekturze i mechanizmie wyszukiwania muzyki w Spotify. Użytkownicy narzekali, że się gubią w gąszczu zmian i trudno im znaleźć taki rodzaj muzyki, który im odpowiada. Jednak zmiany, które wprowadzono na platformie streamingowej, miały swoje uzasadnienie biznesowe. Właściciele i administratorzy, jak się uważa, musieli coś zmienić, bo od 2021 roku Spotify było na poważnym minusie.

W 2023 Spotify wrócił do rentowności. Po raz pierwszy od dwóch lat. Po drodze zredukował zatrudnienie o 17 proc. i znacznie ograniczył wydatki marketingowe. Podniósł również ceny. Przez lata miesięczna subskrypcja Spotify wynosiła równowartość 9,99 dolara. Tak było do lipca 2023 r., kiedy to ceny wzrosły o dolara na głównych rynkach. W 2024 r. ceny ponownie wzrosły o dolara. Cięcia kosztów i podwyżki cen przyniosły efekty. Liczba aktywnych użytkowników miesięcznie wzrosła w 2023 r. o 14 proc. rok do roku do 626 milionów. Liczba abonentów wzrosła o 12 proc. do 246 milionów.



1. Symboliczna wizualizacja końca Spotify

A zatem wyglądałoby na to, że Spotify wcale nie umiera, a nawet wręcz przeciwnie – rozwija się. Sceptycy spojrzeli na dane jednak wnikliwiej. Wynika z nich także to, że liczba użytkowników miesięcznych wzrosła bardziej niż liczba abonentów, co wskazuje na spadek współczynnika przechodzenia z wersji bezpłatnej na płatną (tzw. konwersji). Całkowite przychody wzrosły bardziej niż liczba subskrybentów, co wskazuje, że za wyniki odpowiadają przede wszystkim podwyżki cen dające wzrost średniego przychodu na użytkownika. Czyli Spotify wydaje się nie tyle rozwijać, ile zwiększać obciążenie finansowe dla istniejących użytkowników. To nie wydaje się na pierwszy rzut oka perspektywiczna recepta na wzrost. Ale może kierownictwo, co robi.



Spotify to nie tylko muzyka, ale również podcasty. Miały zróżnicować ofertę, ale w przeszłości były finansowym obciążeniem dla firmy. Zmniejszono więc ich koszty, m.in. dzięki mniejszym wydatkom na wyłączność podcastów (tj. płacenie gwiazdom takim jak Joe Rogan wynagrodzeń za występowanie tylko na jednej platformie). Jak powiedział współzałożyciel i prezes Spotify, Daniel Ek (2): „[podcasty] były obciążeniem – teraz są dla nas kolejnym źródłem zysków”.

Krytycy przy okazji zwrócili szybko uwagę, że skoro podcasty okazują się źródłem zysków dla firmy, pojawia się wrażenie, że trudniej jest jej pozyskiwać wartościową muzykę. Do krytyki ogólnej, którą formułowało wielu użytkowników, że w obecnym interfejsie Spotify około połowy miejsca na stronie powitalnej to nie tyle to w sensie stylu muzyki, co użytkownik słucha, lubi i mógłby lubić na podstawie zapisanych upodobań, ile rekomendacje generowane przez platformę. W przypadku strony z podcastami rekomendacje są oceniane jako oderwane od upodobań użytkownika.

Nie to, czego szukamy

Wielu użytkowników zwraca uwagę, że interfejs Spotify jest teraz przeładowany rekomendacjami i w ogóle wszystkim. Mają poczucie nadmiaru i zagubienia. Zwracają uwagę, że kiedyś doświadczenie użytkownika było znacznie bardziej spersonalizowane. Widział i mógł wyszukiwać przede wszystkim to, co lubi i czego słucha. Dziś czuje się jak w sklepie pełnym różnorodnego towaru, w którym ma problem ze znalezieniem tego, o co mu chodzi.

Ponadto wielu zauważa, że algorytm serwisu podaje im utwory, które już wielokrotnie odrzucali i zaznaczali, że im się nie podobają. Lista piosenek, które „naszym zdaniem mogą Ci się spodobać”, jest w Spotify długa. Odrzucanie pozycji powoduje jedynie podsuvanie kolejnych z tej samej listy, a nie zmianę profilu listy czy jej samej. Ta technika nie reaguje, w ocenie krytyków, na to, co przez swoje odrzucenia komunikuje użytkownik, tylko uparcie stara mu się wcisnąć kolejne rzeczy, których nie chce.

W artykule w „New Yorker” Kyle Chayka cytuje Jarretta Fullera, projektanta i profesora na Uniwersytecie Karoliny Północnej, który mówi: „W ostatniej dekadzie... podejście do projektowania skoncentrowane na użytkowniku zostało zastąpione podejściem skoncentrowanym na korporacji”. Zamiast optymalizować doświadczenia użytkownika, optymalizuje się je pod kątem osiągania zysków. Jak dodaje Fuller, jeśli Spotify uda się zmienić nas



2. Daniel Ek, współzałożyciel Spotify

wszystkich w biernych słuchaczy, to nie będzie miało znaczenia, jakie treści platforma licencjonuje.

Typowa deklaracja pożegnania ze Spotify składa się mniej więcej z takich deklaracji i komentarzy:

- *Żegnaj Spotify! Dlaczego rezygnuję z mojej dziesięcioletniej subskrypcji premium.*
- *Jedyne, co trzymało mnie przy Spotify, to moje playlisty i historia odtwarzania. Dzięki temu algorytm dobrze mnie znał. Ale teraz mam dość.*
- *Kropkę, która przełaziła czarę goryczy, było ostatnie otwarcie aplikacji na komputerze i widok kolejnej zmiany interfejsu użytkownika.*
- *Interfejs użytkownika jest chaotyczny. Znalezienie i odtworzenie muzyki wymaga zbyt wielu kliknięć.*
- *Ciągle promują określone utwory.*
- *Nadal nie ma bezstratnego dźwięku.*
- *Bezsensowne powiadomienia o wyprzedanych koncertach.*
- *Funkcja odtwarzania losowego nadal działa fatalnie.*
- *Zbyt duży nacisk na podcasty i filmy, które mnie nie interesują.*
- *Nie mogę już szybko i łatwo znaleźć muzyki, której chcę słuchać. Mam dość tego, że zmuszają mnie do oglądania i słuchania rzeczy, które mnie nie interesują.*

Precz z algorytmami i kradzieżą!

Od niedawna hasło „Śmierć Spotify” (ang. „Death to Spotify”) przestało być tylko wyrazem narzekania niezadowolonych i sfrustrowanych użytkowników. „Death to Spotify” funkcjonuje teraz jako międzynarodowa kampania, w ramach której muzycy, wytwórnie

muzyczne, DJ-e i inne grupy prowadzą bojkot, twierdząc, że Spotify nie tylko stosuje nieuczciwe praktyki płatnicze i algorytmy, ale także powoli niszczy nasze zainteresowanie dobrą muzyką.

Bojkoty Spotify to nie nowość. Swoje akcje przeciw serwisowi wszczynali znani artyści, Neil Young i Taylor Swift. „Death to Spotify” ma jednak charakter bardziej oddolnej kampanii. Hasło „Śmierć Spotify” pojawiło się po tym, jak dziennikarka Liz Pelly opublikowała w styczniu swoją książkę „Mood Machine: The Rise of Spotify and the Costs of the Perfect Playlist” (z ang. „Maszyna nastrojów: Powstanie Spotify i koszty idealnej playlisty”). W książce tej autorka zwraca uwagę na coś, o czym wielu krytyków wyżej cytowanych pisało już wielokrotnie. Na to, że Spotify zniechęca do słuchania interesującej muzyki, promując playlisty o niskiej wartości.

W Oakland w Kalifornii ogłoszono niedawno serię prelekcji pod hasłem „Śmierć Spotify” (3). W efekcie powstało zgrabne podsumowanie celów ruchu: „Precz z algorytmicznym słuchaniem, precz z kradzieżą tantiem, precz z muzyką generowaną przez sztuczną inteligencję”.

Trzy cele „Death to Spotify” to:

1. Usunięcie mechanizmu odtwarzania opartego na algorytmach.
2. Koniec z kradzieżą tantiem.
3. Wyeliminowanie muzyki generowanej przez sztuczną inteligencję.

Ruch sprzeciwu nabiera więc form nieco bardziej zorganizowanych niż liczne, ale rozproszone narzekania w Internecie. Tymczasem kierownictwo Spotify podało w październiku 2025 r., że liczba aktywnych użytkowników Spotify wzrosła o 11 proc. rok do roku, osiągając poziom 713 milionów. Przychody również wzrosły ponad oczekiwania, osiągając poziom około 4,9 mld dolarów. W sierpniu 2025 r. Spotify po raz kolejny podniosło ceny abonamentów na wielu głównych rynkach. Posunięcie to, wraz z wprowadzeniem nowych produktów, takich jak komunikator w aplikacji, pomogło zwiększyć obroty w tym kwartale. „Działamy szybciej niż kiedykolwiek, obserwujemy wysokie zaangażowanie użytkowników i dłuższy czas spędzany na platformie. Dziękujemy wszystkim zespołom Spotify na całym świecie”, czytamy w komunikacie Spotify. ■

Mirosław Usidus

3. Zdjęcie z jednej z prelekcji w ramach ruchu Death to Spotify w kalifornijskim Oakland © Instagram





Inżynieria jakości

Każdy widział filmiki z testów zderzeniowych. Te słynne ujęcia, w których samochód uderza w ścianę, a plastikowy ludzik leci do przodu w zwolnionym tempie. Dziś to codzienność, ale pierwszy taki test przeprowadzono w latach 30. ubiegłego wieku. Wtedy zamiast manekinów używano worków z piaskiem, a zamiast komputerów notatek z ołówkiem. Wyniki analizowano na podstawie tego, co udało się zobaczyć gołym okiem. Od tamtej pory świat jakości przeszedł rewolucję. Inżynierowie testują już nie tylko wytrzymałość karoserii, ale też trwałość lakieru na zarysowania, odporność deski rozdzielczej na promienie UV i jakość powłok antyodblaskowych na wyświetlaczach, by nie męczyły wzroku kierowcy. To właśnie robią specjaliści od inżynierii jakości. Ludzie, którzy wierzą, że doskonałość zaczyna się tam, gdzie większość przestaje patrzeć. I tego właśnie uczą studia, które kształcą inżynierów gotowych na walkę o jakość. Zapraszamy na studia.

Inżynieria jakości to kierunek, który łączy technikę, zarządzanie i analitykę danych. Jego absolwenci należą dziś do jednej z najbardziej poszukiwanych grup specjalistów w przemyśle. Choć często występuje pod różnymi nazwami, jak na przykład: zarządzanie jakością i produkcją, inżynieria jakości w procesach i usługach czy inżynieria bezpieczeństwa i jakości, to idea pozostaje ta sama: kształcenie ludzi, którzy potrafią sprawić, by procesy działały bezbłędnie, a produkty spełniały najwyższe standardy. Studia te prowadzone są zarówno w trybie dziennym, jak

i zaocznym. Można je realizować na poziomie inżynierskim (7 semestrów), magisterskim (3–4 semestry) lub podyplomowym. Najpełniejszą ofertę ma Uniwersytet Morski w Gdyni, gdzie kierunek inżynieria jakości można studiować jako samodzielny program, w trybie dziennym i zaocznym, na obu stopniach. Studenci wybierają tam m.in. specjalizacje: inżynieria jakości i zarządzanie produktem, zarządzanie jakością usług oraz usługi żywieniowe i dietetyka. Na politechnikach inżynieria jakości najczęściej funkcjonuje jako ścieżka w ramach inżynierii

i zarządzania produkcją lub pokrewnych kierunków. Tak jest na Politechnice Krakowskiej (inżynieria jakości w produkcji zautomatyzowanej) czy Politechnice Poznańskiej (inżynieria bezpieczeństwa i jakości). Z kolei Akademia Górniczo-Hutnicza w Krakowie proponuje studia podyplomowe z inżynierii jakości oprogramowania, gdzie klasyczne metody kontroli jakości łączą się z analizą procesów informatycznych. Koszt takiego kierunku to ok. 7600 zł, choć inne uczelnie oferują tańsze odpowiedniki. Inżynierię jakości można znaleźć również w szkołach niepublicznych, często w trybie zdalnym. Zanim jednak zdecydujemy się na taką formę, warto sprawdzić uprawnienia uczelni i jakość programu. W tym wypadku „inżynieria jakości” nie powinna być tylko nazwą.

Nabywanie jakości

Rekrutacja na studia z zakresu inżynierii jakości jest umiarkowanie wymagająca. Na politechnikach punkty przyznawane są głównie za matematykę i fizykę, w uczelniach zarządczych za geografię lub język obcy. Studia niestacjonarne są łatwiejsze do rozpoczęcia, natomiast studia dzienne na dużych uczelniach wymagają podstaw ścisłych i systematyczności. Na Politechnice Krakowskiej w rekrutacji na rok akademicki 2025/2026 odnotowano 3,03 kandydata na jedno miejsce. Duże różnice zaczynają się dopiero w toku nauki. Na uczelniach technicznych program jest bardziej rozbudowany i obejmuje pełny zestaw przedmiotów inżynierskich: matematykę, statystykę, fizykę, metrologię, a także liczne laboratoria pomiarowe i projektowe. Na uczelniach o profilu praktycznym większy nacisk kładzie się na systemy zarządzania, audyt i analizę procesów. Do najczęściej spotykanych przedmiotów należą: wstęp do inżynierii jakości, zarządzanie jakością, systemy zapewniania jakości, analiza wyników pomiarów, metrologia, dokumentacja jakości, audytowanie i certyfikacja, Branżowe systemy zarządzania jakością, planowanie jakości oraz kwalitologia, czyli nauka o naturze jakości i metodach jej oceny. Studenci uczą się także rozwiązywać problemy produkcyjne i eliminować błędy w procesach. Poznają metody analizy przyczyn (np. technikę „5 × dlaczego”) i sposoby ich wizualizacji za pomocą wykresów przyczynowo-skutkowych, czyli tzw. diagramów Ishikawy. Uczą się oceniać ryzyko, planować kontrolę poszczególnych etapów produkcji i opracowywać działania korygujące, które wdraża się po wykryciu problemu. Na politechnikach pojawiają się też bardziej techniczne zagadnienia: zaawansowane metody pomiarów geometrycznych, roboty i maszyny pomiarowe,

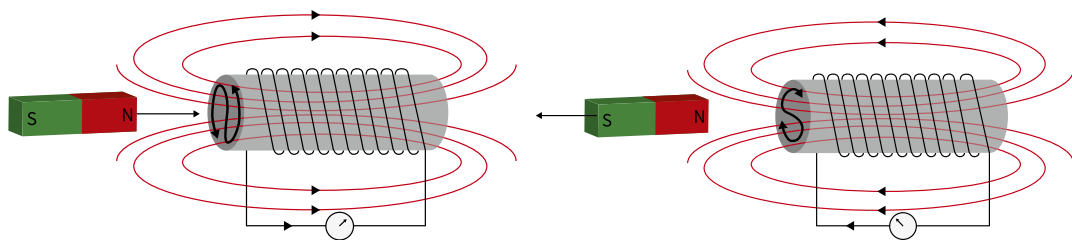
automatyzacja kontroli jakości czy zarządzanie bezpieczeństwem technicznym. Najwięcej trudności studentom sprawiają zajęcia wymagające pracy z danymi, a więc: statystyka, kontrola procesów produkcyjnych i metrologia. To właśnie one uczą precyzji, cierpliwości i logicznego myślenia, cech niezbędnych w zawodzie inżyniera jakości. W praktyce inżynier korzysta z całego zestawu metod doskonalenia procesów. Analizuje dane produkcyjne, planuje eksperymenty, bada przyczyny błędów i wdraża rozwiązania, które zapobiegają ich powtarzaniu. Zna techniki planowania jakości, sposoby oceny niezawodności produktów i narzędzia do raportowania. Coraz częściej pracuje z systemami komputerowymi: SAP, MES czy oprogramowaniem do analizy danych jakościowych i statystycznych. Nowoczesny inżynier jakości to nie tylko kontroler, to także analityk, doradca i łącznik między produkcją, laboratorium a zarządem.

Jakość w pracy

Po ukończeniu studiów absolwent może pracować w firmach produkcyjnych, laboratoriach badawczych, jednostkach certyfikujących, logistyce, sektorze spożywczym, farmaceutycznym czy elektronicznym. Typowe stanowiska to: inżynier jakości, kontroler jakości, audytor wewnętrzny, specjalista ds. systemów zarządzania jakością, pełnomocnik ds. jakości. Wielu absolwentów zdobywa też certyfikaty zawodowe (np. audytor ISO, Six Sigma Green Belt, Black Belt), które zwiększają ich atrakcyjność na rynku pracy i otwierają drzwi do stanowisk kierowniczych. Rynek pracy dla inżynierów jakości jest stabilny. Wraz z rozwojem automatyzacji i sztucznej inteligencji część prostych zadań kontrolnych znika, ale równocześnie pojawiają się nowe role: quality data analyst, digital quality engineer czy process excellence specialist. To dowód, że zawód ewoluuje, zamiast zniknąć.

Inżynieria jakości kształci interdyscyplinarnych specjalistów. Ludzi, którzy potrafią połączyć techniczne myślenie z umiejętnością pracy z ludźmi i danymi. Dobry inżynier jakości nie tylko wykrywa błędy, ale też przewiduje, gdzie mogą się pojawić, zanim jeszcze do nich dojdzie. To kierunek dla osób, które lubią analizować, szukać przyczyn, planować i usprawniać. Dla tych, którzy nie boją się liczb, ale też potrafią współpracować i komunikować się jasno. Inżynieria jakości to nie tylko zawód. To sposób myślenia o świecie, w którym wszystko można zrobić lepiej, jeśli tylko wiemy, jak to zmierzyć. Zapraszamy na studia, bo dobra jakość nigdy nie jest dziełem przypadku. ■

Michał Pacholski



Zjawisko indukcji elektromagnetycznej

W pewnym zbiorze znalazłam bardzo ciekawy przykład zadania, które rozwiązać można tylko wtedy, gdy rozwiązujący wie, czego oczekuje od niego autor. Sens zadania był mniej więcej taki: samolot leciał ze stałą prędkością w ziemskim polu magnetycznym. Podana została wartość tejże prędkości, wartość indukcji pola magnetycznego i rozpiętość skrzydeł samolotu.

Pałapka szkolnego zadania

Celem rozwiązującego było obliczenie siły elektromotorycznej generowanej w trakcie lotu pomiędzy końcówkami skrzydeł, przy czym nie wspomniano ani słowem, aby ich powierzchnia zmieniała się w trakcie lotu na przykład na skutek schowania klap, co zresztą powinno mieć miejsce w trakcie rzeczywistego lotu.

Co na to uczniowie? Uznali zadanie za niezrozumiałe i nierozwiązywalne, nie bacząc na to, że na końcu owej publikacji została podana odpowiedź. Ja potrafię to zadanie rozwiązać w kilka minut, uzyskując zresztą dokładnie ten sam wynik co jego autor. A mimo wszystko przychyliłam się do zdania uczniów. Gdzie zatem trudność i w czym tkwi problem?

Przede wszystkim zacznijmy od tego, w jakiej sytuacji może być generowana siła elektromotoryczna w przypadku przewodnika umieszczonego w polu magnetycznym. Odpowiednie warunki występują wyłącznie wtedy, gdy zrobimy z przewodnika zamknięty obwód (np. ramkę), a strumień pola magnetycznego przechodzący przez powierzchnię ograniczoną tym przewodnikiem będzie się zmieniał.

Zgodnie z treścią zadania nic podobnego nie miało prawa się wydarzyć (stałe pole magnetyczne, stała powierzchnia skrzydeł), a jednak w niewytłumaczalny sposób zaistniało. Do szkolnego zadania

wrócimy za chwilę. Aby odczytać intencje autora, trzeba przede wszystkim dobrze zrozumieć zjawisko indukcji elektromagnetycznej. Z tym niestety część uczniów ma problem.

Odrobina historii

Zjawisko indukcji elektromagnetycznej zostało odkryte w roku 1831 przez Michaela Faradaya, który dowiódł na drodze doświadczalnej, że w zamkniętej pętli przewodnika może powstawać siła elektromotoryczna wywołana zmianami strumienia pola magnetycznego. W praktyce oznacza to, że można wywołać przepływ prądu elektrycznego bez konieczności posiadania zewnętrznego źródła zasilania. Wystarczy umieścić przewodnik w zmiennym polu magnetycznym albo odpowiednio obracać nim względem pola stałego.

Wygenerowana siła elektromotoryczna jest wprost proporcjonalna do zmiany strumienia pola magnetycznego i odwrotnie proporcjonalna do czasu, w jakim ta zmiana nastąpiła, co można opisać wzorem

$$\epsilon = - \frac{\Delta \Phi}{\Delta t}$$

gdzie ϵ oznacza siłę elektromotoryczną wyrażoną w woltach, $\Delta \Phi$ zmianę strumienia pola wyrażoną w weberach, a Δt – przedział czasu, w jakim dokonała się ta zmiana, wyrażony w sekundach. Minus we wzorze oznacza, że siła elektromotoryczna jest

tak skierowana, aby przeciwdziałać przyczynie jej powstawania.

Pole magnetyczne i strumień pola

Aby zagłębić się w to zagadnienie, przypomnijmy sobie kilka pojęć związanych z magnetyzmem. Przede wszystkim, aby zjawisko indukcji magnetycznej zaistniało, musi istnieć pole magnetyczne, czyli obszar, w którym na umieszczone tam ciała działa siła magnetyczna.

Oczywiście nie każdy rodzaj materii ulega oddziaływaniu magnetycznemu, więc nie w każdej sytuacji jesteśmy w stanie stwierdzić obecność pola magnetycznego. Jeśli jednak mamy ferromagnetyk (na przykład opilkę żelaza albo drobne stalowe przedmioty), to możemy określić przebieg linii pola oraz ich gęstość.

Kierunek linii odpowiada kierunkowi wektora natężenia pola, a ich zagęszczenie – wartości natężenia. Gdy na przykład rozsypiemy opilkę w pobliżu magnesu sztabkowego, to okaże się, że linie pola zagęszczają się w bezpośredniej bliskości biegunów, a mniej gęsto biegną w większej odległości od nich. Strumień pola magnetycznego jest proporcjonalny do liczby linii przecinających ustaloną powierzchnię.

W przypadku powierzchni prostopadłej do linii pola możemy strumień policzyć jako $\Phi = BS$, gdzie B jest wartością indukcji magnetycznej, a S jest polem powierzchni, przez którą przechodzą linie pola magnetycznego. Warto w tym miejscu zaznaczyć, że indukcja magnetyczna jest proporcjonalna do natężenia pola i jednocześnie uwzględnia właściwości magnetyczne danego ośrodka, zatem zazwyczaj to ona jest wykorzystywana do opisu zjawisk z dziedziny magnetyzmu, a nie samo natężenie pola magnetycznego.

Zmiany strumienia pola magnetycznego w czasie

Wartość strumienia pola magnetycznego może się zmieniać jeśli zmienia się wartość

indukcji magnetycznej przy stałej powierzchni ramki ($\Delta\Phi = \Delta BS$) albo zmienia się wielkość ramki przy stałej wartości indukcji magnetycznej ($\Delta\Phi = B\Delta S$). Oczywiście obie te wielkości mogą ulegać zmianie równocześnie, jednak nie będziemy się w tej chwili zajmować tego typu przypadkami.

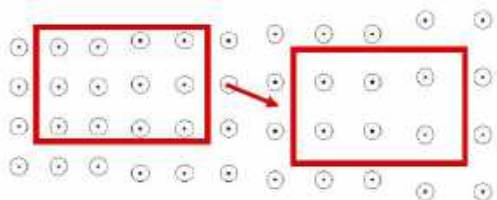
Na **rysunku 1** zaprezentowano przykładową realizację pierwszej z tych możliwości. Ramkę o ustalonym rozmiarze przesuwamy w obszarze niejednorodnego pola magnetycznego. W położeniu początkowym ramka obejmuje dwukrotnie więcej linii pola niż w położeniu końcowym. Zatem strumień pola maleje dwukrotnie. W czasie przesuwania ramki zostaje wygenerowana siła elektromotoryczna powodująca przepływ prądu.

Na **rysunku 2** przedstawiono przykładową realizację drugiego pomysłu na wytworzenie siły elektromotorycznej. W trakcie zmiany rozmiarów ramki zmienia się również liczba linii pola obejmowanych przez przewodnik. Oczywiście może to być ramka pozwalająca na jej rozciąganie, ale również szyny połączone na jednym końcu, po których toczy się metalowa belka. W praktyce rozwiązania mogą być różne i wiele z nich znajduje zastosowanie w technice.

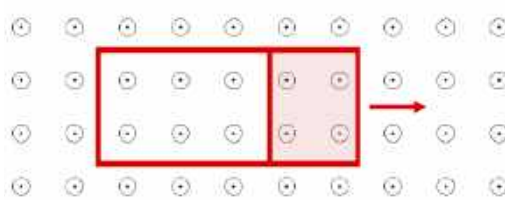
Jak rozwiązać zadanie?

Żeby rozwiązać przytoczone na samym początku zadanie, potrzebna jest przede wszystkim duża doza wyobraźni. Otóż wyobraźmy sobie, że w jakiś sposób skrzydła samolotu są połączone z lotniskiem niewidocznymi, aczkolwiek przewodzącymi prąd drucikami lub szynami. W ten sposób powstaje nam ramka z przewodnika, przedstawiona na **rysunku 2**.

Ponieważ samolot porusza się, powierzchnia tej ramki rośnie wraz z czasem lotu. Jej krótszy bok ma stałą długość równą rozpiętości skrzydeł (oznaczymy tę wielkość jako Y), a dłuższy bok jest równy przebytej drodze, którą nazwiemy Δx . Skoro wartość prędkości jest stała, to $\Delta x = v\Delta t$. Zmiana powierzchni ramki wynosi $\Delta S = Y\Delta x = Yv\Delta t$.



1. Zmiana strumienia pola magnetycznego na skutek ruchu ramki z przewodnika w niejednorodnym polu magnetycznym



2. Zmiana strumienia pola magnetycznego na skutek zmiany rozmiaru ramki z przewodnika w jednorodnym polu magnetycznym. Pokolorowany obszar odpowiada zmianie pola powierzchni ramki



Zatem zmiana strumienia indukcji w czasie lotu jest równa $\Delta\Phi = V\Delta S = Byv\Delta t$.

Teraz podstawiamy wyrażenie na zmianę strumienia indukcji do wzoru na wartość siły elektromotorycznej i w prościutki sposób, do tego w zaledwie paru liniach dostajemy $\epsilon = -Byv$ (skróci się czas w liczniku i mianowniku). Jeśli w poleceniu jest wyraźnie napisane, że mamy obliczyć wartość siły elektromotorycznej, to minus pomijamy, ponieważ określa jej zwrot. Wystarczy już tylko podstawić dane liczbowe i obliczyć wynik.

Na koniec mała uwaga: bez drucików łączących samolot z miejscem startu żadna siła elektromotoryczna nie ma prawa powstać pomiędzy końcówkami skrzydeł, zatem w rzeczywistości nie powstaje. Owszem, powierzchnia samolotu może zostać naładowana elektrycznie i może pojawić się różnica potencjałów pomiędzy różnymi punktami poszycia, ale powstaje ona w trakcie lotu przez chmurę jako efekt tarcia pomiędzy samolotem a kryształkami lodu lub drobinami pyłu wulkanicznego.

Zjawisko to było wielokrotnie obserwowane przez załogi samolotów jako tak zwane „ogień świętego Elma” i zazwyczaj stanowi ono ostrzeżenie o możliwości wystąpienia silnych wyładowań atmosferycznych. W przypadku braku chmur deszczowych oznacza ono, że samolot znalazł się w chmurze pyłu wulkanicznego, co może być bardzo szkodliwe dla silników.

Sprawdź, co potrafisz

W stałym polu magnetycznym o znanej wartości indukcji przesuwamy szybko miedzianą ramkę

z punktu A do punktu B, a następnie z powrotem do punktu A. Czynność powtarzamy wielokrotnie. Wybierz i zaznacz stwierdzenie prawdziwe.

- ☐ A. W ramce nie będzie płynął prąd, ponieważ nie wygenerujemy w ten sposób siły elektromotorycznej.
- ☐ B. W ramce będzie płynął prąd stały, ponieważ kierunek siły elektromotorycznej będzie ten sam w obu fazach ruchu.
- ☐ C. W ramce będzie płynął prąd zmienny, ponieważ kierunek siły elektromotorycznej będzie się zmieniał na przeciwny.
- ☐ D. W ramce będą równocześnie generowane dwie siły elektromotoryczne o przeciwnych kierunkach, więc będą się znosić.

Dla nauczyciela

Zadania tego samego typu jak przytoczone tutaj zadanie z samolotem pojawiają się co jakiś czas na maturze z fizyki. Dlatego warto prześledzić ogólny schemat ich rozwiązania, nawet jeśli na pierwszy rzut oka wydają się mocno uduchowione. Po opanowaniu tego schematu zadania przestają wydawać się nierozwiązywalne, a wręcz okazują się względnie proste.

Zazwyczaj zadania te dotyczą belki poruszającej się po szynach w polu magnetycznym i zdarza się, że oprócz siły elektromotorycznej należy policzyć na przykład natężenie płynącego przez obwód prądu (jeśli znane są wartości oporu elektrycznego poszczególnych elementów). ■

Joanna Borgensztajn

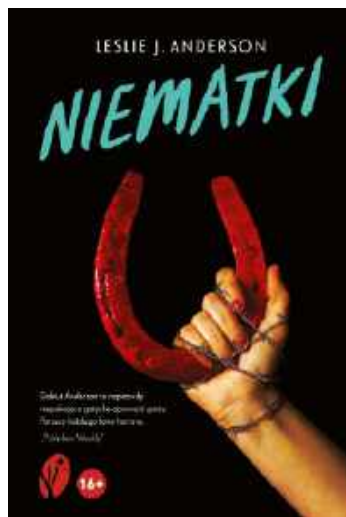
(Uwaga – mamy stałe pole magnetyczne)
Prawidłowa odpowiedź do zadania:

Niematki

Leslie J. Anderson

Wydawnictwo StoryLight, liczba stron: 360, cena sugerowana: 44,99 zł

Dziennikarka trafia do zapadłej miejsciny, gdzie poznaje mroczny sekret skrywany przez kobiety od pokoleń. Carolyn Marshall wciąż próbuje poskładać swoje życie po śmierci męża. Gdy zostaje uwięziona w tragiczny wypadek, jej szef wysyła ją do prowincjonalnego miasteczka Raeford, by zweryfikowała absurdalną pogłoskę: ponoć klacz urodziła zdrowego ludzkiego chłopca. Na miejscu Carolyn zderza się z hermetyczną społecznością, która lepiej traktuje konie – stynie zresztą z ich hodowli – niż ludzi. A kiedy wśród okolicznych pól zostają odnalezione dwa potwornie okaleczone ciała – końskie i ludzkie – staje się dla niej jasne, że za badaną pogłoską kryje się coś więcej. Gdy Carolyn coraz głębiej wnika w życie miasteczka i jego zamkniętych w sobie mieszkańców, zaczyna jej się wydawać, że traci kontakt z rzeczywistością. A może ta nieprawdopodobna historia jest kluczem do mrocznej tajemnicy, która od pokoleń prześladowa kobiety z Raeford?





Chemiczny Nobel 2025, czyli szukanie dziury w całym

Środa 8 października 2025 roku to dzień, w którym Królewska Szwedzka Akademia Nauk ogłosiła listę laureatów Nagrody Nobla w dziedzinie chemii. W 117. edycji (w 8 latach nie przyznano nagrody) wyróżniona została trójka uczonych – Richard Robson, Susumu Kitagawa i Omar M. Yaghi. Uzasadnienie werdyktu brzmiało: „za rozwój struktur metaloorganicznych” (*for the development of metal-organic frameworks*). Podobnie jak w poprzednich latach laureaci otrzymali złote medale, dyplomy i 11 milionów koron szwedzkich do podziału.

Tytułowe „szukanie dziury w całym” jest dość niepochlebnym określeniem, ale w przypadku nagrodzonych ostatnim Noblem osiągnąć stanowiło o ich powodzeniu. Uczeń zaprojektowali szereg nowych materiałów o porowatej strukturze charakteryzujących się unikatowymi właściwościami pochłaniania gazów i cieczy. Zbiorczo nazywane są one **strukturami metaloorganicznymi**, w skrócie **MOF** (ang. *metal-organic frameworks*). Choć materiały te są rzeczywiście wynalazkami z ostatnich dekad, sama idea to jednak...

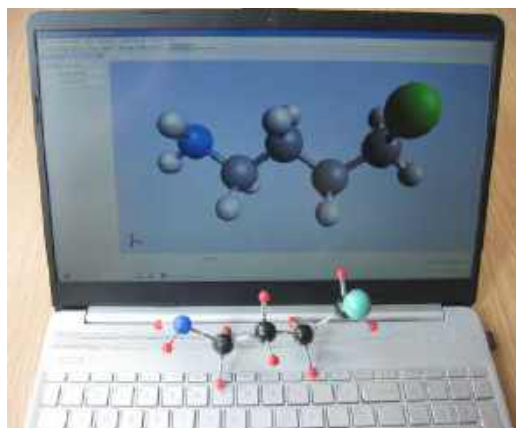
...nic nowego

Porowate materiały pochłaniające znane są chemikom od ponad 200 lat, a praktycznie stosowane od kilku tysięcy. **Węgiel drzewny**, powstający przez zwęglenie drewna bez dostępu powietrza, wykorzystywany był już przez starożytnych do oczyszczania wody i wina, a także w lecznictwie (pisał już o nim ojciec medycyny – Hipokrates – w IV w. p.n.e.). Specjalnie przygotowany węgiel drzewny nosi nazwę **węgla aktywnego** i od dawna ma liczne zastosowania. Przykładami niech będą domowe filtry do wody, oczyszczacze powietrza, wkłady masek przeciwgazowych, a także tabletki przyjmowane przy problemach żołądkowych. Naturalnie występujące glinokrzemiany – **zeolity** – to także od dawna stosowane materiały (wytwarza się również ich sztuczne odpowiedniki). Znajdziesz je w akwaryjnych filtrach do wody, proszkach do prania (zmiękczały wodę) czy też żwirku dla kotów. W przemyśle zarówno zeolity, jak i węgiel aktywny stosuje się w roli pochłaniaczy i katalizatorów.

Oba rodzaje materiałów są tanie, łatwe w produkcji i bezpieczne dla środowiska. Nasuwa ci się zapewne pytanie, co wobec tego prezentują sobą MOF-y, że zasłużyły aż na Nagrodę Nobla?

O pożytkach z używania drewnianych modeli

Chemicy od półtora wieku przedstawiają budowę cząsteczek za pomocą modeli, co umożliwia lepsze wyobrażenie ich struktury przestrzennej. Początkowo były to drewniane kulki przedstawiające atomy połączone patyczkami naśladującymi wiązania chemiczne. Od połowy XX wieku modele konstruuje się z tworzyw sztucznych, a w ostatnich kilkudziesięciu latach istnieją one w pamięci komputerów, zaś



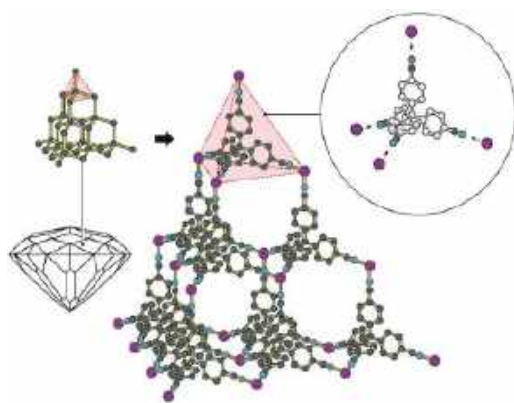
1. Modele cząsteczki związku chemicznego – tradycyjny kulkowy i nowoczesny w formie elektronicznej



wyspecjalizowane oprogramowanie pozwala na dostrzeżenie niuansów budowy molekuł (1). W roku 1974 **Richard Robson** zlecił jednak uniwersyteckiemu warsztatowi przygotowanie modeli z drewna. Oglądając gotowe wyroby, zwrócił uwagę na często pomijany fakt: otwory w drewnianych kulkach, w które wkłada się łączące je patyczki, nie są losowo rozmieszczone, lecz skierowane w ściśle określonych kierunkach przestrzeni tak, jak rzeczywiste wiązania pomiędzy atomami. Spostrzeżenie nasunęło mu pomysł połączenia ze sobą nie tylko atomów, ale całych cząsteczek. Czy i one utworzą strukturę przestrzenną?

Pomysł doczekał się realizacji dopiero po latach. W roku 1989 Robson połączył organiczną cząsteczkę mającą kształt tetraedru (piramida o podstawie trójkąta równobocznego, jej ściany są takimi samymi trójkątami) z jonami miedzi. Miedź z kolei „lubi” otaczać się czterema sąsiadami, również leżącymi w wierzchołkach tetraedru (jon metalu położony jest zaś w środku piramidy). Całość, zgodnie z przewidywaniami Robsona, miała strukturę przestrzenną, taką samą jak w przypadku diamentu. W przeciwieństwie jednak do tego ostatniego, kryształu o zwartej strukturze, otrzymana substancja była porowatej budowy. Przypominała konstrukcję kratownicową, której elementy (cząsteczki organiczne) spojęne były poprzez jony miedzi (2). Wielkość wolnych przestrzeni pomiędzy elementami struktury odpowiadała przeciętnym rozmiarom molekuł i okazało się, że dość łatwo wchodziły one w powstałe luki – analogicznie jak w przypadku węgla aktywnego i zeolitów.

Uczony opublikował wyniki badań, przewidując przyszłość stojącą przed otrzymanymi połączeniami:



2. MOF Robsona (w środku) zawiera wolne przestrzenie i ma budowę podobną do diamentu (z lewej), z prawej szczegóły elementu składowego struktury (jony miedzi oznaczono fioletowymi kulami) (©Johan Jarnestad/The Royal Swedish Academy of Sciences)

przy odpowiednim zaprojektowaniu wnek tak, aby pasowały do konkretnych związków chemicznych, mogły one służyć jako katalizatory reakcji i selektywne pochłaniacze. Chociaż konstrukcje były nietrwałe i łatwo ulegały kolapsowi (zapadały się), to jednak słowa Robsona okazały się prorocze.

O pożytkach z „bezużytecznych” badań

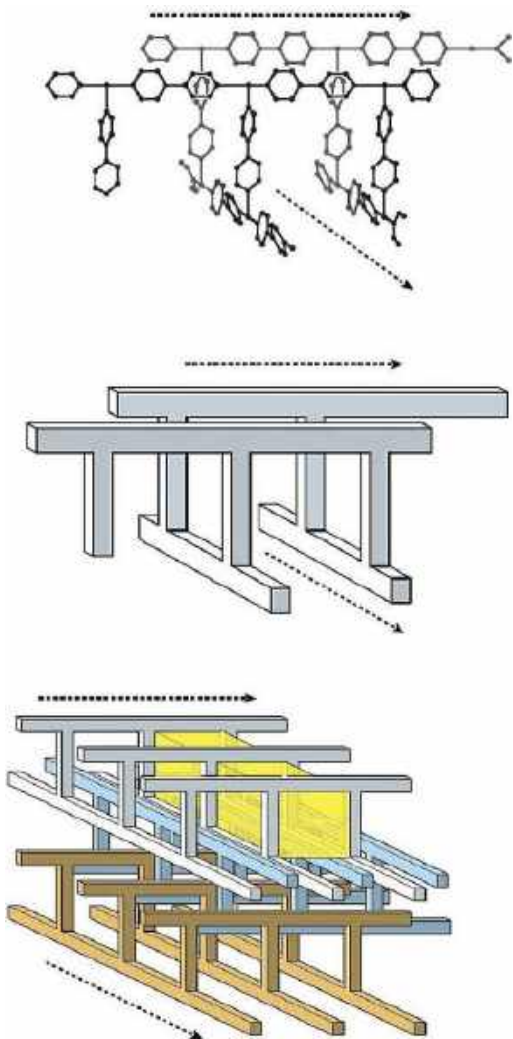
Drugi z noblistów, **Susumo Kitagawa**, uważał, że nawet badania nieprzynoszące natychmiastowych korzyści mogą okazać się bardzo cenne. Ze swymi przekonaniem znalazł się w doborowym towarzystwie. W połowie XIX wieku Michael Faraday zapytany przez ministra skarbu, do czego może przydać się odkryte przez niego zjawisko indukcji elektromagnetycznej, odparł, że jeszcze nie wie, ale zapewnił dygnitarza, że wkrótce będzie z niego ściągał podatki. I rzeczywiście, zjawisko stało się podstawą konstrukcji prądnic, transformatorów, silników i wielu innych urządzeń elektrycznych.

Od roku 1992 Kitagawa syntezował porowate materiały, nie były one jednak stabilne, a urzędnicy odpowiedzialni za finanse nie widzieli celu jego prac. Uczony nie poddawał się i w roku 1997 odniósł sukces – gdy wysuszył otrzymaną substancję, jej struktura nie zapadła się, a wolne przestrzenie można było wypełniać gazami (materiał pochłaniał i uwalniał je pod wpływem zmian temperatury i ciśnienia). Urzędnicy nadal jednak nie chcieli finansować badań, argumentując, że już od dawna chemicy i przemysł mają zeolity o takiej samej funkcjonalności. Po co zatem wydawać pieniądze na coś, co nie będzie użyteczne?

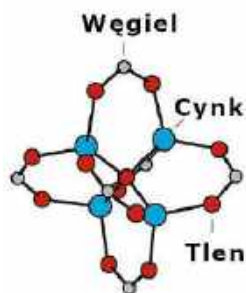
Kitagawa ripostował: w odróżnieniu od twardych krzemianowych zeolitów, nowe materiały mogą być elastyczne (co umożliwiającą giętkie, organiczne kamyki budulcowe), a praktycznie nieskończona liczba związków węgla pozwoli na tworzenie MOF-ów o niespotykanych do tej pory właściwościach. Japoński chemik zdołał przekonać księgowych i jego badania mogły być kontynuowane, jak się później okazało, z sukcesem (3).

O pożytkach z zakradania się do biblioteki

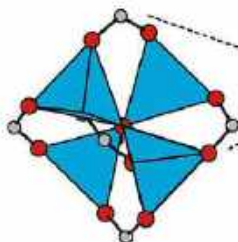
Dzieciństwo **Omara Yaghi** nie było łatwe. Wraz z liczną rodziną wychowywał się w pozbawionym podstawowych wygod domu. Chłopiec nie poddał się i pragnął się kształcić, w czym wspierali go rodzice. Często zakradał się do zamkniętej szkolnej biblioteki, aby czytać książki. Pewnego razu trafił na podręcznik z ilustracjami struktur cząsteczek. Obrazy tak



3. MOF Kitagawy z widocznymi kanałami wewnątrz materiału (©Johan Jarnestad/The Royal Swedish Academy of Sciences)



=



4. Struktura Yaghiego oznaczona symbolem MOF-5 okazała się przełomowa (wolne miejsca w strukturze przedstawione są jako żółte kule) (©Johan Jarnestad/The Royal Swedish Academy of Sciences)

go urzekły, że postanowił studiować chemię. Marzenie zrealizował po emigracji do USA.

Tam też postanowił zmienić sposób syntezy chemicznej. Tradycyjne mieszanie substratów i podgrzewanie naczynia reakcyjnego w celu przyspieszenia przemiany w przypadku związków organicznych zwykle prowadzi do powstania dużej ilości zbędnych produktów ubocznych. Yaghi chciał budować złożone struktury z elementów składowych, zupełnie jak z klocków Lego, które od razu wskazują na właściwe miejsce konstrukcji.

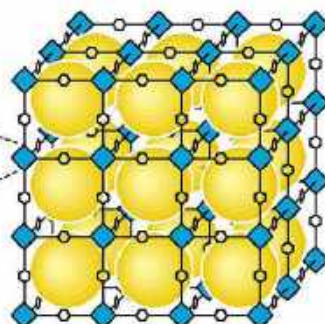
Od roku 1995 uczony tworzył porowate materiały, łącząc, jak wcześniej Robson i Kitagawa, molekuły organiczne za pomocą jonów metali. Okazały się one nadzwyczaj stabilne, wytrzymywały bez uszkodzeń ogrzewanie do temperatury 300...350°C i dawały się wielokrotnie napełniać i opróżniać gazami i cieczami. W artykule opisującym eksperymenty Yaghi użył powszechnie dziś stosowanej nazwy tych materiałów, czyli MOF.

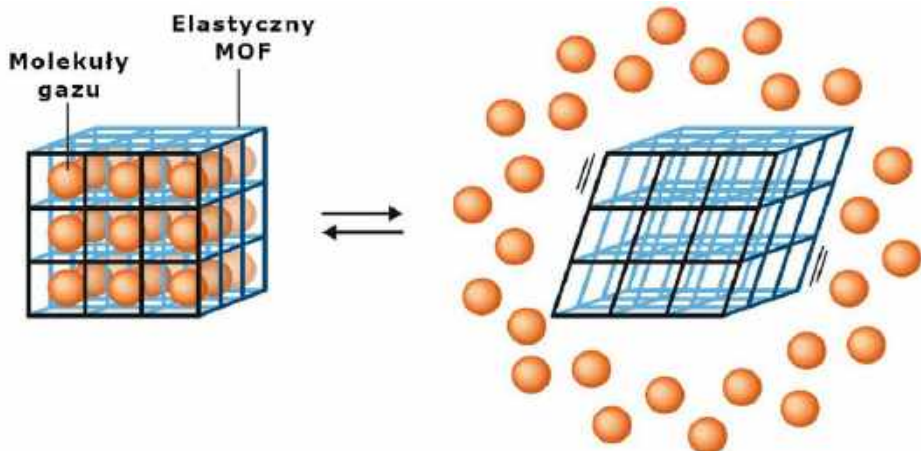
Przełomowy okazał się wytworzony w roku 1999 związek nazwany MOF-5. Materiał jest wyjątkowo stabilny i ma tak rozwiniętą strukturę wewnętrzną, że kilka jego gramów ma powierzchnię porów równą powierzchni boiska piłkarskiego (ponad 7000 m²) (4). W tym czasie konkurencja nie próżnowała – zespół pod kierownictwem Kitagawy zaprezentował pierwszy elastyczny MOF, który zmieniał kształt pod wpływem napełniania i opróżniania, zupełnie jak twoje płuca, gdy oddychasz (5).

Jak działają MOF-y?

Tak samo jak zeolity i węgiel aktywny. Podstawą aktywności jest struktura wszystkich tych materiałów: liczne pory przenikające całą objętość substancji. Ale to nie wystarczy, luki dodatkowo muszą mieć rozmiary porównywalne z rozmiarami cząsteczek

MOF-5



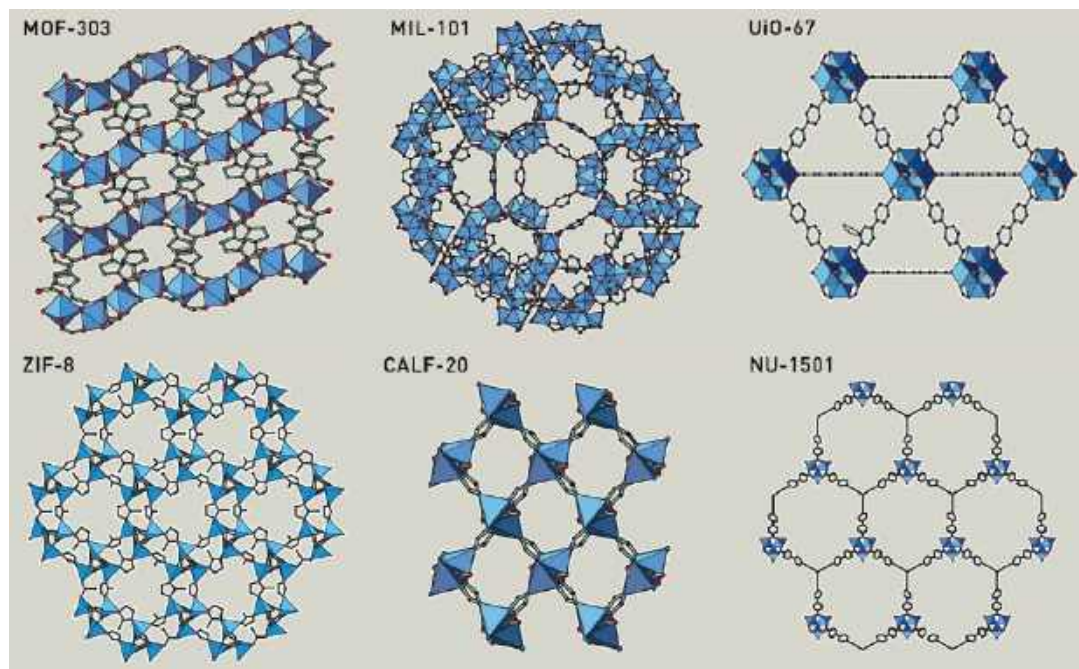


5. Elastyczny MOF zmienia kształt pod wpływem napełniania i opróżniania molekułami gazu (©Johan Jarnestad/The Royal Swedish Academy of Sciences)

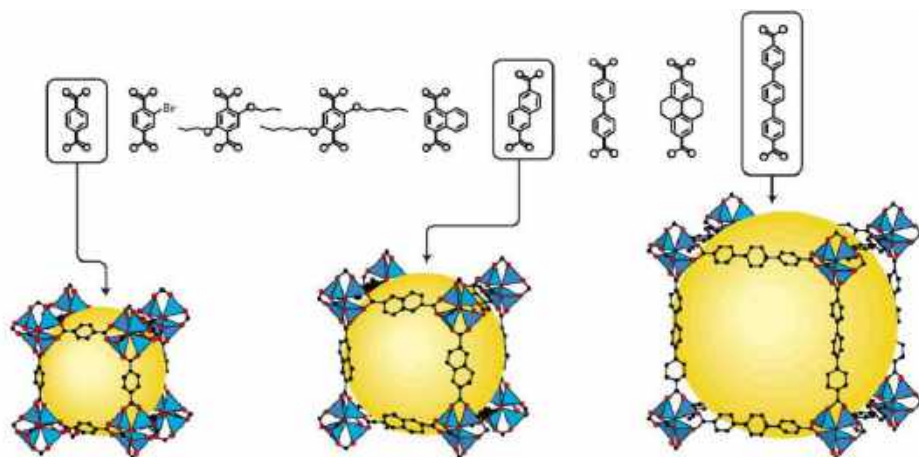
związków chemicznych, wtedy zaczynają działać specyficzne dla mikroświata siły przyciągania i molekuły są zatrzymywane w wolnych przestrzeniach. Przy zmianie warunków (ciśnienie, temperatura) można je jednak stamtąd uwolnić. Jak już wspomniano wyżej, sumaryczna powierzchnia porów dochodzi do kilku tysięcy m² na gram substancji – tak niesamowicie rozwinięta jest wewnętrzna struktura tych materiałów.

Wróćmy do osiągnięć noblistów. Zastosowali oni w praktyce swoje wynalazki. MOF-y

są zdolne do pochłaniania wody nawet z suchego, pustynnego powietrza (potem ją oddają). Mogą także magazynować i uwalniać wodór i metan, co stwarza możliwości łatwego zastosowania tych gazów jako paliw silnikowych (nie trzeba ich skraplać ani przechowywać pod wysokim ciśnieniem). W przemyśle chemicznym i farmaceutycznym pochłaniają toksyczne gazy powstające podczas procesów produkcyjnych oraz pomagają oczyszczać wodę z produktów ropopochodnych i antybiotyków.



6. Struktury różnych MOF-ów (©Johan Jarnestad/The Royal Swedish Academy of Sciences)



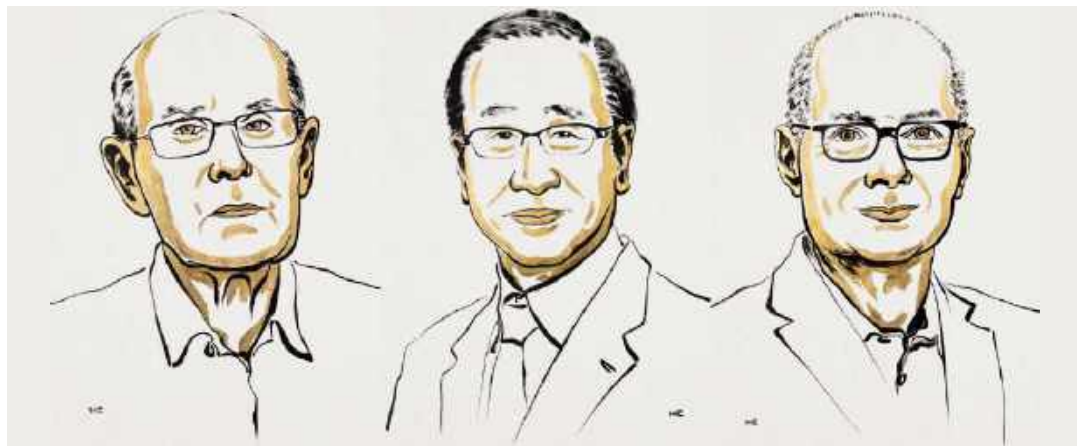
7. Wielkość wolnej przestrzeni (żółta kula) wewnątrz struktury można zmienić, dobierając cząsteczkę organiczną o odpowiednich rozmiarach (©Johan Jarnestad/The Royal Swedish Academy of Sciences)

Podobnie do swych starszych „kuzynów”, zeolitów i węgla aktywnego, mogą być stosowane jako katalizatory, wystarczy tylko dobrać odpowiedni metal do połączenia organicznych „klocków” w porowatą strukturę (6).

Możliwości są zresztą praktycznie nieograniczone – znamy dziesiątki jonów metali i wiele milionów związków organicznych, zatem liczba możliwych struktur jest iście astronomiczna. Do każdego zastosowania, które możemy wymyślić, na pewno znajdzie

się odpowiedni MOF (7). Chemicy i specjaliści zajmujący się inżynierią materiałową będą mieli czym się zajmować. Przemysł również zauważył potencjał tkwiący w nowych materiałach i szeroko inwestuje w prace badawcze. Przyszłość pokaże, czy MOF-y zastąpią stosowane od dawna materiały. Raczej nie, podobnie jak w przypadku wielu innych wynalazków z dziedziny chemii, uzupełnią one tylko ich możliwości. ■

Krzysztof Orlński



8. Laureaci Nagrody Nobla z chemii w roku 2025 (od lewej):

Richard Robson, Susumu Kitagawa, Omar M. Yaghi (Ill. Niklas Elmehed © Nobel Prize Outreach)

Richard Robson (ur. w 1937 w Glusburn, Wielka Brytania) jest profesorem chemii na australijskim uniwersytecie w Melbourne specjalizującym się w chemii polimerów.

Susumu Kitagawa (ur. w 1951 w Kioto) jest profesorem chemii na japońskim uniwersytecie w Kioto, specjalizującym się w chemii związków kompleksowych.

Omar Mwanne Yaghi (ur. w 1965 w Ammanie, Jordania) jest profesorem Uniwersytetu Kalifornijskiego w Berkeley, specjalizującym się w projektowaniu i syntezie materiałów nowych technologii, m.in. MOF-ów (8).



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

Polityków, niezależnie od opcji, jaką prezentują, trzymałbym w pilnie strzeżonym miejscu, żeby nie mogli uciec. Karmił raz dziennie.

Lubi mnie jeden pies z Tulec, rasy beagle”.



O dużych liczbach

Temat tego artykułu przyszedł mi do głowy... po obejrzeniu w telewizji kilku odcinków amerykańskiego programu „Jak działa Wszechświat”. Jest interesujący, ale przeważające tam tematy, podobno zgodne z gustami współczesnego widza: to czarne dziury, kataklizmy, eksplozje i implozje oraz niszcząca wszystko olbrzymia energia i niemalże przemoc w gigantycznej skali. Na szczęście wszystko daleko od Ziemi, nawet nie w naszej Galaktyce. Ale co tam się dzieje, to strach pomyśleć.

Wszyscy wiemy, co to są liczby pierwsze. Zwrot „rozłożyć na czynniki pierwsze” wszedł do języka potocznego. Liczba pierwsza nie jest przez nic podzielna – z wyjątkiem tego, przez co musi być podzielna każda liczba: przez 1 i przez siebie samą. Liczby 1 nie zaliczamy do pierwszych, choć jest podzielna przez 1 i samą siebie (czyli też 1). Ale wykluczamy ją ze szlachetnej rodziny liczb pierwszych. Początkowe liczby pierwsze to 2, 3, 5, 7, 11, 13, 17, 19, 23, 29 i tak dalej.

Napisałem „i tak dalej”. Co to znaczy? O, właśnie. Już Euklides wiedział, że liczb pierwszych jest niekończące wiele. To proste, lecz piękne rozumowanie. Przyjmijmy, że mamy skończoną listę wszystkich liczb pierwszych. Mnożymy je wszystkie i dodajemy 1. Ta nowa liczba nie jest podzielna przez żadną z naszej listy, a więc jest nową liczbą pierwszą – wbrew założeniu, że wypisaliśmy wszystkie. Wiemy zatem tylko, że listę wszystkich liczb pierwszych ma w swoim komputerze dobry Pan Bóg. Pokazuje nam od czasu do czasu jej fragment, ale nie daje poznać, jak jest zbudowana.

Niektóre duże liczby pierwsze wyglądają ciekawie: 3737373737373, 999999996993, 9797979797979 i 1000008000001. Ta ostatnia jest *palindromiczna* – czytana od tyłu jest tą samą liczbą. Znamy wszystkie liczby pierwsze aż do 10^{12} – jest ich ponad 36 miliardów, wciąż znacznie więcej niż ludzi na Ziemi, ale więcej niż szacunkowa liczba wszystkich

Homo sapiens, od początku naszego gatunku. Mianowicie, według Population Reference Bureau (PRB) z 2022 roku jest to około 117 miliardów ludzi. Tyłu ludzi żyło na Ziemi.

Cóż to jest 117 miliardów? To zaledwie 117 z dziewięcioma zerami. Znacznie więcej jest komarów – i to każdego roku; podobno 12 000 na jednego człowieka. Atomów we Wszechświecie jest „tylko” 10^{79} – jedynka i siedemdziesiąt dziewięć zer.

Z bardzo dużymi liczbami mamy do czynienia na co dzień – tylko ich nie zauważamy. W Dzień Niepodległości, 11 listopada, poszedłem na spacer do Puszczy Kampinoskiej, a potem wstąpiłem do otwartej tego dnia kawiarni. Zamówiłem kawę i rogalika i dostałem rachunek numer

5476 3BE7C9 F9F2ED 30024A 26545C 8CC688 106C5B

Zapisany jest jak widać w układzie szesnastkowym, to jest takim, gdzie w podstawie numeracji mamy zero i 15 cyfr. Tradycyjnie oznaczamy je tak: A=10, B=11, C=12, D=13, E=14, F=15. W dziesiętkowym systemie mój rachunek ma numer

46459122280181401594143903048201823029275998
461565979004246715

A w binarnym

1010100 01110110 00111011 11100111
10011111 0010111011010011 00000000
00100100 10100010 01100101

0100010111001000 11001100 01101000
10000110 00100000 01100110 110001011011

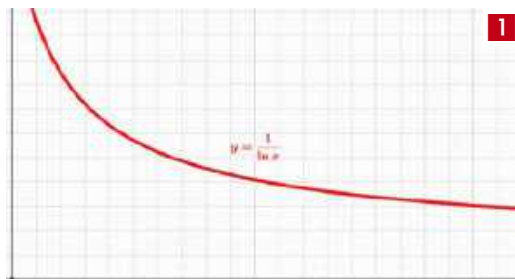
Jak wiemy, Stanisław Lem nie dostał Nagrody Nobla z literatury. Należała mu się, jak mało komu – ale widocznie tylko moim skromnym zdaniem. Nie jestem Szwedzką Akademią Literatury. W jego powieści „Głos Pana” jest fascynujący wątek. Dostajemy (my, Ziemianie) wiadomość z głębi kosmosu. Jest zakodowana w postaci ciągu zer i jedynek i „na pewno” nie jest pochodzenia naturalnego. „Ktoś” ją musiał wysłać. Cywilizacja nadawców już najprawdopodobniej nie istnieje. Co chcieli nam przekazać? Czy da się odczytać komunikat?

Lem pisze wprost, że każdy komunikat da się odczytać – jeżeli to jest komunikat! Nie będę rzecz jasna streszczać powieści, zawierającej – jak to u Lema – głębokie rozważania filozoficzne, ubrane w żarty, aluzje do życia psychicznego komputerów i lekką formę.

Tu jednak przypomniałem sobie „Głos Pana” Lema i puściłem wodze fantazji. Czy ktoś kiedyś odczyta z tego ciągu zer i jedynek, że pewnego dnia, w takiej a takiej kawiarni, taki a taki człowiek zamówił kawę z rogalikiem? Oczywiście, że nie – ale wyobraziłem to jeszcze dalej, że w tym ciągu zakodowane jest całe moje życie!

To tylko żarty. Wrócę do początkowego wątku, że gdzieś daleko od Ziemi dzieją się astronomiczne cuda-niewidy i że z liczbami pierwszymi jest podobnie. Najbliższa okolica jest przebadana na wszystkie strony. Co dalej, za jakąś umowną granicą?

W lipcu tego roku odbędzie się w Pensylwanii w USA kolejny, 30. Kongres Międzynarodowej Unii Matematycznej. Wspomnę, że dziewiętnasty był w 1983 roku w Warszawie. Na Kongresie przyznaje się najważniejszą nagrodę, jaką może otrzymać młody matematyk: Medal Fieldsa. Mówi się, że jest to „matematyczny Nobel”, chociaż finansowo wypada to nader skromnie: 15 000 dolarów kanadyjskich. Piłkarzom z wiodących reprezentacji za tę sumę nie opłaca się nawet kopnąć piłki. Ponadto jest to istotnie nagroda dla młodych – laureat nie może mieć więcej niż 40 lat. Od kilkudziesięciu lat przyznaje się cztery medale. W 2022 roku jedną z laureatek była Ukrainka Maryna Wiazowska (druga kobieta w historii). Wtedy zresztą Kongres miał się odbyć w Petersburgu, ale w związku z napaścią Rosji na Ukrainę przeniesiono go „do chmury” i tylko wręczenie medali było normalnie, stacjonarnie, w Helsinkach. Doceńmy, że pierwszy wykład po otrzymaniu medalu Maryna Wiazowska wygłosiła w Warszawie. Sądzę,



że był to również gest symboliczny za naszą pomoc dla jej rodaków – wtedy szczerą i spontaniczną.

Jednym z czworga laureatów był wtedy Brytyjczyk James Maynard, zajmujący się teorią liczb, w szczególności rozmieszczeniem liczb pierwszych. Wyobraźmy sobie, że wędrujemy po osi liczbowej. Zwykle liczby są zaznaczone zwykłymi małymi punkcikami, ale w liczbach pierwszych jest migająca lampka. Jak gęsto są te lampki? Czy dostatecznie oświetlają prostą? Na początku osi są nawet całkiem gęsto, potem coraz rzadziej. Wiemy mniej więcej jak. Gęstość tych lampek (czyli gęstość rozmieszczenia liczb pierwszych) maleje logarytmicznie (1).

Jest to tylko „prawda statystyczna”. Z wyjątkiem liczb 2 i 3 żadne dwie liczby pierwsze nie mogą stać bezpośrednio obok siebie, ale daleko od zera i jedynki można natknąć się na dwie lampki świecące niemal jedna obok drugiej, w odległości 2: na przykład 999999191 i 999999193 – a niewiele dalej dwie kolejne obok siebie: 1000037 i 1000039.

Mogą zdarzyć się długie ciemne przedziały. Po parze liczb bliźniaczych 1319 i 1312 następną liczbą pierwszą jest wprawdzie już 1327, ale potem dopiero 1361. To jeszcze nic. Po liczbie 1968188556641 występuje aż 601 kolejnych liczb złożonych. Przypominam, że liczba, która nie jest pierwsza, nazywa się złożoną (z wyjątkiem liczby 1, która nie jest ani pierwsza, ani złożona). Największy znany dziś odstęp między kolejnymi liczbami pierwszymi to 1113106 (ponad milion) i występuje po liczbie pierwszej, która ma 18662 cyfr i ze zrozumiałych względów jej nie podam. Zresztą, do czego byłąby nam potrzebna?

Jednym z obszarów badań Jamesa Maynarda były właśnie owe odstępstwa czy luki między kolejnymi liczbami. Wykazał on, że duże luki zdarzają się znacznie częściej, niż moglibyśmy przypuszczać, a dokładnie: niż to wynika z rozkładu statystycznego. To jest tak, jakby wykazał, że w podróży do krańca Wszechświata lecimy w zupełnie pustej przestrzeni międzygalaktycznej częściej, niż moglibyśmy się spodziewać. Tylko od czasu do czasu migną nam za oknem jakieś odległe galaktyki. Trochę może nam smutno

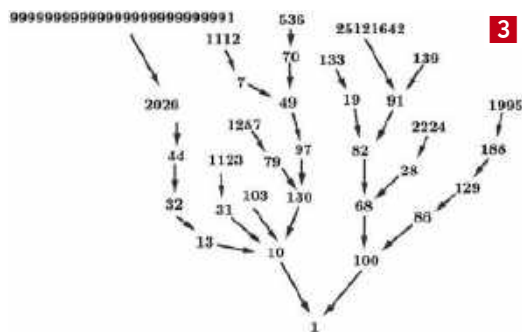
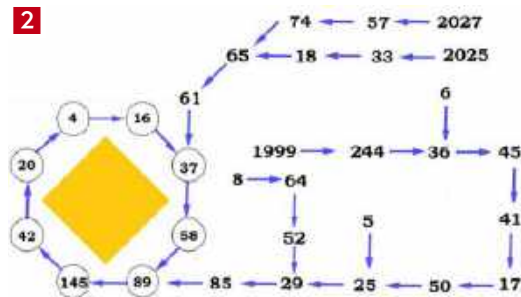
z tego powodu: wyglądamy za okno, a tu nic, ani gwiazdki! Podobno w przestrzeniach międzygalaktycznych mamy taką próżnię, że trudno ją sobie wyobrazić: co najwyżej jeden atom na metr sześcienny. Swoją drogą ciekawe, skąd tam się zaplątał? I czy też mu nie smutno?

Spora liczba Czytelników zadaje sobie pytanie: „komu potrzebne są takie badania?” Trochę to jest tak, jak ogólnie z kulturą. Dobre wykonanie *Koncertu e-moll* ma swoją wartość niezależnie od cen giełdowych na, powiedzmy, buraki cukrowe. Po drugie jednak: właśnie taka teoria liczb przydaje się i w kryptografii wojskowej, i na co dzień, gdy chcemy zadbać o swoje cyberbezpieczeństwo. Ale to oddzielny temat. Pozostańmy przy matematyce.

Rok 2025 zęgałem ze smutkiem, a to dlatego, że pasował do zadania: „mam n lat w roku n^2 ”. Istotnie, osoba urodzona w 1980 roku miała 45 lat w roku $2025=45^2$. Podobny warunek spełniała i moja babcia: urodzona w 1892 roku miała 44 lata w roku $1936=44^2$, a zadanie znałem w wersji dotyczącej milionera amerykańskiego Morgana, który urodził się w roku 1806 i miał 43 lata w roku śmierci Fryderyka Chopina: $1849=43^2$. Następną taką zależność nastąpi nieprędko. Sporo ludzi urodzi się w 2070 roku. Będą mieli, będzie miał 46 lat w roku $2116=46^2$.

Ale na szczęście rok 2026 też jest wyjątkowy, a to z powodu znanego zadania Hugona Steinhausa (1887–1972, matematyka lwowskiego, a potem wrocławskiego). Napiszmy jakąś liczbę naturalną. Niech będzie nią dzień, miesiąc i rok urodzenia Alberta Einsteina: 14 marca 1879=14031879. Podnieśmy każdą cyfrę do kwadratu i dodajmy wszystko: $1^2+4^2+0^2+3^2+1^2+8^2+7^2+9^2=1+16+0+9+1+64+49+81=221$. Z otrzymaną liczbą postąpmy tak samo: $221 \Rightarrow 2^2+2^2+1^2=9$ i dalej tak samo: $9^2=81 \Rightarrow 64+1=65$. Mówiąc matematycznie, iterujemy (czyli powtarzamy) proces: „oblicz sumę kwadratów cyfr”.

Od daty urodzin Alberta Einsteina doszliśmy do 65. Spójrzmy teraz na **rysunek 2**. Dostaniemy dalej 61, 37



i już wpadamy w pętlę, z której się nie wydostaniemy. Tak będzie dla „prawie każdej” liczby startowej. „Prawie każdej” – bo są wyjątki. Jednym z nich jest data urodzin Isaaca Newtona, 25 grudnia 1642 roku. Zobaczmy, co będzie: $25121642 \Rightarrow 2^2 + 5^2 + 1^2 + 2^2 + 1^2 + 6^2 + 4^2 + 2^2 = 91 \Rightarrow 9^2 + 1^2 = 81 + 1 = 82 \Rightarrow 8^2 + 2^2 = 64 + 4 = 68 \Rightarrow 36 + 64 = 100 \Rightarrow 1$ Można wykazać, że zawsze tak będzie: na ogół wpadniemy w pętlę, ale niekiedy dojdziemy do jedynki i z niej już się nie wydostaniemy, jak z czarnej dziury (3). Taką strukturę, jak na rysunku 3, nazywamy w matematyce drzewem. Istotnie, przypomina rozrośniętą lipę. Dokładne określenie to: drzewo jest grafem bez cykli. A czy wiecie, drodzy Czytelnicy, jak matematycy nazywają zbiór drzew? Nie zgadniecie: las!

A do czego dojdziemy, gdy liczbą startową będzie 2026? Zobaczymy $2026 \Rightarrow 2^2 + 0^2 + 2^2 + 6^2 = 44 \Rightarrow 4^2 + 4^2 = 32 \Rightarrow 3^2 + 2^2 = 13 \Rightarrow 1^2 + 3^2 = 10 \Rightarrow 1$. Czarna dziura! Czy to znaczy, że rok, który się właśnie zaczął, będzie ponury? Nie. W serialu astronomicznym, o którym wspominałem, czarne dziury są pozytywnymi bohaterami, bo z jednej strony niszczą życie (siły przyływowe rozerwą każdy obiekt), ale również przyczyniły się do powstania życia, bo stabilizują galaktyki i umożliwiają spokojne rozsiadanie pierwiastków po Wszechświecie. Niech żyją czarne dziury (byle tylko z dala od nas).

W zadaniach z pętlami i drzewami jest coś więcej niż tylko ciekawostki. Pętle i drzewa to podstawowe składowe programów komputerowych. To już temat na inną opowieść.

Autokomentarz. Rysunki 2 i 3 są mojego (M.Sz.) autorstwa. W zmienionej wersji wykorzystałem je w książce „91 gier matematycznych”, wyd. Oficyna Edukacyjna „Krzysztof Pazdro” (2025).

* * *

Nie ma teorii liczb bez Srinivasa Ramanujana (1887–1920). Do niego najlepiej pasuje powiedzenie Hugona Steinhausa: *Geniusz? Gen i już!* Jak można zgadnąć z nazwiska, Ramanujan był Hindusem. Pochodził z biednej rodziny, żyjącej w pewnym okresie w skrajnej nędzy. Nie kończył szkół i tracił

przyznane stypendia, bo nie interesował się niczym poza matematyką. W 1912 roku przyjaciele pomogli mu zdobyć posadę kancelisty z pensją 20 funtów rocznie. W pracy dawał sobie radę i mógł samodzielnie zajmować się matematyką, mówiąc, że natchnienie podsyła mu bogini Namagiri. W 1913 roku odważył się wysłać swoje zapiski do Anglii, do znanego matematyka Godfreya Harolda Hardy'ego. Ten osłupiał z wrażenia – wyniki Ramanujana były oryginalne i głębokie, świadczące o nieprawdopodobnym talencie autora. Przekonał swoich ksenofobicznych kolegów, żeby natychmiast ściągnąć Ramanujana do Anglii i stworzyć mu godne warunki do rozwoju. Przez kilka lat ich współpraca dawała świetne wyniki. Niestety, Ramanujanowi nie służył angielski klimat i niektóre zwyczaje (na przykład skąpo ogrzewane pomieszczenia). Szybko zaczął chorować i już w 1920 roku zmarł. Zostawił tyle artykułów i notatek, że do dziś (2026) nie wszystko z jego spuścizny jest opracowane. Przytaczam wzór, na widok którego Godfrey Harold Hardy osłupiał.

$$\frac{1}{\pi} = \frac{2\sqrt{2}}{9801} \sum_{k=0}^{\infty} \frac{(4k!)(1103+26390k)}{(k!)^4 \cdot 396^{4k}}$$

Nie trzeba być matematykiem, żeby zrozumieć, że taki wzór jako produkt myśli biednego hinduskiego samouka musi być czymś nie tylko wyjątkowym, ale „nadzwyczaj wyjątkowym”. Jako wybitny matematyk, Hardy docenił nie tyle skomplikowany wygląd tego wzoru, ale jego głębsze znaczenie matematyczne. Bardzo dobrze przybliży on liczbę π . Sprawdziłem dziś na swoim domowym laptopie: już cztery wyrazy tego szeregu (tej sumy) dają 32 cyfry rozwinięcia!

Z pobytem Ramanujana w Anglii związanych jest kilka anegdot. Pierwsza dotyczy jego niefortunego przejścia przez trawnik w Trinity College. Po prostu poszedł za kimś na skróty. Wybuchła afera. Prawo do tego mieli bowiem tylko członkowie kolegium – a on nie od razu nim został, a poza tym był Hindusem, a więc kimś gorszej rasy. Dyskutowano, jak ukarać go za ten nieobyczajny wybryk, może zastosować klasyczny dzisiaj push-back, czyli wsadzić na statek i wysłać z powrotem do Madrasu. Na szczęście Hardy był wpływowym członkiem kolegium i powiedział, że taki geniusz może chodzić, gdzie chce. Podobno, gdy Ramanujan został już pełnoprawnym członkiem, demonstracyjnie przechadzał się po trawniku, wywołując zresztą złość niektórych źle do niego nastawionych uczonych lordów.

Jednak najsłynniejsza historia z Ramanujanem związana jest z jego pobytem w szpitalu i liczbą 1729. Można powiedzieć, że jest to klasyczny

„suchar matematyczny”. Godfrey Hardy odwiedził chorego przyjaciela i współpracownika. Rozmowa się nie kleiła, w pewnej chwili Hardy, chcąc znaleźć temat do konwersacji, powiedział: „Wiesz, przyjechałem tu taksówką, bo już dałem wolne mojemu szoferowi”. Ramanujan ożywił się. „Tak? A jaki był jej numer boczny?”. „Oj, mało ciekawy, 1729”. „Jak to mało ciekawy? Przecież to najmniejsza liczba naturalna, którą można rozłożyć na sumę sześciątów na dwa różne sposoby!”

No i zaczęło się. Istotnie, $1729 = 1^3 + 12^3 = 9^3 + 10^3$. Nawet w epoce przedkomputerowej udało się znaleźć kilka następnych liczb o podobnej własności, to znaczy będących sumą sześciątów na dwa sposoby. Następną po 1729 jest $4104 = 2^3 + 16^3 = 9^3 + 15^3$, potem idzie $13832 = 2^3 + 24^3 = 18^3 + 20^3$, a dziś możemy łatwo tę listę przedłużać i przedłużać, chociaż w pewnym momencie komputer się nam zadławi. Na 30 000. miejscu tej listy znajduje się 520 890 296 211, ale samo znalezienie takiego podwójnego rozkładu wcale nie jest łatwe.

Trudno powiedzieć, kto pierwszy spojrział na anegdotkę Ramanujana z innej strony i zaczął poszukiwać liczb, które dadzą się przedstawić w postaci sumy sześciątów dwóch liczb, ale na więcej niż dwa sposoby. Zagadnienie chwyciło. Okazało się trudne, obrosło już w całą teorię, a w komputerach w kilku uniwersytetach przepalają się lampy. Tak, tak, wiem, że w komputerze nie ma lamp – to tylko taka figura retoryczna! Otóż gdy mówimy „liczba taksówkowa numer n ”, to mamy na myśli najmniejszą liczbę, która jest sumą dwóch sześciątów na n sposobów. Ponieważ słowo „taksówka” zaczyna się na T (po angielsku też: taxi-cab), więc n -tą taką liczbę oznaczamy przez $T(n)$. W ten sposób liczba Ramanujana to $T(2)$. Ramanujan nie był zresztą odkrywcą tej własności. Znana była już w XVII wieku. Liczbę $T(3)$ odkrył J. Leech w 1957 roku: $87\,539\,319 = 167^3 + 436^3 = 228^3 + 423^3 = 255^3 + 414^3$. Następną, czwartą, poznaliśmy dopiero w 1991 roku i jest nią $6\,963\,472\,309\,248$, piątą znalazł w 1999 roku David W. Wilson, $T(5) = 48\,988\,659\,276\,962\,496$. Kilka lat trwały poszukiwania szóstej. Kandydatką była $T(6) = 24\,153\,319\,581\,254\,312\,065\,344$. W 2008 roku Uwe Hollerbach napisał szczerze: „my computers screwed-up”, co w łagodnym tłumaczeniu brzmi „moje komputery wysiadły ze zmęczenia”, ale w końcu odpoczęły, a Uwe sprawdził program, którym się posługiwał, usunął błędy i wykazał, że $T(6)$ jest właśnie takie. Jest zwyczajem w teorii liczb, że wynik jest uznany dopiero wtedy, gdy zostanie potwierdzony przez kogoś innego. Tak się stało i mamy już (dopiero?) liczbę $T(6)$.

Dalej już niewiele wiadomo. Warto zauważyć, że jednym z pionierów poszukiwań jest Jarosław



Liczba Ramanujana na taksówce (źródło: Internet)

Wróblewski z Uniwersytetu Wrocławskiego. Jego osiągnięciem (we współpracy z Christianem Boyerem) jest oszacowanie liczb taksówkowych aż do $T(22)$. Konkretnie, $T(22)$ na pewno istnieje i ma nie więcej niż 160 cyfr. Są to już jednak obliczenia dla małych zielonych ludzików z Marsa... albo dla komputera kwantowego, który podobno Chińczycy już budują. Czy życzyć im powodzenia? No cóż, postępu nie da się zatrzymać i tylko można się zastanawiać, czy to jest prawdziwy postęp, czy tylko technologiczny?

Na konferencji Stowarzyszenia na Rzecz Edukacji Matematycznej, we wrześniu 2025 roku, dostałem od kolegi z Częstochowy niesamowity prezent: suwak logarytmiczny. Pokazałem go studentom informatyki. Byli zdumieni i zachwyceni, choć lekko oszołomieni: „na czymś takim naprawdę dało się coś obliczyć?”. Nauczyłem ich podstawowych zasad suwaka. Ale nie będę wymagał na egzaminie. ■

Michał Szurek

Historia ludzkości na haju. Skandaliczne, ale prawdziwe spojrzenie na historię „pod wpływem”

Sam Kelly

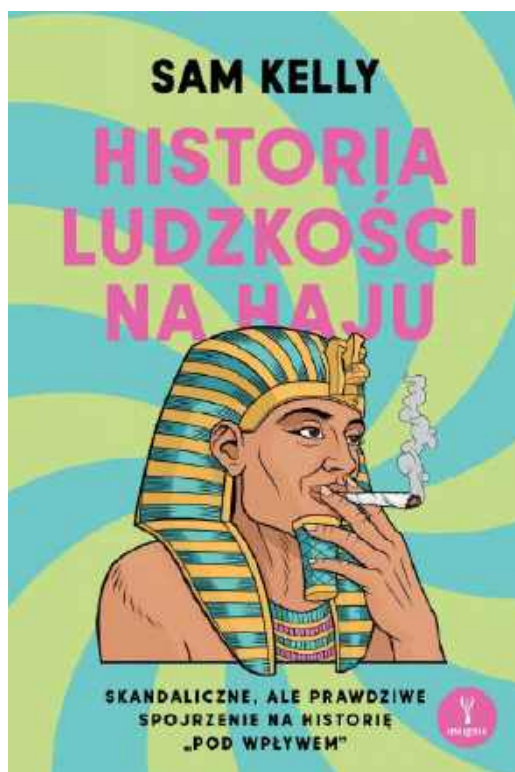
Wydawnictwo Insignis, liczba stron: 400,
cena na okładce: 49,99 zł

Porywająca, pełna humoru i absolutnie prawdziwa książka o słynnych postaciach historycznych i ich życiu „pod wpływem”. Poznajcie zaskakujące fakty o Ramzesie II, Jerzym Waszyngtonie, Vincencie van Goghu, Beatlesach i wielu innych ikonach dziejów i ich używkach!

„Ta książka w niezwykle sposób ożywia znane postacie i epoki. Po jej lekturze już nigdy nie spojrzycie na historię tak samo” – Harlan Coben. Wiedzieliście, że Aleksander Wielki był notorycznym pijakiem, Szekspir ponoć nie stronił od marihuany, a Steve Jobs uważał, że LSD pomogło mu poszerzyć wyobraźnię? A może obito się wam o uszy, że Zygmunt Freud tak uwielbiał kokainę, że nieustannie ją zażywał i przepisywał swoim pacjentom?

W swojej książce Sam Kelly opowie wam o tym, o czym raczej nie dowiedzieliście się od swoich nauczycieli historii, ujawni naprawdę dziwaczne aspekty życia znanych postaci i okraszy to wszystko błyskotliwym, humorystycznym komentarzem. Zacznie od starożytnej Grecji („Wyrocznia delficka w halucynogennych oparach”), a skończy na współczesności („Carl Sagan na astronomicznym haju”).

Przekonacie się, że na kartach historii jest mnóstwo używek i osób, które z nich korzystały. Ta książka to pasjonująca i rozyrywkowa podróż przez ich upadki i wzloty... znaczy: odloty. Zdziwicie się, jak olbrzymi wpływ na historię ludzkości miały rozmaite substancje.





Szkoła Wynalazców

dozwolone do lat 15

Mieście temat w sam raz dla was. Jest oczywistym faktem, że język polski gnie. Młodzi ludzie nie potrafią napisać pisma do urzędu, do dziewczyny, do rodziców z pobytu na wycieczce, itp. królują słowa angielskie i angielsko-podobne. Poza tym „skrótowce” typu: na ra, spoko, oki, dil (!!!). No, ale jeśli młodzież nie czyta, to skąd ma czerpać wzorce dobrej, poprawnej polszczyzny? Opracować system i metodę, zachęcającą młodzież do czytelnictwa.

Oczekiwaliśmy propozycji realnych dających szansę na pobudzenie ciekawości i chęci poczytania książek – wciąż najdoskonalszego źródła wiedzy i języka.

Roman Tarło: zasadą harcerskiego wychowywania było: „uczymy wzorem” i hasło to – jak wynika z historii – było skuteczne. Zdaniem Romana młodzież powinna zobaczyć ludzi cieszących się autorytetem – czytających książki. Powinni to robić rodzice i to zaczynając od najwcześniejszych lat: najpierw czytać dziecku jakieś ciekawe bajki i czytając, przerywać dzienną „porcję” tekstu na jakimś miejscu, w którym akcja się rozwija, ale następny odcinek będzie... jutro.

Współcześni rodzice nie bardzo się sprawdzają w roli osób, zachęcających do czytania, bo... sami nie czytają. I tu oczywiście działa hasło „uczymy wzorem”. Nadziejemy jeszcze mogą być dziadkowie – przedstawiciele tych pokoleń, które czytały. Dziadkowie mogą również wystąpić w roli słuchaczy tekstów, które mogą im czytać wnuki.

Jerzy Nowak: przyczyną może być błędna na ogół polityka w realizacji nauczania języka polskiego. Nikt nie wmówi młodzieży 14–17-letniej, że „Ferdynand” to ciekawa książka, że „Kordian” również, itd. Spis lektur powinien być propozycją, a nie nakazem. Ta propozycja powinna zawierać książki od dawna znane, chętnie czytane, np. książki przygodowe Arkadego Fiedlera, także „Trylogia” H. Sienkiewicza, czy seria „Tomków” Alfreda Szklarskiego. Dobór książek do listy propozycji powinien być starannie opracowany i skonsultowany z młodzieżą.

Być może problem polega na nieciekawym doborze lektur obowiązkowych. Seria o przygodach Harry’ego Pottera jest książką czytaną nawet pod kołdrą przy świetle latarki (tak!)

Seria książek o „Ani z Zielonego Wzgórza” i dalsze to do dziś bestsellery, głównie dla dziewczynek, chociaż chłopcy też chętnie je czytają.

Stefan Kwiatkowski: można zorganizować konkursy, polegające na ogłoszeniu listy kilku książek, a po upływie odpowiednio długiego czasu przeprowadzić zawody ze znajomości tych książek. Zadania konkursowe nie mogą być zbyt banalne, powinny być ciekawe, np. „Czym powinien zajmować się malarz, który namalował portret Oleńki, który nie bardzo podobał się Kmicicowi?” Odp: „piece mu malować, nie one specjały, którymi oczy pasę”. I tym podobne.

Być może wprowadzenie elementu rywalizacji przy odpowiednim doborze książek, skonsultowanym z młodzieżą, dałoby dobry rezultat.

Czytelnictwo jest ważnym elementem nie tylko kultury, ale także gospodarki, polityki i nawet technologii. Umiejętność sprawnego porozumiewania się, precyzyjnego wyrażania myśli, to elementy, które można wykształcić dzięki czytaniu. Wdrożenie młodzieży do sięgania po książkę to zadanie trudne, ale może jednak możliwe do rozwiązania. Propozycje kolegów są interesujące, ale wydają się nieco słabe w sensie oddziaływania na młodych ludzi. Wszelkowiedza mediów: TV, komputerów i smartfonów, to potężny czynnik odciągający młodzież od lektury. Może jednak należałoby wysłać ekipę do Finlandii i zbadać na miejscu, na czym polega ich sukces czytelniczy?

Kolegom gratuluję i zachęcam do rozwiązywania dalszych zadań.

Nowe zadanie

Ruch kołowy i rowerowy – mamy tu na myśli rowery i hulajnogi elektryczne – zgęszcza się i rośnie liczba wypadków z udziałem użytkowników tych nowinek technicznych. Użytkownik leży i nie bardzo wiadomo, czy jest groźnie ranny, czy tylko poobtłukiwany nieszkodliwie. Wiadomo, że najważniejszą częścią ciała jest głowa – chroniona kaskiem. Ale skąd wiadomo, czy uderzenie głową w asfalt jest groźne,



czy wszystko kończy się na strachu. Przydałby się kask taki, który jakoś by sygnalizował siłę uderzenia np. zmianą koloru. W takim razie zadaniem waszym będzie: Opracować ideę kasku rowerowego, który sygnalizowałby groźne dla zdrowia uderzenie głową w jezdnię, np. zmieniając kolor lub wysuwając jakąś plaketkę barwną, itp.

Oczywiście ciśnie się natychmiast pomysł jakiegoś gadżetu elektronicznego. Ale głównym zadaniem

sygnalizatora jest reagowanie na przyspieszenie towarzyszące uderzeniu. To zaś może zrobić zwykła sprężyna obciążona ciężarkiem, która ugnie się więcej lub mniej w zależności od siły zderzenia. Ta zasada prowadzi do prostych rozwiązań, których ideę macie opracować i wysłać do redakcji MT. Wszystkim życzymy dobrych pomysłów, bezpiecznej jazdy na rowerze i hulajnodze i przypominamy o terminie nadsyłania propozycji – do końca lutego br.

Klub Wynalazców

bez ograniczeń wieku

Waszym zadaniem było opracować przepis na dowolne danie narodowe kuchni japońskiej lub danie, będące modyfikacją oryginału z uwagi na produkty.

Wychodząc od powiedzenia: „Odkrycie nowej potrawy daje więcej szczęścia ludzkości niż odkrycie nowej gwiazdy” (Anthelme Brillat-Savarin), nie możemy lekceważyć przepisów na ciekawe i oryginalne dania, które noszą często nazwy od nazwisk ich twórców, jak np. zrazy à la Nelson, zupa cebulowa Przypkowska, tort Sachera i wiele innych. Mieliście szansę zasłać modyfikacjami kuchni japońskiej na polskich produktach.

Ryszard Bogucki: dziś sprawa jest dziecinnie prosta: sięgamy do ChatGTP i mamy to, o co chodzi np. przepis na japońskie curry z ryżem: oryginalny przepis japoński:

Oryginalne składniki:

- gotowy blok curry roux (np. Golden Curry)
- sos Worcestershire
- jabłko starte do sosu
- ryż japoński

Dostępne w Polsce zamienniki:

- Roux – można zrobić samemu: podsmaż masło z mąką, dodaj curry, kurkumę, odrobinę kakao i sos sojowy.
- Sos Worcestershire – można zastąpić octem balsamicznym z łyżeczką sosu sojowego i szczyptą cukru.
- Ryż japoński – dowolny ryż kleisty lub arborio.

Technologia nieco niejasna, ale każdy, kto ma elementarnej pojęcie o gotowaniu, będzie wiedział, jak zrobić to curry.

ChatGTP staje się uniwersalnym nauczycielem wszystkiego! Ciekawe, jak daleko to dojdzie za 10–20 i 50 lat!

Miłosz Warecki – pisze: dzisiaj takich rzeczy nie trzeba wymyślać. Odkąd jest ChatGTP sprawa jest prosta. Bierzymy dowolny przepis japoński, np. na onigiri – kulki ryżowe z nadzieniem

Wersja klasyczna z polskimi składnikami

- 1 szklanka ryżu do sushi ⇒ zamiennik: ryż krótkoziarnisty (ryż do risotto)
- 1 i ¼ szklanki wody
- 1 łyżka octu ryżowego ⇒ zamiennik: ocet jabłkowy lub winny (rozcieńczony wodą)
- ½ łyżeczki cukru
- ¼ łyżeczki soli

Nadzienie:

Klasycznie: grillowany łosoś, tuńczyk z majonezem lub umeboshi (śliwka japońska)

Zamienniki polskie:

- pieczony łosoś lub wędzona makrela
- tuńczyk z majonezem i odrobiną chrzanu
- śliwka w occie lub marynowana żurawina (zamiast umeboshi)
- arkusze nori (suszone wodorosty) ⇒ zamiennik: cienki liść sałaty, szpinaku lub glonów wakame w płatkach namoczonych w wodzie

1. Gotowanie ryżu: Ryż przepłucz kilka razy w zimnej wodzie, aż woda będzie prawie przezroczysta. Zalej 1 i ¼ szklanki wody, przykryj i gotuj na małym ogniu ok. 15 minut, aż wchłonie wodę. Odstaw pod przykryciem na kolejne 10 minut.

2. Doprawienie ryżu: Wymieszaj ocet, cukier i sól, dodaj do gorącego ryżu i delikatnie wymieszaj, żeby go nie roznieść.

3. Formowanie: Zwilż dłonie w wodzie z odrobiną soli. Weź porcję ryżu, zrób wgłębienie, włóż łyżeczkę nadzienia i uformuj kulkę lub trójkąt.

4. Zawinięcie: Owiń dolną część onigiri paskiem nori (lub liściem sałaty/szpinaku).

5. Podanie: Onigiri można jeść ciepłe lub schłodzone. Świetnie smakują jako przekąska lub lunch.
Jak widać, gotowanie jest sztuką i proces przyrządzania potraw wcale nie różni się złożonością i wymaganiami od procesów technicznych.

Obu kolegom gratuluję twórczego skorzystania z ChatGPT i czekamy, co będzie dalej!

Nowe zadanie

Wracamy na solidny grunt techniki. Mamy w mieszkaniach dzbanuszki z wkładami filtrującymi, mamy też niekiedy układ kilku filtrów zainstalowanych na dopływie wody do mieszkania lub domku.

Są to urządzenia profesjonalne i może z wyjątkiem dzbanków – dość drogie. A przecież istnieją inne systemy filtrowania, wśród nich takie, które wykorzystują energię przepływającej przez nie wody. Właśnie idea takiego filtra to temat dla was: zaproponować ideę filtra do wody pitnej, wykorzystującego energię przepływającej przez niego wody do oddzielania cząstek zanieczyszczeń od czystej wody.

Woda przepływając przez filtr niesie energię, która może zrobić wiele. Może np. napędzać jakiś mechanizm wykorzystujący siły odśrodkowe, lub filtry mechaniczne, samooczyszczające się, itp.

Możliwości jest wiele i przy wyobraźni i odrobinie pomysłowości można co najmniej kilka koncepcji takiego oczyszczania wody wymyślić. Wszystkim życzymy dobrych pomysłów i przypominamy o terminie nadsyłania propozycji: koniec lutego br.

Vademecum Młodego Wynalazcy

Zadziwiającym zjawiskiem, nasilającym się w ostatnich 20–30 latach, jest spadek wyobraźni przestrzennej, zwłaszcza u młodzieży. Jeszcze 30 lat temu, odbierając prace konstruktorskie, np. projekty reduktora dwustopniowego z przekładniami walcowymi, chcąc sprawdzić samodzielność pracy ucznia, wystarczyło naszkicować na rysunku złożeniowym reduktora linię oznaczającą przekrój przez reduktor i poprosić o naszkicowanie tego przekroju. Ogromna większość uczniów zadanie to wykonywała poprawnie, oczywiście z wyjątkiem tych, którym ten reduktor rysowały „zaprzysiężone siły”. Dzisiaj takiego pytania nie zadaje nikt, bo ogromna większość uczniów wyłożyłaby się na nim kompletnie. Problem jest poważny, bo przecież programy 3D nie są w stanie zastąpić w pełni koncepcyjnej pracy umysłu, który musi „widzieć” projektowany mechanizm, zanim go narysuje. Geometria przestrzenna jest zmorą bardzo wielu uczniów szkół średnich, bo ona wymaga właśnie widzenia bryły przestrzennie. Inaczej nie da się zadania rozwiązać.

W 1962 lub 63 roku na maturze z matematyki było zadanie, na którym padło bardzo wielu uczniów. Treść bardzo prosta, ale kryjąca podstępną właściwość bryły przestrzennej.

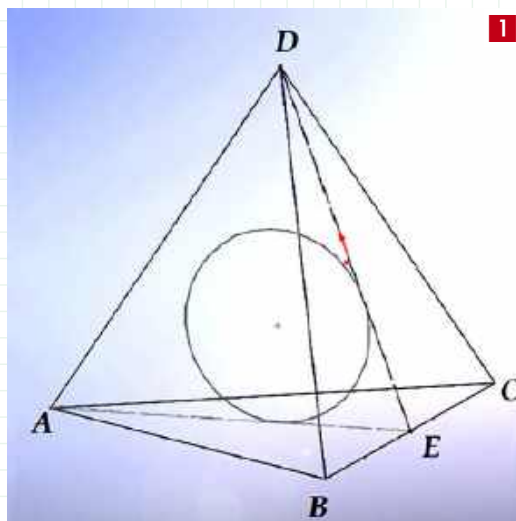
A oto treść tego zadania:

W czworościan foremny wpisano kulę, styczną do wszystkich czterech ścian. Różnica średnicy

kuli i krawędzi czworościanu wynosi 3 cm. Obliczyć długość krawędzi czworościanu i średnicę kuli.

Sytuację tę przedstawia **rysunek 1**.

Przekrój niezbędny do rozwiązania tego zadania poprowadzono przez wysokości ścian: podstawy i jednej ze ścian bocznych, przez punkty A E D. Wynikiem przekroju jest koło wpisane w trójkąt (nierównoboczny!!!) A E D tak, że koło, czyli przekrój kuli, nie styka się z krawędzią AD. Środek kuli, w przekroju – koła, leży na wysokości czworościanu. To trzeba





wiedzieć, a dalej to już prosta trygonometria. Wszyscy, którzy tego zadania nie rozwiązyali, założyli, że koło w przekroju jest styczne do krawędzi AD. Był to oczywiście, elementarny błąd.

Zadaniem zadawanym towarzyszyko przy popijaniu kawy jest taki oto problem: mając sześć patyczków (zapalek), zbudować cztery trójkąty.

Zadania tego na płaszczyźnie nie da się rozwiązać. Dopiero „po długich i ciężkich zmaganiach” niektórzy dochodzą do wniosku, że trzeba patyczki ułożyć jako krawędzie czworoscianu foremnego i wtedy już wszyscy widzimy, że da się z sześciu patyczków ułożyć cztery trójkąty równoboczne (2).



2

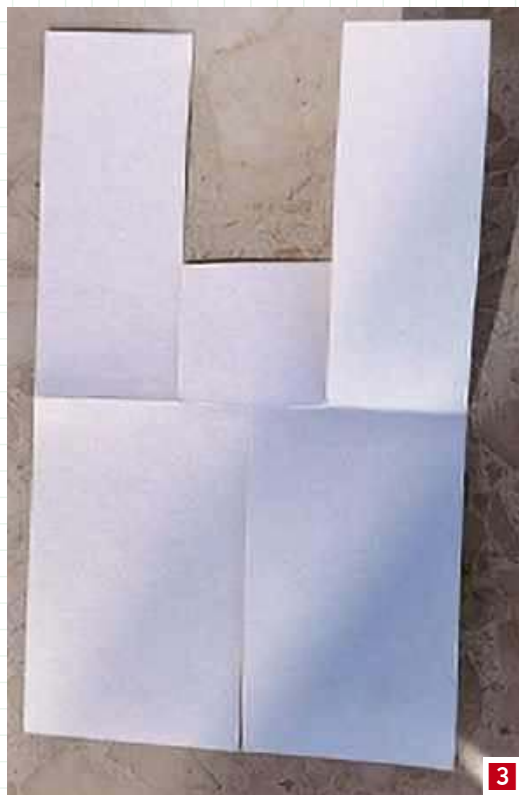
Pokonanie przyzwyczajenia do „płaskiego” widzenia rzeczywistości jest – jak się okazuje – dość trudne.

Jest też taka dziwna zagadka, która przysparza sporo trudności osobom nienawykłym do widzenia przestrzennego, zwłaszcza gdy do tej przestrzeni włączymy ruch. Najpierw należy wykonać jeden prosty element (3). Element należy wstępnie zagiąć, tworząc „główne zagięcie” – tak jak to widać na rysunku 3.

Należy robić to w „tajemnicy” i pilnować, żeby nikt nie podpatrywał. Następnie obracamy szeroki pas razem z wąskim paskiem o około 90°, przeciwnie do ruchu wskazówek zegara, aby przyjął położenie równoległe do wąskiego pasa po przeciwnej stronie głównego zagięcia, równocześnie wąski pas zajmuje pierwotną pozycję paska szerokiego. Trzeci element – skrócony – ustawiamy w pozycji pionowej. W rezultacie powstaje figurka, przypominająca domek z kominem (4).

„Domek” ustawiamy na stole i proponujemy zebranym kolegom wykonanie takiej samej figurki. Dajemy im papier i nożyczki.

Jeden warunek jest bardzo ważny: figurki nie wolno brać do ręki; można ją jedynie oglądać z każdej



3

strony bez dotykania rękami. Zdziwiające jest, jak wiele osób nie potrafi wykonać takiej figurki. Uważają, że „kominek” jest przyklejony i że w ogóle bez kleju tego się zrobić nie da. Po ujawnieniu „tajemnicy” domku ogólna konsternacja: to takie proste!

I czworoscian, i domek to problemy nieco złośliwe, ale proste. Inaczej sytuacja wygląda przy konstruowaniu narzędzi takich jak: wykrojniki, tłoczniaki, i narzędzi kombinowanych jak np. przyrząd wycinający wykrój z blachy i w dalszym ciągu tego samego suwu prasy kształtujący z tej blachy jakiś przestrzenny obiekt. Sytuacja się komplikuje przy projektowaniu narzędzi wielotaktowych. Ma to miejsce przy produkcji elementów z taśmy z blachy stalowej. Wprowadza się taśmę do przyrządu, który najpierw wykonuje pierwszy zbieg np. przetłoczenie dla uzyskania kształtu „rondelka”, pas blachy przesuwają się do przodu i w następnym takcie prasa wycina gotowy detal. Oczywiście tych taktów może być więcej.

Projektowanie takich narzędzi zaczyna się w głowie konstruktora, który musi sobie wyobrazić, jak ten przyrząd ma działać, niejako „widzieć go” w myśli i dopiero, gdy wszystko „gra”, może zacząć rysować. Oczywiście proces tłoczenia na prasie



to zespół kolejnych zabiegów, dość dynamiczny. Nie chodzi o jeden zabieg, ale sekwencję zabiegów, połączonych taśmą blachy.

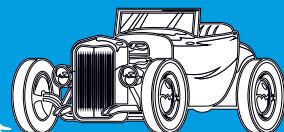
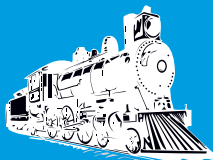
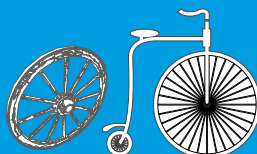
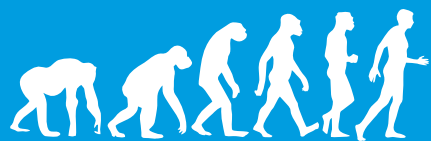
Ciekawym zadaniem są konstrukcje maszyn do pakowania różnych przedmiotów. Jednym z trudniejszych zadań, na jakie trafiłem w swojej praktyce zawodowej, było zaprojektowanie maszyny pakującej okrągłe wkłady do dzbanuszków filtrujących wodę pitną, do pudełek po 2, 4, 6, 8, 9 i 12 sztuk, czyli maszyna musiała realizować 6 programów dla różnej liczby filtrów, pakowanych do różnych pudełek. Ogólny schemat działania maszyny był następujący: maszyna powinna pobierać pudełka dostarczane w formie płasko złożonej, następnie powinna je rozłożyć i wkładać pomiędzy dwie ścianki transportera, który przenosił puste jeszcze i otwarte z obu stron pudełko na stanowisko zaklejania denka, następna stacja to wkładanie przez zespół popychaczy pneumatycznych odpowiedniej liczby wkładów do pudełka, później następowało zaklejanie pudełka i koniec. Maszyna – jak już wspomniano – powinna pakować różne liczby wkładów do różnych pudełek. Czyli odległości ścianek transportera należałoby każdorazowo zmieniać przy zmianie programu pakowania. Ścianek transportera było ok. 80, zmiana metodą odkręcania i przykręcania dla dostosowania ich odległości do szerokości kolejnego asortymentu pudełek byłaby nie do przyjęcia. Drugi problem to sposób pobierania pudełek z magazynka. Pudełko musiało być pobierane tak, aby przy okazji rozkładało się do postaci prostopadłościenną rury z otwartymi obu końcami. Nie dało się tego zrobić w przestrzeni pomiędzy górnym i dolnym biegiem łańcucha transportera

i trzeba było sięgać po pudełko spod dolnego łańcucha, trafiając w luki pomiędzy ściankami. W opisie wygląda to dość prosto. Ale to trzeba było zaprojektować i wykonać! Ruchome elementy maszyny musiały się mijać bezkolizyjnie, z dość sporą szybkością. Pracowałem wtedy w programie 2D Auto CAD. Łatwo zrozumieć, że należało mocno wytężyć wyobraźnię przestrzenną, żeby w ogóle wziąć się do takiej roboty. Program 3D Solid Works, na którym później pracowałem, wcale nie zwalnia od obowiązku ćwiczenia wyobraźni przestrzennej. W Solidzie można zobaczyć przestrzenny rysunek, nawet złożonej maszyny, można ją nawet „uruchomić”, ale istota projektu musi powstać w głowie konstruktora. Wynika stąd wniosek, że ćwiczenie wyobraźni przestrzennej jest obowiązkiem każdego kandydata na inżyniera konstruktora, tym bardziej na wynalazcę.

Jak pracować nad wyrobieniem umiejętności widzenia przestrzennego? W zasadzie metodyka jest dość prosta i niemal oczywista. Na początek dobrze jest wykonywać szkice odręczne prostych kształtów np. mebli domowych, w perspektywie „kawalerskiej”, później dobrze jest przejść do szkicowania budynków i ćwiczenia perspektywy. Tu cenną pomocą mogą być reprodukcje obrazów Warszawy, wykonywanych z ogromną precyzją przez Bernarda Bellotto (Canaletto). Jego obrazy posłużyły przy odbudowie warszawskiego Starego Miasta.

Równolegle warto ćwiczyć „prace ręczne”, czyli majsterkowanie, w najprostszych formach. Dziś młodzież ćwiczy głównie „tresowanie” myszki komputerowej. Oczywiście umiejętność posłużenia się technikami informatycznymi jest dziś obowiązkiem, ale należy pamiętać, że dowolny projekt musi powstać najpierw w głowie, a później można go przenieść do komputera. Zbytnią wiarą w nieomyślność komputera, kalkulatora czy smartfona prowadzi do uzależnienia się i do sytuacji wręcz patologicznych, jak opisywana przez prof. Docenka sytuacja na III roku Wydziału Fizyki Uniwersytetu Piotra i Marii Curie w Paryżu, gdzie student, który miał obliczyć promień kuli ziemskiej na podstawie danych astronomicznych, napisał wynik: $R=10$ mm. Na pytanie profesora: „czy ty zdawałeś sobie sprawę z tego, co napisałeś?”, student beztrząsco odpowiedział: „ale mnie tak wyszło z kalkulatora” (!!!). I o tym przypadku trzeba pamiętać i nie dać się znieślić żadnym maszynom. ■

**Prezes Klubu Wynalazców
Champion TRIZ
Jan Boratyński**



czasy prehistoryczne

Tkaniny

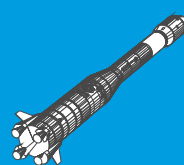
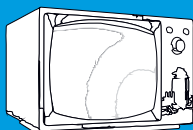
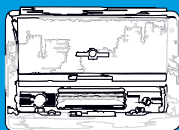
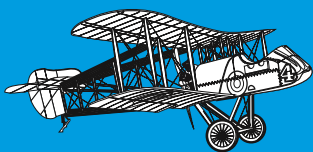
Nie ma zgody w świecie nauki co do tego, kiedy ludzie zaczęli nosić ubrania, a szacunki różnych ekspertów są bardzo zróżnicowane, od 40 tys. do nawet trzech milionów lat temu. Są dowody oparte na znaleziskach w Maroku na to, że ubrania były produkowane już 90–120 tys. lat temu. Najstarsze artefakty barwionych włókien lnianych znaleziono w prehistorycznej jaskini w Gruzji i datuje się je na 36 tys. lat. Figurka sprzed 25 tys. lat, zwana „Wenus z Lespugue”, znaleziona w południowej Francji, przedstawia spódnicę z tkaniny lub skręconego włókna (1). Niektóre inne rzeźby z Europy były ozdobione nakryciami głowy lub pasami materiału. Ok. 10 tys. lat p.n.e. ludzie zaczęli hodować owce na mięso. W miarę hodowli opartej na doborze ich wełna stała się bardziej użyteczna, a około 5 tys. lat p.n.e. była już wystarczająco dobra, aby można ją było prząść.

ok. 8000–5000 p.n.e.

Włókna, z których wytwarza się przędzę, a następnie splata się ją w siatkę, pętelki, dzierga lub tka, tworząc tkaniny, pojawiły się na Bliskim Wschodzie pod koniec epoki kamienia łupanego. Najwcześniejsze znane tkaniny z Bliskiego Wschodu to prawdopodobnie tkaniny lniane używane do owijania zmarłych. Zostały one odkryte podczas wykopalisk neolitycznych w Çatalhöyük w Turcji (2). Istnieją dowody na uprawę lnu od około ósmego tysiąca p.n.e. na Bliskim Wschodzie. W Szwajcarii i na terenach dzisiejszej Gruzji znaleziono pozostałości tkanin lnianych sprzed siedmiu tysięcy lat. Prawdziwym imperium lnianym stał się starożytny Egipt. Z włókna lnianego wyrobiano tam tkaniny, w które następnie owijano mumie. Wytwarzano również sznury, liny, sieci rybackie i żagle. Malowidła egipskie przedstawiają egipskich mężczyzn noszących lniane spódnice i kobiety w wąskich sukienkach z różnymi rodzajami kieszek i kurtek. W Egipcie płótno lniane ceniono za lekkość i przewiewność, przydatne cechy w gorącym klimacie. Lniane żagle pomagały w ekspansji morskiej na Morzu Śródziemnym.

3300–1300 p.n.e.

Archeolodzy prowadzący wykopaliska stanowisk cywilizacji doliny Indusu znaleźli skręcone nici bawełniane. Figurki z terakoty odkryte w Mehrgarh przedstawiają postać mężczyzny noszącego coś, co powszechnie interpretuje się jako turban. Figurka z Mohenjo Daro, nazwana „Królem Kapłanem” (3), przedstawia postać ubraną w szal w kwiatowe wzory. Jak dotąd jest to jedyna rzeźba z doliny Indusu, która przedstawia ubranie w tak szczegółowy sposób. Mieszkańcy Harappy mogli używać naturalnych barwników do farbowania tkanin. Badania pokazują, że uprawiano tam rośliny, z której pozyskuje się barwnik indygo. Bawełna była uprawiana i przetwarzana na włókna także w starożytnych Indiach, Egipcie i Peru, które następnie tkano na proste materiały. Rzymianie sprowadzali bawełniane tkaniny z Indii, traktując je jako luksusowy towar. W średniowieczu bawełna trafiła do Europy w większych ilościach. Renesans przyniósł wzrost zainteresowania tym materiałem, a rewolucja przemysłowa w XVIII i XIX wieku zrewolucjonizowała jego produkcję. Bawełnę zaczęto produkować na masową skalę.



ok. 5000–3000 p.n.e.

W Chinach rozpoczęto produkcję jedwabiu (4), który stał się symbolem luksusu i wyrafinowania. Według legendy, to cesarzowa Leizu odkryła proces pozyskiwania włókna jedwabnego z kokonów jedwabników. W stanowiskach archeologicznych kultury Yangshao w Xia w prowincji Shanxi, gdzie znaleziono kokon jedwabnika domowego przecięty nożem, datowany na okres między 5–3 tys. lat p.n.e. Fragmenty prymitywnych krosien znaleziono w stanowiskach kultury Hemudu w Yuyao w prowincji Zhejiang, datowanych na około 4000 p.n.e. Skrawki jedwabiu znaleziono w stanowisku kultury Liangzhu w Qianshanyang w Huzhou w prowincji Zhejiang, datowanym na 2700 r. p.n.e. Produkcja jedwabiu przez wiele stuleci była pilnie strzeżonym sekretem – za ujawnienie sekretu groziła kara śmierci. Dopiero w pierwszych wiekach naszej ery technologia dotarła do Korei, Japonii, a później do Persji i Europy dzięki Jedwabnemu Szlakowi. W średniowieczu miasta włoskie, takie jak Wenecja, Florencja czy Lukka, stały się ośrodkami europejskiej produkcji jedwabiu, a tkaniny trafiały na dwory królewskie i do bogatych kupców. Pozazińska hodowla jedwabników jest udokumentowana na terenie bizantyńskiej Syrii już w V w.

średniowiecze

1589

XVIII w.

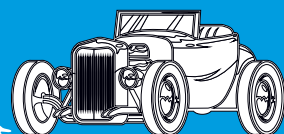
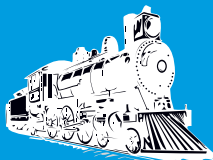
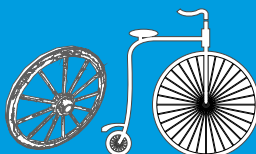
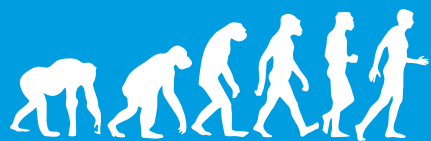
Europa staje się centrum produkcji tkanin, zwłaszcza wełny. Miasta takie jak Florencja czy Gandawa słynęły z wysokiej jakości wyrobów wełnianych. W tym okresie rozwijały się również techniki tkackie, a rzemieślnicy doskonalili swoje umiejętności, tworząc skomplikowane wzory i zdobienia.

William Lee wynajduje w Anglii maszynę do dziania wyrobów pończoszniczych (5). Jej zastosowanie było ważnym etapem mechanizacji przemysłu tekstylnego i odegrało ważną rolę we wczesnej historii rewolucji przemysłowej. Została dostosowana do dziania bawełny i wykonywania ściągaczy, a do 1800 roku – do produkcji koronek.

W Lyonie powstaje cała seria innowacyjnych urządzeń tkackich. Pierwsze z nich zbudowane było około 1725 roku według pomysłu tkacza Basila Bouchona. Do sterowania nićmi osnowy służyła kartonowa wzornica w kształcie perforowanego rulonu. Kolejne udoskonalenie było dziełem tkacza Falcona z roku 1742. Wprowadził on wzornicę w formie połączonych ze sobą tekturowych kart. Kolejnym ważnym historycznie urządzeniem była, powstała około 1740 roku, maszyna skonstruowana przez Jacques'a Vaucansonde. Poprawiono w niej system powrotu igieł, z którym były problemy w wynalazku Falcona. Urządzenie Vaucansona wykorzystywało wzornicę bębnową (6).



1. Replika Wenusa z Lespugue; 2. Fragment materiału tekstylnego z Çatalhöyük; 3. Figurka z Mohenjo Daro; 4. Tkanina jedwabna z Mawangdui w Changsha w Chinach, pochodząca z II wieku p.n.e.; 5. Replika muzealna jednej z wczesnych maszyn do dziania wyrobów pończoszniczych; 6. Rekonstrukcja wynalazku Jacques'a Vaucansonde



1738–85

Przyrząd lub maszyna określane jako krosno do tkania tekstyliów znane jest w prostych formach od czasów prehistorycznych w większości Europy, na Bliskim Wschodzie i w Afryce Północnej. Przez wieki produkcję materiałów tkanych zdominowały dwa główne typy krosien, z obciążeniem osnowy i krosno dwuramiennie. Długość belki w tej maszynie decydowała o szerokości tkaniny tkanej na krosnie i mogła wynosić nawet 2–3 metry. Duże tkaniny na ubrania były najprawdopodobniej produkowane na krosnach z obciążeniem osnowy w prehistorycznych czasach w Europie Środkowej, o czym świadczą liczne znaleziska obciążników krosien z prehistorycznych osad. W XVIII w. innowacje przyspieszyły. Krosno tkackie z ręcznym mechanizmem przerzutowym skonstruował w 1738 roku John Kay. Pierwsze krosno mechaniczne powstało za sprawą Edmunda Cartwrighta w 1785 r. (7) i stało się od razu przyczyną niepokojów społecznych i protestów ludzi tracących z tego powodu pracę, co było połączone z niszczeniem maszyn. Wynalazek ten umożliwił masową produkcję materiałów. Mechanizacja przyczyniła się do obniżenia kosztów produkcji i zwiększenia dostępności tkanin dla szerszej populacji. Bawełna stała się surowcem dominującym, zwłaszcza w Stanach Zjednoczonych, gdzie jej uprawa i przetwórstwo napędzały gospodarkę.

1793

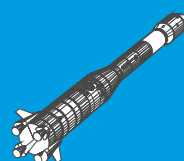
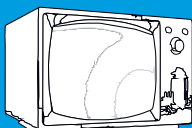
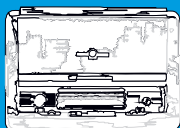
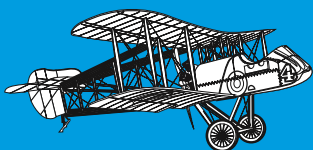
Eli Whitney wynajduje odziarniarkę bawełny, aby przyspieszyć proces usuwania nasion z włókien bawełny. Zaprojektowane przez niego urządzenie naśladowało ręczną pracę. Zamiast ręki trzymającej roślinę, wynalazca opracował sito z drucianym podłużnym rusztowaniem, przytrzymujące ziarna. Obok sita obracał się bęben z drobnymi haczykami, które, podobnie jak grzebień, odrywały włókna bawełny.

1805

Produkcję tkanin o bogatych wzorach umożliwił wynalazek Josepha Marie Jacquarda. To od jego nazwiska pochodzi technika wykorzystywana do dziś przy wykonywaniu tkanin nazywanych żakardowymi. Zapoznał się z konstrukcją Vaucansona. Dzięki wykorzystaniu połączonych kart perforowanych udoskonalił jego pomysł. Karta perforowana, poprzez system otworów, unosiła bądź nie odpowiednio nici osnowy, pozwalając tym samym na powstanie wzoru. Była to zasada binarna „otwór – dziąta, brak otworu – nie dziąta”, 0-1, a więc pierwsze programowanie w dziejach. Z pomysłem Jacquarda zapoznał się sam Napoleon Bonaparte, który wraz z cesarzową Józefiną odwiedził Lyon. Cesarz przyznał miastu patent na krosno żakardowe, a jego wynalazca otrzymał dożywotnią emeryturę w wysokości trzech tysięcy franków. Odkrycie Jacquarda przyspieszyło produkcję tkanin wzorzystych, powodując ich upowszechnienie i spadek cen. Wynalazek z 1805 roku do dziś stanowi podstawę tworzenia tkanin wzorzystych w nieco zmodyfikowanych wersjach.

1883

Brytyjski chemik, Joseph Swan, wyprodukował pierwsze włókno syntetyczne, które nazywał „sztucznym jedwabiem” i zaprezentował je na międzynarodowej wystawie w Londynie. Z tego osiągnięcia wywodzi się cała gama włókien sztucznych. Wkrótce, już w XX w., nadchodzi dynamiczny rozwój syntetycznych włókien, nylonu, poliestru czy akrylu. Materiały te charakteryzowały się wytrzymałością, elastycznością i łatwością w pielęgnacji.



1935

lata 40. XX w.

1941

1958

1979

Nylon został wynaleziony przez Wallace'a Carothersa pracującego dla koncernu DuPont w Wilmington w USA. W 1938 roku wprowadzono na rynek pierwsze produkty z nylonu. Nylonowe pończochy pojawiły się w 1940 roku. Włókno nylonowe było wkrótce używane zamiast jedwabiu do produkcji spadochronów wojskowych podczas II wojny światowej. Po wojnie popularność nylonu rozlała się na cały świat (10). Producenci odzieży szybko przyjęli włókna syntetyczne, często stosując mieszanki różnych włókien w celu uzyskania optymalnych właściwości i tworząc nowe materiały.

Firma DuPont opracowuje włókna akrylowe, syntetyczne włókna, tworzone z poli-meru (poliakrylonitrylu). Akryl swoją fakturą przypomina wełnę i najczęściej jest stosowany jako jej zamiennik ze względu na niższą cenę, większą wytrzymałość i przyjemniejszą w dotyku fakturę.

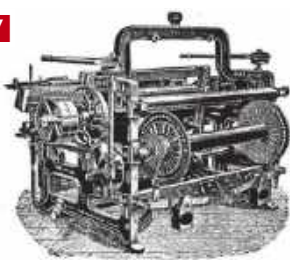
J.R. Winfield i J.T. Dickson w Wielkiej Brytanii jako pierwsi pomyślnie opracowali włókno poliestrowe w laboratorium przy użyciu kwasu tereftalowego i glikolu etylenowego jako surowców i nazwali je terylen. W 1953 roku Stany Zjednoczone wyprodukowały włókno poliestrowe o nazwie handlowej dacron. W kolejnych latach włókno poliestrowe rozpowszechniło się na całym świecie.

Spandex (lycra, elastan i inne wersje), elastyczne włókno syntetyczne (elastomer poliuretanowy), wynalezione przez Josepha Shiversa w laboratoriach DuPont.

Polar, rodzaj dzianiny wykonanej z PET i innych tworzyw sztucznych, która charakteryzuje się hydrofobowością i dobrą izolacją ciepła, może mieć właściwości termoizolacyjne podobne jak wełna, a w przypadku zamoczenia nie traci ich; został opracowany w przedsiębiorstwie Malden Mills. Nazwa „polar” wywodzi się od materiału produkowanego przez to przedsiębiorstwo – Polartec. Wykorzystywany jest m.in. do wytwarzania odzieży sportowej. Dzięki paroprzepuszczalności pozwala na odprowadzanie potu. Obecnie produkowanych jest wiele rodzajów polarów, które, w zależności od zastosowania, różnią się gramaturą.



7



10

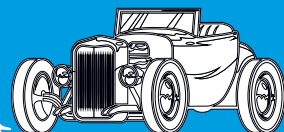
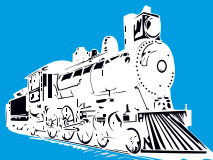
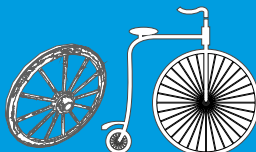
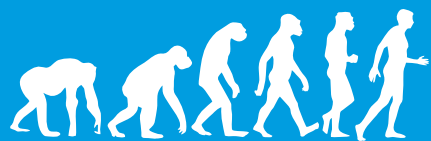


8



9

7. Edmund Cartwright i wynalezione przez niego krosno; 8. Rekonstrukcja maszyny sterowanej dziurkowanymi kartami Josepha Marie Jacquarda; 9. Jedna z pierwszych próbek 'sztucznego jedwabiu' Josepha Swana © Brytyjskie Muzeum Nauki; 10. Tkanina nylonowa sprawdzana w fabryce w Szwecji w 1954 roku



Klasyfikacja tkanin

Tkaniny można podzielić ze względu na różne kryteria. Jednym z nich jest przeznaczenie użytkowe. Właściwości użytkowe tkanin są uzależnione od surowca, z jakiego zostały wykonane, grubości nici, splotu oraz gęstości osnowy i wątku. Ponieważ każdy z surowców wymaga nieco innego wyposażenia tkalni, poszczególne zakłady produkcyjne wytwarzają zwykle tylko jeden typ tkanin. Tkaniny z takich surowców, jak len, konopie, bawełna, można produkować przy użyciu takich samych maszyn. Natomiast tkactwo wełny lub jedwabiu wymaga innego wyposażenia zarówno tkalni, działu przygotowawczego, jak i działu wykończenia tkanin.

Oto asortymenty tkanin według kryterium surowca i przeznaczenia:

1. Lniane i konopne

- pościelowe
- bieliźniane
- odzieżowe
- stołowe (obrusowe)
- dekoracyjne
- specjalne (techniczne)

2. Bawełniane

- pościelowe
- bieliźniane
- odzieżowe
- stołowe (obrusowe)
- dekoracyjne
- specjalne (techniczne)

3. Jutowe

4. Wełniane i wełnopodobne

- ubraniowe
- płaszczowe
- sukienkowe
- kostiumowe
- mundurowe
- koce
- szale, chustki

5. Jedwabne

- pościelowe
- bieliźniane
- odzieżowe
- stołowe (obrusowe)
- dekoracyjne

Tkaniny klasyfikuje się także na wiele innych sposobów, np. pod względem splotu tkackiego, zasadniczego, skośnego, atlasowego, różnego rodzaju splotów wzorzystych, żakardowych i innych. Można też tkaniny dzielić na jednobarwne, barwne tkane, drukowane lub barwione. Ze względu na strukturę wyróżnia się tkaniny jedno- i wielowarstwowe, zaś pod względem faktury – gładkie, tkaniny z okrywą różnego rodzaju, z reliefem. Osobną grupę tworzą tkaniny ludowe i artystyczne. Wytwarzane najczęściej w sposób chałupniczy, ręcznie, na prostych, drewnianych warsztatach tkackich lub nawet na drewnianych ramach. ■

M.U.



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

Pomysł miesiąca 1/2026

Pomysł inteligentnej szuflady jest interesujący. Wydaje się, że, aby działało to dobrze i wygodnie, potrzebne byłoby zaangażowanie sztucznej inteligencji, która na podstawie wypowiadanych instrukcji wskazywałaby poszukiwane przedmioty za pomocą rozpoznawania obrazu.

Autorem pomysłu jest Jan Kaleta

1 Jan Kaleta – proponuje opracowanie inteligentnej szuflady, która „wie”, co się w niej znajduje i po otwarciu sygnalizuje błyskiem miniaturowej lampki LED przedmiot, który jest nam potrzebny.

Faktem jest, że szuflady domowych młodych technik są pełne najrozmaitszych „skarbów”; zawsze troskliwie przechowywanych, „bo mogą się przydać”. Z czasem zmienia się to w ogromny śmietnik, w którym trudno jest znaleźć konkretną rzecz. Rzecz jest trudna, gdyż każdy obiekt powinien mieć jakiś marker, którego aktywacja powodowałaby świecenie przypisanego do niego diody. Ale kto wie...

2 Mateusz Grzywacz – proponuje pilne opracowanie modułowej apteczki, która pilnowałaby czasów zażywania leków oraz okresu ich trwałości, przydatności do spożycia. Obecna sytuacja jest nie do zaakceptowania: daty na pudełkach są niewyraźne, z czasem się rozmazują itd. A przecież lek to jednak nieobojętny dla organizmu środek i pewne reguły muszą być przestrzegane.

Bardzo ciekawy temat. Taka apteczkę byłaby przydatna, zwłaszcza dla osób w podeszłym wieku, które często mają do zażycia spore ilości różnych medykamentów, zwłaszcza suplementów i te pigułki często się im mylą.

3 Beata Konopka – nadesłała propozycję skonstruowania lodówki, która miałaby ruchome półki zmieniające położenie w miarę potrzeb; np. gdy się chce włożyć wysoki garnek „jednym ruchem” półki by się przeorganizowały, robiąc miejsce dla wysokiego obiektu.

Bardzo ciekawa propozycja i w zasadzie już dziś możliwa do realizacji. Jak zwykle: należałoby zbadać rynek i ocenić zapotrzebowanie na tak zmechanizowaną lodówkę, bo przecież producent musi zarobić, a konstrukcja takiej lodówki byłaby jednak dość skomplikowana, a więc i droga.

4 Waldemar Świtoń – uważa, że najwyższy czas zlikwidować niebezpieczeństwo wykipienia płynnych potraw. Garnek z wbudowanym systemem monitorowania temperatury, czasu i odgłosu gotowanej potrawy załatwiłby temat. Przy krytycznym poziomie temperatury, czasu i odgłosu gotującego się np. rosołu,

garnek mogły dać głośny sygnał akustyczny albo powinien wyłączyć zasilanie kuchenki.

Funkcję taką spełnia dostępny w handlu „termomix” – garnek, w którym można gotować i nawet piec różne rzeczy. Garnek ten ma dość rozbudowaną elektronikę i naprawdę ułatwia życie zwłaszcza osobom samotnym, zmuszonym do samodzielnego przygotowywania posiłków.

5 Marcin Włodarski – nadesłał apel do wynalazców i konstruktorów o opracowanie jakiegoś prostego systemu zapinania pasów bezpieczeństwa w samochodach osobowych. Pas dość łatwo wyciąga się ze zwijaka i jego końcówkę należy – jak wiadomo – wetknąć w szczelinę zaczepu dolnego. Zadanie to jest okropne i np. dla kobiet i starszych panów, zwłaszcza przy fotelu przesuniętym, do przodu niemal niewykonalne, a w każdym razie niemiłe trudne. Są podobno samochody z automatycznym załatwianiem tej sprawy, ale tu chodzi o coś prostego, łatwego dostępnego i skutecznego.

To prawda, wpinanie końcówki pasa w dolny zaczep bywa trudne. Niektóre samochody wyraźnie podkreślają swoje japońskie pochodzenie (Japończycy są średnio niżsi od Europejczyków), pojazdy są niskie i w konsekwencji usytuowanie dolnego zaczepu jest niewygodne dla osób wyższego wzrostu. Sprawa wydaje się dość łatwa i chyba tylko niechęć konstruktorów do zajmowania się tak „prymitywnym” problemem jest powodem, dla którego spora część samochodów ma tę konstruktorską skazę.

6 Michał Rutka – proponuje opracowanie koca, który miałby jedną stronę ogrzewającą, a drugą chłodzącą. Taki koc byłby bardzo przydatny w turystyce i wielu innych sytuacjach. Zwykle koce zazwyczaj są używane do ogrzania się podczas chłodnej nocy pod namiotem. A w dzień, gdy słońce zbyt mocno dogrzewa? Taki koc odwróciłoby się na drugą stronę i już.

Na pewno dałoby się wykonać koc jako dwubarwny: z jednej strony srebrzyście biały – odbijający promieniowanie słońca, a z drugiej chyba czarny, ale ten kolor w nocy pod namiotem nie za bardzo się przyda. ■



Obudowę, którą można uznać za początek tej przygody, przedstawiono niemal 100 lat temu w Ameryce. Nie nazwano jej jednak linią transmisyjną, lecz obudową labiryntową. Jej działanie miało przesunąć fazę od tylnej strony membrany, aby wychodziła na zewnątrz w fazie zgodnej z przednią stroną membrany.

W labiryntach, część 3

Taki efekt nie jest jednak możliwy w szerokim zakresie częstotliwości, długość fal zmienia się przecież wraz z częstotliwością, więc fale różnych częstotliwości wychodzą w różnych fazach, promieniowanie labiryntu dodaje się i odejmuje z promieniowaniem przedniej strony membrany, co tworzy charakterystykę pełną wzmocnień i osłabień. W dodatku w labiryncie powstają rezonanse związane z jego całkowitą długością, a także długościami jego poszczególnych fragmentów (labirynt jest w obudowie „zwinięty”). Takie rozwiązanie nie zdobyło wtedy dużej popularności, ale zostało wykorzystane przez A.R. Bailey’a w artykule „A Non-resonant Loudspeaker Enclosure Design” w „Wireless World”, w 1965 roku.

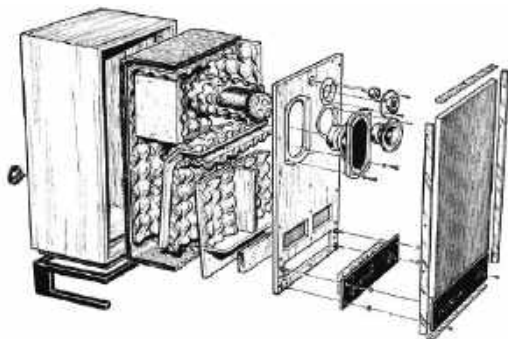
Mija dokładnie 60 lat, od kiedy A.R. Bailey przedstawił koncepcję linii transmisyjnej.

Bailey został powszechnie uznany za jej twórcę, i chociaż podwaliny położono wcześniej, to A.R. Bailey poważnie zmienił założenia.

Bailey powoływał się na wcześniejsze prace o obudowie labiryntowej, ale zaproponował zasadniczą modyfikację jej sposobu działania, związaną właśnie z nowym określeniem – linią transmisyjną. Ono z kolei nawiązuje do działania energetycznych linii przesyłowych, w których na bardzo długich odległościach następuje strata energii. Taka koncepcja miała na celu stworzenie głośnikowi warunków idealnych, przypominających teoretycznie nieskończenie wielką odgradę, czyli odizolowanie promieniowania od tylnej strony membrany (które jest w przeciwnej fazie do promieniowania przedniej, więc ich bezpośrednie spotkanie, gdy nie ma żadnej obudowy ani odgrady, powoduje wygaszenie niskich częstotliwości). Tak działająca obudowa – linia transmisyjna – nie miała w żaden sposób oddziaływać na głośnik ani też emitować energii, i mieć właśnie taką przewagę nad obudową zamkniętą, w której fale, zanim zostaną wytłumione, wielokrotnie się odbijają, uderzając w membranę i zakłócając jej pracę. Pada też argument, że tak działająca linia transmisyjna nie zmienia podstawowych parametrów głośnika, podczas gdy zamknięta,

poprzez skończoną podatność powietrza („poduszkę” powietrzną), podnosi częstotliwość rezonansową i dobroć. To jednak nie musi być wcale szkodliwe, gdy parametry głośnika uwzględniają określone warunki pracy w obudowie.

Bailey sam nie zaprojektował jednak konkretnej konstrukcji, która by go zweryfikowała (a w każdym razie nie natrafiłem na taki trop...). Okazało się, że w praktyce nie można skonstruować idealnej linii transmisyjnej w ścisłym znaczeniu, bowiem dla wytłumienia najniższych fal musiałaby mieć ogromną długość – kilkudziesięciu, a może nawet 100 metrów (w zależności od tego, jak nisko ma tłumić). Z kolei zamknięcie labiryntu, aby nie pozwolić mu niczego promieniować, zamienia go w obudowę... zamkniętą. Chociaż możliwe jest zredukowanie wspomnianych odbić, wymaga to nie tylko ułożenia labiryntu, ale też nadania mu specjalnego kształtu, tymczasem w większości linii/labiryntów niedaleko za głośnikami znajduje się ściana kanału, od której fale odbijają się jak od ściany obudowy zamkniętej i część z nich uderza w membranę.



Pierwsza słynna linia transmisyjna „ery nowożytniej” – IMF Reference Standard Professional Monitor. W roli głośnika niskotonowego nie mniej legendarny „stadion” – KEF B139. Mimo deklaracji producenta, że charakterystyka „sięga od 10 Hz”, pomiary pokazują, że sięga dość równo do 30 Hz, ale przy 20 Hz ma spadek ponad 12 dB

Zorientowano się w tym dość szybko i już słynne, uznawane za klasyczne linie transmisyjne w latach 60. (IMF) w rzeczywistości działały jak częściowo wytłumione labirynty.

Nowe, bardziej realistyczne założenia mówiły o tym, że obudowa taka (wciąż jednak nazywana dumnie linią transmisyjną) ma działać jak labirynt w zakresie najniższych częstotliwości.

Nie powinna zatrzymywać promieniowania od tylnej strony membrany, co jest nawet korzystne, gdyż rezonans ćwierćfalowy odciąża głośnik od dużych amplitud, a w dość szerokim zakresie częstotliwości, dzięki przesunięciu fazy skorelowanej z długością kanału, otwór będzie promieniował w fazie zgodnej z przednią stroną membrany. Natomiast fale wyższych częstotliwości (powyżej ok. 100 Hz), które na skutek coraz bardziej zagęszczających się (na skali częstotliwości) zmian fazy, a więc i relacji fazowych z promieniowaniem przedniej strony membrany, wywoływałyby na charakterystyce wypadkowej wyraźne nierównomierności – zostaną wytłumione.

Wedle takiej koncepcji należało jednak zadbać o możliwie jak najdłuższą linię już nie w celu wytłumienia bardzo długich fal najniższych częstotliwości, tylko dla takiego ich przesunięcia w fazie na tej drodze, aby do jak najniższej częstotliwości wychodziły z kanału przynajmniej w fazie nieodległej od fazy przedniej strony membrany. Naj-najniższe częstotliwości będą jednak wypromieniowane w fazie niewiele przesuniętej względem fazy wprost z tylnej strony membrany, a więc w fazie bardzo przesuniętej względem przedniej strony membrany, i ciśnienie wypadkowe będzie niewielkie.

To także okazało się jednak bardzo trudne do zrealizowania. Im dłuższy labirynt, tym lepsze przetwarzanie najniższych częstotliwości, ale tym większe problemy wyżej. Im niżej sięgniemy, tym proporcjonalnie niżej mamy pierwszy „antyrezonans” (przeciętną fazę promieniowania od głośnika i otworu). Wydawało się, że „selektywne” działanie linii można osiągnąć doбором materiału wytłumiającego – przecież tłumienie większości materiałów wzrasta wraz z częstotliwością, więc wystarczy dobrać odpowiedni rodzaj, ilość i miejsce ulokowania, aby cieszyć się z promieniowania najniższych częstotliwości (i ich wzmocnienia na charakterystyce wypadkowej). Jednak okazuje się, że duża ilość materiału tłumiącego w pierwszym rzędzie tłumi korzystny rezonans ćwierćfalowy, zmniejszając poziom najniższego basu, o który walczymy. Aby wytłumić fale rzędu 100 Hz (i to też nie całkowicie), które już

nam przeszkadzają, trzeba labirynt „zapchać” materiałem tłumiącym, tracąc zbyt wiele w zakresie najniższych częstotliwości.

Kompromis musiał polegać na czymś innym. Projektanci zaczęli kombinować, jak czerpać energię z działania labiryntu w zakresie, w którym jest potrzebna, co jednak nie pozwalało nadmiernie go wytłumiać, a „odcinać” wyższe częstotliwości innymi sposobami. Widzimy je w testowanych konstrukcjach – są to ślepe kanały, działające na zasadzie fal odbitych, albo komory działające jak rezonatory Helmholtza. Warto jednak zwrócić uwagę, że komplikują one układ akustyczny, jeszcze bardziej oddalając go od purystycznego założenia linii transmisyjnej,



Jeden z roboczych prototypów Nautilusa B&W – w sekcji niskotonowej widać zwiniętą, jeszcze ostatecznie niezamkniętą linię transmisyjną, ale jej przekrój na końcu jest tak niewielki, długość tak duża, a cała droga tak silnie wytłumiona, że najwyraźniej chodziło o maksymalne wygaszenie promieniowania tylnej strony membrany zgodnie z oryginalnymi założeniami A.R. Bailey’a, a nie jej wykorzystania i wzmocnienie najniższych częstotliwości. Ostatecznie w projekcie wdrożonym do produkcji linia jest całkowicie zamknięta



zwiększając koszty, a także powiększając obudowę. W praktyce, aby linia transmisyjna nie ustępowała podstawowymi parametrami innym rodzajom obudów, a tym bardziej udowodniła swoją przewagę, staje się bardzo duża, co poważnie nadwyręża jej popularność zarówno wśród profesjonalnych konstruktorów, jak i większości zainteresowanych. Z wyjątkiem najbardziej „zakręconych” na tym punkcie. Z drugiej strony, linie transmisyjne o niewielkiej objętości, zbyt krótkie albo o zbyt małym przekroju kanału, ustępują „zwykłym” obudowom bas-refleks, a nawet zamkniętym, w rozciągnięciu niskich częstotliwości.

Jednym ze sposobów na pozbycie się pasożytniczych rezonansów wyższych częstotliwości jest ograniczenie pracy sekcji niskotonowej do zakresu, w którym działanie linii transmisyjnej przynosi korzyści, czyli w praktyce do ok. 100 Hz. Przed pierwszym antyrezonansem należy głośnik niskotonowy filtrować dolnoprzepustowo, i to z pewnym zapasem oraz dużym nachyleniem, i przekazać przetwarzanie głośnikowi nisko-średniotonowemu. To jednak nie tylko narzuca schemat całej konstrukcji, ale też rodzi pytanie, po co w ogóle stosować linię transmisyjną, jeżeli nie korzysta z niej sekcja przetwarzająca dużą część zakresu niskich tonów? Tylko po to, aby rozszerzyć pasmo? Są na to inne, lepsze sposoby... a przypomnijmy, że ideą linii transmisyjnej miało być „wyczyszczenie” działania głośnika niskotonowego z rezonansów wprowadzanych przez „zwykłe” obudowy. Tymczasem linia transmisyjna wprowadza ich wcale nie mniej, tyle że w inny sposób.

W XX wieku doskonalenie linii transmisyjnych odbywało się wyłącznie na podstawie podstawowych kalkulacji, a następnie żmudnej metody prób i błędów. Konstruktorzy szli różnymi drogami, nie było pełnej zgody nawet co do niektórych zagadnień podstawowych, jakie na przykład powinny być parametry głośników odpowiednich do linii transmisyjnej. Sprawę tę już dawno rozstrzygnięto dla obudów zamkniętych

i bas-refleks, modele dostępne od lat 70. pozwalały obliczyć kształt charakterystyk, dopasowując parametry obudowy do parametrów głośników, a linia transmisyjna nie poddawała się takiej procedurze. Według klasycznych założeń głośniki niskotonowe w niej stosowane miały mieć wysoką dobroć (Qts), ale było wiele projektów (a później nawet większość) z głośnikami o niskiej wartości Qts, jaka bez wątpienia jest odpowiednia do systemów bas-refleks... bowiem działanie linii transmisyjnej może płynnie przejść zarówno w sposób działania obudowy zamkniętej, jak i obudowy z otworem. Przecież każda obudowa z otworem jest też labiryntem...

Rezultaty były bardzo trudne do przewidzenia i nawet po wielu prototypach i modyfikacjach – nie zawsze takie, na jakie miano nadzieję. Mimo to linia transmisyjna zdobyła renomę, a ponieważ firmowe kolumny zwykle były bardzo drogie, hobbyści z zapałem próbowali je kopiować albo co najmniej zainspirowani nimi wymyślali własne projekty. Chociaż szansa, że za pierwszym podejściem zagrają choćby dobrze, była niewielka. Ale zwyciężała nadzieja na wspaniałe rezultaty.

W XXI wieku, dzięki postępowi w analizie działania głośników i obudów, zaczęły pojawiać się lepsze wskazówki, pokazujące zależność między parametrami głośnika, wymiarami linii a charakterystykami końcowymi. Jednak bogactwo zjawisk, trudność w określeniu wszystkich parametrów, a zarazem wrażliwość linii nawet na niewielkie zmiany, np. w rodzaju materiału tłumiącego czy miejsca jego umieszczenia, wciąż powodują, że jest to obudowa o wiele trudniejsza do zaprojektowania niż jakakolwiek inna. I dlatego jest taka piękna...

Pomysłów na zabawę z linią transmisyjną jest wciąż bardzo wiele i czytelnicy łatwo znajdą je w Internecie.

Kiedyś dotarcie do „recept” było znacznie trudniejsze, dzisiaj są dostępne dla każdego. Mimo to linia transmisyjna broni swoich tajemnic i niech nikt się nie łudzi, że znajdzie cudowny przepis. ■

Andrzej Kisiel

AUDIO

przeglądaj,
czytaj
i kup na
www.ulubionykiosk.pl

PRZEGLĄD ELEKTROTECHNICZNY Miedź, mająca przewodność o 15% wyższą od normalnej

Dr. W. P. Devey (General Electric Co., Laboratoria Naukowe) zdołał otrzymać miedź o przewodności o 15% wyższej, niż zwykła. Jest to miedź krystalizowana w pojedynczych dużych kryształach, jakie udało się otrzymać w laboratorjach General Electric Co. w Schenectady. Praktycznego zastosowania odkrycie to w obecnym stanie znaleźć na razie jeszcze nie może, głównie z powodu delikatności kryształów oraz trudności, związanych z ich otrzymywaniem. Odkrycie stanowi jednak ważną zdobycz naukową. Kryształy otrzymuje się przez bardzo powolne nagrzewanie i studzenie czystej miedzi przy zastosowaniu pieców elektrycznych. Im powolniejsze jest studzenie, tem większe otrzymuje się kryształy. Udało się otrzymać kryształy długości 6 cali i średnicy 7/8 cala.

1 stycznia 1926

Guma elektrolityczna

W grudniowym zeszytce Scientific American r. ub. str. 383 znajdujemy opis nader prostego wynalazku, który niezawodnie znajduje szerokie zastosowanie w fabrykacji przewodników, izolowanych gumą. Jest to sposób otrzymywania warstwy gumy na wszelkiego rodzaju metalach i przewodnikach za pomocą elektrolizy. Sposób ten, wynaleziony przez Dr. S. E. Shepparda i przedstawiony na ogólnym zebraniu Amerykańskiego Stowarzyszenia Chemików (American Chemical Society), polega na elektrolizie gumy w stanie koloidalnym w roztworze wodnym z dodaniem niewielkiej ilości amoniaku, prawdopodobnie dla zwiększenia przewodności. Zataczając prąd stały o napięciu około 100 V i natężeniu prądu od 0,0002 A do 0,5 A na 1 cal kwadratu powierzchni elektrody, otrzymamy na anodzie warstwę gumy, którą później można wprost wulkanizować.

1 stycznia 1926

Przebudowa elektrowni w Szarlottenburgu

Do większych robót, które zostały wykonane w Berlinie, po za projektowaną elektrownią w Rummelsburgu, należy również rozbudowa elektrowni w Szarlottenburgu, która należy do „Berliner Elektrizitätswerke A. G.”. Po rozbudowie moc zakładu została powiększona o 16 000 kW do 78 000 kW. Tymczasowo ustawiono 2 zespoły turbinowe (czotowe) na ciśnienie 32 atm. po 8 000 kW i 2 – niskiego ciśnienia 13 atm. po 16 000 kW systemu Siemens Schuckert – M. A. N. – Brünn. Zasilac parą nowe turbiny będzie 12 nowych kotłów po 700 m² pow. ogrz. i 35 atm. ciśnienia. Próby

instalacji dały rezultaty zużycia ciepła 4 000 – 4 200 kal/ kWh. Napięcie stacji 6 000 V, dla przesyłania energii przyjęto napięcie 30 000 V. Zwraca na siebie uwagę szybkość budowy. Od daty zamówienia do puszczania w ruch upłynęło tylko 10 miesięcy, włączając w to 10-cio tygodniowy strajk.

1 stycznia 1926

Urządzenia zabezpieczające w razie wypadku z motorowym

Z rozwojem wagonów tramwajowych z tak zwaną jednoosobową obsługą a zarazem przy zwiększeniu szybkości jazdy, zjawia się konieczność zaopatrywania tramwajów w urządzenia, które w razie nagłego wypadku z motorowym działabyby automatycznie na wyłącznik prądu oraz hamulce wagonowe. Zwykle takie urządzenia związane są z rączką nastawnika, bądź też z pedalem, o który motorowy opiera się podczas jazdy. Ostatnio towarzystwo tramwajowe w Filadelfii wykonało dla swoich wozów takie automatyczne zabezpieczenia, związane z rączką hamulca pneumatycznego. Z chwilą, gdy motorowy wypuści z ręki tę rączkę, prąd zostaje wyłączony, piasecznice wysypują na szyny piasek i jednocześnie wprawiane są w ruch hamulce pneumatyczne. Rączka hamulca pneumatycznego utrzymywana jest stale w położeniu równowagi przez nacisk ręki motorowego oraz dwie sprężyny górną i dolną umieszczoną w głowie wentyla powietrznego. Gdy motorowy odejmie rękę z rączki hamulca, ciśnienie sprężyny dolnej przewyżczy ciśnienie sprężyny górnej i rączka wskakuje w położenie alarmowe; z tego położenia może być ona z powrotem doprowadzona do położenia normalnego przez odpowiedni nacisk i doprowadzenie jej do miejsca skrajnego hamowania. Działanie na wyłącznik prądowe uskutecznia się w sposób następujący: obok głównego cylindra hamulca powietrznego umieszczony jest specjalny niewielki cylinder, który na głowie swojego tłoku ma umieszczony momentalny przerywacz prądowy, działający podczas każdego hamowania. Przerywacz ten jest odpowiednio odsprężynowany za pomocą dwóch sprężyn z brązu fosforowego. Ponowne włączenie prądu odbywa się za pomocą drugiego cylindra powietrznego, utrzymywanego w równowadze za pomocą sprężyny spiralnej. Magnetyczno-elektryczny wentyl piasecznicy podczas alarmowego hamowania działa jako zwykły elektryczny wentyl, który wypuszcza powietrze do piasecznicy.

15 stycznia 1926

Największa przetwornica częstotliwości

Tow. „Brooklyn Co” niedawno zainstalowało ogromną

przetwornicę częstotliwości, mocy 3500 kW, która służy do powiązania pomiędzy sobą 25-cio i 60-cio okresowych sieci tego towarzystwa. Wymiary: podstawa – 14,75 m × 7,2 m, wysokość nad podłogę – 7 m., średnica części wirującej – 4,6 m., średnica wału – 355 mm, waga całkowita – 460 t, waga części wirującej – 280 t. Ciekawe są szczegóły oliwienia czterech łożysk, chłodzonych zapomocą wody, krążącej w kanałach naokoło panewek. Celem ułatwienia rozruchu oliwa wtłacza się pompami do łożysk, podnosi wał, tak że cały wirnik unosi się i spoczywa na cienkiej warstwie smaru, pozostającego pod ciśnieniem 84 kg/cm². Maszyna robi 300 obr./min. Po odłączeniu od sieci biegnie ona z rozprędu jeszcze przez 45 min.

15 stycznia 1926

PRZEGLĄD TECHNICZNY

Stulecie kolei

Wiek XIX-ty przyjęto nazywać wiekiem pary i elektryczności. Jeżeli jednak rozważyć wpływ, jaki na przekształcenie stosunków gospodarczych i kulturalnych świata miały wielkie wynalazki, to właściwiej byłoby nazwać wiek ten Wiekiem Parowozu. Jeszcze w 1825 roku na całym świecie, nie wyłączając Anglii, panowała na lądzie wyłącznie komunikacja kołowa konna. Rydwany dwukółowe znane były Egipcjanom i Asyryjczykom. Wspomina o nich Homer, a świetnym był w starożytnym Rzymie i jego kolonjach stan dróg kołowych, których wyraźne ślady pozostały do dnia dzisiejszego, świadcząc o wysokim ówczesnym rozwoju komunikacji kołowej. W Wiekach Średnich sztuka budowania dróg poszła w zapomnienie. Komunikacja lądowa cofnęła się wstecz do siodła i juku. Dopiero epoka Odrodzenia przywróciła w Europie komunikację kołową. Pierwsza karetka była zbudowana w Anglii w roku 1568 dla królowej Elżbiety. W roku 1625 zjawiała się w Londynie pierwsza dorożka, a w r. 1659 pierwszy dylżans. Dylżans ten, z matami udoskonaleniami, przetrwał do otwarcia pierwszej kolei w dniu 27 września 1825 roku, ale drogi ówczesne były znacznie gorsze od rzymskich. W ten sposób, od starożytnych Egipcjan aż do początku wieku XIX, nie było przełomowego wynalazku w rozwoju komunikacji lądowej, a sposoby przewozu były w gruncie rzeczy te same, a nawet w Wiekach Średnich był dłuższy okres wsteczny. Wystarczy teraz z tym długim, blisko 5000-letnim okresem zestawzić etap, jaki przebyła komunikacja lądowa w krótkim stosunkowo 100-letnim okresie, dzielącym nas od czasu uruchomienia przez

Jerzego Stephensona pierwszej kolei w Anglii, ażeby zrozumieć, jakie doniosłości wydarzeniem w dziejach ludzkości było otwarcie tej pierwszej kolei i jak potężnym czynnikiem kultury i postępu materialnego stał się parowóz. Trzy pierwiastki techniczne przewozu lądowego: zniwelowana droga, nieruchome gładkie szyny i trakcja mechaniczna rozwijały się w Anglii w ciągu szeregu lat już przed Stephensonem, ale oddzielił jedno od drugiego i zależnie od różnych potrzeb miejscowych. Zastąga genjuszu Stephensona było nie tylko udoskonalenie konstrukcyjnie parowozu, ale i stanowcze połączenie tych trzech pierwiastków technicznych w jedną całość organiczną. To zaś pozwoliło mu osiągnąć wybitny postępek dwóch gospodarzy pierwiastków komunikacji: wielkości tadunku i prędkości jego przewozu. Co do ciężaru, postępek dotyczył głównie jednostki tadunku, gdyż powiększeniem ilości zaprzęgów można było i przedtem podołać dowolnej masie ogólnej. Znaczenie nierównie donioslejsze posiadało powiększenie szybkości przewozu. Koń w zaprzęgu nie może dać większej prędkości przeciętnej przewozu jak 10 km/h. I taką prędkością ludzkość zadowalała się w ciągu tysiącleci. Stephenson odrazu powiększył prędkość przeszło dwukrotnie, a użył do tego sposobu, podatnego do dalszego nieograniczonego rozwoju. Już w 75 lat po tem prędkość przewozu na kolejach osiągnęła 100 km/h i przysięgowała pole do dalszego jej rozwoju w automobiliźmie i następnie w lotnictwie. Wynalazek Jerzego Stephensona rozpętał w ludzkości tę manję prędkości komunikacji, której granic w tej chwili przewidzieć niepodobna. Stała się ona głównym bodźcem technicznym i gospodarczym do olbrzymiego rozwoju sieci kolejowej, która w ciągu 100 lat osiągnęła długość 1 250 000 km, to znaczy mogłaby opasać 35 razy kulę ziemską wzdłuż równika. Ten tak olbrzymi wzrost komunikacji kolejowej miał niewątpliwie większy wpływ na rozwój materialnej, a pośrednio i społecznej kultury ludzkości, niż wszystkie inne zdobycze w dziedzinie nauki i techniki, tak liczne w ciągu ubiegłego stulecia. Dlatego dziwić się wypada, że 100-letnia rocznica otwarcia pierwszej kolei, przypadająca d. 27 września r. ub., przeszła prawie niepostrzeżenie w poszczególnych krajach, których rozwój ma tak wiele do zawdzięczenia kolejom, i została upamiętniona uroczystym obchodem jedynie tylko w Anglii, oczywiście kolei.

13 stycznia 1926



pakiet promocyjny
KOCHAM SZACHY
7 e-booków z rabatem
50%



Dla prenumeratorów – 30% rabatu!

Promocja Internetowa – w formularzu zamówienia online zaznacz pole
„Jestem prenumeratorem wydawnictwa AVT, kupuję ze zniżką”
i podaj swój numer prenumeraty.