



nr 5. maj 2023

e-suplement www.mt.com.pl



Tu przejrzysz
i kupisz ten numer

NEWS 24/7
przeglądaj codziennie
na swoim smartfonie

młody **m.technik**

Ciekawi świata są zawsze młodzi

Internet **GPT**

Sieć 3.0 a nawet 4.0

FIZYKA W SZKOLE
Rozpad promieniotwórczy (1)

ISSN 0462-9760 Indeks 365408



cena: **14,90 zł** (w tym 8% VAT)

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



**do 6* wydań
gratis!**

* Cena prenumeraty rocznej wynosi 163,90 zł.
Przy zamówieniu prenumeraty dwuletniej w cenie 268,20 zł
oszczędność wynosi równowartość sześciu wydań „Młodego Technika”

**Wszystkie opcje prenumeraty i e-prenumeraty znajdziesz na stronie
www.UlubionyKiosk.pl**

prenumerata@avt.pl

AVT-Korporacja sp. z o.o., ul. Leszczynowa 11, 03-197 Warszawa
konto 18 1050 1012 1000 0024 3173 1013

eprasa.pl 376593c193



Temat okładowy

Sztuczna inteligencja zmienia internet, jaki znamy, choć w tej chwili nie wiemy jeszcze zbyt dobrze, w jaki sposób go zmieni. Czy to będzie zmiana rzeczywistości, czy też zmiana w stylu – „dużo musi się zmienić, aby wciąż było tak samo” – przekonamy się zapewne już niedługo.

Prenumerata dla szkół i placówek oświatowych 30% taniej!

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

Roczny dostęp online kosztuje 100,00 zł

Zamów na www.UlubionyKiosk.pl/prenumerata/szkolna

Czy rewolucja wypali?

Mamy przecucie, że internet, czy mówiąc może ogólniej, sieć, się zmieni, już właściwie się szybko zmienia, i wszystko to, do czego przyzwyczailiśmy się w ciągu ostatnich ok. dwu dekad, stanie się całkiem inne. Nie wiemy tylko, jaki będzie koniec tej rewolucyjnej drogi i czy przyniesie ona więcej dobrego niż złego. A tego drugiego bez wątpienia sobie nie życzymy.

A co, jeśli internet, jego powstanie, rozwój i globalne upowszechnienie było tylko przygotowaniem do fazy, w którą teraz zaczynamy wchodzić? W tym rozumieniu internet stanowiłby wstępną infrastrukturę, szkielet, a zarazem – wielkie pole uprawy danych, których potrzebuje sztuczna inteligencja, by rozwijać swoją wiedzę i umiejętności. Prawda, że można na to tak spojrzeć?

Piszemy o tym, starając się nie wartościować, nie przesądzać, czy ekspansja algorytmów to coś dobrego, czy złego. Doświadczenie uczy,

że żadne narzędzia techniczne nie są „z natury” ani dobre, ani złe. Siekierą można narąbać drewna, by się ogrzać. Można nią też zabić. Zaś internet, a także AI, to przecież narzędzia. Historia

Internet to tylko przygotowanie gruntu pod AI?

sieci pokazuje, że jej rozwój przynosił zarówno dobre, jak i mniej dobre sposoby wykorzystania narzędzi. Z siecią wypełnioną AI nie powinno być inaczej.

Na pewno powinniśmy zadbać o bezpieczeństwo, zabezpieczyć nowe techniki przed ludźmi, którzy nie mają dobrych intencji i lepiej, by nie mieli wpływu na to, jak działają narzędzia. Jednak bądźmy realistami. Wiele rzeczy przyjmuje niekorzystny obrót, nawet jeśli wszyscy mają dobre intencje. Pewnych konsekwencji nie da się przewidzieć.

W opisach i prognozach nowych rozwiązań technicznych często jest wiele przesady. I znów – zarówno w dobrym, jak i w złym sensie. Dziś wiemy, że np. przewidywania co do tempa ekspansji mobilnej sieci 5G były sporo na wyrost. Miała być rewolucja i ją ktoś nie wypaliła.

Czy wypali zapowiadany pod wrażeniem oszałamiających możliwości modeli AI przewrót technologiczny w internecie i nie tylko? Wydaje się, że na odpowiedź nie będziemy bardzo długo czekać, bo zmiany zachodzą szybko.

Mirosław Usidus

Spis treści

Temat numeru: Internet GPT.

Sieć 3.0 a nawet 4.0

- 26 • Bezpieczeństwo sieciowe w erze sztucznej inteligencji. Automatyzująca się cyberwojna
- 32 • Co się dzieje z 5G? W sieci rozczarowania
- 36 • Najnowsza internetowa wojna na górze. Premium czy klasa ekonomiczna?
- 43 • Sprzęt kolejnej rewolucji. Krzemowe woły robocze

Technika

- 8 Info Zoom
- 16 Dodaj do obserwowanych
- 17 Robotyka: Repozytorium Robotyki – cyfrowe udostępnianie zasobów nauki z obszaru robotyki
- Horyzonty mgłą spowite
- 18 • Generatywna AI ma brudny i kosztowny sekret. Nieokiełznany apetyt na energię
- 21 • Destrukcja nanotechnologiczna. Czy replikatory opanują Ziemię?
- 23 • Upadek nauki łamiącej utarte schematy? Mniej przełomów i innowacji
- 47 Raport MT: Metawersum po przejściach. Rzeczywistość wirtualnie nieumarła

m.technik

- 58 Mobilne aplikacje. Test aplikacji: Na łono natury

Szkoła

- 60 Chemia inna niż w szkole: Chemia jajka (2)
- 64 Koniec i co dalej: Pokój w kosmosie. Zimna wojna stopniowo rozgrzewana do gorącej
- 67 Matematyka z ludzką twarzą: Cudowne, z lekką bezsensowne zadania
- 72 Fizyka bez granic: Czy substancja promieniotwórcza straci z czasem swoje właściwości? (1)
- 74 Edukacja przez szachy: Ruy López de Segura – hiszpański książd, czołowy szachista XVI wieku
- Klub i Szkoła Wynalazców
- 78 • Szkoła Wynalazców, dozwolone do lat 15
- 79 • Klub Wynalazców, bez ograniczeń wieku
- 80 • Vademecum Młodego Wynalazcy
- 83 Pomysły genialne, zwariowane i takie sobie
- Na warsztacie
- 84 • Elektronika dla Ciebie: Symetryczny zasilacz warsztatowy ±1,25 V...±25 V/1,5 A
- 86 • Jak to się robi? Planszówki małego konstruktora
- 89 • Metr wikinga!
- Odkryj historię wynalazków
- 94 • Radary i radiolokacja
- 98 • Klasyfikacja radarów

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



Czy w branży internetowej czeka nas nowe rozdanie. Na razie wydaje się, że to wciąż gra w tym samym gronie gigantów branży internetowo-technologicznej, które znamy od lat, bo tylko ci najwięksi mają środki, by rozwijać niezwykle kosztowne projekty oparte na modelach i generatorach AI. To jednak może się zmienić...

W tym wydaniu MT m.in.:

- **Horyzonty mgłą spowite: Brudne sekrety sztucznej inteligencji.** Wyścig w budowaniu wysokowydajnych, napędzanych sztuczną inteligencją wyszukiwarek według wszelkich znaków na niebie i ziemi będzie wymagał dramatycznego wzrostu potrzebnej mocy obliczeniowej, a wraz z nią ogromnego skoku w zużyciu energii przez firmy technologiczne.
- **Definitywny koniec pokoju w kosmosie.** Eksperci obawiają się, że zimną wojnę, która od dekad toczona jest na orbitach wokół Ziemi zastąpić wkrótce może wojna „gorąca”, czyli otwarty konflikt militarnym w przestrzeni kosmicznej.
- Test aplikacji: Ze smartfonem na łono natury

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

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

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

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

• Redaktor Wydania:
Wojciech Marciniak

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

• Konsultacja graficzna:
Małgorzata Jabłońska

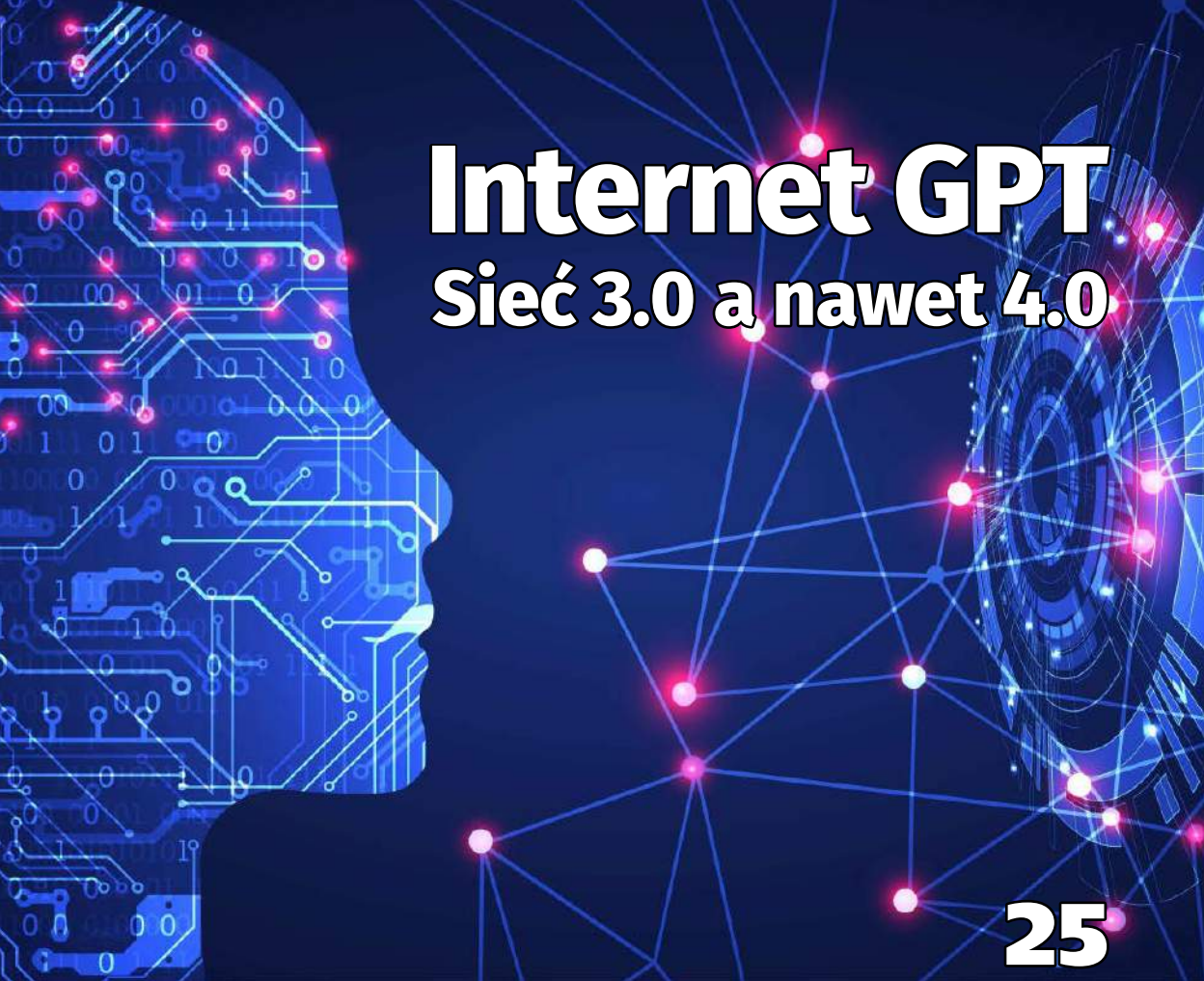
• Dział Reklamy:
e-mail: reklama@mt.com.pl

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

• Prenumerata w Wydawnictwie AVT
www.ulubionyKiosk.pl
tel. 22 257 84 22 (godz. 10:00–14:00)
e-mail: prenumerata@avt.pl

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

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



Internet GPT

Sieć 3.0 a nawet 4.0

25

RAPORT



**Rzeczywistość
wirtualnie nieumarła**

47

Listy miesiąca

Energia z atomu

Opisywany niedawno w „Młodym Techniku” temat energetyki atomowej jest jednym z najbardziej kontrowersyjnych zagadnień związanych z produkcją energii. Współcześnie na całym świecie wiele krajów stawia na atomowe źródła energii, a jednocześnie wiele innych krajów decyduje się na ich wycofanie. Warto zastanowić się, czy energetyka atomowa jest odpowiedzią na potrzeby rozwijającego się świata, czy też przestarzałą technologią.

Energetyka atomowa oferuje wiele korzyści, takich jak niskie koszty produkcji energii, brak emisji gazów cieplarnianych oraz wysoka wydajność. Jednakże istnieją również poważne zagrożenia, takie jak awarie i wypadki, problem składowania odpadów radioaktywnych oraz ryzyko proliferacji broni jądrowej. Mimo że wypadki takie jak Czarnobyl czy Fukushima są rzadkie, to ich skutki są tragiczne i mają wpływ na zdrowie oraz życie ludzi oraz środowisko naturalne.

W dzisiejszych czasach mamy wiele alternatywnych źródeł energii, takich jak energia wiatrowa, słoneczna czy geotermalna, które są bezpieczniejsze i bardziej ekologiczne niż energia atomowa. Jednakże energia atomowa pozostaje jednym z najtańszych źródeł energii i w wielu krajach nadal jest ważnym źródłem energii. Dlatego też należy znaleźć sposób na produkcję energii atomowej, która będzie bezpieczniejsza i bardziej zrównoważona dla środowiska naturalnego.

Wydaje się, że kluczowym wyzwaniem dla przyszłości energetyki atomowej jest bezpieczeństwo. Wymaga to zwiększenia inwestycji w badania i rozwój nowych technologii, które będą bezpieczniejsze dla ludzi i środowiska naturalnego. Należy również wprowadzić bardziej rygorystyczne regulacje, aby zapewnić, że energia atomowa jest produkowana w sposób bezpieczny i zrównoważony dla środowiska.

Energetyka atomowa jest tematem, który budzi wiele kontrowersji i wciąż pozostaje aktualny. Warto przemyśleć, jak możemy wykorzystać tę technologię w sposób bezpieczniejszy i bardziej zrównoważony dla środowiska. Jednocześnie, należy kontynuować rozwój alternatywnych źródeł energii, aby sprostać wyzwaniom zmian klimatycznych i chronić naszą planetę.

Ewelina Chałupiec, Bieńków

Energetyka atomowa

Pragnę zwrócić uwagę na temat energetyki atomowej, która od wielu lat jest przedmiotem dyskusji i kontrowersji w społeczeństwie, a „Młody Technik” poświęcił jej dużą część wydania w lutym. Choć wiele osób uważa ją za niebezpieczną i ryzykowną, to ja uważam, że jest to ważne źródło energii, które może przynieść wiele korzyści dla ludzkości.

Po pierwsze, energetyka atomowa może dostarczyć energię elektryczną na dużą skalę, co jest szczególnie istotne w krajach o rozwijającej się gospodarce. Dzięki temu mieszkańcy tych krajów będą mieli dostęp do taniej i niezawodnej energii, co przyczyni się do poprawy ich jakości życia. Ponadto energetyka atomowa jest czystą formą energii, która nie emituje gazów cieplarnianych, co ma pozytywny wpływ na środowisko naturalne.

Po drugie, w przypadku wypadków w elektrowniach atomowych, jak te, które miały miejsce w Czarnobylu i Fukushima, przeprowadzone zostały liczne badania, które wykazały, że długoterminowe skutki dla zdrowia były znacznie mniejsze, niż się pierwotnie przypuszczało. Dzięki poprawie technologii i bezpieczeństwa w elektrowniach atomowych ryzyko takich wypadków zostało znacznie zmniejszone.

Nie ulega wątpliwości, że energetyka atomowa niesie ze sobą pewne ryzyko, jednak uważam, że poziom tego ryzyka jest obecnie niski w porównaniu do potencjalnych korzyści, jakie może przynieść. Dlatego też, moim zdaniem, powinniśmy kontynuować rozwój tej technologii i inwestować w nią, aby zapewnić ludziom na całym świecie niezawodne źródło energii.

Maria Czuk, Rydułtowy



Synteza termojądrowa – szanse i bariery

Synteza termojądrowa, czyli proces wytwarzania energii przez łączenie jąder atomowych, jest jednym z najważniejszych wyzwań współczesnej nauki. Wielu naukowców uważa, że synteza termojądrowa może być kluczem do rozwiązania problemów związanych z brakiem zrównoważonych źródeł energii.

Jednakże praktyczne wykorzystanie energii syntezy termojądrowej jest nadal bardzo trudne do osiągnięcia. Wymaga to osiągnięcia bardzo wysokich temperatur i ciśnień, które dotychczas nie zostały osiągnięte w żadnym eksperymencie. Dodatkowo, magazynowanie energii syntezy termojądrowej również stanowi wyzwanie, ponieważ wymaga ona bardzo skomplikowanych i kosztownych instalacji.

Pomimo tych problemów, naukowcy na całym świecie kontynuują prace nad opanowaniem energii syntezy termojądrowej. Wiele krajów, w tym Stany Zjednoczone, Unia Europejska, Chiny i Japonia, prowadzi badania nad tą technologią. Wiele z tych projektów jest finansowanych przez rządy i organizacje międzynarodowe, które widzą w energii syntezy termojądrowej potencjał jako źródło energii bezemisyjnej i, praktycznie rzecz biorąc, nieskończonej.

Jednakże nadal istnieją poważne problemy, które muszą zostać rozwiązane, zanim energia syntezy termojądrowej będzie mogła stać się praktycznym źródłem energii. Należy znaleźć sposób na osiągnięcie odpowiednich temperatur i ciśnień, a także na wydajne magazynowanie energii. Dodatkowo, należy rozważyć aspekty bezpieczeństwa, ponieważ energia syntezy termojądrowej jest bardzo mocno związana z bronią jądrową.

Opanowanie energii syntezy termojądrowej jest jednym z największych wyzwań nauki. Pomimo postępów w badaniach, praktyczne wykorzystanie tej technologii pozostaje wciąż odległe. Niemniej jednak, wiele krajów i organizacji wciąż inwestuje w badania i rozwój tej technologii, co może przynieść pozytywne efekty w przyszłości.

I tego sobie życzymy.



Romuald Nowak, Aleksandrów Kujawski

AI ożywia zmarłych

Szanowni Redaktorzy,

Chciałbym poruszyć kontrowersyjny temat związany z algorytmami sztucznej inteligencji ożywiającymi osoby, które zmarły, o czym pisałoście niedawno.

W ostatnich latach postępy w dziedzinie sztucznej inteligencji umożliwiły stworzenie algorytmów, które pozwalają na tworzenie wirtualnych postaci zmarłych osób. Te postacie wykorzystują sztuczną inteligencję do naśladowania stylu mówienia, mimiki i zachowań zmarłej osoby, co ma na celu przypomnienie o niej i ułatwienie radzenia sobie z jej utratą.

Jednakże algorytmy ożywiające osoby, które zmarły, budzą wiele kontrowersji. Niektórzy uważają, że są one bardzo cenne i mogą pomóc ludziom w radzeniu sobie z traumą związaną ze stratą bliskiej osoby. Inni natomiast uważają, że stwarzają one problemy etyczne i mogą prowadzić do złudzenia, że zmarła osoba nadal jest obecna w naszym życiu.

Jeden z największych problemów związanych z algorytmami ożywiającymi osoby, które zmarły, to kwestia prywatności. Aby stworzyć wirtualną postać zmarłej osoby, konieczne jest posiadanie dużej ilości danych o niej, takich jak nagrania wideo, dźwięku i zdjęcia. Takie dane są często przechowywane przez rodziny zmarłych osób i ich dostępność i wykorzystanie bez zgody zmarłej osoby lub jej bliskich może stanowić naruszenie prywatności.

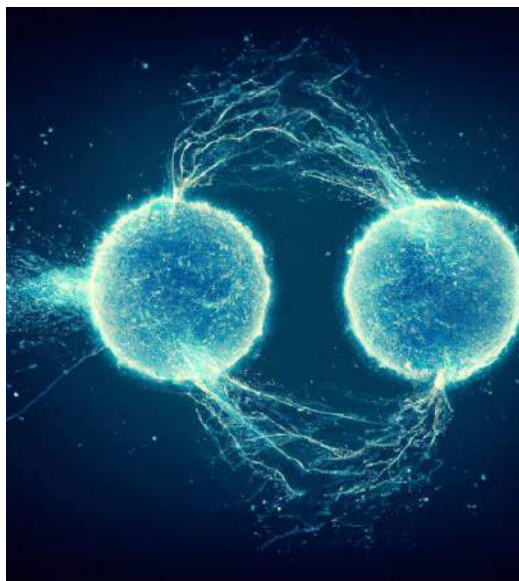


Dodatkowo, algorytmy ożywiające osoby, które zmarły, mogą prowadzić do zakłócenia żałoby. Wirtualna postać zmarłej osoby może stwarzać złudzenie, że ta osoba nadal jest obecna w naszym życiu, co może prowadzić do trudności w zaakceptowaniu faktycznej straty i zakończeniu procesu żałoby.

Podsumowując, algorytmy ożywiające osoby, które zmarły, stanowią kontrowersyjną dziedzinę sztucznej inteligencji. Choć mogą one pomóc w radzeniu sobie z traumą związaną z utratą bliskiej osoby, stwarzają również problemy związane z prywatnością i utrudnieniem procesu żałoby. Należy więc dokładnie rozważyć zalety i wady tej technologii przed jej wykorzystaniem.

Z wyrazami szacunku,

Marian Kończak z Lubartowa



ŹRÓDŁA ENERGII

Atmosferyczny wodór na elektryczność

Jak wynika z publikacji na temat badań naukowców z Uniwersytetu Monash w Melbourne w Australii, która ukazała się w czasopiśmie „Nature”, udało im się odkryć enzym pozyskany z bakterii glebowych *Mycobacterium smegmati*, który przekształca śladowe ilości wodoru w atmosferze w energię elektryczną. Zdaniem badaczy odkrycie to mogłoby znaleźć zastosowanie w ogniwach paliwowych.

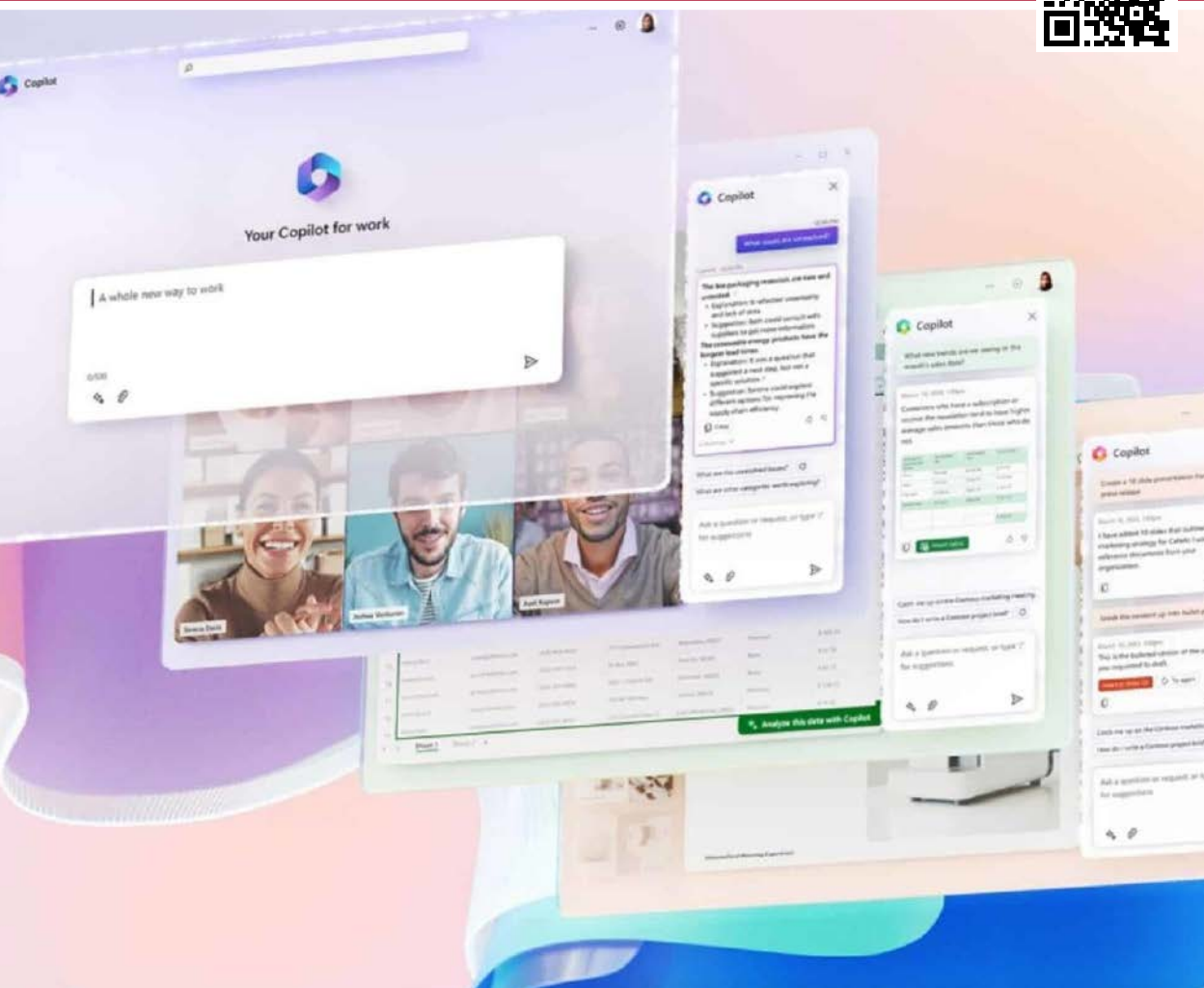
Uczeni nazwali, wyodrębnioną za pomocą zaawansowanych technik mapowania molekularnego, substancję huc. „Huc jest wyjątkowo wydajny”, mówi w komunikacie Rhys Grinter, szef zespołu badawczego na Uniwersytecie Monash. „W przeciwieństwie do wszystkich innych znanych enzymów i katalizatorów chemicznych, pochłania wodór nawet na poziomach mniejszych niż atmosferyczne – już od zawartości 0,00005 procent w powietrzu, którym oddychamy”.

Zdaniem badaczy, ilość prądu wytwarzanego przez huc jest stabilna i wystarczająca do zasilania i ładowania akumulatorów w niewielkich urządzeniach elektrycznych. Dostarczenie większych ilości wodoru znacząco zwiększa produkcję energii elektrycznej. Australijczycy uważają, że jest to szansa na zastosowanie ogniw paliwowych w szerokiej gamie urządzeń, od smartwatchy po samochody. ■



Udostępnienie przez OpenAI nowej wersji swojego tzw. „wielkiego modelu językowego” (ang. „Large Language Model”, LLM) o nazwie GPT-4 było o tyle niespodzianką, że nikt nie wiedział dokładnie, kiedy to się stanie. Gdy więc wydano aktualizację modelu stojącego za jego popularnym ChatGPT, wszyscy byli nieco zaskoczeni. Autorzy nazywają GPT-4 swoim „najbardziej zaawansowanym systemem, produkującym bezpieczniejsze i bardziej użyteczne odpowiedzi”. Został zaprojektowany jako podstawowy silnik zasilający chatboty i wszelkiego rodzaju inne systemy, od wyszukiwarek po asystentów personalnych i korepetytorów online.

GPT-4 rozwija podstawowe możliwości oprogramowania ChatGPT, umożliwiając rozwiązywanie trudniejszych problemów z większą dokładnością dzięki szerszej wiedzy ogólnej i umiejętnościom radzenia sobie z wyzwaniami wyższej klasy. Wprowadza również nowe funkcje, przede wszystkim przetwarzanie obrazów jako danych wejściowych ze szczegółowym rozpoznawaniem zawartości plików graficznych. By poznać możliwości „czwórki”, trzeba być płatnym subskrybentem w planie abonamentowym ChatGPT Plus. Nowy model jest też udostępniony jako programistyczne API dla deweloperów. OpenAI, która zatrudnia około 375 pracowników, jednak jest wspierana miliardami dolarów inwestycji ze strony Microsoftu i znanych osobistości z sektora technologicznego. Wspominane na początku zaskoczenie wynikało także z tego, że spodziewano się, że premiera GPT-4 będzie połączona z wydarzeniem prezentującym nowe usługi Microsoftu. Ta jednak odbyła się dwa dni później. Zade monstrowano na niej rozwiązanie o nazwie Copilot, która wprowadza funkcje generacji AI w MS Word, Excelu, PowerPointcie i w innych aplikacjach Microsoft Office. Warto zaznaczyć, że, według dostępnych informacji, nowy model GPT-4 zasila już również zintegrowane z wyszukiwarką Bing Microsoftu narzędzie chatu AI, które dostępne było już wcześniej, od stycznia 2023 r.



REWOLUCJA AI

Wejście GPT-4, Google'a i innych potentatów do boju na froncie AI

W tym samym mniej więcej czasie, niemal równolegle z prezentacjami OpenAI i Microsoftu, swoje nowe oferty oparte na własnych modelach AI zaprezentowało Google. Były to aktualizacje usług chmury Google'a i pakietu Google Workspace (Gmail, Google Calendar, Google Drive, Google Docs i Google Meet). Swoją premierę, choć wciąż skierowaną do wybranego grona, miał Bard, narzędzie chatu połączone z generatorem AI, pomyslane

jako konkurencja dla ChatGPT. Próbuje od lat konkurować z Google wyszukiwarka DuckDuckGo w odpowiedzi na serię doniesień o produktach AI uruchomiła DuckAssist, nową funkcję, która generuje odpowiedzi w języku naturalnym na zapytania wyszukiwawcze z wykorzystaniem Wikipedii. Firma Adobe również w tym samym mniej więcej czasie zaprezentowała własny generator obrazów o nazwie FireFly. ■



TECHNIKA MEDYCZNA

Elektroniczny bandaż przyspieszający gojenie o 30%

Zespół specjalistów z Uniwersytetu Northwestern opracował innowacyjny typ opatrunku, elastyczny, rozciągliwy bandaż, który przyspiesza gojenie przez dostarczanie elektrostymulacji bezpośrednio do miejsca rany. W badaniach przeprowadzonych na myszach, nowy opatrunek goił wrzody cukrzycowe 30% szybciej niż u zwierząt bez opatrunku.

Według pracy opisującej jego działanie w periodyku naukowym „Science Advances”, bandaż monitoruje również proces gojenia w sposób aktywny, a następnie rozpuszcza się bez uszczerbku dla organizmu wraz z elektrodami, gdy nie jest już potrzebny. Jest przedstawiany jako pierwszy biodegradowalny opatrunek zapewniający elektroterapię.

Nowe urządzenie może dostarczyć skutecznego narzędzia pacjentom z cukrzycą, u których wrzody mogą prowadzić do różnych powikłań, w tym amputacji kończyn, a nawet śmierci. Realizuje cel szybkiego, pozbawionego niekorzystnych efektów ubocznych leczenia otwartych ran, które są szczególnie podatne na infekcje. ■

15 metrów – na takim dystansie kierowca prowadzący samochód z prędkością 100 km/h ma zamknięte oczy podczas kichania.



BIOTECHNOLOGIE

Biodrukowanie 3D wewnątrz ludzkiego ciała

Innowacyjny zrobotyzowany system biodruku 3D, który potrafi drukować biomateriał komórkowy wewnątrz ludzkiego ciała, czyli w organach wewnętrznych i w tkankach, został opracowany przez naukowców z Uniwersytetu Nowej Południowej Walii w Australii i opisany w czasopiśmie „Advanced Science”. Nowe rozwiązanie pozwala na wykorzystanie technologii biodruku 3D do tworzenia precyzyjnych i niestandardowych implantów medycznych bez konieczności ich chirurgicznego usuwania lub przemieszczania.

Urządzenie prototypowe, nazwane F3DB, jest sterowane z zewnątrz i składa się z długiego, elastycznego ramienia robotycznego, na końcu którego znajduje się zwrotna głowica obrotowa, która „drukuję” biokomponent przez miniaturową, wielokierunkową dyszę. Robot może być wykorzystywany ponadto do wytwarzania implantów z materiałów, które są niemożliwe do uzyskania w tradycyjny sposób.

Obecnie technologia bioprintingu wykorzystywana jest głównie do badań naukowych oraz do opracowywania nowych leków. Wymaga użycia dużych drukarek 3D do tworzenia konstrukcji, które są chirurgicznie wszczepiane do organizmu, co niesie ze sobą odrębne zagrożenia, w tym ryzyko uszkodzenia tkanek i infekcji. Według badaczy stojących za opracowaniem nowej techniki, ich znacznie zredukowany co do rozmiaru i poziomu komplikacji system może być używany do tworzenia implantów wewnętrznych w krótkim czasie, zmniejszając koszty i minimalizując ryzyko powikłań chirurgicznych. Może to potencjalnie zrewolucjonizować leczenie schorzeń wymagających niestandardowych implantów medycznych. ■

HAKOWANIE SZTUCZNEJ INTELIGENCJI

GPT skopiowany za grosze

Grupa badaczy z Uniwersytetu Stanforda zbudowała tzw. duży model językowy (LLM) nie gorszy, jak twierdzą, niż GPT firmy OpenAI. Nie byłaby to sensacyjna wiadomość, bo w końcu obecnie powstaje równych modeli tego rodzaju spora liczba. Jednak specjaliści z kalifornijskiej uczelni twierdzą, że zbudowanie LLM kosztowało ich zaledwie ok. 600 dolarów, czyli grosze, w porównaniu z kosztami ponoszonymi przez zespoły pracujące nad tworzeniem takiej AI. Według komentatorów może to pozwolić tworzyć za niewielkie pieniądze sprawnie działające niezależne formy AI.

Zespół badawczy ze Uniwersytetu Stanforda zaczął od otwartego modelu językowego LLaMA 7B firmy Meta, najmniejszego i najtańszego z kilku dostępnych modeli LLaMA. Wstępnie wyszkolony na bilionie „tokenów”, ten stosunkowo niewielki model językowy ma pewne możliwości, ale w większości zadań pozostaje w tyle za ChatGPT. Główny koszt, a w istocie główna przewaga konkurencyjna modeli GPT, pochodzi w dużej mierze z ogromnej ilości czasu i siły roboczej, jaką OpenAI włożyła w post-training swojego modelu. Mając już gotowy model LLaMA 7B, zespół Stanforda poprosił GPT o wybranie 175 par

instrukcji i wyjść napisanych przez człowieka i rozpoczęcie generowania kolejnych w tym samym stylu i formacie, po dwadzieścia na raz. Zostało to zautomatyzowane z wykorzystaniem API firmy OpenAI.

W krótkim czasie zespół miał około 52 tysiące przykładowych rozmów do wykorzystania w postęcytreningu modelu LLaMA. Wygenerowanie tych masowych danych treningowych kosztowało mniej niż 500 dolarów. Następnie wykorzystali te dane do dostrojenia modelu LLaMA – proces ten trwał około trzech godzin na ośmiu 80-gigabajtowych komputerach A100 przetwarzających dane w chmurze. Kosztowało to mniej niż 100 USD. Następnie przetestowali uzyskany w ten sposób model, który nazwali Alpaca, porównując go z modelem językowym ChatGPT w różnych dyscyplinach, w tym w pisaniu e-maili, mediach społecznościowych i narzędziach zwiększających produktywność. Alpaca wygrała 90 z tych testów, GPT – 89. Jak zauważają, zachowanie modelu dorównuje GPT-3.5, jednak przy dostępie do GPT-4 prawdopodobnie osiągnie wyniki porównywalne z ogłoszonym właśnie potężnym modelem OpenAI. Kilka dni po demonstracji modelu Alpaca pojawiła się informacja o jego zamknięciu ze względu na koszty i bezpieczeństwo. ■

5,5 stopnia w skali Richtera – do tak silnego trzęsienia ziemi porównywany był wybuch w składzie amunicji w amerykańskim New Jersey w 1916 r., wywołany przez niemieckich dywersantów.

138 m rozpiętości (wg innych pomiarów 120 m) ma największy na świecie naturalny skalny most – Xian Ren Qiao, czyli Most Bajkowy, znajdujący się nad rzeką Buliu w chińskiej prowincji Guangxi.



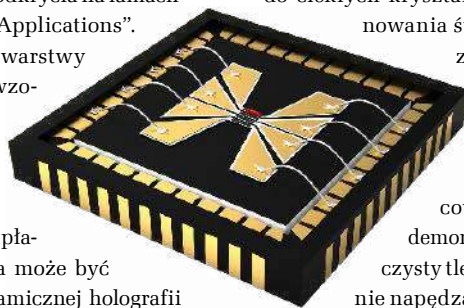
NANOTECHNOLOGIA

Metapowierzchnie w wyświetlaczach

Naukowcy z uczelni brytyjskich i australijskich opracowali technikę wyświetlania opartą na metapowierzchniach z krzemu, opisując swoje odkrycia na łamach periodyku „Light: Science & Applications”.

Metapowierzchnie to cienkie warstwy materiału z nanoskalowymi wzorami, które mogą manipulować światłem w nietypowy sposób. Uważa się, że rozwiązania tego rodzaju mogłyby zastąpić warstwy LCD w telewizorach z płaskim ekranem. Technologia ta może być również wykorzystana do dynamicznej holografii wirtualnej, w technologiach LiDAR oraz do produkcji cieńszych płaskich paneli o rozdzielczości stukrotnie wyższej niż obecne ekrany oparte na LCD, przy jednocześnie zmniejszeniu zużycia energii o połowę.

„Możliwości konwencjonalnych wyświetlaczy osiągnęły swój szczyt i jest mało prawdopodobne, aby w przyszłości uległy znaczącej



poprawie ze względu na liczne ograniczenia”, stwierdza w publikacji Dragomir Neshev, profesor fizyki z Narodowego Uniwersytetu Australii. „Obecnie trwają poszukiwania techniki w pełni półprzewodnikowych płaskich wyświetlaczy o wysokiej rozdzielczości i szybkiej częstotliwości odświeżania. Zaprojektowaliśmy i opracowaliśmy metapowierzchniowe piksele, które, w przeciwieństwie do ciekłych kryształów, nie wymagają do funkcjonowania światła spolaryzowanego, co pozwoli zmniejszyć o połowę zużycie energii przez ekrany”.

Aby kontrolować poszczególne piksele z dużą szybkością modulacji, opracowana przez badaczy platforma demonstracyjna wykorzystuje przezroczysty tlenek przewodzący jako elektrycznie napędzany grzejnik, który może szybko zmieniać właściwości optyczne krzemowych komórek metapowierzchniowych, mających sto razy mniejszą grubość od komórek ciekłokrystalicznych. Rezultatem tej metody jest czas reakcji poniżej jednej milisekundy, co według publikacji dziesięciokrotnie przekracza granicę wykrywalności dla ludzkiego oka. ■

E-MOBILNOŚĆ

Elektryczny rower bez łańcucha

Firmy Schaeffler i Heinzmann, niemieccy dostawcy branży motoryzacyjnej, opracowały nowy układ napędowy do rowerów elektrycznych typu „ride-by-wire”, który nazwali Free Drive. System eliminuje potrzebę stosowania tradycyjnego łańcucha, kasety i przerzutek, zamiast tego wykorzystując silnik elektryczny do bezpośredniego napędzania tylnego koła roweru. Free Drive ma być bardziej wydajny, mniej awaryjny i zapewniać płynniejszą jazdę.

Innowacja polega na tym, że generator elektryczny jest zamontowany na osi obrotu pedałów. Zamiast napędzać łańcuch przenoszący obroty na koło,

rowerzysta bezpośrednio napędza wirnik generatora, który wytwarza energię elektryczną. Ta przewodami przepływa do silnika i jest ponownie zamieniana na energię kinetyczną kół roweru. Wszelka dodatkowa energia elektryczna, która powstaje wskutek pedalowania, ponad to, co jest potrzebne silnikowi, jest magazynowana w akumulatorze, który może być również ładowany z gniazdka.

W tej chwili za najbardziej obiecujące pole zastosowań nowego napędu uważa się roweropodobne pojazdy do przewożenia ładunków, w których stosowanie łańcuchów jest wyzwaniem konstrukcyjnym. System Free Drive pozwala na zaprojektowanie wózków tego typu w inny sposób, obniżając podłogę lub pokład i tworząc znacznie bardziej stabilną platformę, która nie jest ograniczona przez wymagania mechanicznego układu napędowego roweru. Oczywiście rozwiązanie to upraszcza konstrukcję każdego roweru, zaś w e-rowerach wydaje się naturalnym rozwiązaniem. ■





Demonstracja
zmapowanego
konnektomu mózgu
muszki owocowej:
<https://bit.ly/42m8gxL>



BIONIKA

Mapa owadziego mózgu – pierwsza o takim stopniu złożoności

Pierwszą na świecie kompletną, wysokiej rozdzielczości mapę mózgu żywego stworzenia, w tym przypadku larwalnego stadium muszki owocowej z gatunku *Drosophila melanogaster*, opracowali badacze z amerykańskich i brytyjskich uczelni. Jak donosi czasopismo „Science”, to najbardziej złożony i skomplikowany układ połączeń mózgu (konnektom) zwierzęcia, jaki kiedykolwiek powstał.

To zwieńczenie pracy naukowców z uniwersytetu w Cambridge i Uniwersytetu Johnsa Hopkinsa

5 litrów powietrza
może podczas wdechu pobrać
człowiek, choć w trakcie
odpoczynku jest to zwykle nie
więcej niż 5% tej wartości.

i dopiero czwarty kompletny konnektom, jak dotąd. Wcześniej zmapowano znacznie prostsze mózgi mikroskopijnej glisty *Caenorhabditis elegans*, larwy morskiej *Ciona intestinalis* i morskiego robaka *Platynereis dumerilii*, ale te zawierały co najwyżej kilkaset neuronów i kilka tysięcy synaps (połączeń). Zaś wcześniejsze konnektomy muszek owocowych stanowiły jedynie częściowe obrazy.

Wysokiej rozdzielczości konnektom *Drosophila melanogaster*, przedstawia 3016 neuronów i 548 000 połączeń między nimi. Chociaż mikroskop elektronowy uchwycił obraz każdego wycinka, samo obrazowanie trwało dobę dla każdego z 3016 neuronów. Naukowcy pracują obecnie nad mapowaniem mózgu dorosłej muszki owocowej. Obecne narzędzia obliczeniowe mogą prześledzić miliony ścieżek neuronowych, ale nie biliony, które istnieją w ludzkim mózgu. Dlatego specjaliści uważają, że do zmapowania tą techniką naszego konnektomu jest jeszcze długa droga. ■



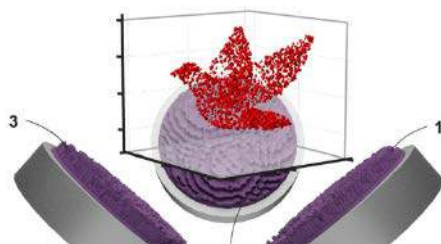
ENERGIA

Fotowoltaika pomiędzy szynami

Szwajcarski start-up Sun-Ways rozpoczyna instalację paneli słonecznych umieszczonych pomiędzy torami kolejowymi w pobliżu stacji Buttes w zachodniej części Szwajcarii. System wymiennych instalacji fotowoltaicznych zintegrowanych z torami powstał we współpracy z politechniką w Lozannie (EPFL).

Szwajcarska firma, z siedzibą w Ecublens, opracowała mechaniczny system instalacji wymiennych paneli słonecznych. Pociąg zaprojektowany przez inną szwajcarską firmę Scheuchzer ma poruszać się wzdłuż szyn, układając panele fotowoltaiczne. Wykorzystuje się tu mechanizm tłokowy do rozwijania paneli o szerokości jednego metra, wstępnie zmontowanych w zakładzie produkcyjnym. Sun-Ways porównuje ten proces do „rozwijania dywanu”. Produkowana przez panele energia trafiać ma do sieci energetycznej. Na wątpliwości dotyczące możliwego pęknięcia, zasypiania śniegiem lub oblodzenia firma ma odpowiedź w postaci sieci czujników monitorujących i systemów topiących śnieg i lód.

Obliczono, że gdyby panele fotowoltaiczne zostały rozmieszczone w torowiskach całej sieci kolejowej Szwajcarii o długości 5317 kilometrów, pokryłyby obszar o wielkości około 760 boisk piłkarskich i produkowałyby, jak szacuje Sun-Ways, jedną terawatogodzinę energii elektrycznej rocznie, co zaspokoiłoby ok. 2 proc. rocznego zapotrzebowania Szwajcarii. ■



BIOMIMETYZM

Obrazy dynamiki grup atomowych

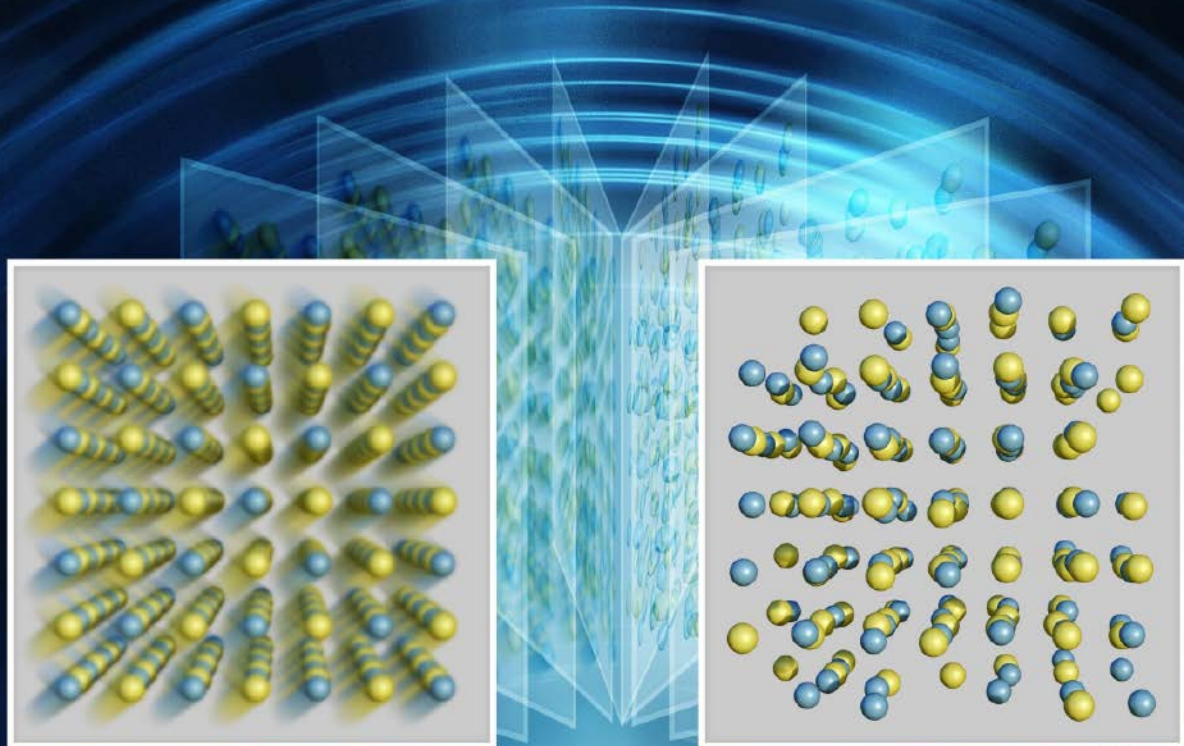
Zespoły badawcze z niemieckiego Instytutu Maxa Plancka zajmującego się badaniami medycznymi oraz uniwersytetu w Heidelbergu opracowały sposób na drukowanie obiektów 3D, w jednym „kawałku”, przy użyciu fal dźwiękowych. Ich technika, która została opisana w czasopiśmie „Science Advances”, wykorzystuje profilowane fale ultradźwiękowe do tworzenia „akustycznych hologramów”, wywierających nacisk na drukowany materiał drukowany, niczym niewidzialna forma.

Grupa wykorzystwała swoją technikę do drukowania obiektów 3D z mikrogranulek materiału, hydrożelowych kulek, a nawet komórek biologicznych – wszystko to całkowicie „bez dotyku”. Jak wyjaśnia w publikacji na temat badań Heiner Kremer, który napisał program drukujący tą techniką, przejście od modelu 3D do roboczej sfery ultradźwiękowej było dość trudne, wymagało zastosowania niestandardowego algorytmu.

Co najważniejsze, fale dźwiękowe nie uszkadzają komórek, co oznacza, że nie są potrzebne żadne chemiczne środki pomocnicze, a w konsekwencji także to, iż technika ta może być szczególnie przydatna w bio-printingu 3D tkanek i kultur komórkowych. Fale ultradźwiękowe mogą także przenikać w głąb tkanek, co pozwala na manipulowanie nimi w sposób mniej inwazyjny. ■

2370°C miała najgorętsza znana skała, której pozostałości znaleźli uczeni w kraterze uderzeniowym (obecnie jeziorze) Mistatin, w Kanadzie.

1 000 000 000 tesli ma pole magnetyczne gwiazdy neutronowej GRO J1008-57, najsilniejsze ze znanych nam we Wszechświecie.



OBRAZOWANIE

Kamera widzi, które atomy tańczą

Uczni z uniwersytetów Columbia w USA i Burgundii we Francji donieśli, że udało im się opracować nowy rodzaj „kamery” z czasem otwarcia migawki wynoszącym 1 bilionową część sekundy, która potrafi dostrzec lokalne zaburzenia w dynamicznym układzie atomów, które są rozmyte przy zastosowaniu wolniejszych migawek.

Nowa metoda, którą nazwali zmienną migawką PDF lub vsPDF (skrót angielskojęzycznego terminu „pair distribution function”), nie działa jak konwencjonalny aparat fotograficzny, lecz wykorzystuje neutrony pochodzące z laboratorium w Oak Ridge (ORNL) Departamentu Energii USA do pomiaru pozycji atomów z czasem otwarcia „przysłony” wynoszącym około jednej pikosekundy, milion milionów (bilion) razy szybciej niż normalne migawki aparatu. Badanie zostało opublikowane w czasopiśmie „Nature Materials”.

„Dzięki tej technice będziemy mogli obserwować materiał i zobaczyć, które atomy tańczą, a które pozostają bez ruchu”, wyjaśnia w komunikacie Simon Billinge z Uniwersytetu Burgundzkiego. Poznanie mechanizmu dynamicznych zaburzeń grup atomów w materiałach może doprowadzić do powstania wydajniejszych energetycznie urządzeń termoelektrycznych, np. półprzewodnikowych chłodziarek i pomp ciepła, a także do lepszego odzyskiwania użytecznej energii z ciepła odpadowego, np. spalin samochodowych i z elektrowni, przez przekształcanie go bezpośrednio w energię elektryczną. ■



Film objaśniający
działanie mechanizmu
obserwacji
dynamicznych
zaburzeń atomowych:
<https://bit.ly/44qRZJT>

4 000 000 lat zajmie wystrzelonej w 1973 roku sondzie Pioneer 11 dotarcie do gwiazdy Lambda w gwiazdozborze Orła.



ARCHITEKTURA

♦ Arabia Saudyjska ogłosiła plany budowy gigantycznej budowli w kształcie sześcienu Mukaab o boku ok. 400 metrów, który ma powstać w centrum nowego osiedla o nazwie Nowa Murabba w stolicy kraju, Rijadzie, i zawierać w sobie skwery z zielenią, hotele, biura i mieszkania, a także mnóstwo prezentacji holograficznych. ♦ W Miami powstaje siedemdziesięciopiętrowy (260 metrów) apartamentowiec nazywany The Residences at 1428 Brickell, który, gdy zostanie ukończony w 2027 roku, ma być pierwszym budynkiem tego typu wyposażonym w okna z fotowoltaicznymi szybami, mającymi zredukować zużycie energii przez budynek. ♦

SZTUCZNA INTELIGENCJA

♦ DARPA doniosła o udanych testach sztucznej inteligencji zainstalowanej w myśliwcu F-16w ramach programu ACE (Air Combat Evolution), który ma na celu zwiększenie zaufania ludzi do sztucznej inteligencji, będącej w stanie kontrolować systemy pokładowe i manewrować samolotem w locie, w sytuacjach bojowych. ♦ W ramach kontraktu pomiędzy Departamentem Obrony USA a firmą RealNetworks z Seattle oprogramowanie firmy, wykorzystujące uczenie maszynowe, wyposażą autonomiczne drony eksploatowane przez siły powietrzne USA w funkcje rozpoznawania twarzy. ♦

ENERGIA

♦ Zespół naukowców pod kierownictwem Po-Chun Hsu z uniwersytetu w Chicago opracował nowy materiał, roztwór wodny na bazie miedzi umieszczony w warstwach metalu i grafenu, zmieniający, dzięki zjawisku elektrochromizmu pod wpływem przyłożonego ładunku elektrycznego, właściwości odbijania i pochłaniania ciepła słonecznego, co w perspektywie można wykorzystać do ogrzewania i klimatyzacji budynków. ♦ Chiński producent turbin wiatrowych Mingyang Smart Energy zaprezentował projekt budowy gigantycznej turbiny wiatrowej MySE 18.X-28X z łopatami o długości 140 metrów, dającymi w sumie średnicę wirnika ponad 280 metrów, który przy średniej rocznej prędko-



ści wiatru 8,5 m/s może wygenerować 80 GWh energii elektrycznej rocznie, co wystarczyć miałyby na potrzeby 96 tys. odbiorców. ♦ Firma o nazwie B2U Storage Solutions ogłosiła, że zainaugurowała w Kalifornii działalność magazynu energii SEPV Sierra o pojemności 25 megawatogodzin, opartego na 1300 akumulatorach, wycofanych z samochodów elektrycznych, przez co, zanim trafią do recyklingu, ich wciąż dostępna pojemność jest użytkowana do przechowywania energii pochodzącej ze źródeł odnawialnych. ♦ Efekt fotowoltaiczny kryształów w ogniwach ferroelektrycznych może być zwiększony tysiąckrotnie, jeśli trzy warstwy tytanianu baru, tytanianu strontu i tytanianu wapnia zostaną naprzemiennie nałożone na siebie we wzorze siatki – podali naukowcy z Uniwersytetu Marcina Lutra w Halle-Wittenberdze po przeprowadzeniu badań. ♦

FIZYKA

♦ Inżynierowie z Uniwersytetu Technologicznego w holenderskim Eindhoven opracowali fotodiode, która, jak twierdzą, charakteryzuje się sprawnością 200%, za czym mogą stać, nie do końca wyjaśnione przez naukowców, efekty kwantowe, powodujące, że sensor reaguje przepływem prądu na różne zakresy światła a nawet na jego brak. ♦ W pracy, opisaną w „Advanced Materials”, naukowcy pod kierownictwem Ruijuan Xu z Uniwersytetu Karoliny Północnej wykazali, że antyferroelektryczny niobian sodu (NaNbO_3) przekształca się w ferroelektryk w miarę jak jego warstwa staje się coraz cieńsza, np. w zakresie od 40 nm do 164 nm zawiera fazy ferroelektryczne w niektórych regionach i fazy antyferroelektryczne w innych, zaś gdy warstwa była cieńsza niż 40 nm, niobian sodu jest już całkowicie ferroelektryczny. ■

M. U.



Repozytorium Robotyki

– cyfrowe udostępnianie zasobów nauki z obszaru robotyki

Robotyka i dziedziny pokrewne rozwijają się obecnie dynamiczniej niż kiedykolwiek dotąd. Nieocenioną pomocą na drodze do kreowania przyszłości tego obszaru jest wiedza zdobyta wcześniej. Z myślą o wsparciu w tym zakresie zarówno świata przemysłu, jak i nauki Przemysłowy Instytut Automatyki i Pomiarów PIAP należący do Sieci Badawczej Łukasiewicz opracował Repozytorium Robotyki. Celem jest zapewnienie dostępu do wiedzy o zdobycach techniki – zarówno historycznych, jak i najnowszych – z zakresu szeroko rozumianej automatyki i robotyki, a co za tym idzie, przyczynienie się do rozwoju nauki i zwiększenia innowacyjności polskiego przemysłu.

Repozytorium Robotyki to unikatowe kompendium, dostępne powszechnie i bezpłatnie. W zdigitalizowanej formie udostępnione są zgromadzone w Łukasiewiczu – PIAP prace naukowe, badawcze i rozwojowe z zakresu robotyki oraz obszarów pokrewnych, takich jak automatyka, pomiary, sztuczna inteligencja, rozpoznawanie obrazów, przetwarzanie mobilne itp. Korzystający z repozytorium mają również dostęp do raportów z badań, czasopism, opisów projektów urządzeń oraz elementów robotycznych, a ponadto do bogatych zasobów biblioteki instytutu Łukasiewicz – PIAP. – Doświadczenie w dziedzinie szeroko rozumianej robotyki gromadzimy od ponad 50 lat. Prowadzone w tym czasie prace, podobnie jak gromadzone biuletyny oraz czasopisma naukowe i branżowe, były skrupulatnie archiwizowane, ale przechowywane w dużej mierze w formie papierowej. To znacznie utrudnia ich znajdowanie oraz ogranicza możliwości korzystania z nich. Dzięki repozytorium te zasoby zyskają drugie życie, tym razem w obiegu cyfrowym i otwartym dla wszystkich zainteresowanych – mówi dr inż. Małgorzata Kaliczyńska, kierownik projektu.

Wsparcie innowacyjności

Łukasiewicz – PIAP może pochwalić się kilkudziesięcioletnim doświadczeniem w prowadzeniu prac badawczo-rozwojowych, a także w tworzeniu rozwiązań z zakresu szeroko pojętej robotyki i wdrażaniu ich w różnych gałęziach przemysłu. Wśród obszarów działań instytutu jest automatyzacja i robotyzacja linii produkcyjnych i fabryk, produkcja robotów mobilnych, antyterrorystycznych i rehabilitacyjnych, tworzenie

rozwiązań z zakresu druku 3D oraz technologii kosmicznych. Dzięki tak szerokiemu spektrum działania instytutu Repozytorium Robotyki ma unikalną wartość historyczną, a jednocześnie może pośrednio przysłużyć się dalszemu rozwojowi robotyki w Polsce. – Polska ma wciąż wiele do nadrobienia w obszarze robotyzacji. Mamy ambicję wpłynąć na zmianę tego stanu. Udostępnianie wiedzy o najnowszych zdobycach techniki z zakresu automatyki i robotyki i ich świadome wykorzystanie na liniach produkcyjnych powinno przyczynić się do zwiększenia innowacyjności polskiego przemysłu. Udostępnimy zarówno wiedzę z początków rozwoju robotyki, jak i tę najnowszą, a internauci, korzystając z zasobów repozytorium, będą mogli analizować rozwój techniki – podkreśla dr inż. Małgorzata Kaliczyńska.

Nie tylko historia

Ideą przyświecającą pracom nad repozytorium jest połączenie przeszłości z teraźniejszością, dlatego stanowi ono źródło wiedzy nie tylko o osiągnięciach archiwalnych, ale również bieżących – będzie stale aktualizowane i wzbogacane o kolejne zasoby. – Obserwując tempo, w jakim ewoluuje technologia IT i możliwości, jakie daje Internet nie można dokładnie przewidzieć, co przyniosą kolejne dwa czy trzy lata, ale pewne jest, że możemy spodziewać się dalszego dynamicznego rozwoju na polu robotyki i dziedzin pokrewnych. Z tego względu planujemy modyfikowanie i rozbudowywanie portalu udostępniającego zasoby bazy wiedzy zgodnie ze zmieniającą się rzeczywistością – podkreśla dr inż. Małgorzata Kaliczyńska.

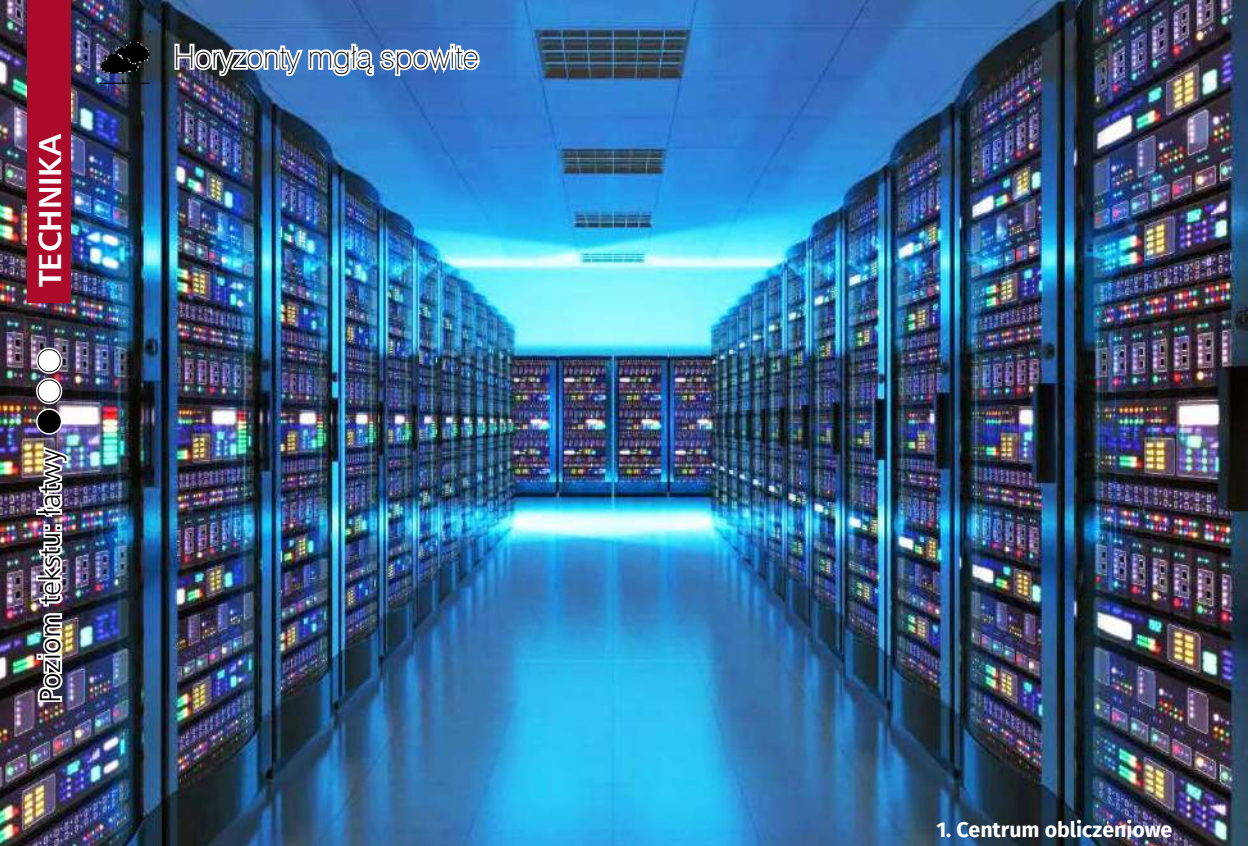
Zapraszamy na stronę roborepo.pl.



Rzeczpospolita
Polska

Unia Europejska
Europejski Fundusz
Rozwoju Regionalnego





1. Centrum obliczeniowe

Generatywna AI ma brudny
i kosztowny sekret

Nieokiełznany apetyt na energię

Obliczenia wykorzystywane do nauki i użytkowania modeli uczenia maszynowego są drogie i pochłaniają mnóstwo energii. Najtęższe głowy zastanawiają się, jak sobie z tym poradzić.

Giganci Big Tech, Google, Microsoft, Meta, Baidu i inni wydali, wydają i będą wydawać ogromne pieniądze na budowę lub zakup generatywnych narzędzi AI, które wykorzystują duże modele językowe (LLM), które odpowiadają na pytania, serwują zestawy informacji, poprawiają teksty, generują kod programistyczny i robią mnóstwo innych rzeczy. Tak naprawdę pełna lista ich zastosowań nie została jeszcze zamknięta.

Niestety to wszystko kosztuje, a jeśli konwersacyjne narzędzia AI zostaną zintegrowane z wyszukiwarkami,

kosztować będzie jeszcze znacznie więcej. Wyścig w budowaniu wysokowydajnych, napędzanych sztuczną inteligencją wyszukiwarek, według wszelkich znaków na niebie i na ziemi, będzie wymagał dramatycznego wzrostu potrzebnej mocy obliczeniowej, a wraz z nią ogromnego skoku w zużyciu energii przez firmy technologiczne.

Szkolenie LLM-ów, takich jak te, które leżą u podstaw ChatGPT firmy OpenAI, który już zasila wyszukiwarkę Bing Microsoftu, oraz jego odpowiednik

ze stajni Google'a nazwany Bard, oznacza przeszukiwanie i obliczanie powiązań w ogromnych zasobach danych. Koszty tych operacji są niebagatelne. Dość powiedzieć, że nie jest przypadkiem, że za LLM-ami stoją największe i najbogatsze firmy.

Chociaż ani OpenAI, ani Google nie ujawniają, jaki jest dokładny koszt obliczeniowy i energetyczny ich produktów, to według analiz zewnętrznych przeprowadzonych przez badaczy niezależnych od tych firm, do szkolenia GPT-3, na którym częściowo opiera się ChatGPT, zużyto 1287 MWh energii elektrycznej. Ale szkolenie to dopiero początek kosztów obliczeniowych i energetycznych. Należy wziąć pod uwagę fakt, że trzeba nie tylko wytrenować system, ale także go uruchomić i obsłużyć miliony użytkowników.

Istnieje również duża różnica pomiędzy wykorzystaniem ChatGPT, który według szacunków banku inwestycyjnego UBS miał w styczniu 2022 i 2023 r. już ok. stu milionów użytkowników dziennie, jako samodzielnego produktu a zintegrowaniem go z Bingiem, który obsługuje pół miliarda wyszukiwań każdego dnia i to bez uwzględnienia wzrostu popularności dzięki dodaniu chatbota.

Martin Bouchard, współzałożyciel kanadyjskiej firmy QScale, uważa, opierając się na analizie planów Microsoftu i Google'a dotyczących wyszukiwania, że dodanie generatywnej AI do tych usług będzie wymagało co najmniej od czterech lub pięciu razy więcej mocy obliczeniowej na samo wyszukiwanie. Zwraca też uwagę, że ChatGPT obecnie (marzec 2023 r.) ma dostęp do danych pochodzących najpóźniej z końca 2021 roku. By sprostać wymaganiom użytkowników wyszukiwarki, będzie musiał dodać wiele nowych informacji i parametrów związanych z dostępem do zasobów internetowych na bieżąco. A to z kolei będzie wymagało znaczących inwestycji w sprzęt, centra danych (1) i infrastrukturę. I oczywiście geometrycznie lub szybciej rosnący apetyt na moc obliczeniową.

Zapotrzebowanie na energię rośnie i bez AI

Według Międzynarodowej Agencji Energii, światowe centra danych już odpowiadają za około jeden procent światowej emisji gazów cieplarnianych. Globalne zużycie energii elektrycznej w centrach danych w 2021 r. wyniosło 220–320 TWh, czyli około 0,9–1,3 proc. globalnego końcowego zapotrzebo-

wania na energię elektryczną. Nie obejmuje to energii wykorzystywanej do wydobywania kryptowalut, która w 2021 r. wynosiła 100–140 TWh.

Chociaż zużycie energii elektrycznej w centrach danych na świecie wzrosło od 2010 roku umiarkowanie, w niektórych mniejszych krajach, w których rozwijają się rynki centrów danych, obserwuje się szybki wzrost. Na przykład zużycie energii elektrycznej w centrach danych w Irlandii wzrosło ponad trzykrotnie od 2015 r. i w 2021 r. będzie stanowić 14 proc.



całkowitego zużycia energii elektrycznej. W Danii przewiduje się, że do 2025 r. zużycie energii w centrach danych wzrośnie trzykrotnie i będzie stanowić około 7 proc. zużycia energii elektrycznej w tym kraju.

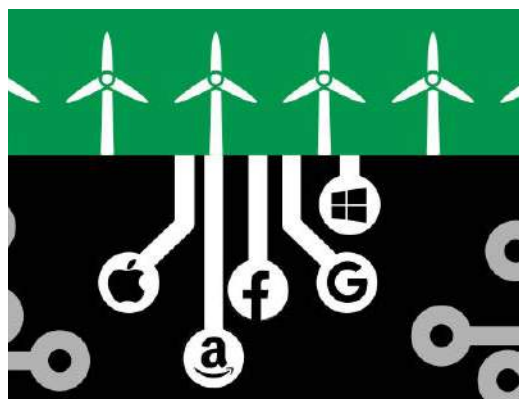
Oczekuje się, że silny wzrost popytu na usługi sieciowe w zakresie transmisji danych będzie się utrzymywał, napędzany tak czy inaczej, niezależnie od rozwoju usług AI, przez usługi tradycyjnie wymagające dużej ilości danych, takie jak strumieniowe przesyłanie wideo, gry w chmurze oraz aplikacje rzeczywistości rozszerzonej i wirtualnej. Przewiduje się również, że ruch danych w sieciach komórkowych będzie nadal szybko rósł, a do 2027 roku wzrośnie czterokrotnie, zaś udział 5G w mobilnym ruchu danych wzrośnie do 60 proc. w tej samej perspektywie czasowej, w porównaniu z 10 proc. w 2021 roku.

Wnioskowanie zużywa więcej energii

W Meta zapotrzebowanie na obliczenia do szkolenia sieci neuronowych (wzrastające o 150 proc. rocznie) i wnioskowania (wzrastające o 105 proc. rocznie) przewyższało w ostatnich latach ogólne zużycie energii w centrum danych (wzrastające o 40 proc. rocznie). Google podaje, że uczenie maszynowe odpowiadało jedynie za 10–15 proc. całkowitego zużycia energii przez centra obliczeniowe firmy, mimo że stanowiło 70–80 proc. całkowitego zapotrzebowania na moc obliczeniową.

Energia elektryczna potrzebna do uruchomienia modelu uczenia maszynowego jest funkcją algorytmu, programu, który go implementuje, liczby procesorów, które uruchamiają program, szybkości i mocy tych procesorów, wydajności centrum danych w dostarczaniu energii i chłodzeniu procesorów oraz miks dostaw energii (odnawialna, gazowa, węglowa, itp.). Napędzanie futurystycznej rewolucji sztucznej inteligencji przez stare dymiące spaliniaki źródła to zabawna, ale mająca wiele wspólnego z rzeczywistością ironia losu (2).

Większość firm wydaje więcej energii na obsługę modelu głębokiej sieci neuronowej (wykonywanie wnioskowania nazywanego po angielsku „inference”) niż na jego szkolenie. NVIDIA oszacowała, że 80–90 proc. obciążenia pracą w systemie uczenia maszynowego stanowi przetwarzanie wnioskowania. Potwierdza to Amazon, który informował, że w jego usługach obliczeniowych AWS 90 proc. zapotrzebowania na moc w chmurze dotyczy wnioskowania. Jeśli całkowita energia w systemach uczenia maszynowego rozkłada się w proporcjach: 10 proc. na szkolenie i w 90 proc. na obsługę, to nawet jeśli model AI wymagałby dwukrotnie większych



3. Big Tech i odnawialne źródła energii

kosztów energii na szkolenie, mógłby zmniejszyć całkowitą emisję dwutlenku węgla, gdyby udało się znacząco zmniejszyć koszt energetyczny w usługach serwowania wniosków.

Tu zaoszczędzić, tam zaoszczędzić, ale co z jakością?

Giganty technologiczne od dawna deklarują dążenie do zmniejszenia emisji i nawet osiągnięcia zerowego lub niższego bilansu emisji (3). Microsoft zobowiązał się do zejścia na ujemny poziom emisji dwutlenku węgla do 2050 roku. Google podjęło ambitniejsze zobowiązania – osiągnięcie zerowej emisji netto w całej swojej działalności już do 2030 roku.

Ślad węglowy i koszty energii związane z integracją AI z wyszukiwaniem można, jak najbardziej, zmniejszyć przez przestawienie centrów danych na czystsze źródła energii, ale to chyba jeszcze nieco potrwa, a rewolucja już teraz puka do drzwi. Podobnie można liczyć również na przeprojektowanie sieci neuronowych, aby stały się bardziej wydajne, zmniejszając tak zwany „czas wnioskowania”, który przekłada się na ilość mocy obliczeniowej wymaganej do pracy algorytmu na nowych danych. Ale znów – to trochę potrwa.

Rzecznik Google'a Jane Park powiedziała magazynowi WIRED, że Google opracowało wersję Barda, która była zasilana lżejszym LLM, co też ma znaczenie energetyczne. „Nasze wyniki pokazują, że połączenie wydajnych modeli, procesorów i centrów danych z czystymi źródłami energii może zmniejszyć ślad węglowy systemu uczenia maszynowego nawet tysiącrotnie”, mówiła Park. Komentatorzy zastanawiają się jednak, jak oszczędności i uszczuplenia w modelach, które znane są z tego, że polegają na jak największych i najbogatszych zasobach danych, wpływają na dokładność i jakość dla użytkownika. ■

Miroslaw Usidus

Destrukcja nanotechnologiczna

Czy replikatory opanują Ziemię?

Rzecz z natury niepozorna, bo mała, niewidoczna. Coraz więcej osób zwraca jednak uwagę, że właśnie to może być największym niebezpieczeństwem.

Na pozór gałąź ta wydaje się obiecywać rozwiązania wszystkich problemów ludzkości – przedłużanie ludzkiego życia w nieskończoność, wyeliminowanie znoju pracy fizycznej, szybka produkcja dowolnej rzeczy, jakiej sobie tylko zapragniemy, co oznacza zarazem, że zbędny stanie się handel międzynarodowy, będziemy w stanie produkować energię w dowolnych ilościach, większość chorób i uszkodzeń ciała będzie łatwo i natychmiastowo leczonych.

Wszystko to dzięki tej jednej technologii, zwanej ogólnie nanotechnologią. Futurolog, autor książki „Engines of Creation” (z ang. „Motory tworzenia”) K. Eric Drexler (1), autor tego terminu, widział w niej przede wszystkim narzędzie kreacji, ale również on zaczął używać innego pojęcia – „motor zniszczenia” na określenie wszystkiego, co opatrujemy przedrostkiem „nano”.



1. K. Eric Drexler

Replikacja w roju aż do końca

W sensie najbardziej chyba ścisłym nanotechnologia dotyczy molekularnych asemblerów, czyli mikroskopijnych maszyn, które mogą manipulować pojedynczymi atomami. Nanotechnologia to maszyny, które zasilane światłem słonecznym lub innymi źródłami energii, mogą wytworzyć cokolwiek ze zwykłych materiałów. W sensie „fizycznym” nie różni się to bardzo od sposobu, w jaki żywe organizmy mogą tworzyć z materiału biologicznego. Żywe organizmy są sterowane przez swoje DNA. Nanomaszyny mogą tworzyć obiekty, według instrukcji w oprogramowaniu, w tym powielać nanomaszyny, takie jak one same lub inne.

Maszyny, o których mowa w nanotechnologicznej dziedzinie, są małe, lekkie i wykorzystują materiały znajdujące się w najbliższym otoczeniu. Uważa się, że nanomaszyny mogłyby być wykorzystywane także militarnie, np. do infiltracji, szpiegowania, a nawet do przejmowania terytoriów przez opanowanie kluczowej infrastruktury, przejmowanie kontroli lub

po prostu niszczenie. Skoro to „motory tworzenia”, to mogą to być zarazem maszyny tworzące kopie samych siebie, w dowolnych ilościach, bez końca, a raczej z końcem w momencie, gdy już nie będzie dostępnego materiału.

Koncepcja samoreplikujących się maszyn została zaproponowana i omówiona przez wielu badaczy, m.in. Homera Jacobsena, Edwarda F. Moore’a, Freemana Dysona, Johna von Neumanna oraz przez wspomnianego K. Erica Drexlera. Przyszły rozwój takiej techniki można sobie wyobrazić jako integralną część projektów obejmujących wydobywanie bogactw naturalnych obcych planet i księżyców, zakładanie fabryk księżycowych, a nawet budowę satelitów słonecznych w kosmosie. John von Neumann pracował nad czymś, co nazwał Universal Builder, samoreplikującą się maszyną, która działałaby w środowisku komórkowym.

Samoreplikująca się maszyna to, jak sama nazwa wskazuje, sztuczny, samoreplikujący się system, który opiera się na konwencjonalnej wielkoskalowej



2. Roje samoreplikujących się nanomaszyn tak jak je widzi generator obrazów AI BlueWillow

technologii i automatyzacji. W literaturze czasami można znaleźć różne terminy, np. nanoroboty lub asemblery. Powielacze nazywano również „maszynami von Neumanna”. Znane jest też w tej gałęzi badań pojęcie „sonda von Neumanna”, oznaczające hipotetyczny pojazd kosmiczny. Byłaby to sonda zdolna do replikacji z surowców znalezionych tam, gdziekolwiek jest wysłana. Pomysł nie pochodzi od Johna von Neumanna, ale jest zastosowaniem koncepcji von Neumanna do eksploracji kosmosu. Koncepcję tę opisał Robert Freitas w 1980 roku.

To właśnie wizja von Neumanna przeniesiona z planu kosmicznego na Ziemię zaczęła z czasem wzbudzać największe emocje. Skoro replikator może korzystać z dostępnego materiału gdziekolwiek w przestrzeni kosmicznej, to może również replikować się zamieniając krok po kroku naszą planetę w masę zreplikowanych nanomaszyn.

Są nie tak daleko idące, jednak wciąż wywołujące dreszczyk grozy wizje. Alarmiści ostrzegają, że działające według opisanych wyżej zasad, nanoreplikatory mogą być zaprogramowane również do kopiowania zaawansowanych typów broni i ulepszania ich przez wykonywanie z coraz lepszych materiałów. Jeśli takie nanomaszyny będą współpracować ze sztuczną inteligencją, to zaczynamy mówić już nie o oprogramowaniu, lecz o systemie, który uczy się i przekracza proste instrukcje, doskonaląc efekty swojego działania.

Jeszcze bardziej przerażająca jest idea wykorzystania w złych celach nanotechnologii medycznych, wpływanie masowo na zdrowie i życie populacji ludzkiej. Najdalej idące obawy dotyczą całości życia ziemskiego, z którym nanomaszyny mogą konkurować i wygrywać w walce o zasoby, energię i budulec. „Rośliny” oparte na nanotechnologii mogłyby zastąpić prawdziwe rośliny. „Bakterie” mogłyby się rozprzestrzeniać jak prawdziwe mikroorganizmy (2), dewastując środowisko lub zagrażając ludzkości nanoinfekcjami.

Realnie, w tej chwili nanotechnologie nie piszą jeszcze tak przerażających scenariuszy, jednak negatywne dla przyrody i środowiska efekty stosowania nanorozwiązań już się dostrzega, np. w wynikach badań uczonych z moskiewskiego uniwersytetu RUDN. Według publikacji, która ukazała się w styczniu 2023 r. w czasopiśmie „Drug and Chemical Toxicology”, tlenki metali stosowane na coraz szerszą skalę w nanotechnologiach trafiają następnie do wody, np. jako jony cynku, zatruwając żyjące tam ryby.

Nanotechnologia nie potrzebuje więc aż „maszyn von Neumanna”, by aktywnie szkodzić życiu na naszej planecie. A to niestety może być dopiero początek. ■

Miroslaw Usidus

Upadek nauki łamiącej utarte schematy?

Mniej przełomów i innowacji

Według publikacji, która ukazała się na początku 2023 roku w magazynie „Nature”, liczba „przełomowych” osiągnięć naukowych spada. Nikt nie wie tak naprawdę, dlaczego tak się dzieje.

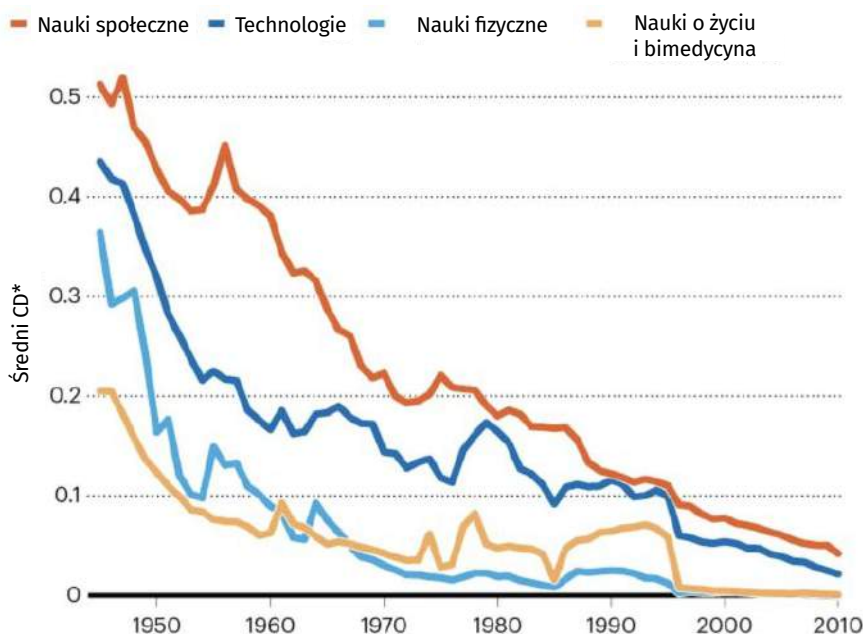
Można to wyrazić nieco precyzyjniej. Choć liczba publikacji o charakterze naukowym rośnie, to odsetek w tej masie publikacji, które sprawiają, że dana dziedzina badań zaczyna zmierzać w nowym kierunku, w ciągu ostatniego półwiecza wyraźnie zmalał. Czy jest to wyraz zjawiska, o którym piszą od lat publicyści, polegającego na domniemanej zapaści w nauce i ludzkiej innowacyjności, po okresie wielkiego rozkwitu w rewolucji przemysłowej, do lat 60.–70. XX wieku?

„Nature” podkreśla, że na przestrzeni ostatnich kilku dekad liczba publikowanych prac badawczych obejmujących różne dziedziny nauki i techniki gwałtownie wzrosła. Jednocześnie ich „przełomowość” spadła. Takie wnioski wynikają z analizy ich „dysrupcyjności”, czyli tego, na ile odbiegają od dotychczasowego dorobku badawczego w danej dziedzinie.

Analiza przeprowadzona na bazie liczącej miliony artykułów wskazuje, że w porównaniu z badaniami

Słabnąca innowacyjność w nauce

Aby zmierzyć, jak dana publikacja naukowa wstrząsa swoją dziedziną badań, badacze zastosowali wskaźnik nazwany indeksem CD, o wartości od 1 dla najbardziej innowacyjnej pracy do –1 dla najmniej przełomowej. Analiza milionów publikacji wskazuje, że przełomowość prac spadała w czasie we wszystkich analizowanych dziedzinach



1. Historia przełomów w nauce

*Average CD₅ is the CD index five years after a paper's publication.



z połowy XX wieku publikacje z XXI wieku częściej rozwijają myśli i wnioski formułowane wcześniej, niż zmieniają radykalnie i zasadniczo kierunek poszukiwań naukowych, co czyni zarazem poprzednie prace nieaktualnymi. Analiza dotycząca wniosków patentowych z lat 1976–2010 wykazała analogiczny mechanizm.

„Dane te sugerują, że coś się zmienia”, uważa Russell Funk, socjolog z Uniwersytetu Minnesoty w Minneapolis, współautor publikacji, która została opublikowana 4 stycznia 2023 r. w magazynie „Nature”. „Nie mamy już do czynienia z takim samym natężeniem przełomowych odkryć, jak kiedyś mieliśmy”.

Jak to w ogóle zbadano? Autorzy publikacji wyjaśniają, że jeśli publikowana praca była wysoce innowacyjna i przełomowa w swojej dziedzinie, to będące jej kontynuacją prace badawcze byłyby rzadziej cytowane w źródłach – zamiast tego cytowałyby samą publikację inicjującą przełom. Wykorzystując dane dotyczące cytowań z 45 milionów artykułów i 3,9 miliona patentów, badacze opracowali probierz „dysrupcyjności”, nazwany indeksem CD, w którym wartości oscylowały od –1 dla najmniej przełomowej pracy do 1 dla najbardziej przełomowej. Po przyjęciu tych założeń i współczynników okazało się, iż średnia wartość indeksu CD spadła (1) o ponad dziewięćdziesiąt proc. w latach 1945–2010 w przypadku prac badawczych i o ponad 78 proc. w latach 1980–2010 w przypadku patentów.

Autorzy przeanalizowali również czasowniki używane najczęściej w artykułach, stwierdzając, że gdy w badaniach z lat pięćdziesiątych częściej używano

słów kojarzących się z tworzeniem lub odkrywaniem, takich jak „wyprodukować” lub „ustalić”, to w badaniach prowadzonych w latach po 2010 r. częściej odnoszono się do przyrostowego postępu, używając określeń takich jak „poprawić” lub „ulepszyć”.

Cytowani przez „Nature” eksperci, choć zgadzają się, że badania te współgrają z ogólnym poczuciem spadku innowacyjności w nauce i technice, ostrzegają przed demonizowaniem tych wyników. „Idealem jest równowaga badań przyrostowych i przełomowych”, zauważa John Walsh z Georgia Institute of Technology w Atlancie. „W świecie, w którym jesteśmy bombardowani komunikatami o ważnych odkryciach, większa liczba replikacji i reprodukcji może nie być taką złą rzeczą”. W jego ocenie te widoczne zmiany w poziomie innowacyjności publikacji mogą w pewnym stopniu wynikać ze zmian zachodzących w samym świecie nauki. Obecnie na świecie jest o wiele więcej badaczy niż w latach 40. XX wieku, co tworzy bardziej konkurencyjne otoczenie, podnosząc stawkę oczekiwań co do badań i patentów. Normą dziś są też duże zespoły, w których z natury istnieje większa skłonność do wyników o charakterze przyrostowym, a nie przełomowym.

Walsh zwraca uwagę na to, że w perspektywie długoterminowej liczba przełomowych prac mogła być w gruncie rzeczy taka sama jak wcześniej, jednak inaczej rozkładają się w czasie, co sprawiać może wrażenie spadku innowacyjności. Jego zdaniem musi to być jeszcze dokładniej zbadane. ■

Miroslaw Usidus

Spowiedź carycy Katarzyny II

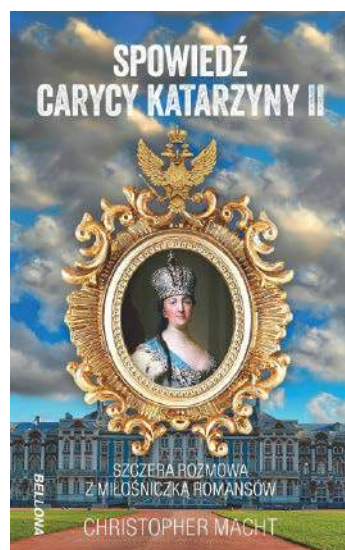
Christopher Macht

Wydawnictwo Bellona, liczba stron: 302, cena: 39,90 zł

Katarzyna II Wielka jest uznawana za jedną z największych postaci, które zasiadły na carskim tronie Rosji! To również kobieta, która kochała seks i nigdy się z tym nie kryła. Gdyby żyła w dzisiejszych czasach, zapewne zrecenzowałaby skandalizującą książkę „50 twarzy Greya” jednym słowem: – Nuda!

Co takiego czyniła wieczorami? Jak wyglądała legendarna Bursztynowa Komnata, w której spędzała czas? I co robiła w swym Pokoju Rozkoszy? Ile jest prawdy w tym, że jeden z jej kochanków przeprowadzał testy mężczyznom, którzy w przyszłości mieli znaleźć się w jej sypialni? I jak do tego wszystkiego odnosił się jej mąż Piotr, człowiek, który również miał zasiąść na rosyjskim tronie? Na te wszystkie pytania władczyni odpowiada w rozmowie ze swoim kochankiem, przyszłym królem Polski Stanisławem Poniatowskim.

Poznajcie zatem historię bodaj największej królowej romanów wszech czasów!





Internet GPT

Sieć 3.0 a nawet 4.0



Announcing

Microsoft Security Copilot

1. Security Copilot Microsoftu

Pod koniec marca 2023 Microsoft zaprezentował narzędzie nazwane „Security Copilot” (1), głosząc, iż wykorzystanie modelu GPT-4 zapoczątkuje „nową erę bezpieczeństwa”, pomagając oszczędzać czas, upraszczać i wyłapywać zagrożenia, które specjalistom ds. bezpieczeństwa mogą przegapić. Nie jest to pierwsze wykorzystanie AI w tej dziedzinie, ale z pewnością pierwsze z potęgą GPT-4 w arsenale.

Bezpieczeństwo sieciowe w erze sztucznej inteligencji

AUTOMATYZUJĄCA SIĘ CYBERWOJNA

Tymczasem rosyjscy hakerzy próbują włamać się do ChatGPT, informował w lutym 2023 r. Check Point, placówka zajmująca się wykrywaniem zagrożeń, w tym np. złośliwego oprogramowania. Docelowo mieliby go wykorzystać do ataków phishingowych. Zdaniem

ośrodka, sondowanie kodu przez rosyjskich hakerów to tylko jeden z wielu przykładów rozszerzającej się liczby podmiotów próbujących uzyskać nielegalny dostęp. Nie jest pewne, czy dla ChatGPT pojawiły się jakiekolwiek podatności typu zero-day, czyli takie, które są nieznane twórcom.

Ataki phishingowe należą do stałego repertuaru działań hakerów i oszustów. Atakujący wysyłają w ich ramach emaila w celu dystrybucji złośliwego oprogramowania, linków phishingowych, przekonują ofiary do przelania pieniędzy. Chatbot AI może zostać użyty do napisania emaili w dowolnych językach, jakich tylko zapragnie atakujący. Cyberprzestępcy nie muszą zatrudniać drogich tłumaczy.

Narzędzia sztucznej inteligencji tworzą więc nowy wektor ataku, jednak z drugiej strony już model stojący za ChatGPT o nazwie GPT-3 sam znajduje więcej luk w zabezpieczeniach niż znane dotychczas oprogramowanie antywirusowe. A kolejna odsłona, GPT-4, jest wielokrotnie lepsza.

Ataki na sterydach

Sztuczna inteligencja już od dłuższego czasu odgrywa ogromną rolę w cyberatakach i należy do arsenału broni hakerskiej. W komunikacie NATO z grudnia 2022 r. nazywa się ją „mieczem obosiecznym” jak też „ogromnym wyzwaniem”. „AI pozwala obrońcom systemów automatycznie skanować sieci i odpierać ataki w takim samym zautomatyzowanym trybie, zamiast ręcznych technik. Ale w drugą stronę to działa oczywiście tak samo”, powiedział dziennikarzom David van Weel, asystent sekretarza generalnego NATO. Na początku grudnia dowódcy wojskowi z ponad trzydziestu krajów (nie wszyscy z nich są członkami NATO) zjechali na poligon cybernetyczny, aby wystawić swoje umiejętności na próbę, jak będą bronić swojego kraju, współpracując z sojusznikami. Stworzono fikcyjne fabuły, a jednym z największych wyzwań w tych cybermanewrach było radzenie sobie z zagrożeniem atakami AI.

Cyberataki, zarówno na infrastrukturę państwową, jak i prywatne firmy, nasiliły się gwałtownie od czasu wybuchu wojny na Ukrainie. Narzędzia oparte na AI mogą być wykorzystywane do lepszego wykrywania i ochrony przed zagrożeniami, ale z drugiej strony cyberprzestępcy mogą wykorzystać technologię do szeroko rozsianych ataków, przed którymi trudniej się bronić, ponieważ jest ich wiele jednocześnie. AI można wykorzystać do próby włamania się do sieci poprzez wykorzystanie danych uwierzytelniających i algorytmów do złamania systemów, mówił cytowany van Weel. Cyberataki AI mogą być wykorzystane nie tylko do wyłączenia infrastruktury,

ale także do wykorzystania informacji, informował Alberto Domingo, dyrektor techniczny ds. cyberprzestrzeni w NATO Allied Command Transformation.

Napędzane przez AI zagrożenia cyberbezpieczeństwa są coraz większym problemem zarówno dla organizacji, jak i osób prywatnych, ponieważ mogą one omijać tradycyjne środki bezpieczeństwa i powodować znaczne szkody.

AI jest obecnie wykorzystywana w różnych formach cyberwojny. Jedną z nich jest tzw. Advanced Persistent Threats (z ang. „zaawansowane trwale zagrożenie”, APT), który ma miejsce, gdy intruz wchodzi do sieci (2) niewykryty i pozostaje w niej przez długi czas w celu kradzieży wrażliwych danych. Sztuczna inteligencja jest w tej technice wykorzystywana do unikania wykrycia i celowania w konkretne organizacje lub osoby. 71 proc. wszystkich detekcji indeksowanych przez CrowdStrike Threat Graph stanowiły włamania przez CrowdStrike Threat Graph stanowiący włamania pozbawione złośliwego oprogramowania.

Inna formą wykorzystującą AI są ataki typu deepfake, korzystające z wygenerowanych, syntetycznych mediów, filmów lub obrazów, w celu podszywania się pod prawdziwe osoby i prowadzenia kampanii oszustw lub dezinformacji. Udowodniono już zagrożenie wynikające z podszywania się pod prezesów firm, innego rodzaju szefów, przywódców politycznych, którzy funkcjonują w sferze publicznej. Politycy pojawiają się w telewizji, wygłaszają przemówienia, w sieci można znaleźć ich nagrania, więc stosunkowo łatwo jest znaleźć nagrania, co pozwala za pomocą już dostępnych narzędzi generować przynajmniej wiarygodnie naśladujący osobę głos a czasem nawet

2. Sieć pełna zagrożeń



Systemy AI w cyberbezpieczeństwie

– przykłady zastosowań

- identyfikacja możliwych zagrożeń
- reagowanie na incydenty cybernetyczne
- systemy bezpieczeństwa domowego
- kamery monitoringu miejskiego i przemysłowego i zapobieganie przestępczości
- wykrywanie oszustw z użyciem kart kredytowych i redukcja ryzyka
- bezpieczeństwo kontroli granicznej
- technologie biometryczne oparte na sztucznej inteligencji
- identyfikacja i obsługa fałszywych opinii klientów
- przeciwdziałanie praniu pieniędzy
- filtry spamu w poczcie elektronicznej
- zabezpieczanie uwierzytelniania
- wykrywanie złośliwego oprogramowania zero-day
- automatyzacja bezpieczeństwa w chmurze
- zwiększanie bezpieczeństwa informacji wrażliwych

obraz. Jeśli pracownik dostaje telefon od szefa firmy, który każe mu coś zrobić, to prawdopodobnie to robi. A ponieważ technologia stojąca za deepfakes wciąż się doskonali, oznacza to, że odróżnienie tego, co jest prawdziwe, od tego, co jest fałszywe, będzie tylko trudniejsze.

Jest też oczywiście złośliwe oprogramowanie napędzane przez AI, które samodzielnie dostosowuje swój sposób działania w odpowiedzi na sytuację i konfigurację systemową, unika wykrycia i utrudnia obronę. Wspominano powyżej phishing, który dzięki algorytmom pozwala atakującym tworzyć bardziej przekonujące emaile i inne wiadomości wyłudzające informacje. Eksperci ds. bezpieczeństwa zauważyli, że generowane przez AI emaile phishingowe mają wyższe wskaźniki otwarcia, co przekłada się na pożądane przez oszustów akcje, np. klikanie w linki lub otwieranie złośliwych załączników. Działa to też w teoretycznie profesjonalnych sieciach społecznościowych. „Jeśli chcesz wygenerować wiarygodne, biznesowe bla-bla w serwisie LinkedIn, aby wyglądało to na prawdziwą osobę z branży, która próbuje nawiązać kontakty, ChatGPT świetnie się do tego nadaje”, zauważyła w jednej z analiz Kelly Shortridge, ekspert ds. cyberbezpieczeństwa w firmie Fastly.

Ataki typu Distributed Denial of Service (DDoS) wykorzystują z kolei AI do identyfikacji i wykorzystania luk w sieci, pozwalając atakującemu na wzmocnienie skali i wpływu ataku. Sztuczna inteligencja może być również wykorzystywana do projektowania złośliwego oprogramowania, które stale się zmienia, aby

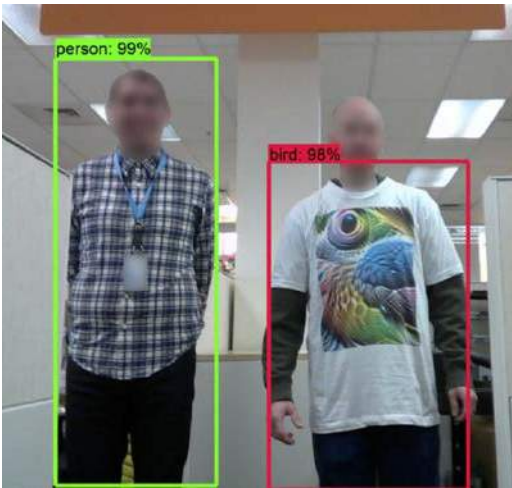
uniknąć wykrycia przez zautomatyzowane narzędzia obronne. Korzystając z AI, złośliwi aktorzy są w stanie przeprowadzić nowe ataki, stworzone poprzez analizę podatności organizacji za pomocą oprogramowania szpiegowskiego.

Programy sztucznej inteligencji są także same w sobie podatne na ataki. Algorytmy uczenia maszynowego mogą zostać zmienione w niepożądany sposób przez manipulacje danymi, zatrutowanie danych, do czego jeszcze wrócimy. Może to mieć szkodliwe a nawet katastrofalne konsekwencje dla tych, którzy polegają na inteligencji systemu. Oczywiście możliwe są też bardziej tradycyjne problemy z cyberbezpieczeństwem, np. luki w kodzie, błędy w oprogramowaniu.

Cyberprzestępcy mogą, jak się przypuszcza, wykorzystać postępy w uczeniu maszynowym do opracowania samoprogramującego się inteligentnego złośliwego oprogramowania, które mogłoby aktualizować się poprzez automatyczne reagowanie na formy cyberobrony, które napotyka, aby mieć największą szansę na skuteczność. „Zobaczymy ataki złośliwego oprogramowania, operacje ransomware i kampanie phishingowe prowadzone całkowicie automatycznie na bazie uczenia maszynowego. Jeszcze tego nie zrobiono, ale nie byłoby to wcale trudne do zrobienia”, ostrzega Mikko Hyppönen z WithSecure. „Takie programy mogą zmienić kod, uczynić go bardziej trudniejszym do zrozumienia, sprawić, by za każdym razem był inny, może próbować tworzyć niewykrywalne wersje. Wszystko to jest technicznie wykonalne, po prostu jeszcze tego nie widzieliśmy – i myślę, że zobaczymy”.

Inteligencja, która daje się zmanipulować

Oszukanie AI, aby myślała, że kot jest psem lub, jak zademonstrowali badacze, panda jest gibonem, lub człowiek – ptakiem (3), jest stosunkowo niewielkim problemem. Kilka lat temu niektórzy badacze pokazali, w jaki sposób mogą stworzyć trójwymiarowe obiekty przeciwne, które mogą oszukać sieć neuronową, by myślała, że żółw jest karabinem. Nie potrzeba jednak wiele wyobraźni, aby wymyślić scenariusze ataku na systemy zarządzane przez AI, w których małe pomyłki lub destrukcyjne działanie aplikowane wprost mogą prowadzić do niebezpiecznych a nawet tragicznych konsekwencji, np. sytuacji, w której „inteligentnie” skanujący otoczenie samochód myli piesze z pojazdem, lub jego system zostaje oszukany tak, by zlekceważyć znaki stopu. Co, gdyby oszukany tak samo został skaner medyczny napędzany przez algorytmu, wystawiając błędną diagnozę? Co by



3. Przykład systemu AI mylącego człowieka z ptakiem

było, gdyby zautomatyzowany system bezpieczeństwa obiektu lub terenu został zmanipulowany tak, aby wpuścić niewłaściwą osobę, a może nawet nie rozpoznać, że w ogóle była tam jakaś osoba?

Ponieważ wszyscy polegamy na automatycznych systemach podejmujących decyzje o ogromnych potencjalnych konsekwencjach, musimy mieć pewność, że systemy AI nie dadzą się nabrać na podejmowanie złych lub nawet niebezpiecznych decyzji. Blokada miasta lub przerwanie świadczenia podstawowych usług to tylko niektóre z najbardziej widocznych problemów, które mogą wynikać z awarii systemów napędzanych przez AI. A „włamanie” to tylko jedna z możliwości.



4. Symbol projektu GARD agencji DARPA

Okazuje się bowiem, że AI poddaje się różnym, mniej oczywistym manipulacjom. Pierwsi użytkownicy Bing Chat napędzanego przez ChatGPT odkryli, że stosunkowo łatwo jest użyć tzw. ataku „prompt injection”, by skłonić chatbota do ujawnienia zasad rządzących jego zachowaniem i jego nazwy kodowej (Sydney). Okazało się, że bot wyklócał się użytkownikami, wpadając w dziwaczne i niepokojące rozmowy. Nic więc dziwnego, że Microsoft zmodyfikował ustawienia czatu Bing, aby wykluczyć takie sytuacje

Z myślą o takich i innych problemach DARPA uruchomiło projekt o nazwie Guaranteeing AI Robustness Against Deception (GARD), którego celem jest uzyskanie pewności, że AI i algorytmy są opracowywane w sposób chroniący je przed próbami manipulacji, podstępów lub innej formy ataku. W dalszej kolejności chodzi o to, by algorytmy AI były broniące przed atakami także wtedy, gdy technologia stanie się bardziej zaawansowana i bardziej dostępna. DARPA współpracuje z wieloma firmami technologicznymi, w tym z IBM i Google. Jednym z kluczowych elementów GARD jest

5. Zatrucie danych



Armory, wirtualna platforma dostępna na GitHubie, która służy jako poligon doświadczalny dla naukowców potrzebujących powtarzalnych, skalowalnych i solidnych ocen obrony przed atakami. Firma IBM opracowała wcześniej Adversarial Robustness Toolbox (ART), zestaw narzędzi dla programistów i naukowców do obrony ich modeli uczenia maszynowego i aplikacji przed zagrożeniami, który jest również dostępny do pobrania z GitHub i stał się głównym elementem projektu GARD.

Zatrute dane

„Zagrożeniem, które najczęściej pojawia się w literaturze akademickiej, jest bezpośrednia modyfikacja obrazu lub wideo. Panda, która wygląda jak panda, ale została sklasyfikowana jako autobus szkolny – tego typu rzeczy”, komentuje z serwisie ZDNet David Slater z Two Six Technologies, firmy, która i jest zaangażowana w projekt GARD.

Ale ta bezpośrednia modyfikacja to tylko niezbyt jeszcze rozpowszechniona forma innego, być może, większego zagrożenia cybernetycznego, jakim jest zatrucie danych (5), polegające na tym, że dane treningowe użyte do stworzenia modelu AI są zmieniane przez atakujących, aby wpływać na decyzje i działania, które podejmuje AI. „Zatrutowanie danych może być jednym z najpotężniejszych zagrożeń i czymś, czym powinniśmy się znacznie bardziej przejmować. Nie potrzeba wyrafinowanego przeciwnika, aby to pociągnąć. Jeśli uda się zatruczyć te modele, a następnie zostaną one szeroko wykorzystane w dalszej części łańcucha użytkowania, wpływ jest zwielokrotniony – a zatrucie jest bardzo trudne do wykrycia i poradzenia sobie z nim, gdy już znajdzie się w modelu”, mówi Slater.

Jeśli algorytm jest trenowany w zamkniętym środowisku, powinien być – w teorii – dość dobrze chroniony przed zatruciem, chyba że hakerzy zdołają się do niego włamać. Ale większy problem pojawia się, gdy AI jest szkolona na zbiorze danych z domeny publicznej, zwłaszcza jeśli jest powszechnie wiadomo, że tak jest. Jednym z przykładów zatrucia jest bot sztucznej inteligencji Microsoftu, Tay. Microsoft wysłał Tay na Twittera, aby weszła w interakcję i uczyła się od ludzi, dzięki czemu mogła nauczyć się używać języka naturalnego i mówić jak ludzie. Ale w ciągu zaledwie kilku godzin, ludzie „zatruli” Tay, który zaczął wygłaszać kontrowersyjne opinie. To może jedynie ciekawostka, jednak gdy weźmiemy pod uwagę sytuację, w której algorytm uczy się ważnych informacji, takich jak dane medyczne, i zostanie zatruty, skutek może być katastrofalny.

Zero zaufania

AI oprócz ekosystemu ataku tworzy również ekosystem obrony. Są wyspecjalizowane firmy, które specjalizują się w wykorzystaniu AI w walce z cyberzagrożeniami, np. Tessian, która wykorzystuje możliwości AI, aby zapobiegać włamaniom do systemów poczty elektronicznej, atakom typu phishing oraz utracie danych w wyniku szkodliwych wiadomości e-mail. Przygotowuje zaawansowane „inteligentne” filtry poczty elektronicznej, które pomagają eliminować niechciane i podejrzane działania w przechodzących i wychodzących wiadomościach e-mail. Oprogramowanie firmy zawiera pulpit nawigacyjny działający w czasie rzeczywistym do monitorowania stanu infrastruktury sieciowej. Inna firma, VMwareInc, wykorzystuje technologie uczenia maszynowego do identyfikacji podejrzanej aktywności w systemach m.in. chmur obliczeniowych, zanim się rozprzestrzeni. Zatem są to niejako systemy wczesnego ostrzegania oparte na AI.

Monitorowanie i analizowanie sieci o dużej skali jest dla człowieka czasochłonne lub dość skomplikowane. Dzięki AI, analizowanie danych pochodzących z różnych źródeł staje się bardziej wydajne i szybsze, co pozwalać ma na szybkie wykrywanie podatności i zagrożeń, zanim nastąpi atak. Systemy wykrywania włamań (IDS) oparte na sztucznej inteligencji wykrywają każde nietypowe lub złośliwe działanie w strumieniu normalnego ruchu w sieci. Platformy operacji IT (AIOps) wykorzystują analitykę big data i uczenie maszynowe do wykrywania problemów poprzez analizę dużych ilości danych i przewidywanie w celu zapobiegania przyszłym problemom.

W przeszłości rozwiązania bezpieczeństwa były w przeważającej mierze reaktywne. Zidentyfikowana nowa próbka złośliwego oprogramowania była analizowana i dodawana do list złośliwego software'u firm zajmujących się bezpieczeństwem cybernetycznym. Branża nadal stosuje to podejście, ale działa bardziej proaktywnie, zwłaszcza w odniesieniu do zagrożeń związanych z inżynierią społeczną (phishing).

Zagrożenia bezpieczeństwa, złośliwe oprogramowanie i taktyka atakujących zwykle jest ewolucją wcześniejszych ataków i złośliwego oprogramowania. Prawdziwie nowych zagrożeń pojawia się każdego roku stosunkowo niewiele. Większość cyberprzestępców to nie tyle twórcy, ile użytkownicy pakietów malware-as-a-service lub ponowni użytkownicy istniejących złośliwych kodów. W niedawnym badaniu ewolucji botnetu Sysrv, widać było, że większość nowych szczepów złośliwego oprogramowania była



6. Sztuczna inteligencja w służbie cyberbezpieczeństwa

kombinacjami i rekombinacjami innych istniejących złośliwych kodów. Być może jednym z najcenniejszych osiągnięć w zakresie cyberbezpieczeństwa AI jest możliwość ostrzegania użytkowników przed wejściem na podejrzane strony internetowe, w tym strony phishingowe. Ponieważ ataki socjotechniczne zwykle powodują największe szkody oraz utratę prywatności i finansów konsumentów, wykorzystanie AI do zapobiegania nowatorskim atakom, zanim pojawią się one w branżowych bazach danych, jest niezwykle ważne.

Precyzyjna identyfikacja, analiza i ocena zagrożeń oraz rekomendacje dla odkrytych zagrożeń dzięki AI pozwalają rozwijać zautomatyzowane modele bezpieczeństwa. IBM podaje, że przyjęcie AI i automatyzacji w bezpieczeństwie pozwala zaoszczędzić ponad 14 tygodni na czasie wykrywania i reagowania na zagrożenia, pomagając zmniejszyć ogólne koszty. W 2022 roku średni koszt naruszenia danych wynosił globalnie 4,35 mln dolarów. Jednak, jak obliczył IBM, firmy z w pełni wdrożonym programem AI i automatyzacją zabezpieczeń zaoszczędziły 3,05 mln dolarów dzięki szybkiemu wykrywaniu i czasowi reakcji na zagrożenia.

AI monitoruje dane transakcyjne w sieci organizacji i chroni je przed możliwymi zagrożeniami. AI wykorzystuje analizę behawioralną do identyfikacji złośliwych działań w czasie rzeczywistym i oferuje natychmiastowe reakcje na zagrożenia. AI wyznacza punkty odniesienia dla ruchu sieciowego w firmie i wykorzystuje go do oceny i ochrony sieci. AI uczy się w końcu i z czasem stale poprawia swoje zabezpieczenia. AI lepiej wykorzystuje swój potencjał,

gdy jest zintegrowana z szerszymi ramami bezpieczeństwa określanymi jako „zero zaufania” (ang. „zero trust”), zaprojektowanymi tak, aby traktować każdą tożsamość użytkownika jako nową granicę bezpieczeństwa. Kluczowym elementem tej polityki jest widoczność i monitorowanie w czasie rzeczywistym całej aktywności w sieci. Algorytmy AI i uczenie maszynowe pozwalają zdefiniować role lub profile bezpieczeństwa dla każdego użytkownika w czasie rzeczywistym na podstawie jego zachowania i wzorców. Analizując kilka zmiennych, w tym między innymi gdzie i kiedy użytkownicy próbują się zalogować, typ urządzenia i konfigurację, systemy te mogą wykrywać anomalie i identyfikować potencjalne zagrożenia w czasie rzeczywistym.

Pomimo swoich zalet w dziedzinie cyberbezpieczeństwa, AI ma również pewne ograniczenia. Rozwiązania tego typu są drogie w implementacji, co sprawia, że są słabo dostępne dla mniejszych podmiotów. Ponadto AI wymaga od firm przyjęcia określonych przypadków użycia, aby funkcjonować dokładnie, co może być dość trudnym wymogiem.

Przewiduje się, że globalny rynek cyberbezpieczeństwa opartego na AI wzrośnie z 17,4 mld USD w 2022 r. do 102,78 mld USD w 2023 r. Firma badawcza Precedence Research, od której pochodzą te szacunki, przewiduje, że najszybciej rozwijające się obszary AI dotyczyć będą walki z oszustwami, identyfikacji emaili phishingowych i złośliwych linków oraz identyfikację nadużyć w zakresie poświadczonych dostępu uprzywilejowanego. ■

Miroslaw Usidus



1. Sieć 5G

Rozczarowanie użytkowników, którym obiecano, że mobilna sieć będzie „śmigać” jak szalona. Prawie zupełny brak wszystkich tych niezwykłych biznesowych zastosowań i usług, które zapowiadano. Ograniczony wciąż zasięg i wciąż marginalne występowania „prawdziwej” 5G, czyli stand-alone opartej na oddzielnej, innej niż LTE, infrastrukturze, i na zakresie fal milimetrowych.

Co się dzieje z 5G?

W SIECI ROZCZAROWANIA

Dochodzą do tego wszystkiego kwestie nieco bardziej polityczne. W 2020 roku np. brytyjskie sieci komórkowe otrzymały nakaz usunięcia zestawu wykonanego przez chińską firmę Huawei z całej infrastruktury 5G, po tym jak uznano go za zagrożenie dla bezpieczeństwa narodowego. Wcześniej sankcje nałożyły stany Zjednoczone.

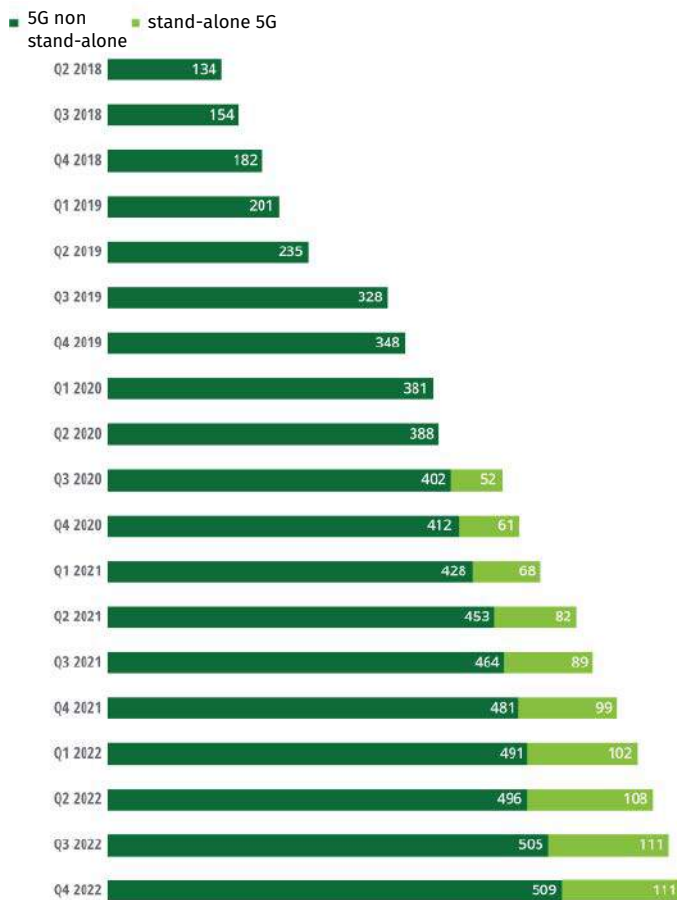
Na razie bez „zabójczej aplikacji”

Szum wokół 5G (1) można dziś uznać za typowy przykład rozdętej bańki nadmiernych oczekiwań.

Balon został nadmuchany ponad miarę głównie przez marketingowców. W czasie pandemii do szumu marketingowego dołączyła szerzona w mediach internetowych dezinformacja na temat rzekomego zabójczego wpływu nadajników i fal w wykorzystywanej przez 5G części widma na zdrowie, co posuwało się niekiedy do sugerowania, że 5G aktywuje koronawirusa.

W różowych wizjach marketingowców wielkich firm sieć 5G miała wyeliminować potrzebę inwestowania w drogie sieci światłowodowe, połączyć autonomiczne samochody, także w ruchu, pozwolić lekarzom wykonywać operacje zdalnie w dowolnym punkcie na Ziemi, wyzwolić eksplozję inteligentnych fabryk, które m.in. miały reindustrializować kraje zachodnie.

Piąta generacja komunikacji mobilnej, znana również jako 5G lub 5G NR (New Radio), jest następczynią 4G LTE, a wcześniej 3G i 2G. Do najważniejszych obietnic 5G należy zwielokrotniona prędkość transferu danych, a także znacznie zmniejszona latencja. Chociaż obecnie 5G wykorzystuje przeważnie te same technologie co 4G LTE i może korzystać z tych samych pasm częstotliwości, istnieją pewne istotne



Źródło: Global Mobile Suppliers Association (GSA), 5G standalone June 2022 summary, dostęp 31 sierpnia 2022 r.

2. Liczba operatorów na świecie inwestujących w rozwój sieci 5G

różnice. Smartfon 5G może teraz doświadczyć wyższej przepustowości niż smartfon 4G przy użyciu tej samej anteny i częstotliwości. Dodatkowo, 5G używa dwóch nowych pasm częstotliwości – 700 MHz (694–790 MHz) i 3,6 GHz (3,4–3,8 GHz). Używać ma także tzw. fal milimetrowych, w paśmie 26 GHz obejmuje zakres od 24,25 GHz do 27,5 GHz. Te mają niewielki zasięg (do 500 metrów), ale za to wielką przepustowość, dlatego nazywa się je „radiowymi światłowodami”.

Wdrażanie 5G rozpoczęło się już w wielu państwach, w tym w Stanach Zjednoczonych, Korei Południowej, Chinach i Wielkiej Brytanii. Niektóre rozpoczęły nawet testy z sieciami fal milimetrowych. Według brytyjskiego badania przeprowadzonego przez firmę Uswitch, rozwój 5G nie zadowala jednak konsumentów. Tylko połowa użytkowników twierdzi, że zauważyła szybsze prędkości lub poprawę stabilności połączenia od czasu uaktualnienia do piątej generacji sieci komórkowych. A jeden na sześciu użytkowników uważa,

że technologia nie spełnia oczekiwań.

Z drugiej strony 5G jest najszybszym wdrożeniem z dotychczas wprowadzanych generacji sieci komórkowych. Do końca 2021 roku w Wielkiej Brytanii, która może być przykładem rozwiniętego kraju zachodniego, 25 proc. populacji zostało objęte zasięgiem 5G. Ten wynik osiągnięto sześć miesięcy krócej niż w przypadku 4G. Sieć nadal jednak nie jest równomiernie rozpowszechniona. Większość wień znajduje się na terenach miejskich. W Wielkiej Brytanii tylko 17 proc. mieszkańców wsi – trzy razy mniej niż mieszkańców miast – twierdzi, że udało im się połączyć z sygnałem 5G.

Większość użytkowników 5G wciąż robi mniej więcej to samo co w 4G, czyli ogląda filmy, przegląda strony internetowe i korzysta z mediów społecznościowych. Trudno mówić o nowych aplikacjach, obiecanych przez 5G, wirtualnej rzeczywistości i grach w chmurze. Dostawcy tych ostatnich gier zresztą za-

chęcają do grania „w chmurach” tylko na stałym łączu, jeśli połączenia ma być bezproblemowe i stabilne.

Sieć 4G dała usługi Ubera, streaming Netflixu. 5G wciąż nie ma swojego „zabójczego” przypadku użycia. W USA 5G rzeczywiście przyniósł pewne zainteresowanie domowym wysokoprzepustowym internetem bezprzewodowym. Verizon i T-Mobile oferują w wybranych obszarach tzw. internet domowy 5G. I to byłoby na razie tyle.

5G non-stand-alone (NSA) było powszechnie uważane za tymczasową prowizorkę dla globalnej branży bezprzewodowej. Wprowadzony w pierwszej partii specyfikacji 5G przez 3GPP, NSA wykorzystuje nową radiową sieć dostępową 5G z istniejącym rdzeniem sieci 4G LTE. Ponieważ było to szybsze i łatwiejsze do uruchomienia, większość operatorów sieci na świecie użyła tej wersji technologii, aby wejść na rynek.

Niektórzy operatorzy pracują nad przejściem z 5G NSA na SA (2). W USA, T-Mobile uruchomił SA 5G

Uruchomienia komercyjnych usług 5G 'standalone' eMMB (operatorzy)

Region	2020	2021	2022
Ameryka Płn.	T-Mobile USA	Rogers – Kanada	Bell Mobility – Kanada AT&T Wireless – USA DISH Wireless – USA Verizon – USA
EMEA (Europa, Bliski Wschód, Afryka)		Telefónica – Niemcy Vodafone – Niemcy stc – Kuwejt Zain – Arabia Saudyjska Vodafone – W. Brytania	Drei Austria stc – Bahrajn Telekom – Niemcy
APAC (Azja-Pacyfik)	China Mobile China Mobile – HK China Telecom AIS – Tajlandia TOT – Tajlandia True – Tajlandia	TPG Telecom – Australia NTT DOCOMO – Japonia Softbank – Japonia Smart – Filipiny M1 – Singapur SingTel – Singapur StarHub – Singapur KT – Korea Płd. Taiwan Mobile	Optus – Australia China Broadnet Reliance Jio – Indie KDDI – Japonia China Telecom – Makau Globe Telecom – Filipiny
CALA (Ameryka Łacińska, Karaiby)			Claro – Brazylia TIM – Brazylia Vivo – Brazylia

Źródło: Dell'Oro Group

3. Rozwój sieci 5G o najwyższej wydajności

już w 2020 roku. Inny amerykański operator, Verizon, ogłosił, że jego standalone 5G „umożliwi dynamiczne przydzielanie zasobów, określane jako plasterkowanie sieci”, zautomatyzowane zmiany konfiguracji sieci, w tym możliwość skalowania w górę lub w dół pojemności funkcji sieciowych w celu zapewnienia „poziomów usług i zasobów sieciowych dla każdego przypadku użycia”. Miało to pozwolić firmom rozpocząć eksplorację zaawansowanych przypadków użycia 5G, takich jak autonomiczne pojazdy, robotyka precyzyjna, usługi inspekcji i dostaw z wykorzystaniem dronów oraz sterowane przez sztuczną inteligencję systemy bezpieczeństwa, kontroli jakości i konserwacji predykccyjnej.

Biznes wciąż w to tak naprawdę nie wchodzi

Jeśli ktoś miałby cofnąć się w czasie, powiedzmy o pięć lat, średnia prędkość transferu danych doświadczana w telefonie komórkowym jest prawdopodobnie o 50 proc. wyższa dzisiaj. Nie jest to jednak zasługa tylko wdrażania 5G. Po pierwsze wyższe prędkości mogą być po prostu lepszym niż poprzednio połączeniem 4G LTE. Poprawa jakości połączeń może być w rzeczywistości efektem wielu czynników, ulepszeń,

i unowocześnienia rozwiązań. Zdaniem trzeźwo myślących ekspertów, nadal jesteśmy we wczesnej fazie ery 5G, która ma przed sobą dziesięć lat rozwoju. Sieć 4G dziś też jest zupełnie inna niż ta na początku wdrażania czwartej generacji w około 2010 roku.

Pewne jest, że słabo wygląda, i to na całym świecie, wdrażanie sieci 5G w zakresie fal milimetrowych. Nawet w krajach najbardziej zaawansowanych w upowszechnieniu sieci 5G, pasmo to ma niewielkie rozpowszechnienie, a infrastruktura obsługująca to pasmo to wciąż wyjątki, a nie masowe wdrożenia (3).

Konsumenci, nie mając żadnego konkretnego powodu, nie palą się do dodatkowych opłat za korzystanie z usług 5G. Zresztą operatorzy i inne firmy mobilnego biznesu chyba nawet tego nie oczekują, twierdząc zwykle, że najlepsze długoterminowe możliwości uzyskania przychodów z 5G będą pochodzić z rynku przedsiębiorstw. Tutaj dokonywane są prawdziwe inwestycje – nie tylko po stronie sieci, ale także w mobilne edge computing, a także budowanie infrastruktury sprzedaży i rozwiązań dla przedsiębiorstw. Jednak budowanie nowego rynku bezprzewodowego dla przedsiębiorstw to długoterminowa gra. Ponadto ta pula będzie dzielona przez większą i inną grupę graczy niż ta, którą widzieliśmy

w przeszłości w branży bezprzewodowej. Nie widać perspektywy znaczących przychodów z rynku przedsiębiorstw przez co najmniej dwa do trzech lat. Funkcje takie jak niskie opóźnienia, plasterkowanie sieci i niektóre z funkcji Massive IoT dopiero teraz zaczynają być dostępne.

Cyfrowa produkcja i technologia „cyfrowych bliźniaków”, które były przedstawiane jako główne aplikacje korporacyjne w sieci 5G, na razie w Europie i Ameryce Północnej mają niewiele wspólnego z rozwojem sieci piątej generacji. Zdaniem ekspertów wynika to też z niezdolności firm do nawiązania odpowiednich relacji z operatorami telekomunikacyjnymi.

Zgiełk wokół 5G wygast, zatem czas na... 6G

Jeśli wrócimy do tego, co rzeczywiście zawierała specyfikacja sieci 5G, która została przyjęta przez międzynarodowe organy normalizacyjne, to widać, że obiecywane w marketingowym zgiełku cuda wcale z niej nie wynikały. Specyfikacje zakładały osiągnięcie pewnych parametrów – po pierwsze uzyskanie prędkości komórkowych powyżej 100 Mb/s, umożliwienie obsługi większej liczby użytkowników jednocześnie w danej lokalizacji komórkowej, umożliwienie telefonowi komórkowemu korzystania z dwóch różnych pasm w tym samym czasie oraz umożliwienie użytkownikowi połączenia z więcej niż jedną lokalizacją w sieci komórkowej, jeśli była taka potrzeba. Podstawowym celem specyfikacji 5G było wyeliminowanie zatorów w lokalizacjach, gdzie występuje wiele osób jednocześnie korzystających z sieci komórkowej. Dokładna analiza specyfikacji 5G wskazuje, że nie chodziło o rewolucję, ale raczej o naturalną ewolucję sieci komórkowej w celu lepszego dostosowania techniki komórkowej do świata, w którym każdy ma telefon komórkowy. Chociaż 5G zapewnia szybsze, stabilniejsze i mniej opóźnione połączenia, oczywiście



4. Skok do 6G

nie jest to, jak obiecywali giganci telekomunikacyjni, „radykalna zmiana sposobu życia i pracy” dzięki 5G.

Największym beneficjentem szumu marketingowego wydają się po czasie producenci telefonów, którzy przekonali klientów, że muszą mieć telefony 5G, choć nie było i nie ma jasnego po temu powodu. Kupujący sprzęt są ogólnie zadowoleni, że prędkości wzrosły. Jednak badania wykazują, że tylko znikomy odsetek jest skłonny zapłacić więcej za szybsze prędkości komórkowe.

Gdy okazało się, że temat 5G nieco się wypalił, pojawiły się sygnały nowego szumu, tym razem wokół... 6G (4). Tym razem ma to być „prawdziwa rewolucja”. Prezes Nokii Pekka Lundmark mówił już w maju 2022 r., że wierzy, iż sieć 6G będzie dostępna do 2030 roku, choć to już nie smartfon byłby jej głównym interfejsem. Cytowany przez serwis BGR.in Lundmark snuje wizję urządzeń rozszerzonej rzeczywistości (AR) i wirtualnej rzeczywistości (VR) w przyszłej sieci 6G. Na Światowym Forum Ekonomicznym w Davos, Lundmark mówił: „Wiele z tych rzeczy będzie wbudowanych bezpośrednio w nasze ciała”. Wobec nieco rozczerwawiającej rzeczywistości 5G branża ucieka jak widać w „rozszerzoną” i „wirtualną” rzeczywistość 6G. ■

Mirosław Usidus

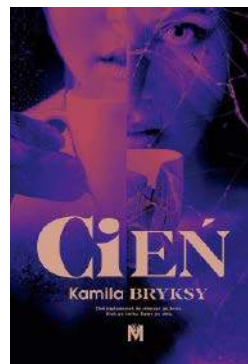
Cień

Kamila Brykсы

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

Ktoś zaplanował, że zniszczy jej życie. Krok po kroku. Dzień po dniu

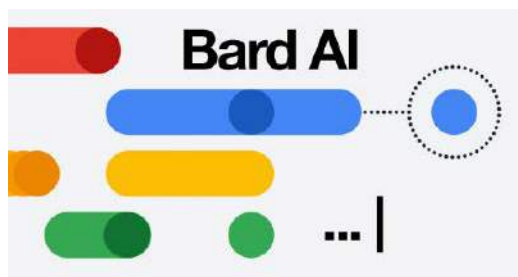
Malwina kończy długoletni związek i próbuje poradzić sobie z problemami finansowymi. Kiedy zaczyna układać swoje życie na nowo, niespodziewanie ktoś postanawia doprowadzić ją do obłędu. Jedno jest pewne: kimkolwiek jest, jej prześladowca nie cofnie się przed niczym.





Najnowsza internetowa wojna na górze

PREMIUM CZY KLASA EKONOMICZNA?



2. Logotyp chatbota Bard

jednak jedynie dla wyselekcjonowanej grupy użytkowników testowych. Pierwsze publikowane w serwisach internetowych wyniki testów „odpowiedzi Google’a na ChatGPT” sugerowały, że potrafi mniej niż usługi OpenAI i Microsoftu.

Microsoft jest obecnie związany biznesowo z OpenAI, która tworzy kolejne wersje modelu GPT (w marcu z dużym hałasem wprowadzono GPT-4). Zainwestował w nią miliard dolarów w 2019 roku, a w ostatnich miesiącach kolejne dziesięć miliardów dolarów, i ma ambitne plany rozwoju kolejnych

produktów wykorzystujących modele AI, nie tylko wyszukiwarkę, ale także oprogramowanie Windows. W marcu 2023 r. dodał do wyszukiwarki Bing kolejny po usłudze chatu model tekstowo-obrazowy AI, zasilany przez DALL-E 2 również autorstwa OpenAI. Microsoft poczynił już wcześniej znaczące inwestycje w sztuczną inteligencję i jest głównym graczem w ekosystemie AI poprzez swoją platformę chmurową Azure. Opracował również szereg produktów AI, takich jak Cognitive Services, które zapewniają wstępnie zbudowane modele AI dla wizji, mowy, języka i podejmowania decyzji.

Warto przypomnieć, że zręby konwersacyjnych interfejsów wspomaganych przez sztuczną inteligencję w wyszukiwaniu informacji i korzystaniu z aplikacji użytkowych zostały nakreślone przez szefa Microsoftu, Satyę Nadellę jeszcze w 2016 r., w wywiadzie dla serwisu „The Verge”. Zatem zapowiedź wprowadzenia zasilanych przez model AI okien dialogowych w Outlooku, Wordzie, Teamsa czy Excelem, to nie tyle nowe pomysły podjęte pod wpływem osiągnięć OpenAI i sukcesu ChatGPT, lecz kontynuacja dłuższej strategii.

Google ma tylko dwa lata?

A co z Google'em? Twórca poczty Gmail, Paul Buchheit, mówił w ostatnim czasie w mediach, że Google dzieli zaledwie rok lub dwa lata od „całkowitej zapaści”. Powodem ma być rewolucja chatbotów AI, które „zniszczą Google w taki sam sposób, jak wyszukiwarka internetowa zabiła ogłoszenia drobne”. Większość przychodów Google'a pochodzi ze sprzedaży reklam w jego własnej wyszukiwarce. W 2021 roku, około 80 proc. z 257,64 mld dolarów przychodów Google pochodziło z reklam. „AI wyeliminuje mechanizm Search Engine Result Page (SERP), czyli to, na czym zarabia Google najwięcej. Nawet jeśli utrzymają się w nurcie rozwoju AI, to nie mogą jej w pełni wdrożyć bez zniszczenia najbardziej wartościowej części własnego biznesu”, komentował Buchheit. Jego zdaniem, postępy w tej dziedzinie z pewnością zmienią sposób, w jaki ludzie uzyskują dostęp do informacji.

Już w marcowym wydaniu MT pisaliśmy o „czerwonym alarmie” w Google z powodu premiery ChatGPT i wszystkiego, co nastąpiło w konsekwencji tego wydarzenia. „The New York Times” donosił w styczniu 2023 r., że firma wezwała współzałożycieli Google'a Larry'ego Page'a i Sergeya Brina, by pomogli w rozwiązaniu problemu. W tym samym czasie spółka nadrzędna Google'a, Alphabet, zmniejszyła zatrudnienie (o 12 tysięcy pracowników), chociaż to akurat ostatnio typowe dla wielkich. Dziesięć tysięcy ludzi zwolnił Microsoft, Facebook też masowo zwalnia.

Nie można jednak powiedzieć, że Google zaspęło kwestię AI. Od ponad dekady rozwija przecież pionierskie projekty w dziedzinie sztucznej inteligencji i to jego zespół opatentował najprężniejszą obecnie na tym obszarze technikę generatorów opartych na wielkich modelach językowych (LLM). Stworzony przez Google'a model LaMDA (Language Model for Dialogue Applications) miał swoją premierę i sporo medialnego rozgłosu już w pierwszej połowie 2022 r. Już w 2012 roku Google zatrudniło Raya Kurzweila, informatyka i wizjonera „osobliwości” technologicznej, do pracy nad modelami przetwarzania języka. Rok później Google kupiło brytyjską firmę DeepMind, która miała (i wciąż ma) na celu stworzenie sztucznej inteligencji ogólnej. Jednak uczeni zatrudnieni w Google i eksperci w dziedzinie techniki powstrzymali projekty, które dziś określilibyśmy mianem prekursorów ChatGPT, ze względu na obawy etyczne wokół masowej inwigilacji. Google zobowiązało się do samoograniczeń w dziedzinie AI. W 2018 roku Google wycofało projekt dotyczący sztucznej inteligencji w projektach wojskowych na skutek buntu pracowników.

W zeszłym roku Deepmind zademonstrował GATO, wykazujący się niezwykłą płynnością w wielu zadaniach. Ta sama sieć może grać na Atari, podpisywać obrazy, czatować, układać klocki za pomocą ramienia robotycznego i wiele więcej, decydując na podstawie kontekstu. Na konto Google idą też postępy w modelowaniu białek, raka, matematyki i wiele innych. AlphaFold stworzony przez DeepMind dokładnie przewiduje modele 3D struktur białkowych (3) i przyspiesza badania w prawie każdej dziedzinie biologii.

Google ma listę programów AI, które planuje zaoferować twórcom oprogramowania i innym firmom, w tym technologii tworzenia obrazów, które mogłyby zwiększyć przychody do działu Cloud. Są też narzędzia, które pomogą innym firmom tworzyć własne

3. Struktury protein



prototypy AI w przeglądarkach internetowych, zwane MakerSuite, które będą miały dwie wersje „Pro”. W maju 2023 r. Google ma, według pogłosek, zapowiedzieć również narzędzie ułatwiające budowanie aplikacji na smartfony z systemem Android, o nazwie Colab + Android Studio, które będzie generować, uzupełniać i naprawiać kod. We wstępnych zapowiedziach mowa też o innym narzędziu do generowania i uzupełniania kodu, o nazwie PaLM-Coder 2. Warto dodać, iż specjaliści Google'a stoją za techniką AutoML-Zero, która służy do samodzielnego, bez ludzkiej interwencji, odkrywania algorytmów uczenia maszynowego. Opisuje się ją jako podłoże dla samostnej ewolucji maszyn, które same opracowują nowe algorytmy pozwalające im się rozwijać i wzbogacać o nowe możliwości.

Wyszukiwarkowy dominator buduje też od pewnego czasu AI tłumaczącą na tysiąc języków. W tej chwili model obsługuje ponad sto różnych języków. Google ogłosiło ów „Universal Speech Model” (USM) w listopadzie 2022. USM ma dwa miliardy parametrów i jest przeszkolony na 28 miliardach zdań pochodzących z ponad trzystu języków. Jego funkcje mają pozwalać na automatyczne rozpoznawanie mowy lub automatyczne generowanie napisów w filmach. Jeśli system będzie w stanie rejestrować i tłumaczyć w locie mowę w czasie rzeczywistym, to będzie coś, czego oczekuje się od lat i dekad. Co zgryźliwsi dodają – byleby tylko AI potrafiła też rozpoznać, czego nie tłumaczyć, bo ostatnią rzeczą, jakiej potrzebujemy, jest zalew tłumaczeń rozmów w tle i wtłaczanie tego wszystkiego do naszych uszu.

Uważa się, że pełna odpowiedź na pytanie, jak Google zamierza postąpić w obliczu rewolucji AI, powinna być znana po dorocznej konferencji deweloperów tej firmy, I/O 2023, w maju tego roku.

Na razie, pomimo wszystkich wymienionych wyżej osiągnięć i projektów, utrzymuje się wrażenie, że Google jest o krok do tyłu. W marcu OpenAI ogłosiła uruchomienie API GPT dla programistów oraz Whisper, modelu transkrypcji mowy, dodając również obsługę wtyczek do ChatGPT, co pozwoli AI na dostęp do internetu, bieżących i aktualnych informacji (zamiast, jak do tej pory, bazy danych obejmującej okres do końca 2021 roku). Wtyczki będą mogły asystować w zakupach lub w innych działaniach użytkowników w internecie, czasami nawet działając za użytkownika (co wzbudza kontrowersje). Kto nie widzi w tym zagrożenia dla cesarstwa Google'a, ten chyba nic nie widzi.

Co ciekawe, jak podał niedawno „The Wall Street Journal”, inżynierowie Google'a, Daniel De Freitas i Noam Shazeer, opracowali przed

laty konwersacyjnego chatbota AI, którego nazwali Meena z myślą o wykorzystaniu go w wyszukiwaniach sieciowych. Jednak wówczas kierownictwo Google'a nie pozwoliło udostępnić go publicznie ze względu na obawy dotyczące bezpieczeństwa. Zarząd udaremnił podejmowane przez inżynierów próby wysłania bota do zewnętrznych badaczy, dodania funkcji czatu do Asystenta Google i uruchomienia publicznego demo. Twórcy opuścili Google'a pod koniec 2021 roku, aby założyć własną firmę.

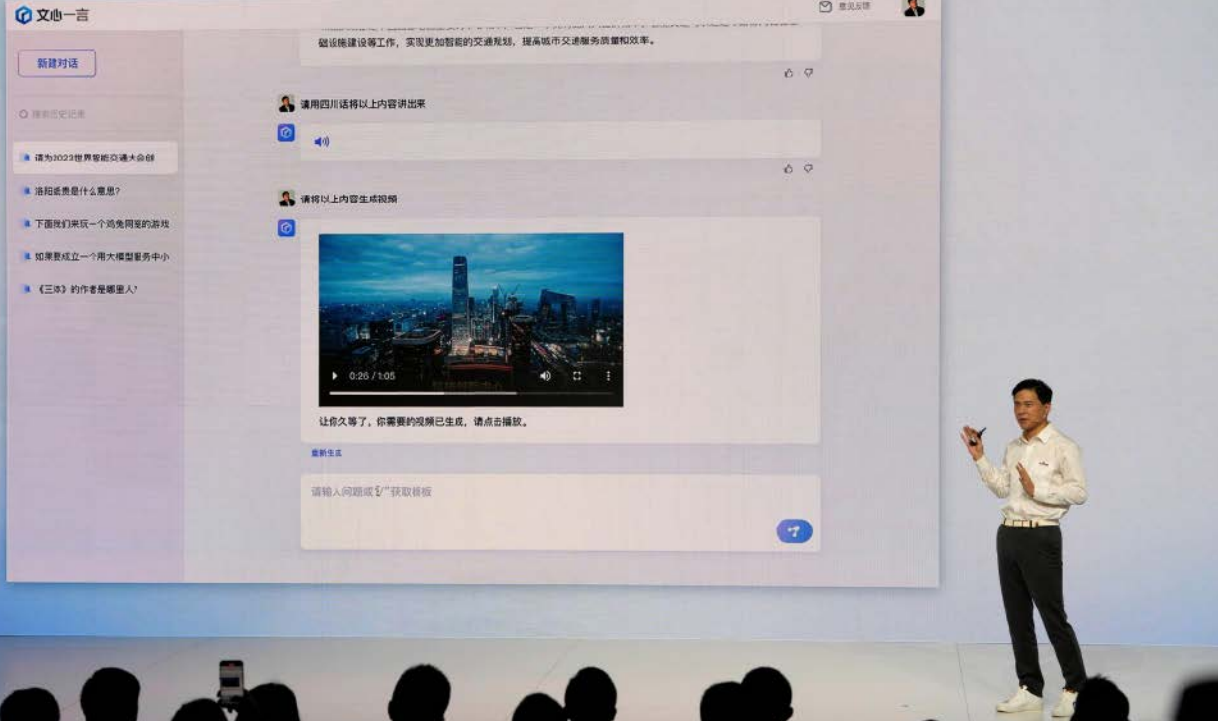
Teraz, według doniesień, Google inwestuje 300 milionów dolarów w Anthropic, start-up złożony z weteranów OpenAI, twórcy ChatGPT. Założony jeszcze w 2021 roku Anthropic stworzył coś podobnego do ChatGPT o nazwie Claude (w momencie przygotowania tego wydania MT jeszcze niedostępny publicznie). Jednak nie jest jasne, czy Google w ogóle planuje zintegrować Claude'a ze swoimi usługami, tak jak zrobił to, robi i zamierza robić Microsoft. Nie ma także pełnej jasności co do planów Google'a wobec wzmiankowanego Barda. Wiadomo, że Google testuje zmiany na swojej stronie głównej, z dodaniem być może jakiejś formy konwersacyjnego interfejsu.

Do ogólnego zamieszania dochodzą inne wydarzenia, takie np. jak wywiad szefa DeepMind (powiązanego z Google), Demisa Hassabis, udzielony magazynowi „Time”, w którym wspomniał o swoich zamiarach uruchomienia podobnego do GPT modelu w tym roku. Miałby to być model Sparrow, opisany w pracy badawczej DeepMind opublikowanej pod koniec 2022 roku. Został zaprojektowany tak, aby rozmawiać z użytkownikiem, odpowiadać na pytania i przeszukiwać Internet, używając Google'a do informowania o swoich odpowiedziach.

Wielkie poruszenie

Oczywiście także wielu innych potentatów technologicznych podejmuje działania i projekty, które mają w ich zamierzeniu pozwolić im utrzymać się na powierzchni pośród burzy zmian, która zaczęła się jeszcze w 2022 r.

Np. porównywalny na swoim rynku z Google chiński gigant Baidu zaprezentował w marcu 2023 r. nowy duży model językowy o nazwie Ernie Bot (4), który potrafi rozwiązywać zadania matematyczne, pisać teksty marketingowe, odpowiadać na pytania dotyczące chińskiej literatury i generować odpowiedzi multimedialne. Ernie Bot to skrót od angielskojęzycznej nazwy „Enhanced Representation from Knowledge Integration” i według doniesień z Kraju Środką nie jest jedynym tego typu projektem na tamtejszym rynku. Jedną z największych firm e-commerce, chińska



4. Prezentacja Ernie Bota firmy Baidu

Alibaba, opracowała szereg produktów AI, w tym napędzany przez model językowy chatbot do obsługi klienta, który może obsługiwać miliony zapytań klientów dziennie. Tencent, chiński potentat na rynku gier, również zaproponował chatbota AI do obsługi klienta.

Amazon ze swoim funkcjonującym od lat inteligentnym asystentem Alexa nie musi niczego w tej dziedzinie udowadniać a i tak pod wpływem ostatnich wydarzeń informuje o pracach nad podobnym do ChatGPT narzędziem oraz o współpracy z konkurencyjnych wobec OpenAI start-upem Stability AI, który m.in. stworzył generator obrazów Stable Diffusion.

Pośród ogólnego wielkiego poruszenia warto zwrócić uwagę, że bierności nie zachowują także nieco mniejsi, ale wciąż ważni w rozgrywce na rynku techniki internetowej i komputerowej, gracze. Próbująca coraz skuteczniej w ostatnich latach rywalizować z Google wyszukiwarka DuckDuckGo uruchomiła opartą na sztucznej inteligencji usługę DuckAssist, Adobe wprowadziła oparty na własnych zasobach i danych generator obrazów FireFly.

Ekspert i komentatorzy zwracają uwagę na brak wyraźnych sygnałów przejmowania się tym, co się dzieje ze strony Apple. Apple podobnie jak Amazon od lat rozwija opartego na AI asystenta o nazwie Siri. Ma liczny zespół specjalistów, ale jakoś nie wrywa się do tworzenia produktu, który miałby być bezpośrednią odpowiedzią na ChatGPT. Zdaniem wielu publicystów jest tak zapewne dlatego, że jako producent sprzętu i systemu nie musi. Twórcy rozwiązań prawdopodobnie

sami przyjdą do Apple z prośbą o zintegrowanie ich rozwiązań z produktami oznaczonymi jabłuszkiem. Zresztą, niewykluczone również, że Apple zaskoczy wszystkich już jesienią nową ofertą, która przyćmi generatory i modele innych firm.

Mającą długą historię projektów (ale również niepowodzeń) firma Meta, posiadaczka przede wszystkim Facebooka, w lutym 2023 r. zademonstrowała duży model językowy, który może działać na pojedynczym procesorze graficznym, co otwiera drzwi do wydajności podobnej do ChatGPT na sprzęcie klasy konsumenckiej. LLaMA-13B podobno przewyższa możliwości ChatGPT, mimo że ma znacznie mniejsze rozmiary. Meta opracowała całą kolekcję modeli językowych LLaMA o rozmiarach od siedmiu do 65 miliardów parametrów. Dla porównania, model GPT-3 OpenAI – podstawowy model ChatGPT – ma 175 mld. Meta twierdzi, że wersja o wielkości 13 mld może być uruchomiona na pojedynczym procesorze graficznym Nvidii A100, który jest stosunkowo dostępny, kosztując kilka dolarów za godzinę wynajmu na platformach chmurowych, w licznych benchmarkach dla modeli językowych AI przewyższa kilkanaście razy większy modelu GPT-3 firmy OpenAI. LLaMA nie ma jednak formy chatbota. Jest to pakiet open source, do którego dostęp może uzyskać każdy. Polityka ta, jak twierdzi firma, ma na celu „demokratyzację dostępu” do AI i pobudzenie badań tej problematyki. Meta opracowała również szereg innych produktów AI, choćby DeepFace, oprogramowanie

do rozpoznawania twarzy, które może rozpoznawać twarze z dużą dokładnością.

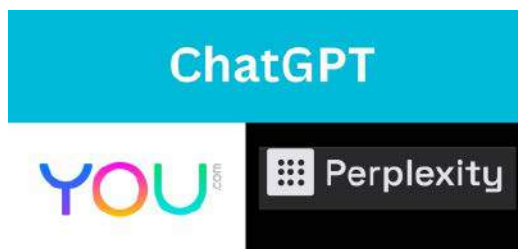
Parametr jest zmienną, którą model uczenia maszynowego wykorzystuje do przewidywania lub klasyfikacji na podstawie danych wejściowych. Liczba parametrów w modelu językowym jest kluczowym czynnikiem wpływającym na jego wydajność, przy czym większe modele są zazwyczaj w stanie poradzić sobie z bardziej złożonymi zadaniami i wytworzyć bardziej spójne dane wyjściowe. Większa liczba parametrów zajmuje jednak więcej miejsca i wymaga więcej zasobów obliczeniowych do uruchomienia. Jeśli więc model może osiągnąć takie same wyniki jak inny model z mniejszą liczbą parametrów, stanowi to znaczący wzrost wydajności. „Uważam, że w ciągu roku lub dwóch będziemy uruchamiać modele językowe o znacznej części możliwości ChatGPT na własnych (najwyższej klasy) telefonach komórkowych i laptopach”, napisał niezależny badacz AI Simon Willison w serwisie Mastodon, analizując znaczenie nowych modeli AI firmy Meta.

Udostępnienie w modelu open source jest możliwe, gdyż Meta szkoliła swoje modele LLaMA, używając publicznie dostępnych zbiorów danych, takich jak Common Crawl, Wikipedia i C4. Byłoby to nowością w branży, w której do tej pory wielcy rywale w wyścigu o AI zachowywali najpotężniejsze zasoby AI dla siebie. „W przeciwieństwie do modeli Chinchilla, PaLM czy GPT-3, używamy wyłącznie publicznie dostępnych zbiorów danych, dzięki czemu nasza praca jest możliwa do odtworzenia, gdy większość innych modeli opiera się na danych, które nie są publicznie dostępne lub nie są udokumentowane”, napisał na Twitterze członek projektu Guillaume Lample.

Specjaliści z Meta snują wizję „własnej”, „lokalnej” sztucznej inteligencji hostowanej na pojedynczym komputerze, w firmowym systemie lub w chmurze. Przyszłe wersje LLaMA mogłyby być, według tych planów, szkolone na danych i np. e-mailach pojedynczych użytkowników lub niewielkich grup. Pozwoliłoby to przygotować model na bardzo zindywidualizowane potrzeby i chronić wrażliwe dane osobowe, tajemnice firmy itd.

Do akcji wkraczają start-upy

Gdy giganci technologiczni przygotowują swoje odpowiedzi na zagrożenie ChatGPT, kilka start-upów uruchomiło wyszukiwarki z interfejsami czatu podobnymi do bota. Należą do nich You.com, Perplexity AI (5) i Neeva. Zbudowane przez nich narzędzia ilustrują zarówno potencjał, jak i wyzwania związane z adaptacją technologii w stylu ChatGPT



5. Nowi konkurenci na rynku chatbotów

do wyszukiwania. You.com, założona przez Richarda Sochera, eksperta w dziedzinie języka i AI, może udzielać odpowiedzi poprzez interfejs czatu. Odpowiedzi są opatrzone cytatami, które mogą pomóc użytkownikowi w ustaleniu pochodzenia danej informacji. Jednak model czasami łączy źródła, które nie pasują do siebie.

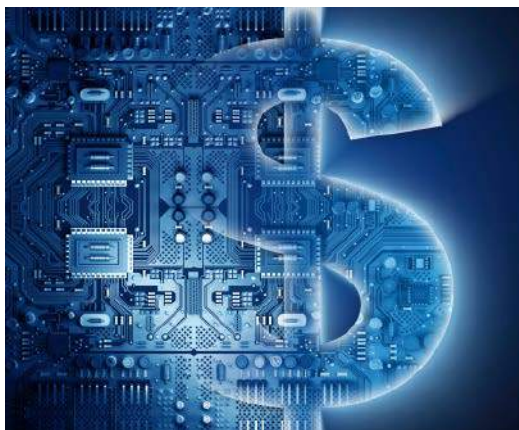
Aravind Srinivas, założyciel i dyrektor generalny start-upu wyszukiwania Perplexity AI, który wcześniej pracował w OpenAI, mówi, że wyzwanie związane z aktualizacją systemu podobnego do ChatGPT o najnowsze informacje oznacza, że muszą one być połączone z czymś innym. „Same nigdy nie będą w stanie być dobrymi wyszukiwarkami”, mówi magazynowi „Wired”.

Saam Motamedi, inwestor w Greylock Partners, który zainwestował w Neeva, firmę zajmującą się wyszukiwaniem opartym na sztucznej inteligencji, zauważa, że nie jest jasne, jak interfejsy czatu pogodzone mogą być z modelem przychodów z wyszukiwarek, czyli reklamą. Google i Bing używają zapytań użytkowników do serwowania reklam, które pojawiają się na górze listy wyników. Motamedi podejrzewa, że nowe formy reklamy mogą wymagać pojawienia się dla interfejsów wyszukiwania w stylu czatu, aby być opłacalne, ale nie jest oczywiste, co to będzie. Jego Neeva pobiera opłatę abonamentową za nieograniczone wyszukiwanie bez reklam.

Drogie szkolenie, jeszcze droższe wnioskowanie

I tak doszliśmy do najważniejszego problemu nurtującego Google'a i innych w kontekście wejściach nowych interfejsów sztucznej inteligencji. Potentaci stali się potentatami, ponieważ znaleźli sposób zarabiania na swoich produktach (6). Jak zarabiać na nowych produktach podobnych do ChatGPT, opartych na modelach AI? Już wiadomo, że Microsoft najprawdopodobniej użyje ChatGPT z MS Office. Ma to sens, ponieważ jest to stosunkowo mało ryzykowny sposób na początek i nie jest prawdopodobne, że zrazi klientów.

Jednak to podejście o niskim ryzyku ma również minusy. Ponieważ MS ma właściwie monopol



6. Jak zarobić na algorytmach?

na segment Office, dodawanie kolejnych funkcji raczej nie sprawi, że więcej osób skorzysta z usług Microsoftu. Praktycznie każdy, kto potrzebuje Teams, Worda czy Excela, i tak go ma, a większość firm nie ma tak naprawdę żadnej alternatywy. Tak więc dodatkowe przychody uzyskane z integracji nie będą raczej duże.

Grupa globalnych gigantów ma dostęp do ogromnych ilości danych, zasobów i wiedzy specjalistycznej, co daje im przewagę konkurencyjną w ekosystemie sztucznej inteligencji. Należy jednak pamiętać, że ekosystem AI stale się rozwija i mogą pojawić się nowi gracze, którzy mogą zakłócić dominację tych globalnych gigantów. Dodatkowo, ramy prawne i względy etyczne również mogą wpłynąć na to, co się stanie. Tak czy inaczej nad siecią napędzaną AI wisi fatum ogromnych kosztów.

Zanim pojawił się ChatGPT OpenAI, niewielki start-up o nazwie Latitude zdołał przyciągnąć uwagę wielu użytkowników swoją grą AI Dungeon (7), która

pozwałała wykorzystywać sztuczną inteligencję do tworzenia fantastycznych opowieści na podstawie podpowiedzi (tzw. promptów). Jednak w miarę wzrostu popularności AI Dungeon koszty utrzymania tekstowej gry fabularnej zaczęły gwałtownie rosnąć. Oprogramowanie do generowania tekstów w AI Dungeon było zasilane technologią językową GPT oferowaną przez OpenAI. „Żartowaliśmy, że mamy pracowników ludzkich i mamy pracowników AI, a na każdego z nich wydawaliśmy mniej więcej tyle samo”, mówił szef firmy Nick Walton w mediach. „Wydawaliśmy setki tysięcy dolarów miesięcznie na AI, a nie jesteśmy dużą firmą, więc był to bardzo wielki koszt”. Do końca 2021 roku Latitude przełączył się z korzystania z oprogramowania GPT na tańsze oprogramowanie językowe oferowane przez startup AI21 Labs i bezpłatne rozwiązania open source. Rachunki Latitude za „korzystanie ze sztucznej inteligencji” spadły kilkakrotnie. Firma zaczęła też pobierać od graczy opłaty abonamentowe

Historia ta odsłania nieprzyjemną prawdę, która kryje się za boomem generatywnych technik AI. Koszty rozwoju i utrzymania oprogramowania mogą być ogromne, zarówno dla firm, które rozwijają modele językowe, jak i tych, które wykorzystują AI do zasilania własnych produktów. Wysoki jest koszt zarówno początkowy, czyli szkolenia, jak też stały, czyli „wnioskowania”, czyli mówiąc w uproszczeniu, obsługi bieżących zapytań, za którymi stoją miliardy obliczeń. Większość szkoleń i wnioskowania odbywa się obecnie na procesorach graficznych (GPU), które nie są tanie. Podstawowy układ Nvidia dla centrów obliczeniowych kosztuje mniej więcej 10 tys. dolarów. Analitycy szacują, że proces szkolenia dużego modelu językowego, takiego jak GPT-3 OpenAI, może kosztować ponad



7. Gra AI Dungeon

cztery miliony dolarów. Koszty innych modeli nie są znacząco niższe. np. największy model LLaMA firmy Meta wykorzystał 2048 jednostek GPU Nvidia A100 do wyszkolenia 1,4 mld tokenów (750 słów to około tysiąca tokenów), co zajęło jakieś 21 dni. Szkolenie zajęło około miliona godzin pracy GPU. Przy cenach usług serwerowych AWS Amazona kosztowałoby to ponad 2,4 miliona dolarów.

Jednak wyszkolenie modelu to dopiero początek wydatków. Pamiętajmy o pozycji w kosztach określonej jako „wnioskowanie”. W przypadku produktu tak popularnego jak ChatGPT, który według szacunków firmy inwestycyjnej UBS osiągnął w styczniu sto milionów miesięcznych aktywnych użytkowników, przetworzenie milionów promptów, które ludzie wprowadzili do oprogramowania tylko w pierwszym miesiącu, mogło kosztować OpenAI 40 milionów dolarów. Koszty jeszcze bardziej rosną, gdy te narzędzia są używane miliardy razy dziennie. Analitycy finansowi szacują, że chatbot Bing AI Microsoftu, który jest zasilany przez model OpenAI ChatGPT, aby udzielać odpowiedzi wszystkim użytkownikom Binga na każde żądanie, potrzebuje infrastruktury wartej nawet cztery miliardy dolarów.

Jeden z prezesów spółki macierzystej Google, Alphabet, John Hennessy powiedział Reutersowi, że prowadzenie rozmowy z dużym modelem językowym może kosztować dziesięć razy więcej

niż standardowe wyszukiwanie po słowach kluczowych. SemiAnalysis, firma badawcza i konsultingowa skupiona na technologii chipów, oszacowała, że dodanie AI w stylu ChatGPT do wyszukiwania może kosztować Alphabet trzy miliardy dolarów, co jest kwotą umiarkowaną, bo uwzględnia oszczędności wynikające z użycia własnych chipów Google'a, Tensor Processing Units (TPU). Jednak rachuby te nie uwzględniają możliwego spadku dochodów z reklamy w wyszukiwarce w rozumieniu tradycyjnym, której popularność może szybko spadać.

Zawsze też jako prosta opcja zarabiania jest pobieranie opłat za dostęp. Dostęp do OpenAI Plus bez opóźnień i ograniczeń, a od niedawna do najnowszego modelu GPT-4, kosztuje dwadzieścia dolarów miesięcznie. Wtedy użytkownik powinien raczej mieć realny powód biznesowy, by używać tego narzędzia. Nie wydaje się to droga do masowego udostępnienia technik sztucznej inteligencji. Alphabet, czyli Google, bada opcje stosowania prostszych i tańszych rozwiązań AI do prostszych zadań. Taki właśnie „mniejszy model” LLM LaMDA zasilać ma chatbota Bard, co oznacza potrzebę „znacząco mniejszej mocy obliczeniowej” i umożliwia kierowanie usługi do większej liczby użytkowników. Użytkownicy jednak mogą nie widzieć wyraźnego powodu, by używać AI „klasy ekonomicznej”. I koło się zamyka. ■

Mirosław Usidus

Następca tronu. Opowieść z Elfhame

Holly Black

Wydawnictwo Jaguar, liczba stron: 400, cena: 52,90 zł

Zbiegła królowa. Niechętny książę. I zadanie, które może zniszczyć ich oboje. Pierwszy tom nowej, mrocznej i urzekającej dylogii od Holly Black.

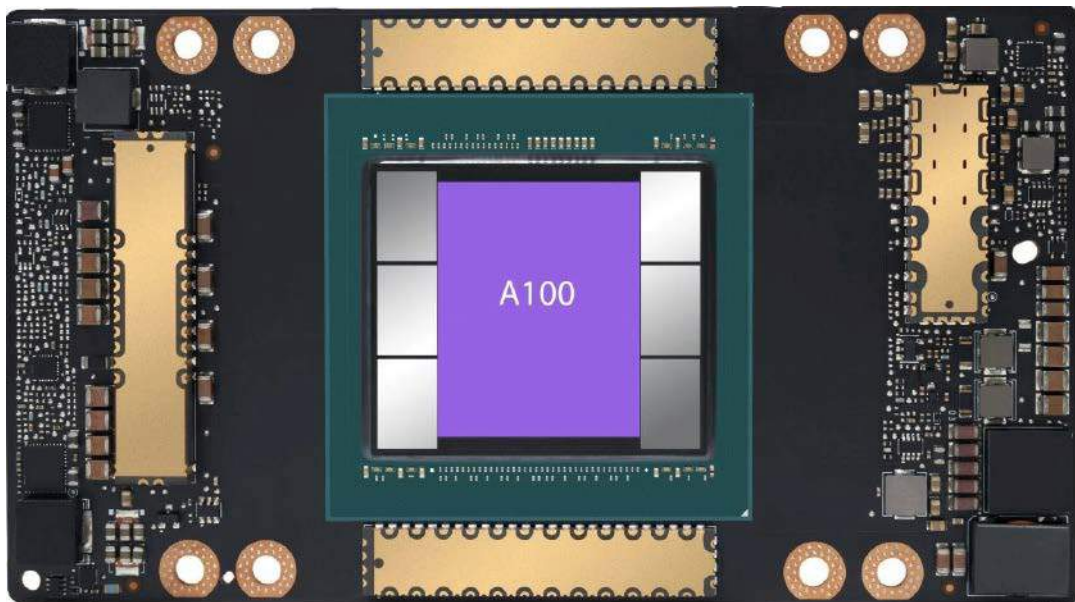
Od wydarzeń opisanych przez Holly Black w bestsellerowej trylogii *Okrutny książę – Zły król – Królowa niczego* upłynęło osiem lat. W Elfhame miłościwie panuje Najwyższy Król Cardan, a pochodząca z ludzkiego rodu Pani Judy sprawuje rządy razem z nim. Uważany przez wszystkich za następcę tronu Dąb wyrósł tymczasem na krzepkiego młodzieńca, który zamiast gnuśnieć na niekończących się biesiadach i opływać w dostatki, podejmuje śmiertelnie niebezpieczną misję.

Stawka w tej grze jest przerażająco wysoka, chodzi bowiem o ocalenie Elfhame w ostatecznym starciu z Nore, zdetronizowaną władczynią Zębatego Dworu. Podlega jej armia wynaturzonych potworów skleconych z wyjątkowo mrocznej magii, patyków oraz ludzkiego truchła.

W niebezpiecznej wyprawie Dębowi towarzyszy dzika Suren o piekielnie ostrych zębach, córka lodowej Pani Nore i zarazem jedyna osoba, której rozkazów nienawistna władczyni Północy musi słuchać.

Czy znany z krwawych czynów ojczym Dęba okaże się sprzymierzeńcem, czy nieprzyjacielem? Czy książę zdoła ocalić elfowy świat przed ponurym losem?





1. GPU Nvidia A100

Chipy zasilające szkolenie i działanie sztucznej inteligencji zostały uznane za towar na tyle strategicznie wrażliwy, że rząd USA nałożył sankcje i ograniczenia na eksport najbardziej zaawansowanego hardware'u do Chin.

Sprzęt kolejnej rewolucji

KRZEMOWE WOŁY ROBOCZE

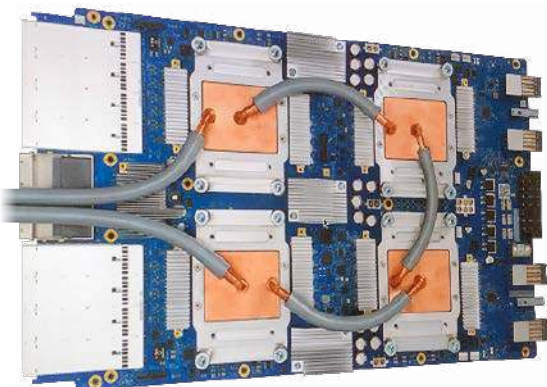
To właśnie było powodem stworzenia przez amerykańską firmę Nvidia Corp, która jest wiodącym projektantem i producentem układów półprzewodnikowych na rynku AI, zmodyfikowanych wersji swoich produktów, których specyfikacje pozwalają na eksport do Chin. Regulacje wprowadzone przez amerykańskie władze spowodowały, że Nvidia nie mogła sprzedawać tam swoich najbardziej zaawansowanych chipów A100 (1) i H100. Zaprojektowane przez Nvidię niedawno A800 i H800 mogą bez przeszkód wejść na chiński rynek. H800 jest wykorzystywany przez Alibaba Group,

Baidu i Tencent, chińskich potentatów rozwijających projekty oparte na uczeniu maszynowym.

Cała powyższa sytuacja z wymuszonymi przez sankcje „zubożonymi” chipami Nvidii dla Chińczyków uwytknęła rolę, jaką w trwającej rewolucji AI odgrywa sprzęt, wydajne, szybkie, procesory, które potrafią sprostać wysokim wymaganiom obliczeniowym.

Tajemnice wydajności sprzętu Nvidii

Nie jest specjalną tajemnicą, że to karty graficzne (GPU), w których specjalizuje się Nvidia, są bardziej wydajną od tradycyjnych procesorów opcją w dużych i skomplikowanych procesach obliczeniowych, stojących za algorytmami AI, od rozpoznawania pieszych przez autonomiczne pojazdy, przez diagnostykę obrazową w medycynie, po superkomputery i deep learning. Przewaga GPU nie polega jedynie na przetwarzaniu równoległym, co jest najczęściej podawanym powodem. Prawdziwa przyczyna jest nieco prostsza i dotyczy przepustowości pamięci. CPU potrafią pobrać (fetching) bardzo szybko niewielkie ilości danych. W scenariuszu kładącym nacisk na szybkość



2. Tensor Processing Unit 3.0

niezależnie od rozmiaru, GPU, mające znacznie większe opóźnienia, działają wolniej. Jednak gdy w grę wchodzi odczytywanie dużych pakietów danych, GPU radzi sobie nieporównywalnie lepiej, osiągając nawet 750 GB/s przepustowości. To sporo w porównaniu do najlepszych procesorów CPU, które wspierają maksymalnie 50 GB/s szybkości odczytu z pamięci.

GPU składają się z tysięcy rdzeni, a każde kolejne wczytanie danych jest w nich szybsze dzięki równoległości wątków. Aby rozwiązać zadanie polegające na obsłudze dużych ilości danych, trzeba czekać na ich pierwsze wczytanie (fetching). Każde kolejne będzie znacznie szybsze, gdyż to proces zrzucania (unloading) zabierający najwięcej czasu jest kolejkowy. GPU są w stanie przechowywać informacje w pamięci cache oraz plikach rejestru w celu ponownego użycia, co sprawia, że są bardziej wydajne w procesach głębokiej nauki maszynowej.

Pierwsza część procesu polega na pobraniu danych z dysku lub pamięci RAM i wgraniu ich na pamięć karty lub cache L1 (pamięć operacyjna) i rejestr. Rejestry są połączone bezpośrednio z jednostką obliczeniową, dla GPU jest to procesor strumieniowy, a dla CPU rdzeń. To tutaj mają miejsce wszystkie obliczenia. Zwykle dąży się do tego, by L1 i rejestr były możliwie blisko jednostki obliczeniowej i jednocześnie były na tyle małe, żeby można było uzyskać do nich szybki dostęp. Karty graficzne są wydajne, gdyż każda ich jednostka obliczeniowa ma pakiet rejestrów – nawet trzydzieści razy większych od tych w CPU, a jednocześnie dwukrotnie szybszych. Rezultatem jest nawet 14 MB zarezerwowanych dla pamięci rejestru pracujących z prędkością 80 TB/s. Dla porównania – przeciętna pamięć cache L1 procesora CPU pracuje z prędkością nie większą niż 5 TB/s, a rejestr rzadko ma więcej niż 128 KB i operuje z prędkością 10, może 20 TB/s.

Pozostałe pięć procent

O wiele mniej rozpowszechniony niż GPU Nvidii, ale znany w świecie AI, jest Tensor Processing Unit (TPU), układ scalony opracowany przez Google'a do uczenia maszynowego sieci neuronowych, wykorzystujący autorskie oprogramowanie Google TensorFlow (2). Google zaczął używać TPU wewnętrznie w 2015 roku, a w 2018 roku udostępnił je do użytku stronom trzecim, zarówno jako część swojej infrastruktury chmury, jak też oferując mniejszą wersję układu na sprzedaż. W porównaniu z procesorem graficznym, TPU są przeznaczone do dużej ilości obliczeń o niskiej precyzji (np. zaledwie 8-bitowej) z większą liczbą operacji wejścia/wyjścia na dżul. Różne typy procesorów nadają się do różnych typów modeli uczenia maszynowego. TPU są dobrze przystosowane do neuronowych sieci konwolucyjnych (CNN), podczas gdy GPU mają zalety dla niektórych w pełni połączonych sieci neuronowych, a CPU mogą mieć zalety dla sieci rekurencyjnych (RNN).

Jednostka przetwarzania tensorowego została ogłoszona w maju 2016 roku na Google I/O, kiedy to firma powiedziała, że TPU była już używana wewnątrz ich centrów danych przez ponad rok. Od 2017 roku Google nadal wykorzystywał procesory i układy graficzne do innych rodzajów uczenia maszynowego. Między innymi miały być wykorzystane w szkoleniu algorytmu AlphaGo przed rozgrywką z arcymistrzem Lee Sedolem, a także w systemie AlphaZero, który produkował programy do gry w szachy. Google wykorzystał również TPU do przetwarzania tekstu w Google Street View. W usłudze zdjęć Google pojedyncza TPU może przetwarzać ponad sto milionów zdjęć dziennie. Jest ona również wykorzystywana w RankBrain, którego Google używa do dostarczania wyników wyszukiwania. W październiku 2019 roku Google zapowiedział smartfon Pixel 4, który zawiera Edge TPU o nazwie Pixel Neural Core. W 2021 roku kolejna generacja tej jednostki obliczeniowej pojawiła się w linii smartfonów Pixel 6.

W kontekście „sprzętu do AI” najczęściej mówi się o kartach graficznych Nvidii lub TPU Google'a, ale nie są to jedyne konstrukcje na rynku, które mogą obsługiwać obliczenia wymagane przez modele i generatory. Na przykład firma AMD zaprezentowała w styczniu 2023 r. procesory AMD Ryzen z serii 7040, które są pierwszymi jednostkami x86 z wbudowanym akceleratorem sztucznej inteligencji. Układy są produkowane w 4-nanometrowym procesie technologicznym, oferując najmocniejszy na świecie zintegrowany układ graficzny, bazujący na RDNA 3.

Sztuczna inteligencja i jej rozwój opiera się na sprzęcie, głównie kartach graficznych. Z drugiej



3. Procesor Hailo

strony AI wykorzystywane jest jako pomoc w konstrukcji hardware'u, czego przykładem jest zaprezentowana niedawno przez izraelską firmę Hailo rodzina procesorów wizyjnych (VPU, Vision Processor Units) Hailo-15 (3). Są one przeznaczone do bezpośredniego zamontowania w kamerach nadzorujących inteligentne miasta, centra handlowe, budynki, hale przemysłowe itp., pozwalając na bieżącą analizę rejestrowanych wydarzeń w tzw. przetwarzaniu krawędziowym (edge computing).

A narzut energetyczny rośnie

Za olśniewającymi możliwościami generatorów tekstu i obrazu skrywa się jednak pewna niewygodna i co tu kryć, dość przykra, tajemnica. Za każdym razem, gdy DALL-E tworzy obraz lub ChatGPT generuje poezję, po stronie serwerowni zachodzi mnóstwo obliczeń, operacji tzw. wnioskowania, które wymagają sporo energii elektrycznej. Już od lat wiadomo o rosnącym udziale centrów danych w zużycia energii na świecie. Prognozy przewidują wzrost tego zużycia z 3 proc. w 2017 roku do 4,5 proc. do 2025 roku. W różnych krajach podaje się różne przewidywania. Chiny przewidują, że ich centra danych zużyją ponad 400 miliardów kWh energii elektrycznej w 2030 roku – 4 proc. całkowitego zużycia energii elektrycznej w tym kraju. Być może tworzone przed 2022 r. prognozy nie uwzględniają trwającej rewolucji AI.

Uwzględnić za to zdają się już nowe przedrostki w Międzynarodowym Układzie Jednostek, które zostały przyjęte w listopadzie 2022 r. To: ronto i quecto dla małych liczb oraz ronna i quetta dla bardzo dużych liczb, przydatne do opisu ilości danych na serwerach internetowych. Nowe prefiksy dla największych i najmniejszych liczb świata zostały potwierdzone w głosowaniu na Generalnej Konferencji Miar i Wąg (CGPM) w Wersalu we Francji. Przewiduje się, że do 2025 roku ilość informacji na świecie sięgnie 175 zettabajtów.

Wracając do problemu wydajności – to stały repertuar dyskusji o CPU i GPU, zaś w kontekście rozwoju AI – rozważań o energochłonności procesu wnioskowania versus szkolenia modelu. Ma to odniesienia do walki o „przewyścienie prawa Moore'a” i fizycznych ograniczeń w pakowaniu większej ilości tranzystorów na matrycach o większych rozmiarach. Jest jeszcze kwestia połączeń pomiędzy układami jednostek obliczeniowych. Tradycyjnie taniej było produkować procesory i układy pamięci oddzielnie, a przez wiele lat prędkość zegara procesora była kluczowym czynnikiem decydującym o wydajności. Obecnie to połączenia między układami scalonymi hamują rozwój mocy obliczeniowej. „Kiedy pamięć i przetwarzanie są oddzielone, połączenie komunikacyjne, które łączy te dwie domeny, staje się głównym wąskim gardłem systemu”, wyjaśnia w jednej z wypowiedzi Jeff Shainline z amerykańskiego Narodowego Instytutu Standardów i Technologii (NIST) i podobne opinie formułuje wielu innych specjalistów.

System AI wykorzystuje różne rodzaje obliczeń podczas szkolenia modelu. W skrócie proces ten zaczyna się od załadowania ogromnych ilości danych, które następnie są przetwarzane w iteracjach nauki maszynowej. Tysiące rdzeni GPU przetwarzają ogromne zbiory danych, takich jak obrazy lub wideo. Jeśli wyniki szkolenia mają być osiągnięte szybko, to najprostszym, ale drogim sposobem jest wynajęcie tylu jednostek GPU w chmurze, ile potrzeba (i – na ile stać zarządzających projektem). Procesy wnioskowania AI wymagają mniejszych mocy obliczeniowych na początku, jednak potem ogromna liczba obliczeń i przewidywań potrzebnych do podjęcia decyzji wymaga znacznie więcej energii niż szkolenie w dłuższej perspektywie.

Jednym z kierunków poszukiwania oszczędności w tej dziedzinie są metody próbkowania i tzw. pipeliningu w przyspieszaniu głębokiej nauki przez zmniejszanie ilości przetwarzanych danych. Do grupy takich rozwiązań należy SALIENT (for SAMpling, SLicing, and data movemeNT), opracowany przez naukowców z Massachusetts Institute of Technology i IBM w celu przewyściania hardware'owych wąskich gardeł. Podejście to ma radykalnie zmniejszyć wymagania dotyczące uruchamiania sieci neuronowych na dużych zbiorach danych. Niespodzianką nie jest jednak to, że te metody ograniczają dokładność i precyzję, czyli nie mogą zostać zastosowane do zadań np. w medycynie czy bezpieczeństwie.

Apple, Nvidia, Intel i AMD zapowiedziały budowę procesorów ze specjalne dla AI tworzonymi

rozwiązaniami wbudowanymi w tradycyjne procesory. Amazon dla swoich usług AWS tworzy nowy procesor Inferentia2. Ale te rozwiązania wciąż wykorzystują tradycyjną architekturę von Neumanna procesorów, zintegrowanej pamięci SRAM i zewnętrznej pamięci DRAM, wszystkie wymagają energii elektrycznej do przenoszenia danych do i z pamięci.

Prace nad przeniesieniem obliczeń bliżej pamięci RAM trwają, choć są jeszcze w fazie wstępnej. Obliczenia w pamięci (in-memory computing lub IMC) rozwiązują opisane wyżej problemy, uruchamiając obliczenia macierzy AI bezpośrednio w module pamięci, unikając przesyłania danych przez magistralę pamięci. IMC sprawdza się w przypadku wnioskowania AI, ponieważ obejmuje ono względnie statyczny (ale duży) zbiór danych o wagach, do którego wielokrotnie uzyskuje się dostęp. Zawsze istnieje potrzeba przesłania pewnych danych do i na zewnątrz układu obliczeniowego, jednak IMC eliminuje większość kosztów transferu energii i opóźnień związanych z przemieszczaniem danych poprzez utrzymywanie danych w tej samej jednostce fizycznej, gdzie mogą być efektywnie używane i ponownie wykorzystywane do wielu obliczeń.

Jak powiedział w kwietniu 2022 roku na konferencji GTC, Bill Dally z Nvidii, w jego firmie sztuczna inteligencja uczestniczy już w projektowaniu nowych procesorów. Zatem nie jest wykluczone, że ostatecznie to sama AI zaprojektuje i zoptymalizuje hardware



4. Komputery i AI

na swoje potrzeby, który rozwiązywałby wszystkie wyżej opisane problemy i ograniczenia.

Nvidia, która posiada około 95 proc. rynku chipów AI, przewiduje, że w ciągu dekady lat modele sztucznej inteligencji będą milion razy potężniejsze niż ChatGPT. Taką opinię wygłosił szef firmy, Jen-Hsun Huang, podczas prezentacji na GTC, dodając, że prawdziwe przełomy dopiero przed nami, a modele sztucznej inteligencji będą w stanie wykonywać jeszcze bardziej skomplikowane zadania (4), w tym przewidywanie zachowania ludzi i zwierząt, projektowanie nowych leków, materiałów i inne. ■

Mirosław Usidus

Pegasus: Jak szpieg, którego nosisz w kieszeni, zagraża prywatności, godności i demokracji

Laurent Richard, Sandrine Rigaud

Wydawnictwo Insignia Media, liczba stron: 496, cena: 49,99 zł

Witajcie w świecie powszechnej inwigilacji. Wystarczy odpowiednio zasobny portfel, bo moralność i etyka nie mają żadnego znaczenia. Oto orwellowski koszmar, który dzieje się dziś, na naszych oczach. Nikt nie jest bezpieczny, bo przecież... Každy z nas nosi w kieszeni szpiega.

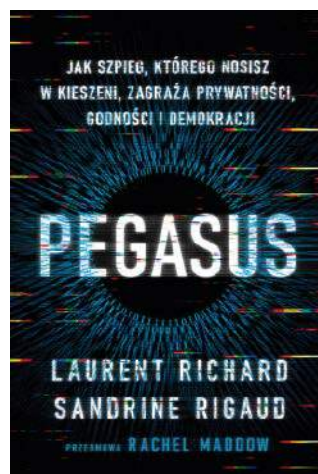
Pegasus – system szpiegowski, reklamowany i sprzedawany przez swoich twórców jako narzędzie do walki z terroryzmem, które w założeniu ma chronić nasze życie. Pegasus skrycie infekuje telefon komórkowy, o czym jego właściciel nie ma najmniejszego pojęcia, a następnie śledzi każdą jego aktywność w czasie rzeczywistym, przejmując kontrolę nad mikrofonem, kamerą i odbiornikiem GPS urządzenia oraz przechwytyując wszystkie filmy, zdjęcia, e-maile, SMS-y, treści i hasła, zaszyfrowane bądź nie.

Obecnie jednak Pegasus jest wykorzystywany nie tylko do walki z terrorystami.

Za jego pomocą śledzeni są dysydenci, reporterzy i wszyscy ci, których autokratyczne rządy chcą szpiegować, a w niektórych przypadkach uciszyć. Pegasus stał się też narzędziem walki politycznej. Jest podstępny i inwazyjny zarazem.

Może działać w twoim telefonie nawet teraz, kiedy to czytasz, zupełnie bez twojej wiedzy: nie zdradzając swojej obecności żadnymi podejrzanymi sygnałami.

Ta książka to mistrzowsko napisana opowieść o działaniu i wykorzystywaniu systemu Pegasus na całym świecie. To skrupulatny i mrozący krew w żyłach zapis śledztwa prowadzonego przez dwójkę nagradzanych dziennikarzy, którzy już wcześniej podejmowali trudne i ważne społecznie tematy. „Pegasus” to pasjonująca historia o tym, jak pewien spektakularny wyciek danych ujawnił zdumiewającą skalę inwigilacji i zatrważające metody, jakimi rządy na całym świecie – zarówno te autorytarne, jak i liberalne – podkopują kluczowe filary demokracji: prywatność, wolność prasy i wolność słowa.





RAPORT

1. Metawersum VR

Metawersum po przejściach

Rzeczywistość wirtualnie nieumarła

Najpierwotniejsza geneza koncepcji świata wirtualnego niknie w odmętach dziejów kultury ludzkiej. Schodząc w przeszłość, dochodzimy do wniosku, że ludzkości od zarania towarzyszy myśl o innym, lepszym, nawet idealnym, oddzielnym od doczesnego świecie. Współczesny raj lub królestwo snu ujęto w techniczne i wizualne ramy wirtualnej rzeczywistości (1), a ostatnio – metawersum.

Metawersum rozumie się zresztą najczęściej nie jako jeden „inny świat”, lecz rodzaj połączonego zespołu, konglomeratu lub systemu światów 3D, który łączyć ma różne wirtualne przestrzenie. Termin został po raz pierwszy użyty przez Neala Stephensona w powieści „Zamieć”. W książce Stephensona przedstawiono rozbudowany świat wirtualny, w którym można prowadzić równoległe życie. W książce „The Metaverse: And How It Will Revolutionize Everything”, Matthew Ball definiuje metawersum jako „masowo skalowaną i interoperacyjną sieć wirtualnych światów 3D renderowanych w czasie rzeczywistym, które mogą być doświadczane synchronicznie i trwale przez docelowo

nieograniczoną liczbę użytkowników z indywidualnym poczuciem obecności i z ciągłością danych, takich jak tożsamość, historia, uprawnienia, przedmioty, komunikacja i płatności”.

W ostatnich latach projekt ogólnie nazywany metawersum stawia sobie za cel zapewnienie ludziom miejsca do kontaktów, interakcji, rozmów, prowadzenia spotkań czy konferencji, korzystania z wirtualnej rozrywki, a także po prostu istnienia w tym świecie. Alternatywna rzeczywistość i przestrzeń do prowadzenia drugiego cyfrowego życia. W metawersum żywych ludzi reprezentują ich awatary, które spędzają czas podobnie jak w świecie fizycznym – pracują, bawią



się, spotykają ze znajomymi i uczestniczą w wydarzeniach sportowych lub kulturalnych. W metawersach, które chyba najlepiej znamy, czyli w światach gier, świat może stanowić odbicie znanej nam rzeczywistości lub przybrać formę świata przeszłości, scenarii science fiction albo fantasy.

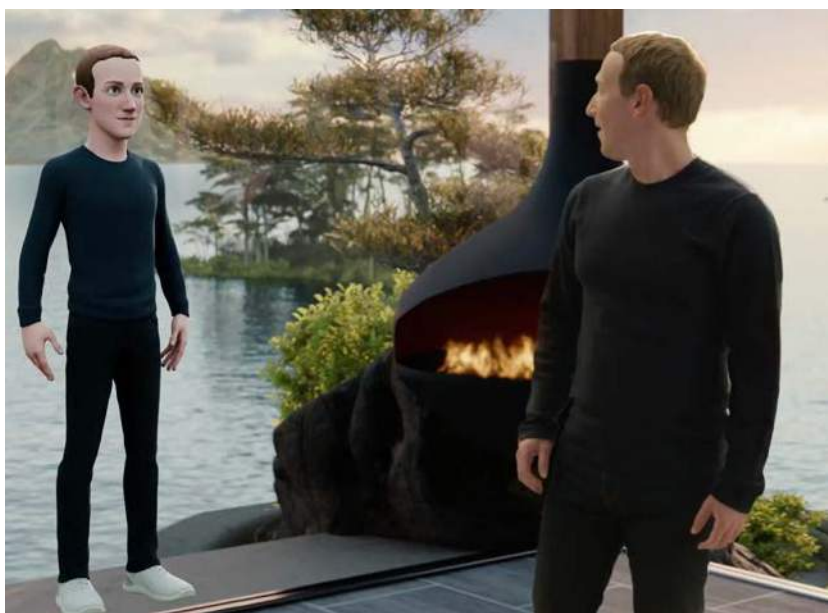
Parę lat temu za sprawą forsowanej przez Marka Zuckerberga zmiany nazwy firmy z Facebook na Meta i postawienia na projekt budowy rzeczywistości 3D, która miała „wciągnąć” użytkowników (2), temat metawersum znalazł się w centrum zainteresowania i dyskusji. Sporo zgiełku, opisywanego zresztą przez MT, towarzyszyło projektom metawersum (lub metawersów) opartych na kryptowalutach i tzw. Web3 (Web 3.0). Projekt Meta, mówiąc najdelikatniej „nie wypalił”, a krach na rynku kryptowalutowym postawił pod znakiem zapytania liczne projekty i cały ekosystem biznesowy oparty na cyfrowych aktywach, NFT, nieruchomościach wirtualnych itd.

Obecna sytuacja nie jest jasna. Nie brakuje skrajnych opinii, mówiących o całkowitym bankructwie idei metawersum. Inni twierdzą, że jesteśmy w fazie spadku po szczycie nadmiernie rozdętych oczekiwań, zgodnie z regułą cykli medialnego szumu firmy Gartner. Innymi słowy, według tego poglądu, metawersum powróci jako projekt w formie dojrzałej, implementowany w sposób spokojniejszy i zrównoważony, stopniowo zdobywając nowe obszary. Nowe pomysły już zresztą zaczęły się pojawiać. Aktywność nie wygasła, a branża technologiczna wciąż jest tematem zainteresowana, o czym świadczą dyskusje podczas Mobile World Congress (MWC), który odbył się na przełomie lutego i marca 2023 r. Określano tam metawersum jako sposób „całkowitego przekształcenia” doświadczenia klienta i innych dziedzin, w tym edukacji, w ciągu najbliższych pięciu lat.

Wydaje się, że obecnie jest dobry moment, aby podsumować pewien etap rozwoju projektu metawersum i dokonać analizy dalszych możliwych kierunków rozwoju.

Według danych Statista, najsilniej w metawersum inwestowały w 2022 r. takie branże jak spółki

technologiczne (67 proc.), media (56 proc.), rozrywka (53 proc.), handel detaliczny (51 proc.) i edukacja (48 proc.). Jak podała firma McKinsey, gospodarka oparta na metawersum może wygenerować do 2030 r. nawet 5 bilionów dolarów przychodów, głównie w handlu elektronicznym, reklamie i w grach. Z kolei Markets and Markets przewiduje, że rynek technik otaczających projekty metawersum dla świata rozrywki, mediów, sztuki, mody i handlu detalicznego osiągnie wartość 426,9 miliarda dolarów do 2027 roku. Można do tych danych dodać przewidywania Gartnera, który oszacował rynek cyfrowych bliźniaków dla przedsiębiorstw na 183 miliardy dolarów do 2031 roku.



2. Mark Zuckerberg i jego awatar w metawersum

Big Tech przesuwą zasoby na AI

Media pisały w pierwszych miesiącach 2023 roku, że Mark Zuckerberg pogrzebał już projekt metawersum. W latach 2021 i 2022 Reality Labs, oddział firmy Meta, zajmujący się projektami metawersum, odnotował skumulowaną stratę w wysokości prawie 24 miliardów dolarów, w tym 13,7 miliarda dolarów w samym ubiegłym roku. Jednak szef nie powiedział – „dość tego, zwijamy biznes”. Powiedział coś innego, mianowicie to, że Meta skoncentruje się teraz na pracy nad generatywną AI. Skoro zasoby firmy przekierowane zostały na konkurencję dla ChatGPT OpenAI i podobnych narzędzi, to dla komentatorów było jasne, że metawersum w Meta idzie w odstawkę. Oczywiście w tym świecie nigdy nic nie dzieje się w trybie aż tak radykalnym.



3. Cyfrowe metawersum przemysłowe

Meta będzie nadal rozwijać pozostałości tego wirtualnego projektu, na przykład zestawy do korzystania z VR, ale skierowane to zostanie do zawężonej grupy odbiorców, głównie graczy.

Metawersum w swoim poważniejszym wydaniu, biznesowym i przemysłowym (3), również przechodzi ostatnio bolesną weryfikację. W lutym 2023 r. Microsoft zlikwidował stuosobowy zespół Industrial Metaverse Core, który utworzył cztery miesiące wcześniej, aby pomóc klientom biznesowym korzystać z metawersum w warunkach przemysłowych. Zespół miał pomóc w tworzeniu interfejsów oprogramowania do obsługi systemów sterowania elektrowniami, robotyką przemysłową i sieciami transportowymi. Było to tłumaczone szerzej zakrojonymi redukcjami personelu w Microsoftzie, który na początku roku zwolnił dziesięć tysięcy pracowników. Może to mieć również, podobnie jak

w przypadku Meta, związek z przesunięciem zainteresowania firmy na projekty AI.

Mówienie, że ChatGPT czy DALL-E 2 firmy OpenAI skradły show metawersum, jest nieporozumieniem. Projektowi metawersum, pomimo szumu medialnego, nieustannie towarzyszyły sceptycyzm i dystans, a nawet przymrużenie oka. Odbiór nowej fali generatywnych modeli AI jest znacznie poważniejszy. Trudno zaprzeczyć, że sztuczna inteligencja przekroczyła pewną barierę możliwości i nie można jej już lekceważyć. Metawersum nie jest traktowane z taką powagą.

Niejedno ma imię

Metawersum nie jest pojęciem rozumianym na jedną tylko modłę. Znany projekt wirtualnych światów oparty na łańcuchach blokowych i kryptowalutach, prezentacje i innego rodzaju kreacje VR, metawersa świata gier i w końcu światy będące wirtualnymi



4. Metawersum Decentraland

**Wybór projektów świata metawersum**

Silks – Gracze-użytkownicy gromadzą tu tokeny NFT odpowiadające koniom wyścigowym. Gdy koń wygrywa wyścigi lub zwiększa wartość, Silks przesyła w czasie rzeczywistym dane na jego temat. To zwiększa wartość NFT odpowiadającego koniowi. Posiadacz NFT otrzymuje również różne nagrody za każdym razem, gdy rzeczywisty koń wygra wyścig w „realu”. Jednakże, aby posiadać konia, użytkownicy muszą najpierw zakupić awatara-dżokeja, który jest częścią metawersum.

Tamaverse – W przestrzeni wirtualnej, znanej jako Tamaverse, użytkownicy hodują zwierzęta domowe, kupując i zwiększając wartość obiektów kolekcjonerskich NFT Tamadoge, będących „wirtualnymi zwierzętami”, pielęgnowanymi przez graczy. Każdy Tamadoge NFT ma unikatowe cechy i umiejętności, co pozwalając graczom rywalizować i wspinać się po tablicy wyników w celu zdobycia nagród w kryptowalucie.

Souls of Nature – Souls of Nature to nowy projekt, który wystartował w 2022 roku. Souls of Nature poświęcone jest dzikim zwierzętom i kwestiom środowiskowym – awatarami graczy są zwierzęta. Zwierzęcy avatar może ewoluować i być ulepszany przez wykonywanie różnych zadań w metawersum, co pozwala gromadzić kryptowalutę, przyczyniając się zarazem do poprawy dobrostanu zwierząt i środowiska, bowiem część dochodów jest przekazywana organizacjom zajmującym się ochroną zwierząt.

Real Estate Investment Club – Real Estate Investment Club to kolekcja NFT stworzonych przez jednego z kreatorów serii „Grand Theft Auto” i „Red Dead Redemption”. NFT służą jako przepustka i avatar do metawersum Real Estate Investment Club, zwanego też REIC – MetaCity, w którym eksperci z dziedziny nieruchomości i finansów mogą gromadzić się, dzielić pomysłami i nawiązywać kontakty towarzyskie. Wewnątrz metawersum gracze mogą spieniężać swoje aktywa i usługi, tworząc wirtualną gospodarkę. W centrum REIC – MetaCity znajduje się REIC MetaHQ, specjalnie zaprojektowany przez zespół programistów, projektantów i architektów obiekt, centralne miejsce spotkań członków Real Estate Investment Club.

Gala Games – Oparty na blockchainie Ethereum zdecentralizowany projekt platformy gier i metawersum, w którym w założeniu sami użytkownicy decydują o kierunku rozwoju wirtualnego w głosowaniach. Jednak wpływ na losy platformy mają jedynie

społecznościami. Najgłośniejsze projekty tego typu, takie jak Decentraland (4), The Sandbox, Somnium Space, Gala Games, ApeCoin, Real Estate Investment Club, mają zwykle silne asocjacje z ekosystemem kryptowalutowym, dlatego tradycyjnie utożsamiano metawersum z tymi projektami. Jednak jest to znaczne uproszczenie.

O metawersum w swoim najpełniejszym wcieleniu nie będzie można mówić, dopóki ktoś nie zmieści techniki hiperrealistycznego wyświetlania 3D w rozmiarze zwykłych okularów, które mogłyby działać przez cały dzień na jednym ładowaniu akumulatora. Metawersum na ekranie smartfona czy laptopa to metawersum niepełne, nawet można by rzec, niepełnosprawne.

Wiele osób łączyło w ostatnich latach metawersum z innym pojęciem, Web3, „nowym WWW”, opartym na rozwiązaniach blockchain, NFT i „przedmiotach cyfrowych” o odmierzanej w kryptowalutach wartości. W ostatnim kwartale ubiegłego roku finansowanie start-upów Web3 spadło o 74 proc. w porównaniu z tym samym kwartałem roku 2021. Pęknięcie bańki spekulacyjnej w tej sferze nie jest powodem, aby spisać na straty projekt metawersum w ogóle. To jest na tyle pojemnym konceptem, że może się urzeczywistnić w sposób, który nie ma z kryptowalutami nic wspólnego. Jednak problemy świata krypto mają wpływ na opinię o wirtualnych projektach.

Przy okazji rozważań o relacjach świata waluty cyfrowej z koncepcjami wirtualnych światów powstaje fundamentalne pytanie – co, jeśli nie chęć wzbogacenia się na kolekcjonowaniu obrazków małp (NFT, Non-Fungible Tokens), ma być realnym motorem skłaniającym ludzi do użytkowania i angażowania się w metawersum. Spotkania i telekonferencje wirtualne? Szczyt koniunktury minął wraz z pandemią. Nie są to rzeczy aż tak atrakcyjne same w sobie, by napędzać projekt. Koncerty? Jak do tej pory, było nieco marketingowych wydarzeń tego rodzaju, ale to wciąż jedynie ciekawostki. Problem metawersum polega na tym, że na razie nie ma do zaoferowania nic, co dawałoby jakąkolwiek wyjątkową, pociągającą wartość, rzeczywiście ważny powód, by z niego korzystać.

Są oczywiście gry. Światy 3D, takie jak Roblox czy Minecraft, mają po latach istnienia ogromne rzesze użytkowników. Ich popularność nie jest silnie powiązana z projektem metawersum w rozumieniu, o którym mówimy. W świecie gier coraz silniejszą pozycję ma oparta na wirtualnej rzeczywistości platforma Quest (Oculus firmy Meta), ale trudno ją określić jako kamień milowy na drodze do metawersum.



5. Zakupy w metawersum

W ciągu kilku ostatnich lat byliśmy świadkami szeregu interesujących marketingowych wypraw znanych firm do metawersum. Nike, Balenciaga i Hyundai to tylko niektóre przykłady marek, które weszły do metawersum, aby eksperymentować i angażować klientów, a także po prostu sprzedawać (5). Ciekawe perspektywy widać w wykorzystaniu wirtualnych technik w treningu sportowym (6). Także organizacje szkolące, motywujące i angażujące swoich pracowników mogą wiele zdziałać w wirtualnych światach. Na przykład globalna firma konsultingowa Accenture już wprowadziła tysiące pracowników do swojego metawersum. Czy to będzie ścieżka do popularyzacji metawersum jako projektu i idei? Okaże się.

6. Wirtualne ćwiczenia fizyczne



posiadacze NFT. W obrębie metawersum możliwy jest handel i wymiana obiektów wirtualnych.

Otherside – Punktem wyjściowym tego metawersum jest znana kolekcja NFT, Bored Ape Yacht Club, która po raz pierwszy weszła na rynek w kwietniu 2021 roku, z 10 tysiącami algorytmicznie utworzonych NFT. Otherside został po raz pierwszy ujawniony w marcu 2022 roku. Ma łączyć mechanikę gier MMORPG i wirtualnych światów opartych na Web3 – twierdzą twórcy. Gracze mogą posiadać „ziemię”, wchodzić w interakcje z różnymi ekosystemami, a posiadane przez nich NFT będą mogły być zamieniane w postacię w grze. Dostępnych jest tu dwieście tysięcy działek „ziemi”, z których każda ma swoją zdefiniowaną lokalizację.

Decentraland – Jest chyba najbardziej znanym metawersum bazującym na kryptowalutach. Najlepiej, jak się zdaje, opisuje go analogia z Minecraftem, choć być może bardziej przypomina znaną grę Sim City. Ludzie mogą eksplorować, wchodzić w interakcje i dokonywać transakcji między sobą w sferze 3D o obszernej mapie. Zapewnia wirtualną ziemię, którą gracze mogą kupić, a następnie rozwijać poprzez umieszczenie infrastruktury w miejscu. Platforma oferuje 90 601 działek, każda o wymiarach 16 metrów na 16 metrów. Każda działka to niefinansowy token (NFT), który jest nabywany za pomocą kryptowaluty. Wersja Beta Decentraland została wydana w 2019 roku. Oficjalny debiut nastąpił w 2020 roku.

The Sandbox – Idea gry Sandbox jest podobna do Decentraland. To platforma, która pozwala użytkownikom swobodnie tworzyć, budować, kupować, wchodzić w interakcje, sprzedawać w wirtualnym świecie. Siła członków społeczności jest podkreślana w tym opartym na blockchainie metawersum przez umożliwienie im bycia jednocześnie twórcami i użytkownikami. W metawersum Sandbox znajduje się 166 464 działek. Gracze mogą tworzyć awatary, aby uzyskać dostęp do wielu gier, miejsc docelowych i ośrodków dostępnych w metawersum The Sandbox. Rodzimym tokenem The Sandbox jest Sandbox (SAND). Metawersum oddane jest w ręce zdecentralizowanej społeczności, co oznacza, że użytkownicy i konsumenci nie będą rozwijać niczego, czego społeczność nie chce.

Somnium – Somnium Space to darmowy i open-source'owy projekt metawersum. Użytkownicy mogą budować domy, kupować cyfrowe grunty i struktury, grać w hiperrealistyczne gry wideo, zakładać firmy, a także organizować koncerty i wydarzenia na żywo. Założyciel i dyrektor generalny Somnium Space, Artur



Sychow, zbudował platformę w 2017 roku i upublicznił ją we wrześniu 2018 roku. Obszar w Somnium jest dostępny za pośrednictwem wszystkich kompatybilnych urządzeń – komputerów PC, sprzętu VR i sieci, a dodatkowo jest również przyjazny dla urządzeń mobilnych. Program wykorzystuje blockchain Ethereum do tokenizacji aktywów, takich jak awatary, ziemia, obiekty noszone i inne. Rodzimą kryptowalutą jest tu CUBE, który może być używany do kupowania wirtualnych produktów, płacenia za towary i usługi w metawersum, płacenia za gry i wydarzenia, np. udział w koncercie, wynajmowania ziemi, a także nagradzania graczy.

Star Atlas – Projekt metawersum, w którym oczekuje się, że gracze będą utrzymywać swoje statki kosmiczne, zapewniając wystarczającą ilość paliwa, amunicji, zestawów narzędzi i żywności, jednym słowem – utrzymać swoją flotę w pełnej sprawności gotowości. Gracze kupują statki wewnątrz gry za kryptowaluty. Metawersum zbudowano na silniku graficznym Unreal Engine 5.

Metahero – Projekt kryptowalutowy, który został uruchomiony w lipcu 2021 roku z celem stworzenia ultrarealistycznego metawersum, pozwalającego użytkownikom skanować rzeczy ze świata rzeczywistego i samych siebie (co tworzy tzw. digital humans) do świata wirtualnego. Metahero konstruuje modele 3D ludzi w domenie cyfrowej za pomocą techniki skanowania firmy Wolf Digital World, z której korzystają firmy produkujące gry, takie jak CD Projekt, odpowiedzialny za Cyberpunk 2077 i serię Wiedźmin.

Bloktopia – 21-kondygnacyjny wieżowiec napędzany przez VR, który działa jako środowisko, w którym użytkownicy mogą się spotykać, prowadzić własne firmy, grać w gry i nie tylko. Ma własny token krypto o nazwie BLOK.

Upland – Platforma handlu i wymiany wirtualnych nieruchomości odwzorowujących rzeczywiste adresy na Ziemi. Użytkownicy mogą posiadać własne firmy, grać i łączyć się z innymi Uplandersami. Używana kryptowaluta – UPX.

RFOX – Ogromna wirtualna przestrzeń, której celem jest zapewnienie w pełni wciągającego doświadczenia wraz z elementami gier, mediów i finansów cyfrowych. Tokenem użytkowym jest tu VFOX.

Portals – Wspólne środowisko miejskie oparte na przeglądarce, które oferuje użytkownikom nieskończone możliwości eksploracji i tworzenia angażujących wirtualnych przestrzeni za pomocą łatwych w użyciu narzędzi.

Wysoki, wirtualny sędzie...

Wbrew minorowym nastrojom wielu komentatorów, w ostatnim czasie pojawia się sporo informacji o świeżych pomysłach na wykorzystanie metawersum.

W połowie lutego kolumbijski sąd przeprowadził pierwszą w historii rozprawę prawną w wirtualnym świecie. Podczas dwugodzinnej sesji, zorganizowanej przez kolumbijski sąd administracyjny w Magdalenie, uczestnicy sporu drogowego pojawili się jako awatary w wirtualnej sali sądowej (7). Awatar sędzi Marii Quinones Triana ubrany był w czarne szaty sędziowskie. „Dawało to większe poczucie realności niż telekonferencja wideo”, oceniła Quiones w wypowiedzi dla agencji Reuters. „To eksperyment akademicki, który ma pokazać, że prowadzenie rozpraw w ten sposób jest możliwe”, dodała. „Jeśli wszyscy uczestnicy się na to zgodzą, to mój sąd może tak procedować”. Zdaniem Juana Davida Gutierrez, profesora Uniwersytetu Rosario w Kolumbii, przed wirtualnymi sądami jest jeszcze długa droga. „Potrzebny jest do tego sprzęt, który posiada bardzo niewiele osób. A to rodzi pytania o równy dostęp do wymiaru sprawiedliwości”, powiedział Reutersowi.

Od dawna uważano, że teoretycznie, przed wypróbowaniem ich na prawdziwych pacjentach, skuteczność nowych metod chirurgicznych w metawersum testować mogą lekarze. Niedawno Uniwersytet George’a Waszyngtona zaadaptował zaawansowane narzędzie wirtualnej rzeczywistości na potrzeby neurochirurgii. Dzięki temu lekarze mogli wirtualnie zbadać mózg i ciało pacjenta przed operacją. Operacje z wykorzystaniem technik wirtualnych są również wykonywane przez urządzenia robotyczne sterowane przez chirurgów (którzy są obecni w zdalnej lokalizacji), co poprawia dokładność procedury i zmniejsza związane z nią ryzyko i komplikacje. Docieramy do idei wykorzystania w medycynie tzw. cyfrowych bliźniaków, które pozwalają lekarzom wstępnie badać, symulować operacje, szkolić się a nawet operować zdalnie z wykorzystaniem cyfrowego odpowiednika organu pacjenta. Metawersum umożliwia wykorzystanie cyfrowych bliźniaków w charakterze bardziej ogólnym, jako poligonu doświadczalnego dla przyszłych technologii, prognozowania cykli powrotu do zdrowia pacjentów i reakcji na leczenie.

Skoro dotarliśmy do cyfrowych bliźniaków, znanych lepiej z przemysłu, to warto wyjaśnić, że są one również częścią metawersum, w tym przypadku biznesowo-przemysłowego. Mają pomagać przedsiębiorstwom uchwycić wirtualne reprezentacje fizycznych kanałów, produktów i operacji, aby usprawnić fizyczne procesy. „Uważamy, że metawersum wpłynie na każdy

aspekt działania każdej firmy, w tym na to, jak wykonywana jest praca, jakie produkty są oferowane, jak produkty są dystrybuowane i jak działają firmy”, prognozował Jason Warnke z firmy doradczej IT Accenture. Jego zdaniem, firmy będą budować swoje własne metawersa biznesowe jako część modelowania, przygotowań i planowania inwestycji i innych działań w świecie rzeczywistym. Accenture niedawno samo zbudowało cyfrowego bliźniaka planowanego biura do analizy różnych układów przestrzeni roboczej, technologii i procesów, co radykalnie zmniejszyło liczbę późniejszych kosztownych błędów i przeróbek.

Niektórzy idą dalej, przewidując, że udoskonalone techniki odwzorowywania obiektów 3D, takie jak lidar, będą

Połączenie bardziej wydajnych technik modelowania, wykorzystujących AI i nowsze algorytmy, z potężniejszymi komputerami umożliwić ma firmom analizowanie coraz większych systemów. Ponadto będą w stanie połączyć cyfrowe bliźniaki z analizy danych w rzeczywistym świecie z cyfrowymi bliźniakami „projektowymi”, opartymi na fizyce ruchu i pracy. To powinno pozwolić stworzyć dokładne, realistyczne modele coraz większych zbiorów aktywów w czasie rzeczywistym.

Metawersum jest ostatecznie przeznaczone nie tylko dla ludzi. Pomoże inżynierom w szkoleniu inteligentnych robotów. Tak uważa Apurva Shah, dyrektor generalny firmy Duality Robotics, produkującej



7. Posiedzenie kolumbijskiego sądu w metawersum

modelować zasoby fizyczne firm. „To pozwoli odtworzyć całe organizacje, wraz z systemami, infrastrukturą i ludźmi, za pomocą cyfrowych bliźniaków w metawersum”, przewiduje Frank Diana z Tata Consultancy Services w wypowiedzi dla mediów. Metawersum zostanie wykorzystane do stworzenia znacznie lepszych doświadczeń w zakresie obsługi klienta poprzez umożliwienie agentom cyfrowego wsparcia bliźniaczego wirtualnego siedzenia obok konsumenta i pomagania mu w rozwiązywaniu problemów, uważa Diana. Na przykład klient, który ma problemy z instalacją nowego kranu, może zwrócić się do agenta obsługi klienta o pomoc w instalacji w rzeczywistości wirtualnej.

oprogramowanie do tworzenia cyfrowych bliźniaków. Wysokiej jakości modele rzeczywistych środowisk, takich jak fabryki, magazyny, place budowy i kopalnie, mogą pomóc w zaplanowaniu pracy i współpracy robotów o różnych konfiguracjach i funkcjach w tej samej przestrzeni.

W branży, nazwijmy to szumnie, wirtualnej, wciąż na wczesnym etapie rozwoju, pojawiają się raz po raz nowe projekty biznesowe. Wiosną ubiegłego roku platforma wirtualnej rzeczywistości Obsess, która posiada ponad dwieście wirtualnych sklepów, nawiązała współpracę z programem SAP.iO, który umożliwia partnerom SAP współpracę z deweloperami w celu





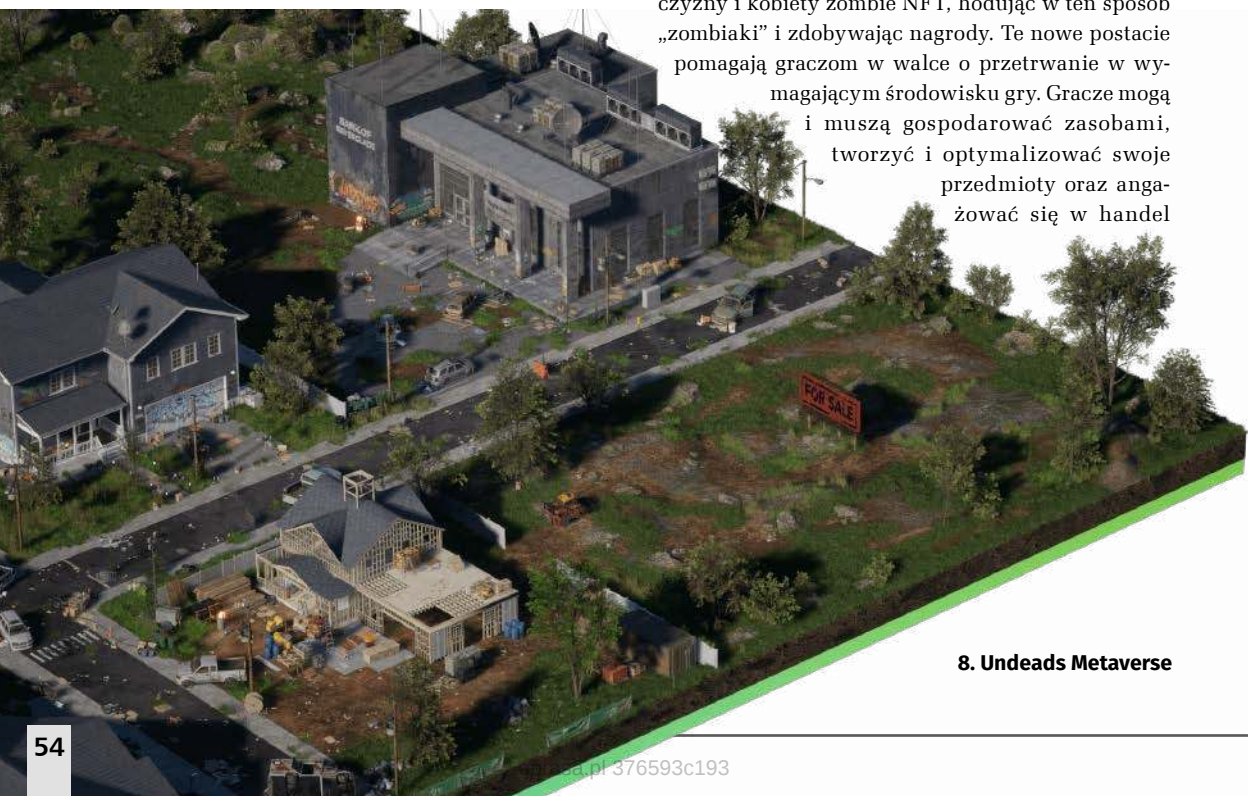
zintegrowania platformy e-commerce i wirtualnych sklepów oraz stworzenia interfejsu, który angażuje młodych konsumentów. Klienci Obsess i SAP mogą tworzyć cyfrowe miejsca zakupów, które mogą różnić się od rzeczywistych sklepów lub być do nich podobne. Konsumentci wchodzą do świata VR i tam mogą angażować się w kontakt ze sprzedawcą. W listopadzie ubiegłego roku Obsess wprowadził Branded Avatars, funkcję, która umożliwia tworzenie środowisk typu metawersum na stronach e-sklepów, oferując kupującym awatary, które mogą uczestniczyć w zakupach lub wydarzeniach w czasie rzeczywistym wraz z awatarami znajomych.

Łatwo sobie wyobrazić wiele produktywnych zastosowań metawersum, od edukacji, przez wirtualne spotkania, po rozrywkę. I większość z nich jest już w trakcie wdrażania w takiej czy innej formie. W tej chwili większość tej działalności toczy się wokół dóbr wirtualnych tworzonych przez marki, które klienci już znają i są wobec nich lojalni, i często koncentruje się wokół awatarów w grach online. Ale są już pewne oznaki, że to się zmienia. Oprócz sprzedaży wirtualnych rzeczy do wykorzystania w metawersum, firmy zaczynają patrzeć na sprzedaż wirtualnych dóbr, które są cyfrowym bliźniakiem przedmiotów w świecie fizycznym. Użytkownik może nabyć prawdziwą wersję, gdy kupi wirtualną dla swojego awatara. To zaczyna dotyczyć nawet rynku nieruchomości. Wirtualne działki czy rezydencje były kwitującym biznesem w „Second

Life” jeszcze w 2007 roku. W metawersum obecnym, w przeciwieństwie do roku 2007, wydaje się, że mniejszy nacisk kładzie się na tworzenie fantastycznych krain, a większy na cyfrowego bliźniaka prawdziwej nieruchomości. W styczniu KB Home, jeden z największych deweloperów w Stanach Zjednoczonych, otworzył swoje podwoje dla społeczności metawersum Decentraland, gdzie potencjalni nabywcy mogą wejść, zbadać i zabawić się opcjami konfiguracji domów, które KB Home sprzedaje w świecie rzeczywistym.

Wśród zombiaków i kryptowalut

W branży gier gałąź wykorzystującą techniki łańcuchów blokowych nazywa się hasłowo GameFi. Wielu deweloperów gier wprowadziło na rynek swoje własne produkcje oparte na blockchainie. Jednym z najnowszych typowych dla tej fali przykładów jest Undeads Metaverse (8), gra survivalowa AAA, która wykorzystuje mechanikę play-to-earn (P2E). Gracze zostają tu przeniesieni do postapokaliptycznego świata, w którym muszą wybrać, czy stanąć po stronie ludzi, czy zombie. Gracze mogą angażować się w konflikt z postaciami zombie i ludzi (obie formy reprezentowane jako kryptowalutowe non-fungible tokens NFT), walcząc o zasoby i terytorium. Każdy z NFT ludzkich lub zombie może być własnością gracza, który ją nabył lub „wyhodował” na platformie jako obiekt o określonej wartości w kryptowalucie. Posiadacze NFT mogą kreować nowe postacie np. przez łączenie w pary mężczyzn i kobiety zombie NFT, hodując w ten sposób „zombiaki” i zdobywając nagrody. Te nowe postacie pomagają graczom w walce o przetrwanie w wymagającym środowisku gry. Gracze mogą i muszą gospodarować zasobami, tworzyć i optymalizować swoje przedmioty oraz angażować się w handel



8. Undeads Metaverse

z innymi graczami. Mogą nawet zostać w metawersum gry przedsiębiorcami, posiadać „ziemię”, pobierać opłaty za każdą interakcję w grze odbywającą się na ich posiadłości. Mogą eksplorować niezbadane terytoria i bronić swoich osad, a jako właściciele ziemi mogą pobierać opłaty za każdą interakcję w grze opartą na lokalizacji gracza w świecie gry.

Ważną innowacją jest to, że gracze mogą wchodzić do Undeads Metaverse bez żadnych typowych dla Web3 barier, przez płynny i bezproblemowy proces wprowadzania do gry, który nie wymaga od graczy natychmiastowego zdobycia NFT, z możliwością generowania tokenów na późniejszym etapie. Oprócz gry, zespół deweloperów buduje VR Social Hub z wieloma grami VR opartymi na doświadczeniu metawersum i wirtualnej rzeczywistości. Gracze mogą spotkać się, aby spędzić czas, nawiązać kontakty towarzyskie, bawić się, wchodzić w interakcje i zarabiać. VR Social Hub oferuje m.in. wędkarstwo sportowe, kręgle, salę bilardową VR i wiele innych aktywności. Środowisko Undeads Metaverse VR jest zasilane na Unreal Engine 5.1 firmy Epic.

Prawdziwym sukcesem na wielką skalę w branży wirtualnych światów nie wydaje się jednak żadna z gier opartych na kryptowalutach, tylko platforma powstała wiele lat temu, gdy o metawersum nie było jeszcze za bardzo mowy. W 2022 r. gracze spędzili prawie pięćdziesiąt miliardów godzin na platformie Roblox – wynikało z raportu opublikowanego przez firmę analityczną SuperData. Sam Roblox poinformował, że jego dzienna suma aktywnych użytkowników wzrosła o 23 proc. do 56 milionów, porównując rok 2022 z poprzednim. Firma powiedziała również, że w styczniu jej liczba dziennych aktywnych użytkowników sięgnęła 65 milionów.

Roblox to gra, która pozwala użytkownikom tworzyć i udostępniać własne gry, a także bawić się w gry stworzone przez innych. Platforma jest szczególnie popularna wśród dzieci i młodzieży, ale zyskuje też coraz większą popularność wśród dorosłych. Według niektórych ekspertów, takie platformy jak Roblox, to załączkowa forma metawersum – wirtualnej przestrzeni, w której ludzie mogą łączyć się, wchodzić w interakcje i tworzyć treści. W oczach wielu liderów technologii popularność Robloxa wśród młodszych użytkowników, w połączeniu ze sposobem funkcjonowania platformy, służy jako model działającego już metawersum.

Co ciekawe platforma, podobnie jak gry oparte na blockchainie, ma własną walutę, Robux. Według Roblox udało się w 2022 r. wygenerować 2,9 miliarda dolarów, sprzedając wewnętrzną walutę platformy. Waluty tej gracze używają np. do zakupu awatarów służących do gier.

Uwaga na dane osobiste

Oprócz znanych kwestii związanych ze sprzętem i sceptycyzmem użytkowników, metawersum boryka się też niestety z innymi poważnymi zastrzeżeniami. Należy do nich wielka obawa o prywatność użytkowników. Według opublikowanych na początku 2023 roku wyników badań zespołu uczonych z Uniwersytetu Kalifornijskiego w Berkeley, prywatność może być w metawersum w ogóle niemożliwa. Prowadzone przez Viveka Naira badanie analizowało wielki zbiór danych o interakcjach użytkowników w wirtualnej rzeczywistości (VR). Najważniejszy z jego wyników wniosek jest taki, że w wirtualnej rzeczywistości potrzeba do identyfikacji użytkownika bardzo niewiele danych, co potencjalnie eliminuje wszelkie szanse na anonimowość. Dzieje się tak dlatego, że najbardziej podstawowy strumień danych potrzebny do interakcji z wirtualnym światem, proste dane o ruchach użytkownika mogą wystarczać do jednoznacznej identyfikacji użytkownika w dużej populacji. Mówiąc „proste dane o ruchu”, badacze mają na myśli trzy najbardziej podstawowe punkty danych śledzone przez systemy rzeczywistości wirtualnej – punkt na głowie użytkownika i po jednym na każdej ręce. Stanowią minimalny zestaw danych wymaganych do umożliwienia użytkownikowi naturalnej interakcji w środowisku wirtualnym.

W badaniu przeanalizowano ponad 2,5 miliona nagrań z VR (w pełni zanonimizowanych) pochodzących od ponad 50 tysięcy użytkowników popularnej aplikacji Beat Saber (9) i stwierdzono, że poszczególnych użytkowników można było jednoznacznie zidentyfikować z ponad 94 proc. dokładnością, wykorzystując zaledwie 100 sekund danych ruchowych, zaś połowa wszystkich użytkowników mogła zostać jednoznacznie rozpoznana przy użyciu zaledwie dwu sekund danych z ruchu. Osiągnięcie takiego poziomu dokładności wymagało zastosowania technik AI, ale ilość danych była niewielka.

Inaczej mówiąc – gdy użytkownik VR zamachnie się wirtualną szablą na lecący w jego kierunku obiekt, dane o jego ruchu mogą bardziej jednoznacznie go identyfikować niż jego odcisk jego palca. Za każdym razem, gdy użytkownik zakłada zestaw gogli, chwytą dwa standardowe kontrolery ręczne i rozpoczyna interakcję w wirtualnym lub rozszerzonym świecie, pozostawia za sobą cyfrowe ślady dokładniejsze niż odciski palców. Ponadto te same dane o ruchu mogą być wykorzystane do dokładnego wnioskowania o wielu specyficznych cechach osobistych użytkowników, w tym m.in. o wzroście, ułożeniu rąk i płci. A to w połączeniu z innymi danymi



powszechnie śledzonymi w środowiskach wirtualnych i rozszerzonych może oznaczać jeszcze dokładniejsze rozpoznanie.

Deweloperzy i firmy wiedzą o tym problemie. Dlatego proponują różne metody anonimizujące i blokujące dane prywatne w światach wirtualnych. Jednym z podejść jest zaciemnianie danych o ruchu, zanim zostaną przesłane ze sprzętu użytkownika do zewnętrznych serwerów. Niestety, oznacza to wprowadzanie szumu do systemu czy chroni prywatność użytkowników, ale zmniejsza również precyzję zręcznościowych aplikacji wymagających umiejętności fizycznych. Dla wielu użytkowników może to być nie do przyjęcia. Alternatywnym podejściem jest wprowadzanie regulacji, przepisów chroniących dane w VR, które uniemożliwiłyby platformom przechowywanie i analizowanie danych o ruchu człowieka w czasie. Jednak takie przepisy byłyby trudne do wyegzekwowania i mogłyby spotkać się z oporem ze strony branży.

Warto dodać, że w miarę jak techniki przechwytywania 3D będą gromadzić coraz większy zakres informacji osobowych i biometrycznych, mogą stać się łatwym kąskiem dla hakerów. Zdarzały się już ataki typu deepfake, wykorzystujące dane, wizualne lub audio, do podszywania się pod konkretnych ludzi, szefów, klientów, przedstawicieli władz lub instytucji, co prowadzić może do wyłudzenia danych lub innych przestępstw. Wiele podmiotów też zaczyna uczulać się na nową generację cyfrowych podróbek lub podróbek drukowanych w 3D.

Bez metanaiwności, proszę

Sceptyków nie brakuje, z którejkolwiek strony się do metawersum podejdziesz. Czy naprawdę jesteśmy na poziomie mocy obliczeniowej i wydajności sprzętu wystarczającej do obsługi VR i AR? Zestawy sprzętowe są wciąż nieopłacalne i choć koszty spadły, wciąż nie są tanie. Za nami wiele falstartów i niezrealizowanych obietnic. Kto np. pamięta Magic Leap? Ta firma była spowitą przez długi czas mgłą tajemniczy gwiazdą świata AR od 2011 roku, dopóki nie pokazała swojego produktu w 2018 roku. Zza mgły tajemniczości wyłonił się mocno niedoskonały produkt, który poniósł rynkową porażkę.

Z drugiej strony, jak widzieliśmy ostatnio w przypadku ChatGPT, duże skoki po stronie rozwiązań technicznych mogą w szybkim tempie zmienić wszystko. Być może metawersum wkrótce będzie miało swój historyczny moment. Na razie fakt, że firmy mające

opinię konserwatywnych i stroniących od ryzyka, np. PwC traktują ten koncept na tyle poważnie, by tworzyć zespoły zajmujące się projektem metawersum, wskazuje na znaczenie tematu.

Specjaliści z branży oprogramowania na ogół przewidują, że wejście metawersum do tzw. głównego nurtu zajmie od trzech do dziesięciu lat. Nie brak opinii, że wdrażanie wirtualnych technik będzie przebiegać etapami, przy czym rozszerzona rzeczywistość (AR) będzie miała charakter przygotowania, etapu pośredniczącego, który pozwoli dotrzeć do coraz szerszych grup użytkowników. Potem zacznie się stopniowe przechodzenie na światy w pełni wirtualne, w miarę rozwoju sprzętu dla VR i rozwiązań haptycznych (interfejsów dotykowych). Po stronie sprzętu i infrastruktury obliczeniowej czy wszystkich zwrócone



9. Gra Beat Saber VR

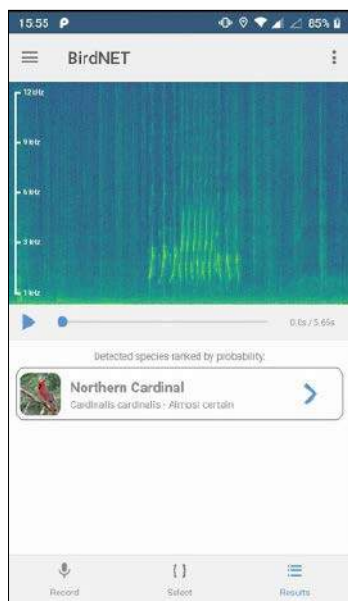
są na Amazon, Apple, Google, Meta i Microsoft. Bowiem, jeśli metawersum ma mieć przyszłość, tylko kolosy Big Tech mają fundusze i zasoby, aby inwestować, rozwijać i wprowadzać z sukcesem na rynek nowe produkty.

Obawy związane z metawersum dotyczą głównie prywatności i nierówności w dostępie do tych rozwiązań w skali społeczeństw i świata. Są też pytania o charakterze politycznym. Komu możemy zaufać, gdy mowa o kontroli, zarządzaniu i administrowaniu metawersum? Jeśli nie będzie ograniczone geograficznie, to pytanie to ma geopolityczny charakter. W jaki sposób będziemy ścigać metazłoczyńców metaprzestępców, metaszpiegów? Założenie, że wszyscy w metawersum będą dobrzy i będą kierować się zasadami etyki, dobra wspólnego, jest cokolwiek naiwne. ■

Miroslaw Usidus



Test aplikacji: Na łono natury



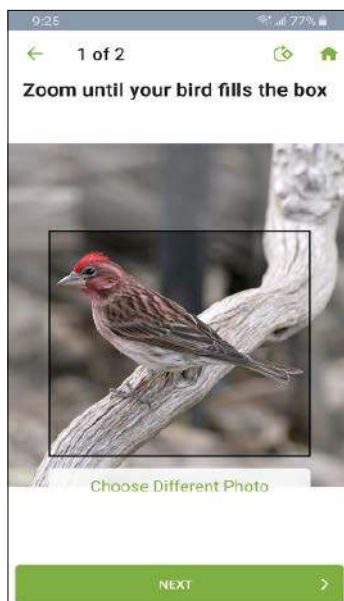
BirdNET

BirdNET to aplikacja do rozpoznawania ptaków oparta na sztucznej inteligencji. Aplikacja ta umożliwia użytkownikom zarejestrowanie dźwięków ptaków, a następnie rozpoznanie gatunku na podstawie analizy akustycznej.

Aplikacja BirdNET jest dostępna na systemy Android i iOS, a jej interfejs ma opinię bardzo intuicyjnego i łatwego w użyciu. Użytkownik może nagrać krótki plik dźwiękowy za pomocą wbudowanego w urządzenie mikrofonu. Aplikacja przetwarza go i przedstawia listę sugestii dotyczących możliwych gatunków ptaków, których śpiew może być zarejestrowany na nagraniu.

Aplikacja BirdNET wykorzystuje zaawansowane algorytmy sztucznej inteligencji do analizy śpiewu ptaków a jej baza danych zawiera ponad tysiąc gatunków z całego świata. Program jest przydatnym narzędziem pracy lub realizacji hobby dla ornitologów, naukowców, miłośników ptaków, i w końcu także dla każdej osoby, która interesuje się przyrodą.

BirdNET	
Producent	Stefan Kahl
Platforma	Android, iOS, Google Chrome
Oceny	Możliwości 8/10
	Łatwość obsługi 9/10
	Ocena ogólna 8,5/10



Merlin Bird ID

Podobnie jak BirdNET pomaga użytkownikom w identyfikacji ptaków, jednak za pomocą nieco większej liczby parametrów niż jedynie głos. Aplikacja ta została stworzona przez Cornell Lab of Ornithology, organizację non-profit z siedzibą w Stanach Zjednoczonych, która zajmuje się badaniami i ochroną ptaków.

Aplikacja Merlin Bird ID umożliwia użytkownikom identyfikację ptaków na podstawie ich wyglądu, zachowania i śpiewu. Aplikacja umożliwia dostęp do baz danych zawierających informacje o tysiącach różnych gatunków ptaków, a także wykorzystuje algorytmy sztucznej inteligencji, aby pomóc użytkownikom w identyfikacji ptaków na podstawie zdjęć lub nagranych dźwięków.

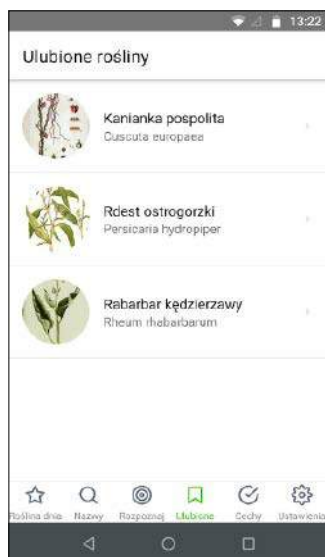
Aplikacja Merlin Bird ID oferuje również funkcję „Explore Birds”, która umożliwia użytkownikom przeglądanie listy ptaków występujących w ich okolicy i uzyskiwanie informacji o nich, takich jak opisy, zdjęcia i mapy występowania. Program ten uznaje się za szczególnie przydatny dla początkujących ornitologów i dla osób, którzy chcą rozwijać swoją wiedzę na temat ptaków.

Merlin Bird ID	
Producent	Cornell Lab of Ornithology
Platforma	Android, iOS
Oceny	Możliwości 9/10
	Łatwość obsługi 9/10
	Ocena ogólna 9/10

Smartfony i ich systemy operacyjne, czyli słówko o platformach

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

- **iOS** – system firmy Apple (tej od komputerów Macintosh), instalowany w urządzeniach iPhone, iPod Touch, iPad;
- **Android** – system firmy Google, niektórzy twierdzą, że wkrótce podbije cały świat. Rzeczywiście, Android jest coraz częściej instalowany w smartfonach m.in. takich firm, jak Huawei, HTC, LG, Motorola, Samsung, Sony Ericsson, ZTE (a także, co oczywiście, w smartfonach firmy Google);
- **Symbian** – system operacyjny open source (czyli bezpłatny i z tzw. otwartym kodem), obecnie najczęściej spotykany w telefonach firmy Nokia. Inne, mniej popularne systemy operacyjne dla telefonów komórkowych, to:
- **Bada** – system rozwijany przez firmę Samsung;
- **Windows Phone** – system firmy Microsoft, następca Windows Mobile, czyli po prostu Windows do urządzeń przenośnych;
- **BlackBerry** – system kanadyjskiej firmy Research In Motion, przeznaczony przede wszystkim do zastosowań biznesowych, instalowany w produkowanych przez nią smartfonach z charakterystyczną, pełną klawiaturą QWERTY. Także w niektórych telefonach innych firm (HTC, Motorola, Nokia, Samsung, Sony Ericsson).



Atlas roślin

Aplikacja o tej dość jasno mówiącej o co chodzi nazwie, została stworzona przez Polskie Towarzystwo Botaniczne, która służy do identyfikacji gatunków roślin w Polsce. Aplikacja ta umożliwia użytkownikom rozpoznanie roślin na podstawie zdjęć lub opisów tekstowych, zawierając również informacje, zdjęcia i mapy występowania poszczególnych gatunków.

Program oferuje rozpoznawanie roślin na podstawie zdjęcia, oznaczanie roślin na podstawie cech, dodawanie sugestii identyfikacyjnych pomagających innym użytkownikom, przeglądanie taksonomii w celu łatwiejszej identyfikacji gatunków, albo po prostu możliwość pozostawienia zdjęcia, aby inni użytkownicy pomogli w określeniu gatunku.

Zdjęcia dostarczane przez użytkowników tworzą... „Atlas roślin” właśnie. Dostępny także w trybie offline. Są klasyfikowane nie tylko naukowo, taksonomicznie, ale również według bardziej potocznych kryteriów, np. na lecznicze, trujące, miododajne, wodne, nadwodne, drzewa i krzewy, trawiaste i według podobnych, niekoniecznie ścisłych ale zgodnych z potocznymi klasyfikacjami reguł.

Atlas roślin

Producent	Atlas roślin
Platforma	Android
Oceny	
Możliwości	6/10
Łatwość obsługi	8/10
Ocena ogólna	7/10



Seek

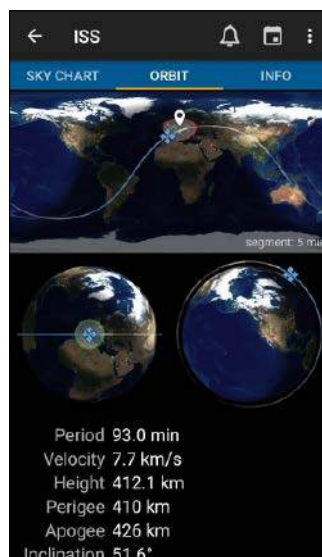
Jedna z najlepszych aplikacji mobilnych, umożliwiających użytkownikom identyfikację różnych gatunków roślin i zwierząt w oparciu o zdjęcia zrobione smartfonem. Aplikacja została stworzona przez organizację non-profit o nazwie iNaturalist, która zajmuje się badaniem i ochroną różnorodności biologicznej.

Działą w prosty sposób – po zrobieniu zdjęcia rośliny lub zwierzęcia, aplikacja automatycznie przeprowadza analizę obrazu i próbuje zidentyfikować gatunek. Aplikacja wykorzystuje technologię sztucznej inteligencji do rozpoznawania obiektów i posiada bazy danych zawierające informacje o setkach tysięcy różnych gatunków.

Seek nie tylko umożliwia użytkownikom identyfikację roślin i zwierząt, ale również zachęca do nauki i eksploracji przyrody. Dzięki aplikacji użytkownicy mogą zdobywać punkty za zidentyfikowanie różnych gatunków i odkrywanie nowych miejsc, a także włączyć się w projekty naukowe, które mają na celu zbieranie danych o różnorodności biologicznej.

Seek

Producent	iNaturalist
Platforma	Android, iOS
Oceny	
Możliwości	9,5/10
Łatwość obsługi	8,5/10
Ocena ogólna	9/10



Heavens-Above

Program umożliwia użytkownikom śledzenie ruchu satelitów, gwiazd, planet, Księżyca i innych ciał niebieskich na niebie. Aplikacja ta pozwala użytkownikom również na sprawdzenie godzin wschodów i zachodów słońca, faz Księżyca, a także na określenie położenia ciał niebieskich na niebie w czasie rzeczywistym. Heavens Above wykorzystuje dane z różnych źródeł, takich jak amerykańska agencja kosmiczna NASA i Europejska Agencja Kosmiczna.

Dzięki niej użytkownicy mogą dowiedzieć się, kiedy i gdzie będzie możliwe zobaczenie przelotu Międzynarodowej Stacji Kosmicznej, czy też jakie gwiazdy i planety będą widoczne na niebie w danym miejscu i czasie. Oferuje również interaktywne mapy nieba, które pozwalają użytkownikom na łatwe odnalezienie interesujących ich ciał niebieskich.

Warto zwrócić uwagę na pewne drobne dodatki, cenione przez wielu miłośników astronomii i nie tylko. Heavens Above serwuje pewne dodatkowe funkcje, takie jak przewidywanie przelotów satelitów na podstawie lokalizacji użytkownika i informacje na temat satelitów, które nie są widoczne na niebie, ale znajdują się w orbicie wokół Ziemi.

Heavens-Above

Producent	Heavens-Above
Platforma	Android
Oceny	
Możliwości	7/10
Łatwość obsługi	9/10
Ocena ogólna	8/10

Chemia jajka (2)

W odcinku z ubiegłego miesiąca przekonałeś się, że nawet skorupka zwykłego jajka jest ciekawym obiektem doświadczeń (czy kolorowe jajka się spodobają?). Bieżąca część artykułu poświęcona będzie chemicznemu spojrzeniu na zawartość, którą ona chroni. Przed tobą białko i żółtko.

Oba składniki, a zwłaszcza żółtko, to rzeczywiście najcenniejsze składniki jajka, wszak z nich musi wyrosnąć nowy organizm. Obie części to również wdzieczny obiekt do domowego eksperymentowania, zapraszam zatem do wspólnej pracy.

O jajku statystycznie

Przeciętne jajko kurze waży 50–60 gramów przy objętości wynoszącej około 50 cm³. Skorupka chroniąca wewnątrz stanowi w przybliżeniu 10% masy całości, z czego, jak już wiesz, 95% to węglan wapnia (z drobną domieszką innych składników nieorganicznych), a reszta – związki białkowe. Znajdujące się pod skorupką białko to w przybliżeniu 60% masy całego jajka. Jest ono około 15% roztworem białek zapewniającym ochronę żółtka przed wstrząsami (widoczne w surowym białku „sznureczki” – chalazy – utrzymują żółtko w położeniu centralnym) oraz rezerwuarem wody. Żółtko, stanowiące około 30% masy całego jajka, to z kolei gęsta mieszanina wody (mniej niż połowa masy) z białkami ($\frac{1}{3}$ suchej masy) i lipidami ($\frac{2}{3}$ suchej masy) oraz związkami nieorganicznymi.

Jajko wybite ze skorupki składa się w $\frac{3}{4}$ z wody, w reszcie zaś po połowie z białek i lipidów. Całe jajko dostarcza około 7 gramów białka o bardzo dobrym zestawie aminokwasów, dużej ilości niezbędnych nienasyconych kwasów tłuszczowych (NNTK), witamin A, B, D, E i K, luteiny korzystnie wpływającej na wzrok, lecytyny i choline poprawiających zdolności uczenia się oraz składników mineralnych (Fe, Mg, Ca, P, Na, K, S) (1).

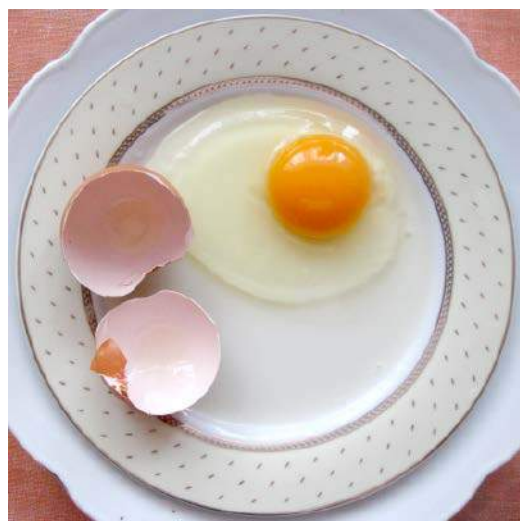
W jajku znajduje się również cholesterol, składnik naszej diety mający od dawna zdecydowanie złą prasę. Obecnie jednak odchodzi się od eliminowania jaj z jadłospisu, a Światowa Organizacja Zdrowia wręcz zaleca, aby zdrowi ludzie jedli od jednego do dwóch jaj dziennie (tygodniowo 7–10). Oczywiście, również w tym przypadku najważniejszy jest rozsądek i zróżnicowana dieta zawierająca możliwe mało przetworzone składniki.

Żółtawe białko

Kolor półpłynnej części jajka pochodzi od ryboflawiny, czyli witaminy B₂, obecnej w dość dużej ilości w białku. Po ugotowaniu barwa zanika i białko rzeczywiście staje się białe. Z białkiem jaja wykonałeś kilka doświadczeń przy okazji eksperymentów z ołowiem (numery 10–12 „Młodego Technika” z ubiegłego roku): otrzymywanie siarczku ołowiu i nieodwracalną denaturację pod wpływem jego soli. Powtórz fragment drugiej próby.

Przygotuj dość stężony roztwór białka jaja, mieszając je z wodą w proporcji 1:2–1:4 (więcej wody). Sporządź również nasycony roztwór soli kuchennej NaCl. Gorąca woda przyspieszy rozpuszczanie soli, ale praktycznie nie zwiększy jej stężenia – rozpuszczalność chlorku sodu jedynie w bardzo niewielkim stopniu zwiększa się wraz ze wzrostem temperatury. Przed próbą ostudź roztwór tak, aby nie „ugotować” białka.

Roztwór białka jest dość gęsty, ciągnący się, w wyraźnie żółtawym kolorze. Zaczynaj dodawać po kropli roztwór soli i cały czas mieszaj zawartość naczynia. Jeżeli spodziewałeś się wyniku podobnego,



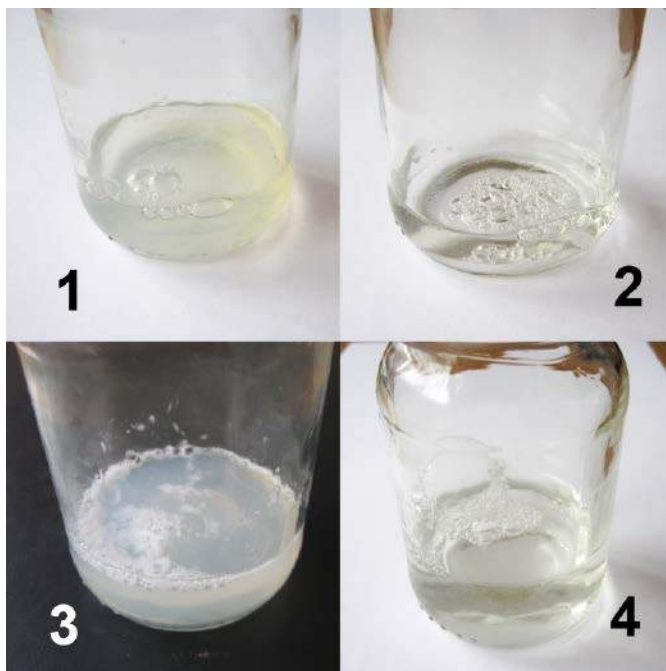
1. Jajka – ciekawy obiekt do eksperymentowania

jak w doświadczeniu z koagulacją białka pod wpływem soli, czyli wytrącania osadu, możesz być zaskoczony. Mieszanina początkowo staje się jaśniejsza, bardziej ruchliwa, mniej gęsta. Dopiero kolejne porcje solanki wywołują oczekiwany efekt – w cieczy pojawiają się kłaczkki białka. Jeżeli teraz dolejesz czystej wody zamiast roztworu NaCl, wytrącone białko rozpuści się.

Cząsteczki białek, ze względu na swoje gigantyczne rozmiary, nie są w stanie utworzyć roztworu właściwego w wodzie, jak np. sól kuchenna czy cukier, lecz **roztwór koloidalny**. Pierwsza porcja solanki wywołała zjawisko określane jako **wsalanie**. Molekuły białka mają liczne fragmenty polarne, w związku z czym przyciągają jony, najczęściej jednego znaku. W ten sposób jednoimiennie naładowane cząsteczki odpychają się od siebie, co zapobiega ich „sklejaniu” i zwiększa rozpuszczalność białka. Dalszy wzrost stężenia jonów działa jednak przeciwnie: otaczają się one cząsteczkami wody „zabieranymi” molekułom białka. Te zaś zbliżają się do siebie, łączą i **koaguluja** (wytrącają z roztworu). Zjawisko to określa się z kolei jako **wysalanie**. Dodatek większej ilości wody powoduje **peptyzację** osadu, czyli ponowne rozpuszczenie białek (2).

Zjawisko wysalania jest wykorzystywane do rozdzielania białek na poszczególne frakcje (rozpuszczanie i wytrącanie zachodzi przy różnych stężeniach soli i odczynie roztworu), ale i w sztuce kulinarnej. Podczas ubijania piany z białek przepisy zalecają dodanie szczypty soli – cząsteczki białek zaadsorbują jony na swojej powierzchni i nie będą miały tendencji do sklejania się, a w efekcie piana stanie się trwalsza. Pamiętaj jednak, aby dodać tylko odrobinę soli, inaczej powstaną kłaczkki wytrąconego białka.

Białko jaja jest wdziecznym obiektem jednej z reakcji charakterystycznych dla białek – **próby biuretowej**. Do białka rozpuszczonego w wodzie dodaj porcję roztworu wodorotlenku sodu NaOH, następnie zaś roztworu siarczanu(VI) miedzi(II) CuSO_4 . W naczyniu powstanie fioletowe zabarwienie, a w przypadku większego stężenia białka – również „kluch” zdenaturowanej substancji organicznej. Reakcja zachodzi w obecności białek i peptydów, a wykrywa obecność wiązań peptydowych położonych obok siebie i przedzielonych co najwyżej jednym atomem węgla. Powstający związek



2. Roztwór białka jaja w wodzie (1) po dodaniu niewielkiej ilości solanki staje się bardziej przezroczysty (wsalanie, 2). Większa ilość roztworu soli powoduje wytrącanie osadu białka (wysalanie, 3). Po dodaniu czystej wody osad rozpuszcza się (peptyzacja, 4)

kompleksowy miedzi(II) ma fioletowe zabarwienie. Próbkę wykorzystuje się m.in. w analityce medycznej do wykrywania białek w płynach ustrojowych (intensywność barwy jest proporcjonalna do ich stężenia) (3).

Natomiast **próba ksantoproteinowa**, druga z reakcji służących do wykrywania białek, nie powiedzie się w przypadku białka jaja kurzego. Reakcja ze stężonym roztworem kwasu azotowego(V) HNO_3 prowadząca do żółtego zabarwienia substancji białkowej (po dodaniu wodorotlenku sodu zmienia się ono



3. Próba biuretowa

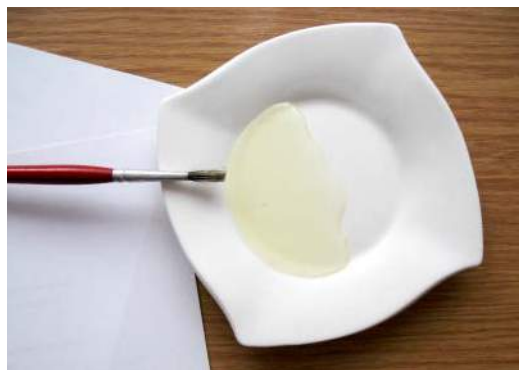
na pomarańczowe) wymaga obecności w niej aminokwasów zawierających pierścień benzenowy: fenylalaniny, tryptofanu i tyrozyny. W **albuminach**, głównym składniku białka jaja, znajduje się bardzo mało tych aminokwasów i wynik próby jest niewiódący (ze spożywczych surowców do jej wykonania bardzo dobrze sprawdzi się biały twaróg).

Z białka jaja można zrobić bezy, ale ma ono także ...

...zastosowania nie tylko kulinarne

Jeżeli zabrakło ci kleju do papieru, możesz użyć białka. Zmieszaj je z niewielką ilością wody, posmaruj kartki i docisnij. Po wyschnięciu klej dobrze trzyma (4). Nie śmiej się z tego zastosowania, białka mają właściwości klejące, w średniowieczu dodawano je do zaprawy murarskiej. Dowodem trwałości tak przygotowanych spoin są katedry, monumentalne budowle od wieków wzbudzające podziw w całej Europie. Klejące właściwości białka jaj wykorzystano oczywiście i w kuchni. Są one koniecznym dodatkiem do wypieków oraz potraw z mięsa mielonego, pełniąc w nich funkcję spoiwa. Smarowanie zaś białkiem powierzchni bułek drożdżowych przed wypiekiem zapewnia im ciemną, chrupiącą skórkę z charakterystycznym połyskiem. Przemiana zachodząca podczas pieczenia nosi nazwę **reakcji Maillarda**, a wiele produktów spożywczych zawdzięcza jej smak i zapach.

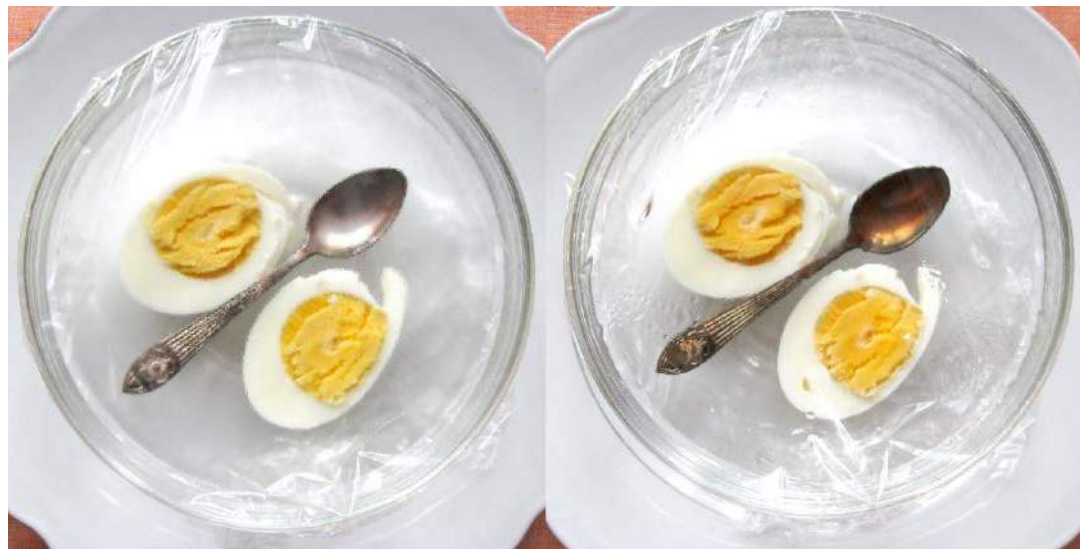
Białko jaja zawiera enzym o nazwie **lizozym**, który niszczy bakterie dostające się przez porowatą skorupkę (w warunkach naturalnych jaja znajdują się w środowisku dalekim od sterylności, stąd konieczność mycia



4. Białko jaja to dobry klej do papieru

ich przed użyciem w kuchni). Białka jaja możesz użyć na oparzoną lub podrażnioną skórę, co złagodzi ból, przyspieszy gojenie i zapewni jałowość. W warunkach polowych (np. biwak czy wycieczka) jako prowizoryczny plaster na ranę sprawdzi się błona z wewnętrznej strony skorupki.

Białkiem jaja wyczyścisz również przedmioty ze skóry. Wręcz przeciwnie działa ono natomiast na wyroby srebrne – te z kolei ciemnieją pod wpływem związków siarki zawartych zarówno w białku, jak i żółtku. Nie powinno być zatem zaskoczeniem, że srebrna czy też posrebrzana łyżeczka używana do jedzenia jajek na miękko po pewnym czasie pokryje się ciemnym nalotem. Efekt możesz wykorzystać celowo do nadania przedmiotowi ze srebra wiekowego wyglądu. Wystarczy powlec go ubitym na pianę białkiem lub nawet pomalować samym białkiem, aby pokrył się warstwą patyny (5).



5. Inny sposób patynowania srebrnej łyżeczki to umieszczenie jej w naczyniu z jajkiem ugotowanym na twardo (pojemnik przykryty folią). Po godzinie łyżeczka jest ciemnoszara

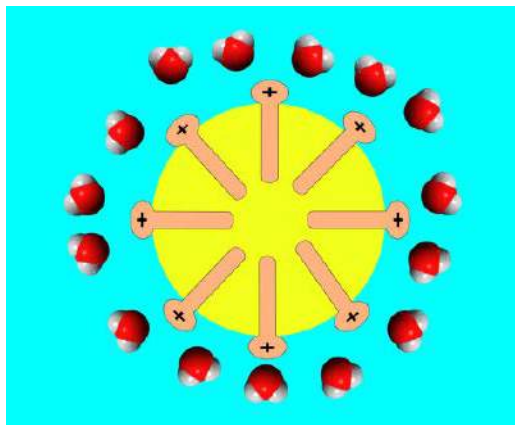


6. Żółtko jaja w roli emulgatora: po prawej widoczna warstwa oleju roślinnego na powierzchni wody, po prawej jednolity roztwór po dodaniu żółtka i wymieszaniu

Najcenniejsza część...

...jaja to żółtko. Oprócz ogromnych wartości odżywczych, jego właściwości fizykochemiczne powodują, że jest szeroko używane w przemyśle spożywczym i domowej kuchni. Niech jednak informacje teoretyczne potwierdzi praktyka. Przygotuj słoik, nieco oleju roślinnego oraz żółtko jaja. To ostatnie możliwie starannie oddziel od białka (możesz pobrać próbkę ze środka żółtka). Do słoika wlej wodę i dodaj pół łyżeczki oleju. Ciepły tłuszcz nie miesza się z wodą i utworzy górną warstwę o żółtawym zabarwieniu. Do naczynia wlej teraz pół łyżeczki żółtka i dobrze wstrząśnij zawartością. Jeżeli żółtko zawierało domieszkę białka, na powierzchni utworzyła się piana. Spójrz na ciecz w słoiku: ma ona bladeżółte zabarwienie, ale tworzy jednolity roztwór, bez oddzielnej warstwy oleju (6).

Zgodnie ze starym, jeszcze alchemicznym sposobem rozpuszczenia się w podobnym) niepolarny olej roślinny nie miesza się z wodą o polarnym charakterze. W doświadczeniu żółtko jaja spełniło funkcję **emulgatora** i umożliwiło połączenie tłuszczu z wodą – żółtawy roztwór, który otrzymałeś po zmieszaniu, to **emulsja typu O/W**, czyli olej rozpuszczony w wodzie. Za takie działanie żółtka odpowiedzialny jest jego składnik o nazwie **lecytyna**. To pochodna zwykłych tłuszczów, mająca budowę podobną do... mydła. Działanie również jest analogiczne: niepolarna część cząsteczki wnika do tłuszczu, natomiast polarna skierowana jest w stronę wody. W efekcie kropelki oleju są rozproszone w całej objętości roztworu, a jednoimienny ładunek na powierzchni zapobiega ich połączeniu (7). W taki sposób żółtko

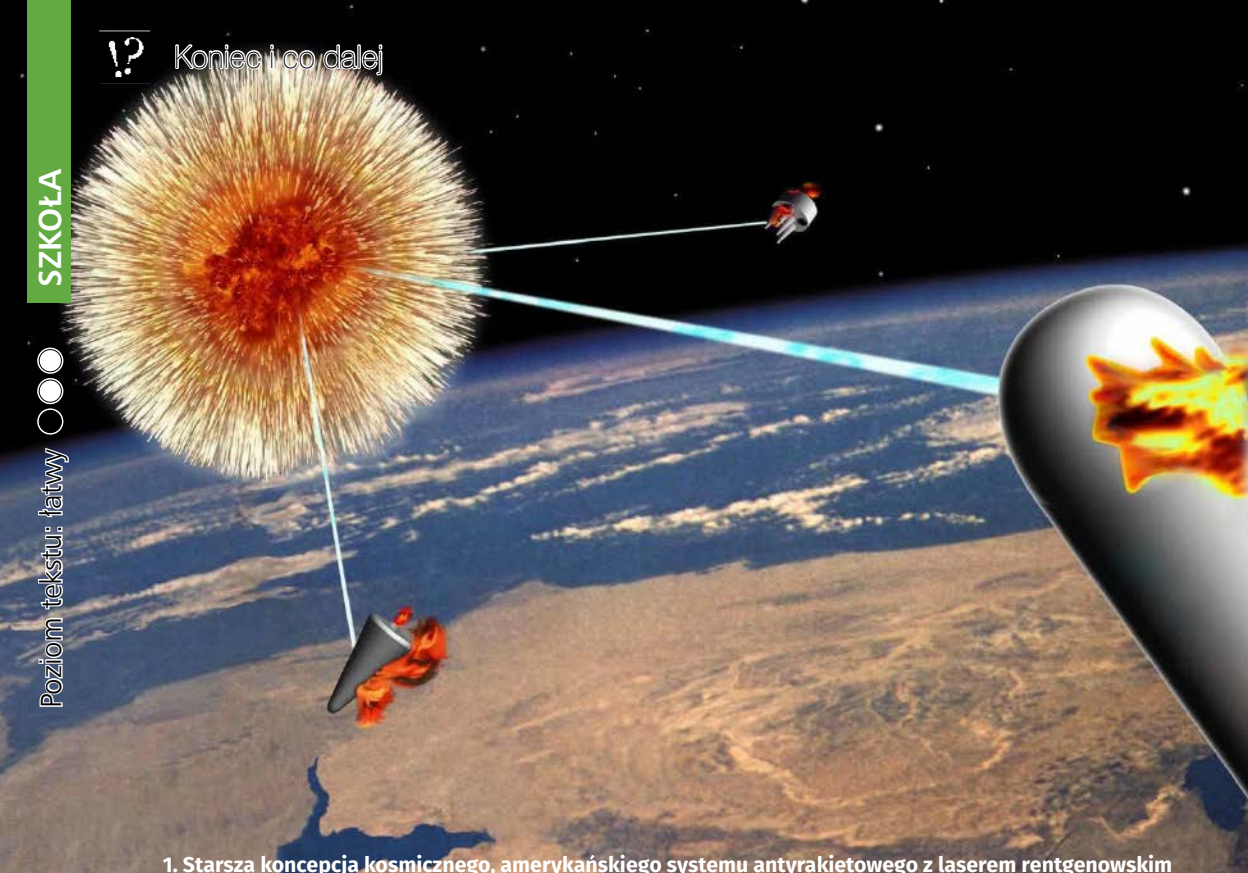


7. Schemat miceli: do kropelki tłuszczu wnika niepolarne fragmenty cząsteczki lecytyny, części polarne naładowane dodatnio wystają na zewnątrz i otaczają się molekułami wody (w stronę lecytyny kieruje się atom tlenu, czyli ujemny koniec cząsteczki wody)

jaja wykorzystywane jest do otrzymywania majonezu, lodów czy też kremów, czyli wszędzie tam, gdzie zachodzi konieczność połączenia tłuszczu z roztworem wodnym w jednolitą masę. Z kolei przykładem niekulinarnej zastosowania żółtka jest wykorzystanie go jako odżywki do włosów.

Jak zatem widzisz, nawet zwykłe jajko może stać się obiektem ciekawych i pouczających doświadczeń oraz wyjaśnić tajniki otaczającego świata. Nie będzie przesadą stwierdzenie, że chemik zawsze znajdzie coś interesującego ze swojej dziedziny. ■

Krzysztof Orlński



1. Starsza koncepcja kosmicznego, amerykańskiego systemu antyrakietowego z laserem rentgenowskim

Pokój w kosmosie

Zimna wojna stopniowo rozgrzewana do gorącej

Rywalizacja i wręcz wyścig zbrojeń mocarstw w kosmosie (1) to rzecz znana od połowy XX wieku. Jednak żadnych działań zbrojnych, „fizycznej” wymiany ciosów, bezpośrednich działań o charakterze militarnym w przestrzeni kosmicznej, jak się wydaje, nie było. Wydarzenia ostatnich lat, zestrzeliwanie satelitów lub manewry wymierzone w statki szpiegowskie, zmieniają sytuację.

Były major Sił Powietrznych USA Even „Jolly” Rogers obawia się gorącej wojny kosmicznej. USA, jak mówi, już teraz „aktywnie konkurują z Rosją i Chinami o swobodę działania i dominację w domenie kosmicznej, i rozwija się ona bardzo szybko”.

Sam jest współtwórcą start-upu True Anomaly, Inc, którego celem jest, jak sam pisał na Twitterze, „rozwiązywanie najtrudniejszych dla amerykańskiej armii

problemów wojny orbitalnej”. True Anomaly planuje wystrzelić w październiku 2023 r. na pokładzie rakiety SpaceX dwa statki kosmiczne o nazwie Jackal (ang. „szakal”), które przeznaczone są do prowadzenia „orbitalnych pościgów” na niskiej orbicie okołoziemskiej (2). Szakale nie będą wyposażone w działa, głowice bojowe czy laserowe, ale będą zdolne do operacji określanych jako „rendez-vous proximity” (RPO),

polegającej na manewrowaniu w pobliżu innych satelitów i czynnej ingerencji w systemy przeciwnika za pomocą układu czujników. True Anomaly twierdzi, że może to pozwolić penetrować systemy nadzoru i uzbrojenia kosmicznych jednostek przeciwnika lub pomóc w przechwyceniu nieprzyjacielskiej komunikacji.

W swojej pierwszej misji, nazwanej Demo-1, szakale będą jedynie szpiegować siebie nawzajem, używając silników odrzutowych, radaru i kamer wielozakresowych, dążąc do zbliżenia się do siebie na odległość nawet kilkuset metrów. Jeśli testy się powiodą, to jak przewiduje Rogers, amerykańskie siły zbrojne będą mogły rozmieścić na orbicie tysiące autonomicznych statków kosmicznych, kontrolowanych przez operatorów-ludzi na Ziemi oraz przez algorytmy AI, „by ścigać nieprzyjaciół, gdziekolwiek lecą”.

„Jeśli poważnie traktujesz zadanie obrony i ochrony domeny kosmicznej, to musisz mieć zdolność do wykonywania manewrów i siłę ognia”, mówi, cytowany przez magazyn „Wired” w lutym 2023 r., Rogers. Tu należałoby wyjaśnić, że choć wojskowi często używają pojęcia „ogień”, który przychodzi na myśl broń kinetyczną, palną, w kontekście kosmicznym zwykle odnosi się do zagłuszania, wojny elektronicznej i cyberataków. Rogers, wyjaśniając wagę tych działań, odwołuje się do potencjalnego zagrożenia, jakie dla formacji wojskowej, np. lotniskowca i bazującej na nim grupy uderzeniowej, stanowić może system satelitarny. Sugeruje więc, że szakale, aktywnie atakując satelity szpiegowskie, stanowiłyby osłonę dla operacji naziemnych (i nawodnych).

Armia amerykańska już od pewnego czasu pracuje nad planami stworzenia systemów przeszkadzających albo wręcz uniemożliwiających działanie satelitów wroga. I wcale nie musi to obejmować przestrzeni kosmicznej. Jeden z projektów, HELEIOS (High-Altitude Extended-Range Long Endurance Intelligence Observation System) miałby być siecią balonów lub pojazdów na baterie słoneczne, które mogą unosić się na wysokości co najmniej 20 km od powierzchni Ziemi. HELEIOS to nie satelity szpiegowskie, ale działają jak satelita. Za jego pomocą armia byłaby w stanie monitorować, przechwytywać i blokować komunikację wroga.

Zbliżyć się, by...

W szerszym ujęciu, stosowanie metody zbliżeniowej RPO w rywalizacji w przestrzeni kosmicznej nie jest pomysłem nowym. W raporcie z września 2022 r. roku Secure World Foundation, prywatna fundacja promująca współpracę zamiast rywalizacji w przestrzeni kosmicznej, wymieniła dziesiątki wojskowych operacji RPO na orbitach geostacjonarnych i niskich



2. Logotyp systemu Jackal i amerykańska flaga

od czasów zimnej wojny. Większość z nich polegała na tym, że amerykańskie, rosyjskie lub chińskie statki kosmiczne podchodziły do satelitów drugiej strony, prawdopodobnie w celu „przyjrzenia się” im lub podsłuchania komunikacji.

Przypadki takie z ostatnich lat relacjonowaliśmy w MT; np. to, że gen. David Thompson, wiceszef operacji kosmicznych US Space Force powiedział w wywiadzie dla „Washington Post” w listopadzie 2021, że w 2019 roku Rosjanie umieścili na orbicie pewne urządzenie „niebezpiecznie blisko” satelity szpiegowskiego USA. Pojazd ten zmienił kurs i przeprowadził próbę broni antysatelitarnej, wyrzucając cel i strzelając do niego pociskiem. Warto przypomnieć też przypadek rosyjskiej sondy Kosmos 2542, która w 2019 r. wykonywała na orbicie zadziwiające, nigdy wcześniej niewidziane, manewry, „przeszkadzające” amerykańskiemu satelicie zwiadowczemu USA 245 w wykonywaniu jego zadań. Eksperci spekulowali, że rosyjski satelita podąża za USA 245 w celu zebrania danych o jego misji. Obserwując satelitę, Kosmos 2542 może być w stanie określić możliwości kamer i czujników amerykańskiej sondy. Za pomocą sondy o częstotliwości radiowej możliwe byłoby nawet nasłuchiwanie słabych sygnałów z USA 245, co mogłoby wskazać Rosjanom, kiedy amerykański satelita robi zdjęcia i jakie dane przetwarza.

Warto przypomnieć, że znacznie wcześniej, w sierpniu 2014 roku, inny rosyjski satelita Kosmos 2499 wykonał serię podobnie zagadkowych manewrów. Tajemnicze ewolucje rosyjskich satelitów nie ograniczają się jedynie do niskiej orbity okołoziemskiej – na orbicie geostacjonarnej statek, oficjalnie kojarzony z telekomunikacyjną konstelacją Łucz, a w istocie będący prawdopodobnie militarnym satelitą SIGINT o nazwie Olimp-K, zbliżał się w 2018 r. do innych satelitów (w tym włoskich i francuskich – nie tylko wojskowych).

Oczywiście, technika RPO może mieć pokojowe zastosowania. Planowane są np. holowniki kosmiczne do naprawy lub przemieszczania uszkodzonych satelitów albo usuwania niebezpiecznych śmieci kosmicznych. Wspomniana fundacja Secure World wspiera projekt Confers, którego celem jest ustalenie dobrowolnych standardów technicznych komercyjnego RPO.

True Anomaly jest jednym z około sześćdziesięciu członków Confers, jednak jej cele nie mają zbyt wiele wspólnego z pokojem.

Atomówką w Starlink?

Już wiadomo, chociażby po ponad roku wojny na Ukrainie, że konstelacje satelitów komunikacyjnych (3) takie jak Starlink firmy SpaceX służą, owszem, do komunikacji i dostarczania internetu tam, gdzie nie ma infrastruktury, ale także świetnie pomagają w realizacji celów wojskowych, zapewniając łączność dromom i systemom śledzenia pocisków i robiąc wiele innych rzeczy. Chińczycy już zastanawiają się, jak takie systemy zwalczać. W październiku 2022 r. naukowcy jednego z chińskich wojskowych laboratoriów nuklearnych (NINT) napisali w publikacji udostępnianej przez „South China Morning Post”, że detonacja ładunku jądrowego o mocy 10 megaton w pobliżu krawędzi przestrzeni kosmicznej mogłaby za pomocą chmury promieniowania szybko uszkodzić lub zniszczyć dużą liczbę satelitów na niskiej orbicie okołoziemskiej. Oczywiście pomysł wydaje się kontrowersyjny nawet dla samych Chińczyków, ale dobitnie wskazuje, że systemy satelitarne są dla rywalizujących z USA mocarstw problemem i że szuka się pomysłów na walkę z nimi.

Do zwalczania rozproszonych konstelacji takich jak Starlink w niewielkim stopniu nadaje się bardziej konwencjonalna naziemna broń antysatelitarna, rakiety ASAT i podobne środki. Satelity, o których mowa, są niewielkie i stosunkowo łatwo wymienialne, co zwiększa koszty ich niszczenia, a atak na kilka satelitów w niewielkim stopniu pogorszyłby możliwości konstelacji. W kwietniu Dave Tremper, dyrektor ds. walki elektronicznej w amerykańskim resorcie obrony, podkreślił odporność komercyjnej konstelacji na rosyjskie zagłuszanie. Derek Tournear, szef amerykańskiej wojskowej Agencji Rozwoju Przestrzeni Kosmicznej (SDA), zauważył, że Rosja nie zestrzeliła żadnego satelity Starlink, mimo wcześniejszych gróźb i posiadania środków, by to zrobić.

Wszystko to dzieje się w kontekście nowych działań prowadzonych przez rząd USA, które zmierzają do wprowadzenia ogólnoświatowego zakazu lub przynajmniej globalnego moratorium na niszczące testy broni antysatelitarnej w przestrzeni kosmicznej, które stanowią wyraźne zagrożenie dla komercyjnych i naukowych zasobów na orbicie.

Kosmiczne plany wojenne

O skomplikowanej fizyce wojny orbitalnej pisaliśmy w „Młodym Techniku” w listopadzie 2021 roku. Wygląda na to, że opisane przez nas wyzwania nie



3. Konstelacja satelitów

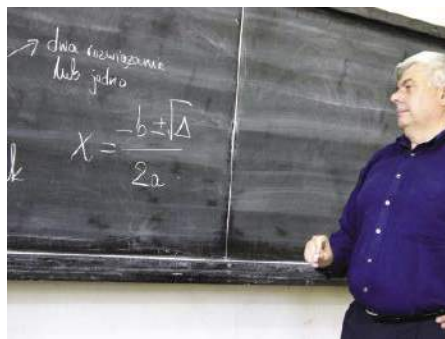
zrażają przygotowujących się do wojny w kosmosie. Techniki walki mogłyby przybrać wiele różnych form, niekoniecznie tzw. kinetycznych ataków; np. w ramach cyberwojny hakerzy mogliby włamać się do satelitów w celu ich demontażu lub przechwycenia. Lasery stacjonujące na ziemi mogą zostać używane do „oślepienia” czujników satelitarnych, co, według nieoficjalnych informacji, od pewnego czasu robią już Chińczycy amerykańskim satelitom. W przypadku bardziej otwartego konfliktu walczące strony mogą usiłować usunąć satelity wroga za pomocą rakiet ASAT albo satelitów manewrujących na orbicie, w atakach „kamikadze” lub za pomocą laserów, pocisków.

W obliczu zagrożeń wojsko USA postanowiło działać. Stąd działania i plany „wzmacniające” i zabezpieczające konstelację amerykańskich satelitów przed możliwym atakiem. Planuje się m.in. wyrzelenie dodatkowych satelitów, które w przypadku zniszczenia krytycznych systemów będą stanowiły kopie zapasowe, a także rozwój mniejszych satelitów, które będą trudniejsze do zaatakowania, jak również satelitów obronnych zdolnych do wykrywania lub nawet zapobiegania zagrożeniom na Ziemi i w kosmosie. Proponuje się także m.in. uaktualnienie GPS do tego, co nazywane jest GPS III – systemu bardziej odpornego na zagłuszanie. GPS był atakowany już niejednokrotnie.

Gdy pytamy o przyszłość „pokoju w kosmosie”, nieuchronnie przychodzi refleksja w postaci pytania, czy pokój tam kiedykolwiek panował. Realisci twierdzą, że nie było go tam nigdy, więc ma co się kończyć. W razie jednak wybuchu gorącej wojny, strać, eksplozji i wybuchów na orbicie, każdy powinien jednak dostrzec różnicę. ■

Miroslaw Usidus

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



Cudowne, z lekka bezsensowne zadania

Na obecnych wiosennych spacerach myślę o następnej książce, ale nie wiem, czy będzie miała powodzenie, bo zbieram różnego rodzaju uduchowione zadania. Nie wiem, jakie zadanie miało ową odpowiedź: „608 arbuźów”. I to oddzielne zadanie: „ułożyć zadanie, do którego odpowiedź brzmi: 608 arbuźów”.

Kilka następnych zadań pochodzi od Lewisa Carolla, autora *Alicji w Krainie Czarów*. Lewis Carroll był matematykiem, znanym także ze swoich abstrakcyjno-nonsensownych idei. Doradzał na przykład posługiwanie się mapami w skali 1:1 i sugerował zmianę oznaczeń $\sin$ i $\cos$ na... jakieś dziwne. W 1891 roku opublikował niewielką książeczkę *Problemy do poduszki*. Zawierała ona 72 zadania matematyczne, które rozwiązywał, gdy nie mógł spać. Już pierwsze zadanie... może nas nieźle uśpić. Ale spróbujmy przetrzymać ogarniającą nas senność.

Zadanie 1. Gdy 1 lipca mój zegar ścienny wybił godzinę ósmą, na zegarku kieszonkowym była 8.04. Poszedłem na spacer do Greenwich i gdy w obserwatorium było południe, mój zegarek pokazał 5.59. Dnia 30 lipca o godzinie 9 rano według zegara ściennego, zegarek ręczny pokazywał 8.57, a tego samego dnia w południe według zegarka ręcznego, na zegarze w Greenwich była już 12.05. Tego dnia o jedenastej wieczorem obydwa zegary (ręczny i ścienny) pokazały tę samą godzinę. Jaką godzinę pokażą zegary w prawdziwe południe (tj. GMT) 31 lipca?

Zadanie to przepisałem z... własnych notatek sprzed wielu lat. Zacząłem rozwiązywać i... doszedłem do wniosku, że właśnie w tej wersji jest to tzw. bardzo dobre zadanie. Treść może i nudnawa, ale...

Rozwiązanie jest ciekawsze od samego zadania. Przeczytajmy kilka razy. Widać, że „coś jest nie tak”. Pierwsze podejrzenie budzi godzina 5.59, gdy obserwatorium Greenwich podaje południe. Po drugie, jeżeli przez dwadzieścia dziewięć dni (od 1 do 30 lipca) były tak niewielkie różnice między zegarami, cóż może zająć w jeden dzień? Po trzecie: pierwszego lipca rano zegarek ścienny spieszy się względem ściennego cztery minuty, 30 lipca rano późni się trzy minuty a potem od razu wyrównują się. Co u licha?

Mamy do czynienia nie tylko z zadaniem matematycznym, ale i łamigłówką: „co się przeinaczyło przy kolejnych kopiowaniach tekstu? Nie umiem już dojść, czy błąd był w oryginalnej książce, czy w moich notatkach, kilkakrotnie przepisywanych. I to właśnie jest właściwe zadanie: jak było naprawdę? Jakie zadanie chciał ułożyć autor?

Na kursach, gdzie uczy się redaktorów książek, egzamin polega między innymi na zredagowaniu tekstu: przeczytania tekstu, wyłapania błędów i chropowatości. Jeżeli na przykład autor napisał, że benzyna była w powszechnym użyciu od końca XVIII wieku, redaktor „musi” to zauważyć i poprawić na „XIX”. Jeżeli napisze, że ci mali Murzyni należą do plemienia Masajów, redaktor musi poprawić na „Pigmejów” – Masajowie to właśnie plemię bardzo wysokie. I tak dalej.

Może i to jest pomysł na nauczanie „inne niż zwykle”. Nie „trzy z”: zakuj, zdaj, zapomnij, tylko „trzy p”: przeczytaj, pomyśl, popraw.

Wróćmy do zadania z zegarami. Nietrudno zgodzić się z tym, że tak zwany chochlik drukarski zamienił 11.59 na 5.59. Co z 30 lipca? Postawmy roboczą hipotezę: autorowi zadania chodziło o trzeci lipca, nie trzydziesty. Zero zostało „gdzieś, kiedyś” dopisane. Sprawdźmy tę hipotezę.

Kryptolodzy mawiali, że każdy szyfr da się złamać, każda zakodowana informacja da się odczytać – pod warunkiem że otrzymany tekst jest rzeczywiście komunikatem, że zawiera jakąś treść. Dzisiaj znamy szyfry, dla których prawdopodobieństwo złamania w czasie, powiedzmy, kilku lat (!) jest bliskie zeru, ale jeśli chodzi o proste kodowanie, pogląd kryptologów jest prawdziwy. Opisuje to z emocją Stanisław Lem (*Głos Pana*).

Załóżmy więc, że Lewis Carrol ułożył dobre, niezłośliwe zadanie. Mamy tylko je zrekonstruować, jak archeolog wazę ze skorupiek.

I może to jest pomysł..., może i to wpłatać w lekcje matematyki? Nie chwał się, że znasz algorytm. Najpierw odkryj, wysupłaj, wyjaw sens. Rozwiązać zadanie... to byle komputer potrafi. Ale znaleźć sens? Wytropić, „co autor miał na myśli”.

No, to do pracy. Przeróbmy 5.59 na 11.59, bo to wynika z kontekstu. Przetestujmy hipotezę, że „30 lipca”, powinno być „3 lipca”. Pierwszego lipca o ósmej rano zegarek ręczny wyprzedzał ścienny o cztery minuty, a 49 godzin później (3 lipca o dziewiątej) spóźniał się o trzy minuty. To nie może być przypadek, czujemy, że jesteśmy na właściwym szlaku, jak na odnalezionym z powrotem szlaku w górach. A zatem

dalej pracujemy przy założeniu, że ręczny spóźnia się względem ściennego o minutę na siedem godzin. Dość dużo, ale w XIX wieku nie było elektroniki, jaka dzisiaj oplątała cały świat.

Następny kłopot czeka nas 3 lipca wieczorem. Zegarek ręczny jakoś dziwnie przyspieszył. Analizujemy tekst: „w południe według zegarka ręcznego, na zegarze w Greenwich była już 12.05”. Zarówno Sherlock Holmes, jak i nawet Hercules Poirot (jak pamiętamy, był Belgiem) pomyślałby wtedy: co robi człowiek, który stwierdza, że jego zegarek późni się aż o pięć minut w stosunku do czasu, którym żyje całe Imperium? To proste. Człowiek ten ustawia zegarek na właściwy czas – popycha zegarek o pięć minut. Minuty sprzed dwóch dni nie uwzględnił. Tego nie było w treści zadania – to musieliśmy sami wyciągnąć z analizy tekstu. To dlatego o jedenastej wieczorem 3 lipca ścienny zrównał się z ręcznym: rano „sam z siebie” opóźniał się o trzy minuty, do wieczora opóźnił się jeszcze o dwie (bo upłynęło dwa razy po siedem godzin), no, ale w południe dostał „doping” pięciominutowy.

Teraz już prosto... Z początkowych danych wynika, że ręczny spóźnia się według Greenwich o 4 minuty na 48 godzin, a więc dwie minuty na dobę. 3 lipca w południe były nastawione tak samo. Zatem 31 lipca, po dwudziestu ośmiu dobach, zegarek ręczny wskaże 11.04... o ile nie będziemy go regulować, a to sugeruje zadanie. Co wskaże ścienny – pozostawiamy Czytelnikowi z komentarzem, że na pewno tak kiepskich zegarków nie było w wiktoriańskiej Anglii... Chyba że wymyślimy inną interpretację zadania.

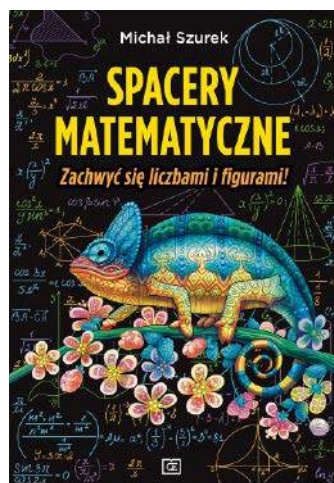
Zadanie 2 (też Lewis Carroll). Podróźni A i B idą tą samą drogą w tym samym kierunku, maszerując codziennie

Spacery matematyczne. Zachwyć się liczbami i figurami!

Michał Szurek

Wydawnictwo: Oficyna Edukacyjna Krzysztof Pazdro, liczba stron: 288, cena: 38,50 zł

Ukazała się niedawno moja książka *Spacery matematyczne*. Zacząłem o niej myśleć w czasie pandemii, gdy z braku możliwości dalszych wyjazdów, spacerowałem w najbliższej okolicy, przemierzając podwarszawskie pola. Na spacerach dobrze się myśli... Czy wszyscy Czytelnicy wiedzą, kim byli perypatetycy? Nie, to nie są fachowcy od montowania parapetów ani nie zwolennicy imprez zwanych parapetówkami. To byli uczniowie Arystotelesa, nazywani tak ze względu na zwyczaj prowadzenia wykładów i dysput filozoficznych w trakcie przechadzek. Takim perypatetykiem był pewien mój kolega, z którym kiedyś współpracowałem w pracy naukowej. Nie wiem, ile razy okrążyliśmy dziewiąte piętro Pałacu Kultury i Nauki w Warszawie, gdzie mieścił się nasz Instytut. Kolega miał lepszą kondycję fizyczną niż ja i pięćdziesiąt kótek nie robiło na nim wrażenia. Teraz sam chodzę na podobne spacery. Dobrze mi się też myśli... w wodzie, na pływalni. Aha, a *Spacery matematyczne* zacząłem zdaniem, zasłyszonym od uczennicy: „Lubię matematykę, bo tylko tam mogę kupić 608 arbużów i nikt nie zapyta, po co mi tyle?”



od 6 rano do 6 wieczorem. W poniedziałki A przechodzi 10 mil, we wtorki 9 mil, w środy 8 mil i tak dalej aż do soboty, gdy przechodzi 5 mil. Podróżny B zaczyna od 2 mil w poniedziałek, następnie we wtorek przebywa 4 mile, w środę sześć, w czwartek osiem, piątek dziesięć, w sobotę dwanaście. W niedzielę odpoczywają. Pewnego poniedziałku rano (o godz. 6) B wyprzedził A o 14 mil. Zakładając, że idą jednostajnym krokiem, wyliczyć, gdzie i kiedy się spotkają.

Komentarz: zadanie jest urocze (co docenimy, rozwiązując) i bawi nierealnym doбором danych. Lewis Carroll mało liczył się z realiami. Na przykład w poniedziałek podróżny B szedł 12 godzin jednostajnym krokiem i przeszedł dwie mile. Mila na sześć godzin to dwieście kilkadziesiąt metrów na godzinę... metrów, nie kilometrów! 12 mil w sobotę to też niewygodowane wymagania: 3 kilometry na godzinę, co prawda jednostajnie i bez przystanku....

Teraz rozwiązanie. Zauważmy, że tygodniowo A przebywał $10+9+8+7+6+5=45$ mil, a B tylko $2+4+6+8+10+12=42$, a zatem A nadrobił 3 mile na tydzień. W poniedziałek wieczorem A był już tylko 6 mil za B. We wtorek wieczorem tylko milę. Co się działo w środę? Do godziny 12 piechur A przeszedł połowę dziennego dystansu, czyli 4 mile, a podróżny też połowę swojej środowej porcji, czyli 3 mile. Wtedy się spotkali.

Ale co było dalej? Do końca dnia A nadrobił milę i tak wystartowali w czwartek rano. Czwartkowa porcja dla A to 7 mil, dla B osiem. A zatem czwartkowy wieczór znów ich połączył. Z treści zadania wynika, że A wypił więcej piwa niż B, bo w piątek włókł się (6 mil przez cały dzień), podczas gdy B „machnął” dziesięć. Nadrobił cztery, a w sobotę dołożył jeszcze siedem. Niedzielnny ranek zastał ich oddległych o 11 mil, a więc przez tydzień A nadrobił trzy mile... i powtarzamy obliczenia. Wynik: podróżny A doganiał B w środę w południe, w czwartek na samym końcu dnia, następnie we wtorek w drugim tygodniu o 15.34, w kolejny poniedziałek o 18.00 (na końcu dnia) i po raz ostatni w piątek ok. 12.26.

Bardziej jako ciekawostkę i żeby „ocalić od zapomnienia” niż do poważnej analizy przytaczam dwa zadania o pociągach z ponad 100-letniego zbiorku Stanisława Michalskiego *Nowy zbiór przykładów i zadań*



algebraicznych, wraz z teorią, wyd. M. Arcta, Warszawa–Poznań–Lwów–Równe–Lublin–Łódź–Wilno. Zbiorek był w użyciu jeszcze przed odzyskaniem niepodległości przez Polskę; mój egzemplarz, wyraźnie przeredagowany (wydanie czwarte), pochodzi z 1922 roku. Autor pisze, że zastosował się do wytycznych Ministerstwa z dnia 18 lipca 1921 roku. Zbiorek kupiłem za złotówkę w sklepie ze złomem (!) w Kazimierzu w 2004 roku. Aha, nie wiem, czy Czytelnicy wiedzą, że nie zrealizowałem chłopięcych marzeń: chciałem być konduktorem kolejowym, a zostałem profesorem matematyki. Nie wszystkie plany udaje nam się zrealizować. Ale oto zadanie:

Zadanie 3. O godzinie 8.40 wyszedł z Warszawy pociąg pospieszny do Poznania via Toruń. O tej samej godzinie wyruszył z Poznania balon sterowy typu „Zeppelin” i poszybował w kierunku Warszawy, lecąc stale ponad torem kolejowym również via Toruń. Szybkość pociągu pospiesznego 72 km, a balonu sterowego 80 km na godzinę. Pociąg stracił wskutek postoju na stacjach 1 godzinę 16 minut, natomiast balon szybował bez przerwy z jednakową szybkością. Odległość pomiędzy Warszawą a Poznaniem wynosi 380 kilometrów. W ile godzin po wzlocie balon znajdzie się nad parowozem pociągu pospiesznego?

Zadanie też urocze, również ze względu na wątki pozamatematyczne. Narzuca się pytanie: dlaczego pociąg pospieszny z Warszawy do Poznania jechał przez Toruń, a nie wprost, przez Konin, Koło i Swarzędz? Z oczywistej przyczyny: w XIX wieku z Warszawy do Poznania jechało się koleją kaliską albo właśnie przez Toruń. Obie trasy były międzynarodowe. Dopiero



w 1921 roku wybudowano linię z Konina do Poznania przez Strzałkowo, co było najkrótszym połączeniem Warszawy z Poznaniem. Nawet dzisiaj, kiedy wyjeżdżam z Poznania pociągiem w kierunku Warszawy, mam wrażenie, że jest to odgałęzienie od linii głównej na Toruń. Zbiorek, z którego biorę zadanie, został napisany kilka lat wcześniej i nie był poprawiany. A czy wszyscy wiedzą, co to był „balon sterowy typu Zeppelin”?

Zadanie ma też niejasną treść. Nie wiadomo, kiedy pociąg stracił ową godzinę i 16 minut. Przyjmijmy, że przed spotkaniem z balonem. Możemy więc założyć, że wyjechał już z Warszawy z tym opóźnieniem. Byłaby wtedy godzina 9.56. Do tego czasu balon przeleciał $101 \frac{1}{3}$ kilometra. I teraz już prosto: dzieliło ich $278 \frac{2}{3}$ kilometra przy szybkości względnej 152 km na godzinę. Trzeba podzielić:

$$\begin{array}{r} 278 + \frac{2}{3} \\ \hline 3 \\ \hline 152 \end{array}$$

...ojoj, jakie to trudne... równa się, o, jak prosto wychodzi: $\frac{11}{6}$ godziny; to znaczy godzina i 50 minut. Teraz też nie jest łatwo: jeżeli od 9.56 upłynęła godzina pięćdziesiąt minut, to która jest? Oczywiście trzeba dodać dwie godziny do 9.56 (będzie 11.56) i odjąć dziesięć minut. Jedenasta czterdzieści sześć.

No, to ja dam Czytelnikowi zadanie, na mój ulubiony temat: *ułoż zadanie*. Zmień jak najmniej treść zadania wyjściowego, by balon nadleciał nad parowóz punktualnie o godzinie 12.00. Możesz oczywiście przesunąć wszystko o te 14 minut. Ale może manipuluj prędkościami i opóźnieniem pociągu?

Zadanie 4. (Michalski, 1922). W poniedziałek o godzinie 12-iej w południe podług czasu miejscowego z miasta A wychodzi pociąg pospieszny, przebywający po 6 mil na godzinę i dąży w kierunku miasta B, leżącego na zachód od A na tymże równoleżniku. Pociąg przybywa do B w czwartek o godzinie 7 minut 26 rano czasu miejscowego (B). Jaka jest odległość między miastami A i B, jeżeli pociąg nie zatrzymuje się na przystankach i co mila, wskutek biegu w kierunku zachodnim, zyskuje 22 sekundy?

Zadanie jest bardzo ciekawe zupełnie niezależnie od swojej niebanalnej treści matematycznej. Daje odskocznik do innych problemów, pozwala wczuć

się w nie tak dawną przeszłość. Zastanówmy się najpierw nad znaczeniem słów „zyskuje 22 sekundy”. Przyjmujemy, że znaczy to, że co mila czas lokalny zmienia się o 22 sekundy. Gdy nie było nawet słowa „globalizacja”, a podróże trwały długo, regiony miały swój własny czas, niezwiązany z przynależnością do stref czasowych. Na przykład przewodnik Walerego Eljasza po Tatrach z 1891 roku podawał rozkład jazdy pociągiem do Tatr z komentarzem: „daty podane są według czasu krakowskiego; czas lwowski różni się od tego o 16 minut”. Czas na terenie całej Polski ujednolicono dopiero po odzyskaniu niepodległości. Zadanie to jest najwyraźniej reliktem z wcześniejszej edycji zbioru... i dobrze, że się zachowało. Jak i w wielu innych cytowanych przeze mnie zadaniach, ich autor mało dba o realia. Nasz pociąg jechał około 19 godzin bez przerwy.

O jakiej mili może być mowa w zadaniu? Były wtedy (pod koniec XIX wieku) na ziemi polskiej trzy mile: geograficzna=7,426 km, austriacka=7,586 km i polska=8,536 km. Najbardziej prawdopodobna (również ze względu na rozwiązanie) jest mila geograficzna.

Rozwiązujemy. Przyjmijmy, że maszynista nie przedstawiał zegarka. Która godzina była na jego zegarku, gdy po tak długiej podróży dotarł do celu? Trzeba do 7,26 dodać $\frac{11m}{30}$ minut, gdzie m jest liczbą przebytych mil. Dlaczego dodać? Wyobraź sobie, że lecisz samolotem do Anglii, gdzie czas jest o godzinę wcześniejszy...

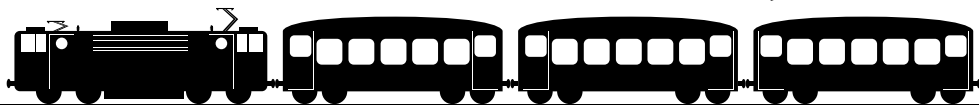
A więc pociąg był w drodze 19 godzin 26 minut + $\frac{11m}{30}$ minut. Lepiej przeliczyć to na minuty: $1166 + \frac{11m}{30}$ minut. Prędkość 6 mil na godzinę to jedna dziesiąta mili na minutę. Droga to prędkość razy czas;

$$1166 + \frac{11m}{30}$$

a zatem $m = \frac{1166 + \frac{11m}{30}}{10}$.

Wyznaczamy stąd $m = \frac{34980}{289}$ mil, czyli nieco ponad 121 mil, ok. 900 km.

Nasuwa się kilka komentarzy. Dzisiaj tak się nie układa zadań – uczeń jest przyzwyczajony, że odpowiedź nie powinna być „dzika”. Można dyskutować z tym poglądem, bo przecież w zadaniach „z życia” na ogół nie ma okrągłych liczb. Po drugie, sprawdźmy jednak rozwiązanie z realiami geograficznymi. Pociąg „zyskał” $\frac{34980 \cdot 22}{289}$ sekund, to jest blisko





trzy czwarte godziny. Zakładamy, że rzecz się dzieje w naszych szerokościach geograficznych. Długość równoleżnika krakowskiego to około 40 000 km razy $\cos(50^\circ)$, czyli około 25700 km (dlaczego? – a, to osobne zadanie!), czyli 900 km odpowiada w przybliżeniu 50 minutom różnicy w czasie lokalnym. Zadanie było zgodne z realiami!

Wykorzystajmy też cytowaną informację z przewodnika Walerego Eljasza, że „czas krakowski różni się od lwowskiego o 16 minut” do obliczenia odległości między tymi dwoma niegdyśniejszymi stolicami Galicji. Skoro 900 km odpowiadało 50 minutom, to 16 minut daje (jak to obliczyć?) około 288 km. A jak jest w rzeczywistości? Możesz sprawdzić? Jak? No wiesz, takie pytanie w XXI wieku?

I na zakończenie jeszcze dwa zadania. Pierwsze znów o pociągach.

Zadanie 5 (Michalski, 1922). Pociąg osobowy idzie z A przez B do C (przypuszczamy, że B i C leżą na prawo od A), zatrzymując się w B siedem minut. W 2 minuty po wyjściu z B pociąg ten spotyka inny pociąg pospieszny, idący na jego spotkanie z 2 razy większą szybkością. Pociąg pospieszny wyszedł ze stacji C w chwili, gdy osobowy znajdował się o 28 km na lewo od B. Wiemy, że pociąg pospieszny przebywa przestrzeń pomiędzy C i B w ciągu $1\frac{1}{2}$ godziny; następnie, gdyby pociąg ten po przybyciu do A i nie zatrzymując się na tej stacji, natychmiast wyruszył z powrotem, to przybyłby do C w 3 minuty po przybyciu do tej stacji osobowego pociągu. Ile kilometrów robi na godzinę każdy pociąg i jaka jest odległość pomiędzy A, B i C?

Wielu z Czytelników pomyśli sobie: „a cóż to za głupie zadanie – nic dziwnego, że matematyka

to ogłupianie ludzi za pomocą cyferek!” Każdy nauczyciel matematyki musi odpowiadać na pytanie: „Po co mam to robić – przecież to mi się w życiu nie przyda!”. Różnie na to odpowiadamy, ale niekiedy trafia do uczniów porównanie, że to po prostu trening intelektualny. Jeżeli nie potrafisz pomyśleć nad takim zadaniem, to może nie jesteś tak dobry intelektualnie, jak mi się zdawało? A sportowcy też robią na treningu ćwiczenia, których potem nie ma na meczu...

Rozwiązania nie podam.

Drugie zadanie pochodzi podobno z college'u w Cambridge, z... 1815 roku – jakże jest ono nieekologiczne. Dziś już nie strzela się do zwierząt dla rozrywki. Wiem, że tacy ludzie są i dziwią się temu.

I tu nie podam rozwiązania, nie chcąc zabierać Czytelnikom przyjemności myślenia nad tym cudośnie bezsensownym zadaniem.

Zadanie 6. Dwaj myśliwi przynieśli z polowania 10 ptaków. Suma kwadratów liczb strzałów, jakie oddał każdy z nich, wyniosła 2880, iloczyn zaś tych samych liczb był 48 razy większy od iloczynu liczby ptaków, jakie ustrzelił każdy z nich. Gdyby B oddał tyle strzałów co A, to ustrzeliłby 5 razy więcej ptaków niż A. Ile ptaków zabił każdy z nich?

Jedna rzecz od razu się narzuca. Pan A był kiepskim myśliwym. A co do zadania – cóż, taka była wtedy matematyka. Można się pośmiać z tych zadań, ale myślenia na pewno uczyły. Zresztą, naukę języków obcych traktowano też jako ćwiczenia umysłowe. Bo i kto 200 lat temu podróżował? Garstka arystokratów. ■

Sierociniec janczarów

Piotr Gajdziński

Wydawnictwo MUZA S.A., liczba stron: 384, cena: 44,90 zł

Do Wolsztyna, niewielkiej miejscowości w zachodniej Wielkopolsce, na zjazd z okazji rocznicy matury przyjeżdża dziennikarz Rafał Terlecki. Paweł Zdanowski, jeden z dawnych szkolnych kolegów, próbuje go zainteresować sprawą brutalnego morderstwa popełnionego w miasteczku dwa lata wcześniej. Ofiarą była szesnastoletnia dziewczyna, której zamordowane ciało znaleziono na terenie ogródków działkowych. Zabójca trafił do więzienia, ale wydaje się, że ta wstrząsająca historia ma szersze tło i sięga czasów przedwojennych. Wkrótce na jaw wychodzą przerażające fakty. W toku dziennikarskiego śledztwa Terlecki odkrywa, że przed kilkudziesięciu laty w Wolsztynie dochodziło do zbrodni tudzież przypominających morderstwo szesnastoletniej Joli. Sytuacja się komplikuje, bo wśród archiwalnych dokumentów Rafał natrafia na dobrze znane mu nazwiska. Jaką rolę w tej historii mógł odegrać jego własny ojciec i ojciec Pawła? Jakie znaczenie dla śledztwa mają losy pewnego opętanego księdza? Prawda okazuje się szokująca.



Czy substancja promieniotwórcza straci z czasem swoje właściwości? (1)

Jak zapewne wiadomo, cała materia w przyrodzie zbudowana jest z atomów. Z powodu różnic w ich budowie, przekładających się na różnice we właściwościach fizycznych i chemicznych, każdy atom możemy zaklasyfikować jako pewien pierwiastek.

Promieniotwórczość w przyrodzie

Niemniej nawet atomy tego samego pierwiastka mogą różnić się między sobą. Dla danego pierwiastka stała pozostaje liczba protonów w jądrze atomowym (a co za tym idzie – również liczba elektronów krążących wokół jądra). Liczba neutronów nie musi w każdym przypadku być identyczna i jeśli dwa atomy tego samego pierwiastka różnią się tą liczbą, nazywamy je izotopami.

Jakkolwiek różne izotopy jednego pierwiastka wykazują w zasadzie identyczne właściwości chemiczne, to już niekoniecznie zachowują się tak samo, jeśli chodzi o ich właściwości fizyczne. Każdy pierwiastek ma przynajmniej jeden izotop stabilny, który nie ulega żadnym rozpadom. Może również mieć kilka izotopów niestabilnych (promieniotwórczych), które

ulegają samoistnym przemianom, emitując bardzo przenikliwe promieniowanie jądrowe.

W wyniku różnego rodzaju przemian jądrowych wydzielą się pewna ilość energii, co znajduje zastosowanie na przykład w elektrowniach atomowych czy niektórych rodzajach silników napędzających pojazdy kosmiczne. Osobnym (i dosyć szerokim zagadnieniem) jest również wykorzystanie promieniotwórczości w terapii nowotworów czy diagnostyce medycznej.

Odrobina historii

Często odkrycie promieniotwórczości przypisuje się Marii Skłodowskiej-Curie. W rzeczywistości zjawisko to zostało odkryte przez Henryka Becquerela w 1896 roku. Skłodowska-Curie odkryła natomiast



Antoine Henri Becquerel (ur. 15 grudnia 1852 w Paryżu, zm. 25 sierpnia 1908 w Le Croisic) – francuski chemik i fizyk, laureat Nagrody Nobla w dziedzinie fizyki w 1903 za odkrycie promieniotwórczości. (źródło: wikipedia.org)

nowy pierwiastek rad, pracując pod okiem Becquerela i pisząc u niego pracę doktorską.

Prawo rozpadu promieniotwórczego

Załóżmy, że mamy pewną próbkę zawierającą w chwili początkowej N_0 jąder pewnego izotopu niestabilnego i chcielibyśmy wiedzieć, czy substancja ta już na zawsze pozostanie promieniotwórcza, czy też z czasem straci swoje właściwości. Wzór opisujący zależność liczby jąder promieniotwórczych w próbce od czasu ma postać $N(t) = N_0 e^{-\lambda t}$, gdzie λ jest pewną stałą, charakterystyczną dla danego izotopu. Wartość $N(t)$ z kolei to liczba jąder pozostająca w próbce po czasie t .

Osoby znające pojęcie funkcji wykładniczej i potrafiące zinterpretować powyższy wzór z pewnością zauważą, iż wyrażenie to zawsze przyjmuje wartość dodatnią. Choć wartość ta maleje z czasem, nigdy nie dochodzi do zera. Czy w takim razie żadna substancja promieniotwórcza nie rozpadnie się do końca i zawsze pozostaną w niej jakieś jądra niestabilne? Gdyby tak było w istocie, mielibyśmy poważny problem. Bo co w takim razie z odpadami z elektrowni jądrowych albo pacjentami, którzy musieli przyjąć pewną dawkę izotopu promieniotwórczego w ramach radioterapii?

Na szczęście wzór ten należy rozumieć jedynie jako opis matematyczny zjawiska rozpadu i pewne przybliżenie rzeczywistości. Uzyskane wyniki (bardzo często są to liczby niecałkowite) odzwierciedlają jedynie prawdopodobieństwo pozostania jakichś jąder promieniotwórczych w próbce, o czym będzie

można się przekonać, wykonując prostą symulację rozpadu jądrowego.

Sprawdź swoją wiedzę

Przed przystąpieniem do przeprowadzenia symulacji sprawdź, co pamiętasz z powyższego tekstu. W tym celu rozwiąż wykreślankę, znajdując pojęcia, które się w nim pojawiły. Szukane słowa mogły zostać zapisane poziomo, pionowo, ukośnie – również w odwrotnej kolejności liter. ■

Joanna Borgensztajn

J	M	N	A	R	T	L	T	E	N	I	S
U	A	A	C	A	U	N	O	T	O	R	P
R	K	D	P	D	A	J	N	U	R	E	K
A	A	I	R	I	K	O	A	S	T	A	L
P	R	A	W	O	R	O	Z	P	A	D	U
C	Y	T	L	T	A	N	D	E	T	A	C
E	Ż	I	U	E	T	T	N	Z	P	M	H
N	P	E	L	R	K	Ł	O	P	O	T	Y
A	N	T	E	A	R	I	A	M	T	O	M
B	U	T	K	P	U	Z	D	R	O	G	N
G	H	B	Ę	I	K	R	Y	M	Z	W	U
K	E	T	S	A	I	W	R	E	I	P	E

Wskazówka: każda litera „o” należy do przynajmniej jednego szukanego hasła.



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

W poprzednim numerze „Młodego Technika” pisałem o włoskich szachistach, którzy po zwycięstwie w 1575 roku na dworze króla Hiszpanii Filipa II, zapoczątkowali dominację włoskiej szkoły szachów w Europie. Pokonali oni wtedy Ruy Lopeza de Segura (1530–1580), słynnego hiszpańskiego mistrza i teoretyka królewskiej gry, uważanego do tychczas za najsilniejszego szachistę na świecie (1).

Ruy López de Segura – hiszpański ksiądz, czołowy szachista XVI wieku

Rodrigo (Ruy) López de Segura urodził się w 1530 roku w małym hiszpańskim mieście Zafra w prowincji Badajoz w regionie Extremadura w rodzinie zamożnych kupców Marii i Hernána López de Segura. Ruy López był wybitnym szachistą, księdzem katolickim, proboszczem parafii La Candelaria w Zafrze (2).

Od najmłodszych lat Ruy López był zakochany w szachach. Jednym z tych, którzy wywarli na niego największy wpływ, był Pedro Damiano, który w 1512 roku w Rzymie wydał *Questo libro e da imparare giocare*

a scachi (Książka do nauki gry w szachy). Podręcznik ten, ośmiokrotnie przedrukowywany, był najpopularniejszą książką szachową w pierwszej połowie XVI wieku. Przez ponad dwadzieścia lat Ruy López uważany był za najlepszego szachistę na świecie. Wielki miłośnik i mecenas szachów, hiszpański król Filip II, zaprosił Lopeza, aby dołączył do jego królewskiego dworu jako doradca i jego osobisty spowiednik. Pierwsze zarejestrowane partie Lopeza pochodzą z 1560 roku, kiedy papież Pius IV wezwał go do Rzymu w sprawach kościelnych. Tam po raz pierwszy



1. Ruy López de Segura, źródło: <https://bit.ly/3VCteq2>



2. Ołtarz główny w zbudowanej w 1546 roku kolegiacie Candelaria w Zafrze, źródło: <https://bit.ly/3M2oSWj>



3. Księga o odkryciach i sztuce gry w szachy autorstwa Ruy Lopeza, źródło: <https://bit.ly/3MojC5i>

López spotkał się z najsilniejszymi włoskimi szachistami i ich pokonał. Według włoskiego szachisty Alessandro Salvio, López podróżował do Włoch również w 1573 roku, gdzie również pokonał najsilniejszych włoskich szachistów.

W 1561 roku Ruy López wydał podręcznik do szachów pod tytułem *Libro de la invención liberal y el arte del juego del Axedrez* (Księga o odkryciach i sztuce gry w szachy) (3).

Książka Lopeza była znacznie dojrzsza od wcześniejszych książek szachowych i jest cennym podręcznikiem analizującym strategię i taktykę gry w szachy. Autor skoncentrował się bardziej na teorii niż na rozwiązywaniu konkretnych problemów szachowych. Była to pierwsza próba usystematyzowania i analitycznej oceny znanych otwarć szachowych. W podręczniku tym znajduje się m.in. analiza otwarcia nazywanego obecnie Ruy López (partia hiszpańska). López nie opracował tego debiutu, który zyskał popularność dopiero znacznie później, ale udoskonalił go i szczegółowo analizował i dlatego w literaturze anglojęzycznej otwarcie to nosi nazwę Ruy López. Do dziś partia hiszpańska, czyli otwarcie Ruy López, jest jedną z najpopularniejszych strategii debiutów w szachach, stosowaną m.in. przez aktualnego mistrza świata Magnusa Carlsena.

W swojej książce, podzielonej na 4 części, López przedstawił sześćdziesiąt sześć partii, z których dwadzieścia cztery zostały zaczerpnięte z książki Damiano z 1512 roku. Najpierw przedstawił kilka



4. Ruy López na znaczkach pocztowych, źródło: <https://bit.ly/3pgVD9e>



5. Spektakl żywych szachów w Zafrze, źródło: <https://bit.ly/42aVTVF>

mitologicznych źródeł gry i omówił jej zalety, zasady i strategię, przeplatając szeregiem cytatów (po łacinie) autorów klasycznych. Druga część koncentruje się na debiutach i jest dziedzictwem Lopeza jako „ojca teorii debiutów”. W trzeciej i czwartej polemizował z analizami Damiano. López rzadko prezentował warianty, które kończyły się matem. Zamiast tego kończył komentarzami, takimi jak np. że czarne muszą stracić swojego hetmana lub że białe mają bardzo dobrą partię.

W 1575 roku wziął udział w pierwszym międzynarodowym turnieju, który odbył się na dworze króla Filipa II. López zajął wtedy dopiero trzecie miejsce w turnieju po przegranej z włoskimi szachistami Leonardo di Bona i Paolo Boi.

Oprócz klasycznych partii López potrafił i lubił grać „na ślepo” (odmiana gry, w której gracze odwracają się od szachownicy i ogłaszają swoje ruchy, nie widząc ani nie dotykając figur). Ruy López uważany jest, przez wielu miłośników królewskiej gry, za pierwszego nieoficjalnego mistrza świata w szachach, następnym był Leonardo di Bona (od 1575 roku).

Ruy López jest obecnie bohaterem znaczków z różnych części świata, takich jak Kuba (1976), Kambodża (1986), Gwinea Bissau (1988), Laos 1988, Malediwy 2019, Mozambik 2019 i Hiszpania (2022) (4).



6. Ruy López na znaczku pocztowym, Hiszpania 2022, źródło: <https://bit.ly/3pgVD9e>



7. Pomnik Ruy Lopeza w Zafrze, źródło: <https://bit.ly/44CJ747>

W Zafrze wszystko emanuje szachami, rozgrywane są spektakle żywych szachów (5), w 2022 z inicjatywy lokalnego klubu szachowego wydano znaczek



9. Partia hiszpańska, pozycja po 5...Ge7 (wariant zamknięty)



10. Partia hiszpańska, pozycja po 5...S:e4 (wariant otwarty)



11. Partia hiszpańska, pozycja po 6...Gb7 (obrona archangielska)



8. Ruy López autorstwa Xulio Formoso, <https://bit.ly/42aVTVF>

pocztowy poświęcony genialnemu szachiście (w tle plansza z partią hiszpańską) (6), stanął jego pomnik w Parku Pokoju w tzw. Rincón de Ruy López (7).

Ruy López – partia hiszpańska

Partia hiszpańska nazywana też Ruy López na cześć XVI-wiecznego hiszpańskiego księdza Ruy Lopeza de Segura to szachowe otwarcie rozpoczynające się od posunięć **1.e4 e5 2.Sf3 Sc6 3.Gb5**. Ideą debiutu jest kontrolowanie centrum za pomocą pionków i szybkie rozwijanie figur. Ruch 3.Gb5 ma na celu wywarcie nacisku na piona e czarnych, a także przygotowanie się do roszady. Czarne mają teraz kilka możliwych odpowiedzi, z których najpopularniejszym jest **3...a6**, znane jako **obrona Morphy'ego**. Po 3...a6 najpopularniejsze jest **4.Ga4**, utrzymujące związek skoczka na c6. Czarne mogą wtedy zagrać **4...Sf6**, rozwijając kolejną figurę i przygotowując się do roszady.

Od tego momentu białe mają wiele możliwości kontynuowania gry, ale najpopularniejsza jest roszada **5.O-O**, na co czarne mogą odpowiedzieć: **5...Ge7** (diagram 9), także przygotowując się do swojej roszady (wariant zamknięty), **5...S:e4** (wariant otwarty,



12. Wariant wymienny w partii hiszpańskiej



13. Wariant wymienny w partii hiszpańskiej, pozycja po 5...Hd4



14. Obrona berlińska w partii hiszpańskiej

diagram 10) lub 5...b5 6.Gb3 Gb7 (obrona archangielska, diagram 11).

Partia hiszpańska może prowadzić do otwartych i złożonych pozycji, w których obie strony mają wiele możliwości kontynuacji gry.

Bobby Fischer często po 3...a6 grał białymi **wariant wymienny** 4.G:c6 (diagram 12) z następnym 4...d:c6 5.O-O f6. Wiele czołowych szachistów uważa obecnie, że wariant wymienny nie stwarza wystarczająco dużo trudności dla czarnych. Należy zauważyć, że pozorna groźba białych wygrania piona e za pomocą 5.Se5? jest iluzoryczna. Czarne mogą odpowiedzieć 5...Hd4 (diagram 13) i odzyskać materiał z lepszą pozycją.

Innym ciekawym wariantem partii hiszpańskiej jest obrona berlińska powstająca po posunięciach 1.e4 e5 2.Sf3 Sc6 3.Gb5 Sf6 (diagram 14). Debiut ten został po raz pierwszy dogłębnie przeanalizowany w XIX wieku i otrzymał swoją nazwę od berlińczyków, którzy badali jego warianty. Obrona berlińska stała się popularna po meczu o mistrzostwo świata PCA w szachach pomiędzy Władimirem Kramnikiem i Garrim Kasparowem w 2000 roku, w którym pretendent Kramnik czterokrotnie zremisował czarnymi, co pomogło mu w efekcie wygrać ten mecz. Obecnie coraz więcej arcymistrzów stosuje obronę berlińską

w grach turniejowych. Ze strategicznego punktu widzenia obrona berlińska jest zwykle używana przez zawodników, którzy chcą uzyskać remis jako czarne lub grać defensywnie i dochodzić do końcówek. Najczęstszym czwartym ruchem białych jest 4.O-O. Często stosowana kontynuacja to 4...S:e4 5.d4 Sd6 6.G:c6 d:c6 7.d:e5 Sf5 8.H:d8+ K:d8 (nazywana berlińską ścianą, diagram 15), która była rozgrywana we wszystkich czterech partiach meczu. Czarne nie mogą już roszować (choć przy braku hetmanów nie jest to tak bardzo ważne) i mają zdublowane piony. Plan dla czarnych to zamknięcie wszystkich dróg do swojego obozu i wykorzystanie przewagi pary gońców.

Partia hiszpańska kryje wiele pułapek, w które mogą wpaść mniej doświadczeni szachiści. Pisałem o wielu z nich w numerach MT 12/2019 – MT 3/2020. A to jeszcze jedna:

1.e4 e5 2.Sf3 Sc6 3.Gb5 Sf6 4.O-O S:e4 5.We1 d5 (diagram 16, błąd, należało odejść skoczkiem na d6 lub na f6) 6.d3 Sd6 7.G:c6 b:c6 8.S:e5 (diagram 17, grozi zdobycie hetmana, najlepsze jest teraz 8...Ge6 9.S:c6 Hd7 10.Se5 z dużą przewagą białych, nie ratuje też 8...Ge7 9.S:c6 Hd7 10.W:e7+). ■

Zadania do samodzielnego rozwiązania i rozwiązanie zadań z MT4/2023 na stronie 82



15. Berlińska ściana, pozycja po 10...K:d8



16. Pułapka w partii hiszpańskiej, pozycja po 5...d5?



17. Pułapka w partii hiszpańskiej, pozycja po 8.S:e5



Szkoła Wynalazców

dozwolone do lat 15

Mieścieście zadanie domowe z gatunku drobnych, praktycznych usprawnień, ułatwiających życie: *Zaproponować bezkolizyjne i wygodne umieszczenie wagi osobowej w małej łazience, zabezpieczające wagę przed zalaniem, a jednocześnie dające możliwość natychmiastowego korzystania.*

Dla wielu osób: sportowców, modelek, aktorek, ludzi chorych np. na cukrzycę, częste ważenie się jest nieodzownym elementem dbałości o zdrowie, wygląd bądź przestrzegania norm sportowych jak np. bokserów. Łazienki mieszkań budowanych w latach 60. do 80. ub. stulecia były małe, wręcz mikroskopijnie małe. Jeżeli w takiej łazience była wanna, pralka, umywalka, pojemnik na brudną bieliznę i odzież do prania, to już na wagę naprawdę nie było miejsca. Zwłaszcza że najlepiej byłoby, żeby waga – dziś niemal wyłącznie elektroniczna – leżała sobie poziomo, na posadzce. Typowa waga łazienkowa ma wymiary 30×30 cm. Nie jest to dużo, ale trzeba przewidzieć pewien luz na osobę ważącą się. W takiej sytuacji jedna z zasad trizowskich mówi: „rozwiązania należy szukać w przestrzeni”. W najprostszym wariancie oznacza to, że wagę należy umieścić pionowo.

Oznacza to jednak pewną niedogodność: wagę trzeba wziąć do ręki, położyć na posadzce w pozycji do ważenia. Waga stojąca pionowo narażona jest na upadek, poza tym musi być zabezpieczona przed oblanie strumieniem wody z prysznicza. Waga musi więc mieć jakiś stelaż. I co to naś czytelnicy?

Zbigniew Rataj – wagę można umieścić w plastikowej, wodoodpornej torbie i powiesić na haku, osadzonym w ścianie łazienki. Torba powinna być zamykana „kłapą” bez zatrasków, itp., żeby dało się jednym ruchem wyjąć wagę i włożyć z powrotem. Na dnie torby powinien być niewielki otwór, żeby woda, która się jednak może do niej dostać – mogła wypłynąć.

Dobra, praktyczna koncepcja. Zwraca uwagę prze-myślenie szczegółów jak: kłapa luźna bez zatrasków i otwór w dnie torby. To rzeczywiście są ważne rzeczy. Problemem może być jedynie próba szybkiego wyjęcia wagi, co może skończyć się zrzućciem całej torby z haka. Być może należałoby zapewnić utrzymanie torby np. dodatkowym hakiem od dołu.

Barbara Kowalska uważa, że najlepsza byłby półeczka w formie koszyczka z nierdzewnych prętów, na tyle głęboka, żeby wykluczała wypadnięcie wagi

po nieuważnym potrąceniu. Półeczka w formie koszyczka zapewniałaby przewiewność i nawet oblanie (oczywiście nie masywnym strumieniem, ale np. prysznicem) nie byłoby groźne, bo woda by ściekla i wyschła.

Bardzo dobra koncepcja. Koszyczek z nierdzewki jest trwalszy niż torba z folii, wygląda „łazienkowo”, czyli dobrze się wpisuje w zestaw akcesoriów i wyposażenia łazienki. Poza tym jest lepszy dla naszego rozleniwionego wygodnictwa, bo nie trzeba by otwierać torby i później zamykać jej kłapą.

Wacław Żurek pisze: wagi mają zazwyczaj płytę szklaną z hartowanego szkła. Taka płyta dobrze współpracuje z przyssawką. Można więc zainstalować na ścianie sporą przyssawkę i z jej pomocą przymocować wagę do ściany w pozycji pionowej. Zdjęcie wagi i założenie to sekundy!

Nie jest to zbyt pewna i stabilna metoda. Wszyscy wiemy, jak wygląda zamocowanie uchwyty do nawigacji samochodowej na szybie – właśnie za pomocą przyssawki, która niekiedy po prostu „puszcza”. Przyssawki przemysłowe mają podciśnienie, wytwarzane za pomocą różnego rodzaju pompek, są więc aktywne. Trzymanie przyssawką taką, jaką proponuje kolega, jest zależne od czystości płyty wagi i nawet niewielkie zanieczyszczenie może spowodować odpadnięcie jej i upadek na wyłożoną płytkami ceramicznymi podłogę. Byłoby to więc bardzo twarde lądowanie!

Wszystkim kolegom i koleżance gratuluję i zapraszam do następnych konkursów!

Nowe zadanie:

Tomek ma babcię, która jest wielką miłośniczką kwiatów doniczkowych. Ma tych kwiatów sporo i jedyny problem, jaki ma, to kwiaty stojące na wysoko umieszczonych półkach i w doniczkach wiszących na hakach, zamocowanych w suficie. Kwiatami – jak wiadomo – trzeba się opiekować: podlewać, dodawać nawóz, itp. Babcia Tomka ma już „swoje lata” i chodzić po stołkach i drabinkach nie może. Potrzebny

jest sposób i jakiś prosty i lekki przyrząd, umożliwiający podlewanie kwiatków wiszących na wysokości ok. 2,2 m. Przyrząd musi być prosty, a więc żadne mechanizmy i pompy nie wchodzi w rachubę. Pomóżcie babci Tomka i rozwiążcie ten problem:

Zaproponować sposób i przyrząd umożliwiający podlewanie kwiatów znajdujących się na wysokości ok. 2.2 m.

Klub Wynalazców

bez ograniczeń wieku

Zadaniem waszym było: Zaproponować poprawkę do systemu ochrony zbiorów bibliotecznych, wykluczając nieuczciwe kombinacje niektórych użytkowników.

Jak podano w założeniach zadania, należało zachować istniejący system sygnalizacji wynoszenia książki, natomiast należało skupić się na sposobie umocowania lub schowania w książce chipu, blaszki lub innego markera, który dawałby sygnał podczas przechodzenia przez bramkę kontrolną. Brak sygnału oznaczać mógłby, że jakaś książka jest wynoszona nielegalnie. Jest znany fakt, że chip lub blaszkę, a także etykietę aktywną magnetycznie można dość łatwo z książki usunąć. Można postąpić odwrotnie: książki legalnie pożyczone mogą mieć np. chip deaktywowany przez obsługę biblioteki, i wtedy książka się „nie odzywa” przy przekraczaniu bramki kontrolnej. Jednakże kluczowym problemem jest zainstalowanie chipu lub blaszki tak, żeby nie dało się go usunąć. No i co na to nasi klubowicze?

Mateusz Frankowski proponuje zamiast blaszki użyć proszku ferromagnetycznego, który może być nanoszony metodą drukowania. Nadruk powinien być na okładce, dobrze widoczny bez otwierania książki. Brak nadruku oznaczałby być może nielegalne usunięcie go z książki. Usunięcie nadruku powinno prowadzić do uszkodzenia okładki – co oczywiście będzie jasnym sygnałem, że coś tu nie jest w porządku.

Zważywszy, że nie chcemy zmieniać aparatury wykrywającej ferromagnetyk – użycie proszku jest niezłym pomysłem, bo przecież proszku nadrukowanego i zabezpieczonego np. folią nie da się usunąć bez podrapania okładki, a to już będzie doskonale widoczne bez szczegółowego przyglądania się książce. Nadruk proszku może mieć postać kodu QR i zawierać informacje o książce i np. o bibliotece.

Marek Borkowski – żeby nie tak łatwo dało się usunąć marker z książki – najprostszym sposobem

Przyrząd musi być lekki (starsza pani to nie sportowiec ciężarowiec), powinien umożliwiać dozowanie wody tak, żeby się nie przelewała przez doniczkę i nie powinien zajmować zbyt wiele miejsca w mieszkaniu.

Wszystkim życzę dobrych pomysłów i przypominam o terminie nadsyłania propozycji: do końca czerwca br.

byłoby osadzenie w okładce nitu, takiego do skóry. Taki nit jest nieusuwalny bez zdemolowania okładki.

Rzeczywiście – jeśli system reaguje na drobny kawałek metalu, to nit może się okazać najlepszym markerem.

Obu kolegom gratuluję i zapraszam do dalszych zmagania z techniką.

Nowe zadanie:

Tym razem zadanie z „rasowej” mechaniki:

Łożysko toczne składa się z dwóch współśrodkowych, stalowych pierścieni, wewnątrz których swobodnie poruszają się elementy toczne: kulki, wałeczki lub baryłki. Łożyska toczne są dziś stosowane niemal we wszystkich mechanizmach. Bywają sytuacje, gdy trzeba zdemontować łożysko bez zniszczenia pierścienia zewnętrznego. Jak to zrobić, skoro pierścień zewnętrzny mocno trzyma kulki wewnątrz? Najprościej jest odpowiedzieć tak, jak to mówią instrukcje wojskowe, dotyczące składania i rozkładania broni: „składanie wykonujemy w porządku odwrotnym do rozkładania”, i już! I wszystko jasne! Niektóre typy łożysk tocznych mają odpowiednie wycięcia w obu pierścieniach (na ogół łożyska wałeczkowe), umożliwiające względnie łatwy ich demontaż. Ogromna większość łożysk to łożyska kulkowe, a one nic takiego nie mają. Co robić? I to jest właśnie zadanie dla was:

Zaproponować sposób demontażu łożyska kulkowego, taki, żeby nadawało się do ponownego montażu; bez uszkodzenia pierścieni.

Oczywiście ani demontaż, ani montaż nie odbędzie się tak „za darmo”. Bliższych szczegółów podać nie mogę, bo przecież nie ja mam to zadanie rozwiązać!

Przypominam o terminie nadsyłania rozwiązań: do końca czerwca br.



Vademecum Młodego Wynalazcy

I znów na chwilę powrócimy do systemu standardowych rozwiązań zadań wynalazczych. Cały system „standardów” to usystematyzowane „ściagi”, przydatne w ogromnej liczbie zadań wynalazczych. Dziś w bardzo wielu firmach wydaje się „standardy” w formie broszury, którą warto mieć pod ręką. W broszurach tych są wolne kartki na wpisanie rozwiązań uzyskanych za pomocą standardów, ale dotyczących problematyki firmy. Z czasem staje się ona encyklopedią różnych spraw, z której można czerpać pomoc dla nowych problemów. Podkreślam po raz kolejny: standardów NIE czytamy jak *W pustyni i w puszczy*, ale dość często PRZEGLĄDAMY, żeby wiedzieć co tam jest.

2.4.8. Dynamizacja fepola

Przypominam: fepole – to minimalny system techniczny tak jak wepole (dwie substancje i pole), z tym że jedna z substancji jest ferromagnetykiem. Jeżeli dany jest system fepolowy, jego sterowalność może być podniesiona dzięki podniesieniu jego dynamiki tj. przejścia do elastycznego, zmieniającego strukturę systemu.

Przykład 1

Urządzenie do kontroli grubości ścianek, pustych wewnątrz detali z nieferromagnetycznego materiału, zawierające przetwornik indukcyjny z systemem pomiarowym i ferromagnetyczny element, rozmieszczone po obu stronach kontrolowanej ścianki. Znamienny tym, że dla podniesienia dokładności pomiaru ferromagnetyczny element wykonano w postaci nadmuchiwanej elastycznej poduszki, pokrytej ferromagnetyczną warstwą.

Przykład 2

Sposób symulacji różnych właściwości gleby na stanowiskach badawczych do badania roboczych organów maszyn rolniczych, polegający na wprowadzeniu do gleby cząsteczek ferromagnetycznych. Znamienny tym, że dla poszerzenia warunków badania sprawności roboczych organów maszyn rolniczych na cząsteczki oddziałuje się regulowanym polem elektromagnetycznym.

Daje to możliwość kształtowania własności mechanicznych gleby, pozwalając na symulację np. oporu mechanicznego gleby przy orce, bronowaniu, kultywacji itp. Można też symulować glebę ciężką, lekką, itd.

Zadanie nr 8

Znana jest metoda szlifowania części form wtryskowych narzędziem w postaci balonu z elastycznego materiału, którego robocza powierzchnia

pokryta jest materiałem ściernym. Szlifowanie odbywa się w warunkach ciągłego docisku narzędzia do obrabianego detalu. Dla uzyskanie równomiernego docisku materiału ściernego, do wnętrza balona wprowadza się zawiesinę materiału ferromagnetycznego, na który oddziałuje się stałym polem magnetycznym. Sposób ten pozwala na uzyskanie równomiernego docisku ścierniwa do powierzchni obrabianej i podniesienie dokładności obróbki. Jednakże jednocześnie na skutek powiększenia powierzchni kontaktu narzędzia z detalem, w strefie skrawania podnosi się temperatura i przyspieszając tępienie się ostrzy ścierniwa, prowadzi do podniesienia chropowatości powierzchni i pogarsza wydajność procesu.

Rozwiązanie zadania nr 8 wg standardów 2.4.3, 2.4.7., 2.4.8:

Stały docisk ścierniwa zastąpić zmiennym, narzędzie wibruje i tarcie się zmniejsza. W tym celu należy wprowadzić dodatkowe zmienne pole magnetyczne, działające na ferrozawiesinę (dynamizacja procesu). Żeby zaś działanie magnetycznego pola na ferrozawiesinę było możliwie maksymalne, cząsteczki ferromagnetyczne produkuje się z materiału o właściwościach magnetostrykcyjnych (kolejne wykorzystanie efektu fizycznego).

2.4.9. Strukturyzacja fepola

Jeśli dany jest system fepolowy, jego efektywność może być podniesiona przez przejście od pól jednorodnych lub mających nieuporządkowaną strukturę do pól niejednorodnych lub mających strukturę uporządkowaną (procesy ciągłe lub przemienne).

Przykład 3

Metoda formowania detali z tworzyw termoplastycznych. W charakterze stempla kształtującego wykorzystuje się ferproszek, na który działa pole termiczne, przekraczające w miejscach najmniejszego odkształcenia punkt Curie.

Jeżeli do substancji wchodzącej w fepole (lub mogącej wejść w fepole) powinna być dodana określona przestrzenna struktura, to proces należy prowadzić w polu o strukturze odpowiadającej żądanej strukturze substancji:

Przykład 4

Sposób uzyskiwania „włosów” na powierzchni termoplastycznego materiału, polegający na wyciąganiu powierzchniowych warstw z następującym później ochłodzeniem. Znamienny tym, że w celu podniesienia wydajności i możliwości sterowania procesem

tworzenia włosów, przed przeprowadzeniem operacji, do powierzchniowych warstw materiału wprowadza się ferrocząsteczki, następnie nagrzewa się materiał do temperatury mięknięcia, a kształtowanie włosów prowadzi się metodą wyciągania ferrocząsteczek za pomocą elektromagnesu z nadbiegunnikami w formie negatywu detalu z „włosami”.

2.4.10. Uzgodnienie rytmiki w fepolu

Jeżeli dany jest „przedfepolowy” lub fepolowy system, jego efektywność można podnieść metodą synchronizacji rytmu wchodzących w jego skład elementów.

Przykład 5

Przy wibromagnetycznej separacji materiału zaproponowano zsynchronizować częstotliwość wibrującego pola magnetycznego z częstotliwością własną vibracji materiału. Powoduje to zmniejszenie sił tarcia wzajemnego cząsteczek materiału i prowadzi do podniesienia efektywności separacji.

Przykład 6

Sposób transportu ferromagnetycznych materiałów sypkich, drobno- i gruboziarnistych metodą poddania ich działaniu vibracji, znamieny tym, że w celu podniesienia szybkości transportu na wibrowany materiał w momencie początkowym fazy odrywania działa się impulsami magnetycznego pola, którego linie sił biegną równolegle do kierunku transportu. Częstotliwość impulsów magnetycznych synchronizuje się odpowiednio z fazą odrywania wibrowanego materiału.

2.4.11. Przejęcie do fepola – wepola z wzajemnie oddziałującymi prądami elektrycznymi

Jeżeli wprowadzanie ferromagnetyków lub magnesowanie jest niemożliwe lub zbyt trudne, należy wykorzystać wzajemne oddziaływanie zewnętrznego pola

elektromagnetycznego z bezpośrednio doprowadzonymi lub bezstykowo indukowanymi prądami lub wzajemne oddziaływanie prądów.

Przykład 7

Metoda urabiania górotworu: dla podniesienia efektywności urobku przepuszcza się przez skałę impulsy elektryczne pomiędzy dwoma równoległymi przewodnikami.

Przykład 8

Sposób chwytania i przytrzymywania metalowych nieferromagnetycznych detali znamieny tym, że w celu podniesienia pewności procesu chwytania i przytrzymywania detalu, przez jego masę w strefie działania pola magnetycznego przepuszcza się prąd elektryczny w kierunku prostopadłym do linii sił pola.

Przykład 9

Sposób zbiórki jagód w uprawach szpalerowych metodą wprawiania w drgania drutów szpalerowych wraz z przymocowanymi do nich odciągami, znamieny tym, że w celu obniżenia nakładów pracy i zmniejszenia uszkodzeń jagód zastosowano magnes o stałym polu, obejmujący biegunami druty szpaleru, przez które przepuszcza się prąd zmienny, a magnes przemieszcza się wzdłuż drutu.

Objaśnienia:

1. Jeżeli fepola – to systemy, do których wprowadzono ferromagnetyczne cząsteczki, to epola – systemy, gdzie zamiast ferromagnetycznych cząsteczek działają (lub wzajemnie oddziałują) prądy.

2. Rozwój epoli – podobnie jak rozwój fepoli – powtarza ogólną linię:

- proste epola,
- kompleksowe epola,
- epola w środowisku zewnętrznym,
- dynamizacja,
- przebudowa strukturalna,
- synchronizacja rytmu.

Legandy i latte

Travis Baldree

Wydawnictwo Insignis, liczba stron: 336, cena: 44,90 zł

Jeszcze niedawno walczyła z potworami... Dziś zapraszana na pyszną kawę. Znużona dziesięcioleciami walk i krwawych wypraw orczyca Viv uznaje, że nadeszła pora zawiesić miecz na kołku. Ma marzenie – pragnie otworzyć pierwszą kawiarnię w mieście Thune. Viv chce zostawić przeszłość za sobą, ale nie powinna mierzyć się z tym w pojedynkę. Na jej drodze staną bowiem rywale – dawni i nowi. Mroczne podziemie Thune może sprawić, że orczyca łatwo będzie znów chwycić za miecz, a przecież obiecała sobie, że więcej tego nie zrobi. Stare porzekadło mówi: „Przed wyruszeniem w drogę należy zebrać drużynę”. A jej członkowie nadejdą z najmniej oczekiwanych kierunków. Nieważne, czy przyciągnie ich pradawna magia, smakowite słodkości czy świeżo zaparzona kawa – staną się dla Viv bardziej nieodzowni, niż sobie wyobrażała. Ma prawdziwą nagrodą w wędrówce są ci, których spotykamy na szlaku.





Materiał z analizy epoli wzrasta, gdy jego analiza wykaże, czy jest celowe wydzielanie standardów epolowych jako oddzielną grupę.

Zadanie 9

Przy rozcinaniu kamieni szlachetnych i czystych kryształów stosuje się do cięcia bardzo cienkie piłki. Im cieńsza piłka, tym mniejsze odpady. Napęd piłki może być różny: ręczny, mechaniczny, elektromagnetyczny itd. Problem polega na zapewnieniu ściśle stałej, co do wartości i kierunku, siły docisku piły do dna szczeliny cięcia. Stałość wartości siły zapewnia jednorodność płaszczyzny cięcia (bez zmętnień, naprężeń termicznych itp.) Zmienność kierunku działania siły to nieuchronność odprysków. Potrzebna idea metody, dającej ściśle stałą siłę docisku piły do materiału.

Rozwiązanie zadania 9

Zadanie rozwiązujemy, korzystając z idei zastosowanej w rozwiązaniu zadania 8. Tu rozwiązano zadanie o przeciwnych warunkach: zapewnienie drgań drutu szpalerowego (taki drut zasadniczo niczym nie różni się od cienkiej piłki do cięcia minerałów) Rozwiązanie zadania 9 jest analogią do zadania 8, tylko że naturalnie wykorzystano nie zmienny, a stały prąd (nie trzeba wzbudzać drgań, lecz je gasić).

Wykorzystanie płynów elektroreologicznych 2.4.12

Szczególną postacią epola – elektoreologiczna zawieszina (wprowadzić np. drobno zmielony kwarcowy

proszek do toluenu) o dającej się sterować lepkości. Jeżeli nie da się zastosować ferromagnetycznej cieczy, można wykorzystać elektoreologiczną cieć.

Przykład 10

Mimośrodowy generator drgań. Masy mimośrodowe rozmieszczone w elektoreologicznej cieczy. Lepkością cieczy można sterować za pomocą pola elektromagnetycznego.

Przykład 11

Elektoreologiczna cieć ze zmienną lepkością, wykorzystana w amortyzatorach pojazdów transportowych.

Przykład 12

Zastosowanie elektrolepkiej zawiesziny w urządzeniu do cięcia materiałów, w charakterze środka mocującego.

Przykład 13

Wąż składający się z wewnętrznej i zewnętrznej warstwy, pomiędzy którymi znajduje się warstwa elektroprowadzących włókien, odizolowanych od siebie warstwą elastycznego materiału izolującego. Znamienny tym, że w celu umożliwienia sterowania sztywnością węża elastyczny materiał izolujący wykonano jako porowaty, przesycony elektoreologiczną zawiesziną.

System standatrdów trzeba poznać i... zaprzyjaźnić się z nim. Najłatwiej dzieje się to, gdy uda się wam rozwiązać parę rzeczywistych problemów z własnej praktyki, czego serdecznie wszystkim czytelnikom życzę.

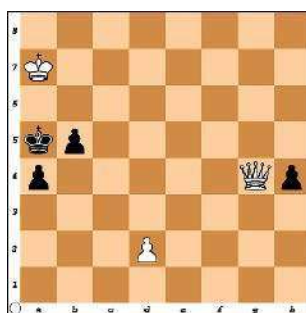
Prezes Klubu Wynalazców

Champion TRIZ

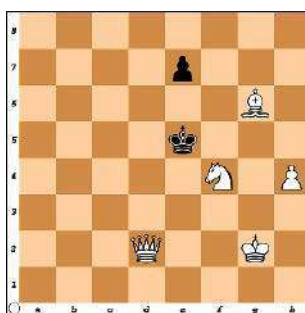
Jan Boratyński

Dokończenie ze strony 77

Zadania do samodzielnego rozwiązania



Zadanie 1
Iszta, 1893
Mat w 2 posunięciach



Zadanie 2
Hoffman, 1900
Mat w 2 posunięciach

Rozwiązanie zadań z MT 4/2023

Zadanie 1

Montvid 1893

Mat w 2 posunięciach

Rozwiązanie: 1. Sh4

Zadanie 2

Potempski 1893

Mat w 2 posunięciach

Rozwiązanie: 1. He6

Archiwalne odcinki o tematyce szachów
<http://bit.ly/2VohMA1>

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 5/2023

Przemawia do nas pomysł przygotowania prostej, przeznaczonej dla niekoniecznie zaawansowanych profesjonalistów aplikacji z zapisem nutowym asystującej przy wykonywaniu utworów przez osoby grające na instrumentach, nawet amatorsko.

Autorem pomysłu jest Zygmunt Łada

Zbigniew Górski ma dziadka, który zajmuje się pszczołami. Dziadek chciał wnuka zainteresować pszczołami, ale skończyło się to fatalnie; Zbyszek przez trzy dni musiał przykładać różne kompresy, ale i tak chodził spuchnięty. W związku z tym proponuje wynalezienie jakiegoś straszaka elektronicznego, który odpędzałby pszczoły od osoby, zagładającej do ula. Tradycyjne odymianie pszczoł jest – zdaniem Zbyszka – „starożytne i kłopotliwe”. Mamy przecież XXI wiek!

Może się uda. Wiadomo, że komary doczekały się już odstraszaczy elektronicznych, może więc dla pszczoł też się uda coś zrobić. Poza tym warto „przeczesać” internet: może coś takiego już jest? Problem właściwie polega głównie na tym, żeby tylko wokół pszczelarza wytworzyć strefę bezpieczną, a niekoniecznie usuwać pszczoły z całej pasieki.

Marek Bielecki proponuje opracowanie łyżki do mieszania np. miodu z wodą. łyżka powinna mieć napęd elektryczny i obracać się naprzemiennie: w lewo i w prawo, co dałoby efekt szybszego wymieszania. Niektóre substancje wyjątkowo opornie się rozpuszczają w wodzie i taka łyżka byłaby ogromnym ułatwieniem przy sporządzaniu różnych mieszanek leczniczych.

Wynalazek dla leniwych, ale ogromna liczba wynalazków ma właściwie taki cel: ograniczyć nasz wysiłek i pozwolić na swobodne uprawianie „dolce far niente”. Faktem jednak jest, że taki np. miód spadziowy bardzo wolno się rozpuszcza, a ogólnie znane blendery mieszają zbyt gwałtownie.

Dariusz Zakrzewski nie lubi okresu jesienno-zimowego, kiedy do samochodu wnosi się błoto, śnieg i liście z drzew, które zalegają na całym parkingu. Dariusz proponuje wyposażenie samochodów w automatycznie wysuwającą się spod nadwozia elektryczną, obracającą się szczotkę, która mogłaby usunąć nieczystości z podszwy buta i przy odpowiednio długim włosie oczyścić również boki. Wnętrze samochodu pozostałoby czyste.

Pomysł jest realny i w zasadzie da się zrealizować. Wnoszenie do wnętrza samochodu śniegu, błota i liści powoduje konieczność uciążliwego czyszczenia wnętrza pojazdu, czego nikt

robić nie lubi. Wydaje się, że wystarczyłaby taka szczotka, zainstalowana przy drzwiach kierowcy, który najczęściej wsiada i wysiada z samochodu.

Zygmunt Łada dziwi się zakorzenionemu od setek lat zwyczajowi zapisywania nut na arkuszach papierowych. Powoduje to konieczność odwracania kart np. dużego koncertu, w czym pomagać musi asystentka. Powoduje to żłośliwe krytyki typu:

„na wczorajszym koncercie najlepiej prezentowała się pani Ania”. W dobie rozwiniętej techniki ekranów i Podów i tabletów dawno już powinno się wyprodukować tablet, w którego pamięci można by zapisać potężną ilość tekstu nutowego, a poza tym układ śledzący postępowanie gry mógłby automatycznie zmieniać stronę lub przesuwając ją na ekranie.

Jest już dostępny program MuseScore, dający ogromne możliwości zapisywania nut, komponowania z odsłuchem, zastępowania brakujących instrumentalistów, itd. Jest to program na komputer, ale zapis nutowy można przekazać na całą orkiestrę, wyposażoną monitory.

Jest to już wyższa szkoła jazdy w stosunku do skromnego pomysłu Zygmunta, ale to raczej dla profesjonalistów i dużych zespołów. Miejsce dla pomysłu Zygmunta byłoby występy solistów, a nawet uczniów – do ćwiczeń nowych utworów.

Marek Lewiński obserwuje sytuacje na parkingu osiedlowym, kiedy to kierowcy „gimnastykują się” i pracownicy próbują zmieścić swój samochód w za ciasnym miejscu do zaparkowania. Ćwicząc więc „kopertki” i katując układ kierowniczy do wielu kolejnych prób.

Sytuację zmieniłoby radykalnie wyposażenie samochodów w przednie i tylne koła skrętne. Taki pojazd mógłby skręcać w miejscu, co wydatnie poprawiłoby jego zdolności manewrowe.

Z parkowaniem jest coraz gorzej i pomysł, który może wydawać się ekstrawagancją, za kilka lat może stać się koniecznością. Szereg pojazdów np. niektóre wojskowe mają taki system i to się sprawdza w trudnych warunkach przestrzennych.



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



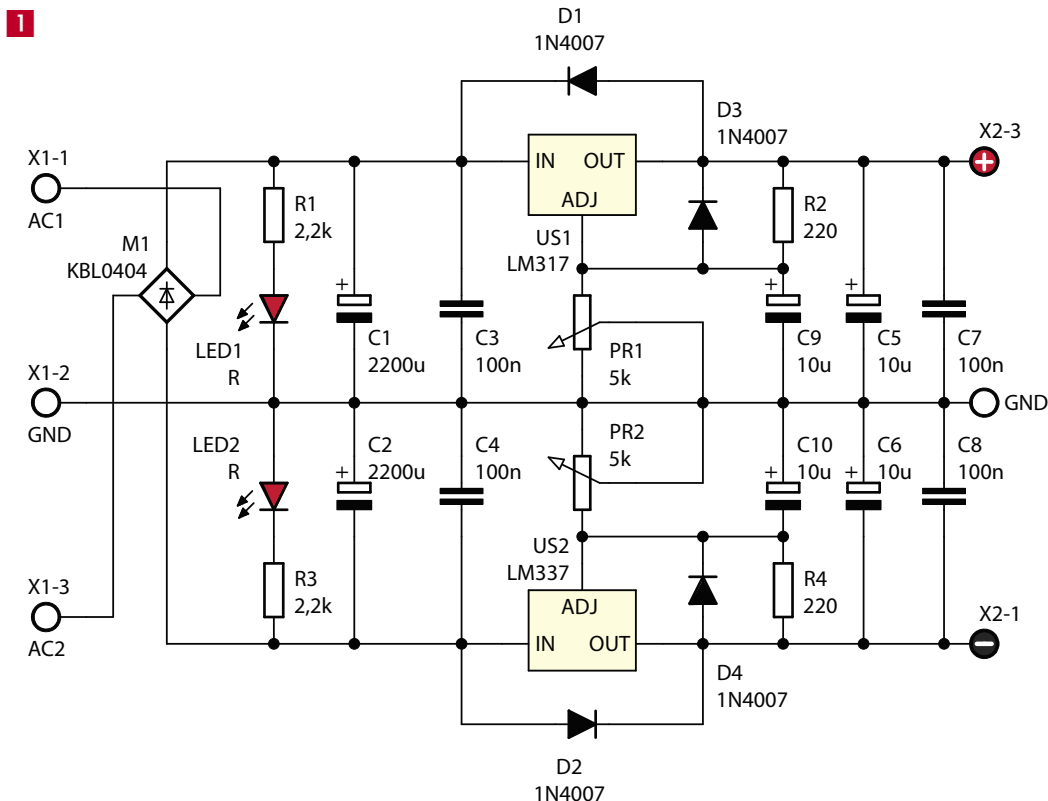
Symetryczny zasilacz warsztatowy $\pm 1,25 \text{ V} \dots \pm 25 \text{ V} / 1,5 \text{ A}$

Jak bardzo przydatnym urządzeniem w pracowni jest zasilacz, nie trzeba przekonywać żadnego praktyka, natomiast początkujący szybko się przekona, że bez niego nic nie da się zdziałać. Bardzo często uruchamiane urządzenia wymagają symetrycznego zasilania, a nie zawsze posiadane zasilacze z pojedynczym napięciem można łączyć w szereg w celu uzyskania potrzebnego napięcia. Prezentowany zasilacz jest więc kolejnym urządzeniem rozbudowującym możliwości naszego warsztatu.

Schemat ideowy przedstawiono na **rysunku 1**. Napięcie referencyjne ustawiane jest za pomocą potencjometrów PR1 i PR2. LM317 jest regulatorem napięcia dodatniego, natomiast LM337 napięcia ujemnego. Układy LM potrzebują zaledwie kilku elementów zewnętrznych oraz mają wbudowane zabezpieczenie

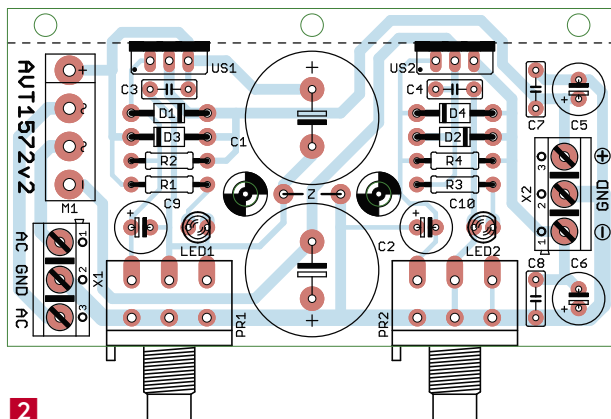
Właściwości

- napięcia wyjściowe dodatnie i ujemne
- napięcia wyjściowe regulowane 1,2...25 VDC
- maksymalny ciągły prąd wyjściowy: 2×1,5 A
- kontrolki napięć wyjściowych – diody LED
- zabezpieczenie przeciwzwarciowe i termiczne
- zalecany transformator: 2×17...19 VAC
- wymiary płytki: 45×81 mm



termiczne i ograniczające prąd przy zwarciu wyjścia do masy. Zakres napięć wyjściowych dla zasilania 2×17...19 VAC wynosi od ±1,25 V do ±25 V. Układy LM317 oraz LM337 mają wbudowane krótkotrwałe zabezpieczenie przeciwzwarciowe oraz termiczne.

Przy doborze transformatora należy zwrócić uwagę na znamionowe napięcie kondensatorów C1, C2. Transformator należy dobrać tak, aby jego napięcie wtórne po wyprostowaniu przez mostek M1 było nie większe od znamionowego napięcia kondensatorów.



Montaż i uruchomienie

Schemat montażowy przedstawiony jest na **rysunku 2**. Montaż należy zacząć od wlutowania zwory „Z”, a ostatnimi montowanymi elementami muszą być kondensatory C1, C2, tuż po przykręceniu układów do radiatora. Układy US1 i US2 należy odizolować od radiatora przekładką mikową lub silikonową. Montaż jest typowy, a układ zmontowany ze sprawnych elementów nie wymaga żadnej regulacji i po dołączeniu transformatora działa natychmiast poprawnie. ■

Wykaz elementów:

R1, R3: 2,2 kΩ	LED1, LED2: czerwona
R2, R4: 220 Ω	3 mm
PR1, PR2: 5 kΩ liniowy	M1: mostek prostowniczy
C1, C2: 2200 uF/35 V	X1, X2: DG301-5.0/3
C3, C4, C7, C8: 100 nF	Podkładki silikonowe
C5, C6, C9, C10: 10 uF/63 V	TO220 ×2 tulejki dystansowe ×2
US1: LM317	Wkręty ×3 radiator
US2: LM337	
D1-D4: 1N4007	



1. Chwaszczyno – z daleka widoczna główna siedziba firmy ALEXANDER (mp – materiały producenta)

Planszówki małego konstruktora

Kiedy byłem małym chłopcem (hej! 😊), marzyłem o tym, żeby pracować jako projektant w fabryce zabawek – nic więc dziwnego, że kiedy pojawiła się możliwość wizyty w tego rodzaju zakładzie i zobaczenia na własne oczy całego procesu produkcji, nie wahałem się ani chwili. Dziś dzielę się tym wyjątkowym doświadczeniem na łamach „Młodego Technika”.

ALEXANDER to jeden z najbardziej znanych w Polsce producentów gier i zestawów kreatywnych. Miałem już okazję poznać kilka produktów tej firmy – w tym modeli konstrukcyjnych (ze zrozumiałych

2. Projektowanie, czyli obrazowanie pomysłu – tu chyba najbardziej technika wiąże naukę ze sztuką! (mp)



względów najbardziej interesujących ze powodów zawodowych). Główną siedzibą zakładu jest Chwaszczyno – niewielka miejscowość na Pomorzu, nieopodal Gdyni. Przy oddalonej od centrum miejscowości ulicy Telewizyjnej znajduje się biurowiec oraz przylegająca do niego większość hal produkcyjnych i magazynowych [1]. Ta rodzima i rodzinna firma od 47 lat od podstaw produkuje własne zabawki, zestawy konstrukcyjne, gry planszowe i edukacyjne, ale także losowe, sprawnościowe, karciane i łamigłówek – nie tylko dla dzieci – także dla dorosłych i seniorów! Szczyci się tym, że niemal wszystkie elementy własnych gier i zestawów są wykonywane na miejscu.

POMYSŁ na dobry produkt, który ma uczyć i bawić, trzeba najpierw znaleźć – ale także i odpowiednio opracować. Na miejscu zajmują się tym firmowe projektantki i projektanci [2]. Ponieważ rodzajów wyrobów jest wiele,



3 4



5 6



7 8



3. Doskonała narzędziownia to nie tylko numerycznie sterowane obrabiarki, ale przede wszystkim doświadczeni fachowcy i zegarmistrzowska precyzja ich pracy (PD) **4.** Tłocznia elementów metalowych – tu do zestawów serii Młody Konstruktor (mp) **5.** By krawędzie metalowych elementów nie kaleczyły rąk młodych budowniczych, przed malowaniem są poddane uszlachetnianiu w wykładarkach wibracyjnych (mp) **6.** Elementy plastikowe są wykonywane w jednej z wielu automatycznych wtryskarek, często uzupełnionych o dodatkowe usprawnienia (PD) **7.** Do produkcji zestawów zawierających elementy z drewna i sklejkę wycinane laserowo wykorzystuje się tu kilkanaście ploterów CNC różnych producentów i wielkości (mp) **8.** Części drewniane (m.in. do zestawów Młody Konstruktor Junior) powstają także w stolarni – automatycznie lub spod ręki firmowych fachowców (mp)

a każdy z nich ma swoją specyfikę (inaczej projektuje się grafiki użytkowe, inaczej modele metalowe, inaczej gry, inaczej zestawy kreatywne czy sprawnościowe), dlatego też i tu są różne specjalizacje. Jak w większości nowoczesnych biur projektowych, obsługa programów 3D jest tu standardem. Mając wybór z wielu

wcześniej opracowanych i przetestowanych akcesoriów, tworzy się tu także zestawy na specjalne zamówienie poszczególnych zleceniodawców – może więc kiedyś i na Twoje, droga Czytelniczko/drogi Czytelniku?...

OPRZYRZĄDOWANIE to bodaj najbardziej skomplikowana, a przy tym absolutnie niezbędna część pracy,



9 10



9. Spakowane w pudełka zestawy trafiają do magazynu wyrobów gotowych – wielkiej hali wysokiego składowania. 10. Osobne, niemałe biuro, zajmuje wzorcownia – tu można obejrzeć najbardziej popularne wyroby firmy – w tym najmiłsze młodym technikom zestawy konstrukcyjne.

pozwalającej przekuć (niekiedy wręcz dosłownie) projekt w gotowy wyrób. Stare majsterskie przysłowie mówi: „Poznasz fachowca po narzędziach jego!” – przyznaję, że to właśnie narzędziownia zrobiła na mnie największe wrażenie podczas zwiedzania zakładu. Tu, z zegarmistrzowską precyzją, najlepsi fachowcy tworzą między innymi formy wtryskowe, sztance i wiele innych elementów linii produkcyjnych [3]. Profesjonalne urządzenia, ale co jeszcze ważniejsze, doświadczeni specjaliści przy okazji pomagają stworzyć także narzędzia do produkcji między innymi dla przemysłu samochodowego.

PRODUKCJA tylu różnych asortymentów elementów, składających się na końcowy wyrób, z oczywistych względów nie może odbywać się w jednej hali – osobne więc pomieszczenia stanowią drukarnie (offsetowa, sitodrukowa, tampondrukowa), intro-ligatorynia, szwalnia, stolarnia [8] czy lakiernia. Inne

pomieszczenia zajmują linie do produkcji elementów metalowych [4–5], inne do plastikowych [6] oraz wypalanych laserowo [7], jeszcze inne przeznaczone są na konfekcjonowanie i pakowanie zestawów. Każdy z wymienionych wyżej rodzajów produkcji musi mieć też swoją doświadczoną, a przy tym ściśle współdziałającą z pozostałymi działami, załogę.

PRZED DROGĄ DO FINALNEGO ODBIORCY złożone zestawy trafiają do wielkiej hali – magazynu produktów gotowych [9] – tylko pojedyncze lądują na półkach firmowej wzorcowni [10]. Ponieważ zakład prowadzi zarówno sprzedaż hurtową, jak i wysyłkowy sklep detaliczny, skuteczna koordynacja zamówień i produkcji jest odrębną sztuką kolejnej firmowej ekipy. ■

Paweł Dejnak

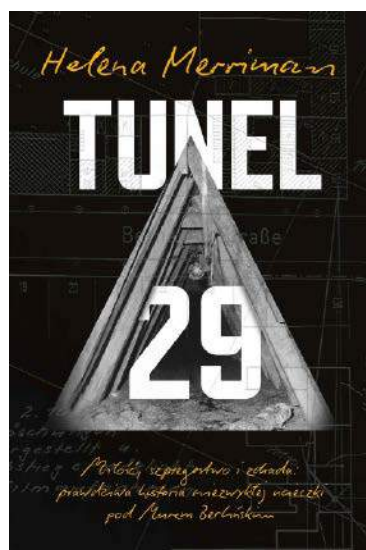
(zdjęcia z archiwum autora oraz materiałów producenta)

Tunel 29. Miłość, szpiegostwo i zdrada: prawdziwa historia niezwykle ucieczki pod murem berlińskim

Helena Merriman

Wydawnictwo Insignis, liczba stron: 456, cena: 49,99 zł

Niedawno uciekł z jednego z najbardziej brutalnych reżimów na świecie. Teraz postanawia przedostać się tam z powrotem. Lato, rok 1962. Joachim Rudolph, student, kopie tunel pod murem berlińskim. Po drugiej stronie – w Berlinie Wschodnim – czekają dziesiątki mężczyzn, kobiet i dzieci. Ci ludzie gotowi są zaryzykować wszystko, nawet życie, byle stamtąd uciec. Ta książka to prawdziwa historia najbardziej niezwykłego tunelu ewakuacyjnego wykopanego pod murem berlińskim, opowiedziana przez autorkę podcastu BBC Radio 4 (trzy miliony słuchaczy, nagrody Best Radio Podcast oraz Best Moment w plebiscycie British Podcast Awards 2020). Opierając się na setkach godzin wywiadów z ocalałymi i tysiącach stron dokumentów Stasi, Helena Merriman w pasjonujący sposób przedstawia historię splatającą losy genialnej grupy studentów kopiących tunel, czarującą rudowłosej kurierki, amerykańskiej sieci informacyjnej, planującej sfilmować śmiałą ucieczkę, i szpiega Stasi, który za graje całemu przedsięwzięciu. „Tunel 29” to poruszająca opowieść o tym, co się dzieje, gdy ludzie tracą wolność, i o tym, do czego są zdolni, by ją odzyskać. To literatura faktu, która wciąga bardziej niż niejedna powieść sensacyjna.





Metr wikinga!

Małe modele wikingów budowaliśmy już razem z „Młodym Technikiem” przy okazji regat rynnowych – ściśle w wydaniu październikowym z 2014 r. naszego miesięcznika. Tym razem Na warsztacie zajmiemy się znacznie okazalszą jego wersją, która świetnie sprawi się jako flagowy model nie tylko dla wodnej drużyny harcerskiej.

Te atrakcyjne wizualnie rodzaje skandynawskich łodzi od lat budujemy na zajęciach modelarskich w wersjach pływających (sklejkowych) oraz papierowych (wystawowych) – zwykle jednak jako niewielkie, kilkunastocentymetrowe jednostki. Czasem jednak zdarza się potrzeba wykonania zdecydowanie większego modelu, który będzie mógł być dumą drużyny lub pracowni modelarstwa szkatułowego. Na taką okoliczność opracowałem sklejkową łódź wikingów o blisko metrowej długości.

Głównym materiałem budowlanym jest tu lekka modelarska sklejka topolowa grubości 3 mm. Mniejsze formatki tego surowca można także uzyskać z demontażu wyrzucanych jednorazowych skrzynek po owocach cytrusowych – dostępne prawie w każdym warzywniaku. Prócz sklejki potrzebne będą drewniane

pręty (najlepiej bukowe) kilku średnic (w przedziale od 6 do 15 mm), materiał na żagiel, trochę bejcy, ew. lakieru do drewna i kolorowych wodoodpornych farb.

Podstawą do rozpoczęcia budowy są odpowiednie rysunki (zestawieniowy [A] oraz elementów sklejkowych [B]) stanowiące część tego artykułu. Po wydrukowaniu ich w docelowych rozmiarach należy skopiować obrysy elementów na sklejkę i wyciąć je starannie (ręcznie lub za pomocą plotera). Szczegóły dalszego montażu zawierają opisy do zdjęć ilustrujących budowę prototypu.

Starannie wykonany model może stanowić cenną dekorację klubu, harcówki lub modelarni – ciesząc oczy i serca wykonawców oraz widzów – czego wszystkim po równo życzę. ■

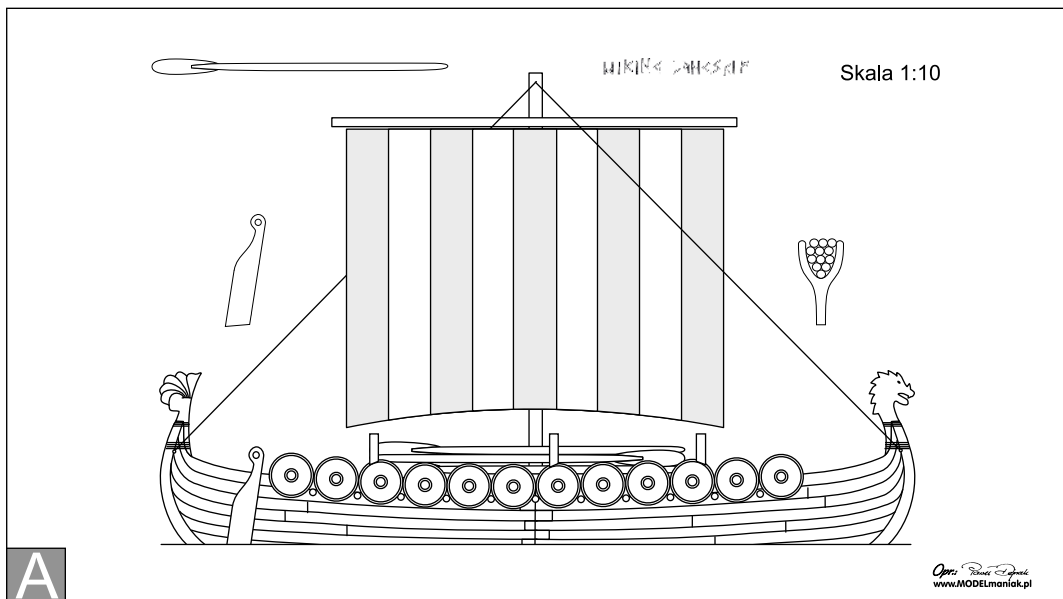
Paweł Dejnak



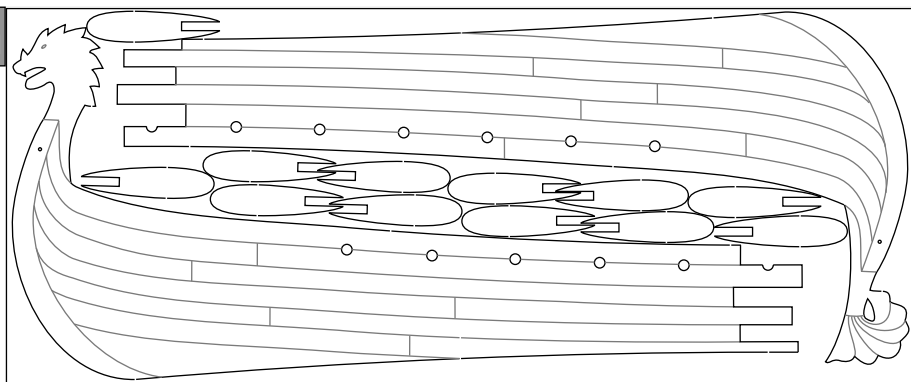
1

1. Z uwagi na swoją awanturniczą historię, ale i ciekawe zdobnictwo, łódzie średniowiecznych skandynawskich piratów i odkrywców cieszą się niestąbną popularnością – także wśród modelarzy. 2. Trzeba przyznać wikingom, że ich łódzie miały naprawdę imponujące możliwości – kiedy wreszcie przekonali się do „niemęskich” żagli (prawdziwi mężczyźni używają wiosel) na wiele wieków przed Kolumbem odkryli Amerykę (nie ma już co do tego najmniejszych wątpliwości!).

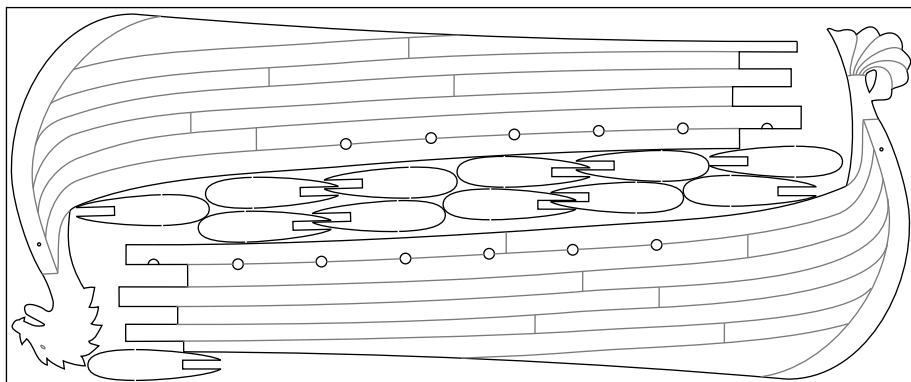
2



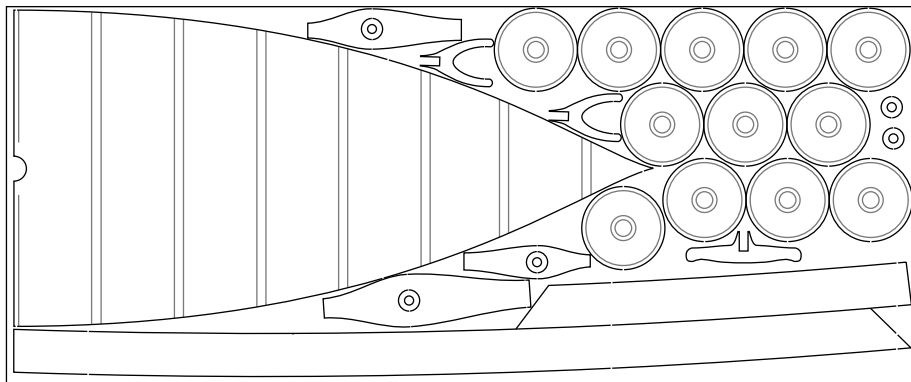
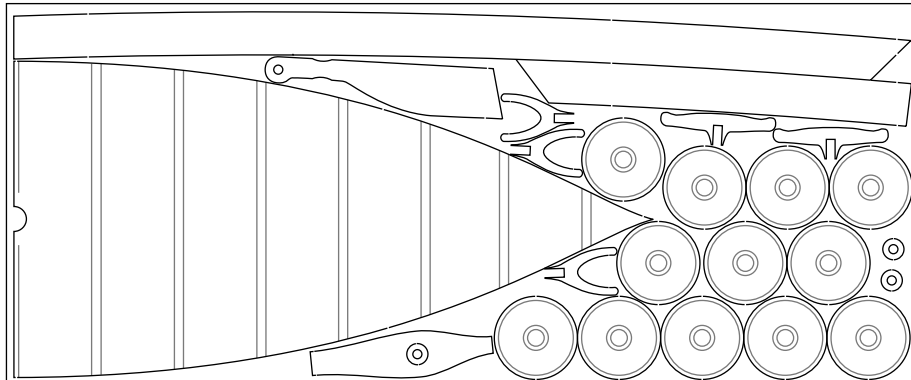
B



Sklejka topolowa 3 mm



Elementy modelu w skali 1:5



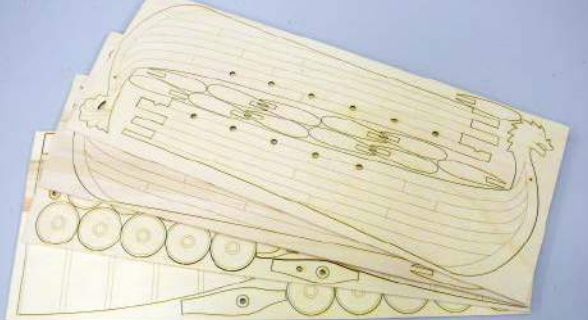


3

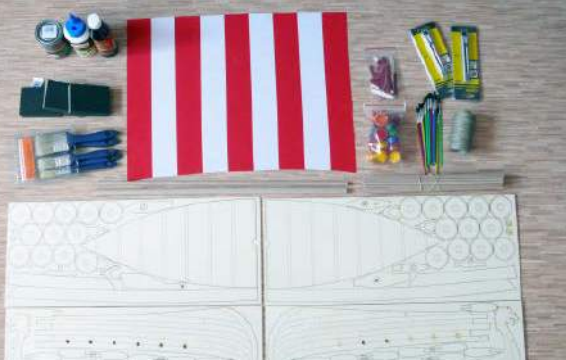
3. Makietowe odwzorowania wymagają wiele pracy i niemało historycznej wiedzy. Nawet jeśli robimy prostsze wersje, warto się im przyjrzeć, by wiedzieć, co upraszczamy. **4.** Punktem wyjścia do dzisiejszego projektu metrowej smoczej łodzi był młodotechnikowy model klasy rynnowej opublikowany w numerze 10 z 2014 r. **5.** Nie ma wystarczająco mocnych dowodów, że żagle miały takie układy kolorów, ale znalazłszy taki właśnie materiał, od razu wiedziałem, że z niego właśnie powstaną żagle nowej pokazowej jednostki!

4 5





6 7



8



9



10



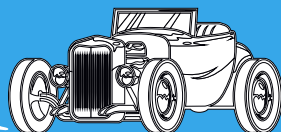
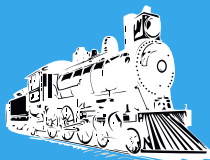
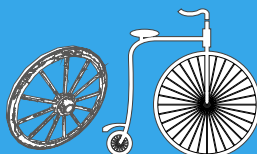
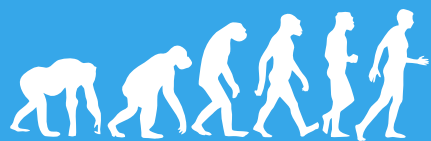
11



12 13



6. Przyznaję, mając dostęp do plików wektorowych wcześniej opracowanego już modelu oraz ploter laserowy pod ręką, najwygodniej było mi wyciąć wszystkie elementy sklejkowe w technologii CNC. Można jednak obejść się również bez niego – sklejka topolowa (zwana też cytrusówką od skrzynek, z których ją pozyskujemy) jest bardzo przyjaznym w obróbce materiałem. **7.** Dodatkowo, do zmontowania modelu potrzebne będą: klej wodoodporny do drewna, bejca, farby akrylowe wodoodporne do ozdobienia tarcz, linki i kilka typowo warsztatowych narzędzi. **8.** Prace nad kadłubem warto zacząć od obustronnego, dwukrotnego pobejcowania sklejkowych elementów. Jeśli miałyby być wystawiony na działanie wody, przyda się także odpowiedni lakier. **9.** Prócz sklejkowych, do modelu potrzebne będą również drewniane pręty na maszt, reję, trzony wiosła itp. Je również warto pokryć bejcą. **10.** Kiedy podsychają burty, można zająć się zdobieniami tarczy – wikingowie malowali je w swoje rodowe barwy. **11.** Wzorów jest dużo, zawsze można wybrać coś dla siebie! **12.** Burty modelu są łączone pośrodku, pokład nadaje im odpowiedni kształt – ale stosowna „perswazja” w postaci ścisków stolarskich jest bardzo przydatna. Stopniowo dodaje się inne elementy wyposażenia – tarcze, gniazdo masztu, podpory masztu czy stojaki na wiosła. **13.** Po postawieniu masztu i wciągnięciu żagla model jest już prawie gotowy – jeszcze tylko olinowanie (choć symboliczne) – i można cieszyć się niezwykłą jednostką. Niech długo służy i cieszy wszystkich wykonawców i ich gości!



Radary i radiolokacja

1842–45

Austriacki naukowiec Christian Doppler publikuje teorię zmiany częstotliwości fali, wywołanej ruchem źródła falowego. Zjawisko to znane jest jako efekt Dopplera. W 1845 roku dopplerowska teoria została potwierdzona doświadczalnie przez holenderskiego fizyka Buys Ballota.

1886

Heinrich Hertz (1) opisuje teoretycznie istnienie fal elektromagnetycznych, które mogą być wykorzystane do przesyłania informacji. Hertz używał urządzenia, które co do zasady działania było podobne do współczesnego radaru impulsowego. Wykazał również, że fale elektromagnetyczne mogą być odbijane przez metalowe przedmioty oraz ulegają refrakcji podczas przejścia przez wykonany z dielektryka pryzmat.

1904

Christian Hülsmeier, niemiecki wynalazca, konstruuje pierwszy radar, zwany wówczas telemobiloskopem (2) lub telemobilofonem, do wykrywania obiektów na wodzie. Urządzenie to działało na zasadzie emisji fal elektromagnetycznych i odbicia ich od obiektów. We wczesnych latach XX wieku zbudował on urządzenie, które dzisiaj można nazwać monostatycznym radarem impulsowym, a było udoskonaloną wersją urządzenia Hertza (termin monostatyczny oznacza, że nadajnik i odbiornik radaru umieszczone są w jednym urządzeniu). Wynalazek Hülsmeyera nie spotkał się jednak z zainteresowaniem i wkrótce o nim zapomniano.

lata 20. XX wieku

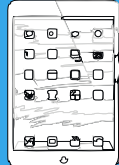
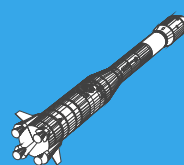
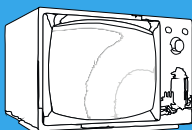
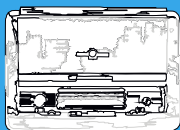
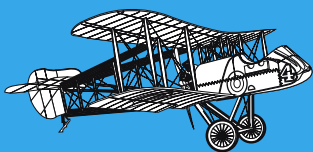
Naukowcy zaczynają eksperymentować z wykorzystaniem efektu Dopplera w radarze. W 1922 roku dwóch Amerykanów, Hoyt Taylor i Leo C. Young, przypadkowo zaobserwowali fluktuacje odbieranego sygnału, gdy statek przepływał między odbiornikiem i nadajnikiem umieszczonymi na przeciwnych brzegach rzeki. Dzisiaj urządzenie takie nazwalibyśmy radarem bistatycznym o fali ciągłej (termin bistatyczny oznacza, że nadajnik i odbiornik radaru umieszczone są w różnych miejscach, oddległych nawet o kilkadziesiąt kilometrów). Pierwsze próby wykrycia ruchu obiektów przy użyciu radaru dopplerowskiego zostały przeprowadzone w 1927 roku.

1921

Amerkański fizyk Albert Wallace Hull (3) wynajduje magnetron, wydajną lampę nadawczą, która przetwarza wejściową energię prądu stałego na energię elektryczną wysokiej częstotliwości. Przetwarzanie energii odbywa się w specjalnie ukształtowanej komorze anodowej umieszczonej w silnym polu magnetycznym.

1935

Brytyjczyk Robert Watson-Watt projektuje pierwszy radar lotniczy o nazwie Chain Home. Był to system, który wykorzystywał emisję fal radiowych i ich odbicie od samolotów, aby określić ich położenie. Zadeklarował on Brytyjskiemu Komitetowi Obrony Powietrznej aparat, za pomocą którego wykrył z odległości 50 kilometrów nadlatujący samolot. W tym samym roku przystąpiono w Wielkiej Brytanii do budowy nadbrzeżnej sieci eksperymentalnych radarów (4), składającej się z pięciu stacji pracujących na falach o długości 25 metrów. Sieć ta pozwalała wykrywać nadlatujące samoloty z odległości 90 kilometrów. W 1937 roku zbudowano w Anglii radar znacznie doskonalszy, pracujący na falach 1,5 metra, przy czym urządzenie miało takie wymiary, że mogło być umieszczane na samolotach. W przededniu II wojny światowej Wielka Brytania miała już sprawną sieć stacji radarowych, pracujących na falach o długości 12 metrów, służących do zdalnego ostrzegania, i sieć stacji pracujących na falach o długości 1,5 metra służących do wykrywania nisko lecących samolotów. Brytyjskie systemy radarowe odegrały kluczową rolę w bitwie o Anglię, umożliwiając wykrywanie nadlatujących samolotów niemieckich.



1936–37

1939

1940

1947–49

1947

Według różnych źródeł wynalazcami ważnego w technice radarowej klustronu było dwóch techników z General Electric, George F. Metcalf i William C. Hahn, albo też bracia Russell i Sigurd Varianowie z Uniwersytetu Stanforda. Klustron to lampa mikrofalowa z modulacją prędkości elektronów, ważny element stacji radiolokacyjnych spełniający funkcję wzmacniacza lub generatora przebiegów mikrofalowych (o częstotliwościach od setek megaherców w górę). Składa się z katody wysyłającej elektrony, zespołu elektrod ogniskujących elektrony w wąską wiązkę, anody przyspieszającej oraz przynajmniej dwóch rezonatorów i kolektora. Był pierwszym praktycznym źródłem fal radiowych o dużej mocy w zakresie mikrofalowym. Przed jego wynalezieniem źródłami były lampa Barkhausena–Kurza i magnetron z dzieloną anodą, ograniczone do małych mocy.

Zastosowanie magnetronu w radiolokacji pozwoliło na znaczne zredukowanie rozmiarów radarów, zwłaszcza anten, co umożliwiało instalowanie ich na pokładach samolotów. Dwóch inżynierów z uniwersytetu w Birmingham, John Randall i Henry Boot, buduje niewielki, ale mocny radar na bazie magnetronu wielownękowego, w który następnie zaczęto wyposażać samoloty B-17. Pozwoliło to później odnajdywać i zwalczać niemieckie łodzie podwodne, w nocy i we mgle.

Amerykańska firma RCA opracowuje pierwszy radar dopplerowski, który potrafi wykrywać ruch obiektów na powierzchni Ziemi.

W 1947 roku w Glastonbury w amerykańskim stanie Connecticut rozpoczęto testowanie radaru do pomiaru prędkości przejeżdżających pojazdów (5). W lutym 1949 roku pojawia się pierwszy na świecie punkt pomiaru prędkości. Oficerowie pracowali w parach – jeden obserwował system radarowy, zbierając wizualne dowody przekroczenia prędkości przez przejeżdżające samochody i wysyłał sygnały radiowe do drugiego funkcjonariusza, podając mu numer rejestracyjny pojazdu, drugi funkcjonariusz czekał na samochód, zatrzymywał kierowcę i wystawiał mu mandat.

Ideę radaru impulsowego z liniową modulacją częstotliwości opracowuje i patentuje Sidney Darlington, amerykański inżynier i wynalazca pracujący w Bell Laboratories.



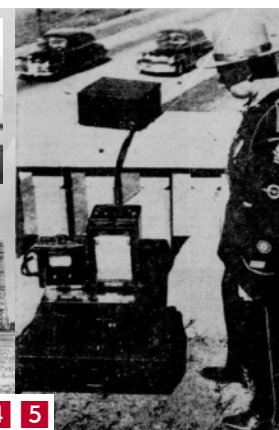
1



3

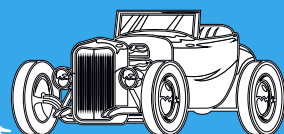
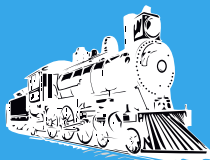
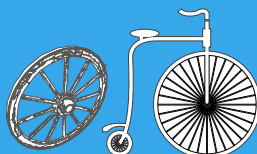
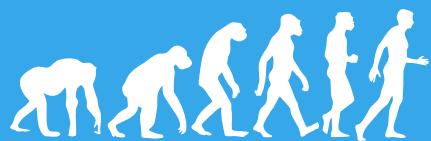


4



5

1. Heinrich Hertz; 2. Telemobiloskop Christiana Huelsmeyera; 3. Albert Wallace Hull ze swoim wynalazkiem; 4. System radarowy Chain Home; 5. Pierwsze testy radaru do pomiaru prędkości samochodów w Glastonbury w stanie Connecticut



lata 50. XX wieku

Prace rozwojowe nad radarem pozahoryzontalnym (OTH), którego pierwszą koncepcję przedstawił amerykański uczoney Robert Morris Page. Prowadzone były w USA, w ZSRR (6), a także w Czechosłowacji.

lata 50.–60. XX wieku

Początki wykorzystania radaru do celów cywilnych, takich jak kontrola ruchu lotniczego, meteorologia, badania obiektów kosmicznych, przemysł i medycyna.

1951–78

Patent na radar z syntetyczną aperturą (SAR) John Wiley złożył na początku lat 50. Kierował zespołem w Goodyear Aircraft Corporation, pracującym nad techniką pozwalającą rozszerzyć i poprawić rozdzielczość obrazów generowanych przez radar. W radarze z syntetyczną aperturą (SAR) zwykłej wielkości sygnał z anteny, mocowanej np. w samolocie, jest przetwarzany w złożonym systemie wykorzystującym m.in. matematyczne przekształcenie Fouriera. Impulsy są wypromieniowywane w poprzecznym paśmie na teren. Powrót sygnałów jest rozłożony w czasie, ze względu na odbicia od elementów znajdujących się w różnych odległościach. Amplituda i faza powrotów są łączone przez procesor sygnału w celu generacji obrazu. W 1974 r. po raz pierwszy zaproponowano koncepcję interferometrii SAR do obserwacji Ziemi. W 1978 roku, wystrzelono na orbitę pierwszy kosmiczny radar typu SAR do obserwacji Ziemi, na pokładzie satelity SEASAT (7).

1957

W USA rozpoczynają się prace nad projektem sieci radarów meteorologicznych, załączkiem Narodowej Sieci Radarów Meteorologicznych. W latach 50. naukowcy rozpoczęli prace nad wykorzystaniem radaru dopplerowskiego w meteorologii, do badania ruchu chmur, opadów i innych zjawisk atmosferycznych. W kolejnych latach radar tego typu zaczął być wykorzystywany w kontrolowaniu ruchu lotniczego, a także w badaniach geologicznych i oceanograficznych.

1967

W Wielkiej Brytanii powstaje pierwszy niewielki radar przenośny, który mógł być łatwo przenoszony przez policjantów na drodze.

1975

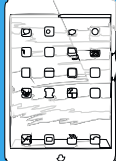
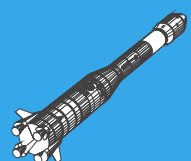
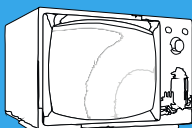
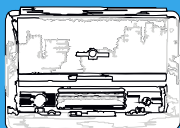
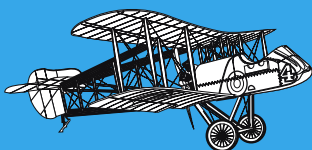
W Japonii opracowany zostaje pierwszy radar meteorologiczny, który wykorzystywał falę 5 cm. Był to radar z reflektorem obracającym się wokół anteny.

lata 80. XX wieku

Wprowadzanie, przede wszystkim w zastosowaniach militarnych, radaru pasywnego, który wykorzystuje fale elektromagnetyczne generowane przez inne źródła, takie jak telewizory i radia, do wykrywania celów. Ten rodzaj radaru stał się ważnym narzędziem wywiadowczym, ponieważ umożliwia wykrycie celów, które nie emitują własnych sygnałów (8).

przełom lat 80. i 90.

Pojawia się radar z fazą skanowania (phased array), który umożliwia szybkie i skuteczne skanowanie dużych obszarów. Ten rodzaj radaru stał się bardzo ważnym narzędziem w obronie powietrznej, umożliwiając szybkie wykrycie i śledzenie lotów nieprzyjacielskich samolotów. Jego następcą jest radar z aktywną fazą skanowania, umożliwiający jeszcze szybsze i bardziej precyzyjne skanowanie, a także wykrywanie celów, które poruszają się z dużą prędkością (9).



lata 90. XX wieku

Pojawiają się pierwsze radary typu stealth, zdolne do wykrywania celów o niskiej sygnaturze radarowej. Celem było zwalczanie samolotów niewykrywalnych przez konwencjonalne radary.

lata 2000 i później

W latach 2000–2010 nastąpił znaczny postęp w technologii mikrofalowej i fotonowej, pozwalający na rozwój radarów o większej czułości i zdolności do detekcji celów. Wprowadzano również nowe techniki przetwarzania sygnału radarowego, w tym systemy cyfrowe i sztuczną inteligencję. Pojawił się radar z bardzo krótką falą (VHF), który umożliwia wykrywanie celów o niskiej wysokości, takich jak drony, i wykrywanie samolotów stealth, a także wykrywanie celów na dużych odległościach. W 2004 roku naukowcy z Caltech zademonstrowali pierwszy zintegrowany, oparty na krzemie odbiornik dla szyku fazowanego anten (phased array) na częstotliwości 24 GHz z ośmioma elementami. Następnie w 2005 r. zademonstrowali nadajnik CMOS 24 GHz dla układu phased array. W 2007 roku badacze z Agencji Zaawansowanych Projektów Badawczych w Obszarze Obronności (DARPA) ogłosili powstanie szesnastoelementowej anteny radaru w szyku fazowanym, która również była zintegrowana ze wszystkimi niezbędnymi obwodami na jednym chipie krzemowym i pracowała na częstotliwości 30–50 GHz.

od 2016 do dziś

Start i rozwój techniki radaru kwantowego. W USA projekty radarów kwantowych finansuje DARPA. W 2016 roku Chiński Uniwersytet Nauk i Technologii opracował prototyp radaru kwantowego ze zdolnością wykrywania celów o niskiej sygnaturze radarowej na odległości do 100 metrów. W 2019 roku naukowcy z brytyjskiego uniwersytetu w Glasgow ogłosili, że opracowali radar kwantowy, który jest w stanie wykryć cel w odległości do 20 metrów. W 2020 roku zespół z australijskiego Uniwersytetu Monash ogłosił, że opracował radar kwantowy, który jest w stanie wykryć obiekty o rozmiarze od 10 do 50 centymetrów na odległości od metrów do kilku kilometrów. Radary kwantowe są jednak nadal w stosunkowo wczesnej fazie eksperymentalnej.



6



7

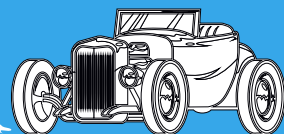
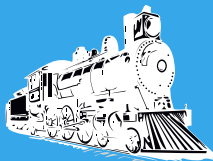
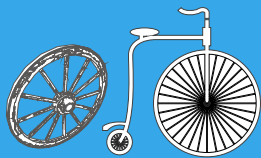
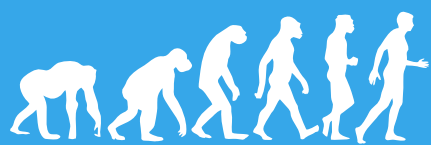


8



9

6. Stacja radarowa OTH w pobliżu Czarnobyla; 7. Obraz powierzchni Ziemi uzyskany za pomocą radaru SAR z satelity SEASAT; 8. System radarów pasywnych Twintvis; 9. Nowoczesny chiński radar typu phased array z domniemanymi możliwościami wykrywania samolotów stealth



ODKRYJ HISTORIĘ WYNALEZKÓW

Klasyfikacja radarów

Radar to skrót angielskiego terminu Radio Detection and Ranging, co można przetłumaczyć na Radiowe Wykrywanie i Pomiar Odległości. Radar wykorzystuje fale elektromagnetyczne do wykrywania obiektów oraz pomiaru ich odległości, prędkości i innych parametrów. Radary można klasyfikować na wiele sposobów, według rozlicznych kryteriów. Pod względem zastosowania można je podzielić na:

- Radary meteorologiczne – służące do monitorowania pogody, w tym wykrywania opadów, burz i tornad, w których analiza odbić, przystosowanego do obserwacji obiektów meteorologicznych radaru pozwala ustalić prognozę pogody, a także strukturę burz;
- Radary nawigacyjne – stosowane w lotnictwie, do nawigacji statków i w samochodach, umożliwiające dokładne określenie położenia i prędkości;
- Radary wojskowe – służące do wykrywania i identyfikacji celów, w tym samolotów, okrętów, pojazdów i ludzi;
- Radary kosmiczne – służące do badania planet, gwiazd i innych obiektów kosmicznych;
- Radary samochodowe, zdalne prędkościomierze, fotoradary – zbiór rozwiązań stosowanych na drogach, zarówno przez kierowców, jak i służby bezpieczeństwa ruchu do określania prędkości, odległości lub innych parametrów w ruchu samochodowym;
- Radary przemysłowe – stosowane w przemyśle, np. w kontrolowaniu jakości produktów, dozowaniu substancji lub pomiarze poziomu cieczy;
- Radary medyczne – używane w medycynie, np. do obrazowania narządów wewnętrznych lub do pomiaru prędkości krwi w naczyniach;
- Radary do zastosowań specjalnych – używane do różnych celów, takich jak wykrywanie min, monitorowanie ruchu pojazdów czy wykrywanie obiektów na dnie morza.

Gdy weźmiemy pod uwagę zasadę fizyczną, spełniane funkcje i techniczne szczegóły konstrukcji, to rysuje się taka klasyfikacja radarów:

- Radary pasywne – działające na zasadzie odbicia fal elektromagnetycznych od obiektów, które nie generują własnego sygnału radarowego, np. statków powietrznych;
- Radary aktywne – generujące własny sygnał radarowy, który jest odbijany od obiektów, umożliwiając ich wykrycie;
- Radary kierunkowe – umożliwiające określenie kierunku, z którego nadchodzi sygnał;
- Radary z syntetyczną aperturą – wykorzystujące matematyczne algorytmy do poprawy rozdzielczości obrazów;



- Radary z fazową aperturą – wykorzystujące anteny o skomplikowanej strukturze (w tzw. szyku fazowym anten), które pozwalają na uzyskanie bardzo dobrej jakości obrazów.

Jeśli zechcemy podzielić radary według sposobu emisji sygnału, to najogólniejszy podział wygląda tak:

- Radary z falą ciągłą – emitują transmitowany sygnał w sposób ciągły, a sygnał echa jest stale odbierany i przetwarzany – ich odmianą są radary dopplerowskie;
- Radary impulsowe – urządzenia radarowe, które emitują krótkie, silne impulsy i odbierają echa pomiędzy nimi (podczas tzw. przerwy spoczynkowej). ■

M.U.

*** Pisownia oryginalna ***



PFRZEGLĄD TECHNICZNY Przymusowe zastosowanie spirytusu dla celów motoro- wych we Francji

Skutkiem wprowadzenia monopolu państwowego na spirytus we Francji (oprócz spirytusu używanego do picia oraz wytwarzanego z owoców i winogron) znalazły się tak znaczne ilości spirytusu w posiadaniu państwa, że zachodzi konieczność znaleźć zastosowanie techniczne dla tego paliwa. W tym celu parlament uchwalił wniosek, na skutek którego pozwolenie na przewóz do kraju paliw cudzoziemskich: benzyny, benzolu i inn. olejów mineralnych, udzielane będzie jedynie pod warunkiem jednoczesnego wykupienia ilości spirytusu równej około 10% importowanego paliwa. Próby jeszcze z czasów wielkiej wojny dowiodły, że dodawanie 10% spirytusu do benzyny i benzolu daje dobre wyniki w silnikach spalinowych. Nassegnanie ceny i % należy, zgodnie z nowym prawem, do państwa, którego władze tym sposobem otrzymują szerokie pełnomocnictwa do wprowadzenia w kraju „paliwa narodowego” (carburant national) które by dzięki swemu składowi i cenie uwzględniło w pierwszej linii interesa gospodarki krajowej, dając jednocześnie najlepsze rezultaty przy spalaniu w silnikach. Wysokość % została ustalona na 10% gdyż ilość spirytusu, otrzymywanego z ziemiopłodów, wynosi obecnie we Francji około 1/10 ilości wwożonych paliw mineralnych.

8 maja 1923

Zagadnienie paliwa narodowego we Francji

Jak już donosiłmy poprzednio w naszym piśmie, sprawa stworzenia t. zw. paliwa narodowego we Francji jest obecnie bardzo aktualna. Z licznych głosów w tej sprawie, spotykanych w pismach technicznych francuskich, podajemy wiadomości ważniejsze. Okazuje się, że już obecnie Francja może pokryć dwie piąte swego zapotrzebowania na paliwo płynne, wykorzystując jedynie swoją własną produkcję alkoholu. P. Potart, w komunikacji, wygłoszonej na posiedzeniu Stow. Inż. Cywil. w dniu 23 marca r. b., dochodzi do wniosku, że na wyprodukowanie pozostałych trzech piątych potrzebnego paliwa płynnego w postaci benzolu, należałoby zużyć rocznie 1500 000 t koksu, oraz pracę mechaniczną, odpowiadającą niezmienniej mocy około 900 000 k. m., co wobec przedwojennego wydobycia we Francji wszelkich rodzajów paliwa stałego, wynoszącego 40 000 000 t, nie stanowi wielkiego procentu. Sprawa zastosowania mieszaniny alkoholu i benzolu została już technicznie rozwiązana. Wątpliwa była jedynie kwestja mieszaniny alkoholu z benzyną, szczególnie ze względu na jej trwałość. P. D. Berthelot, w referacie wygłoszonym na wyżej wspomnianem posiedzeniu Tow. Inż. Cywil., rozpatruje następującą kwestję: 1. Trwałość. Dwa składniki: alkohol i benzyna mają dążność do odzielania się, szczególnie przy niskich temperaturach, ponieważ alkohol zawiera pewien procent wody. Starano się temu zapobiedz przez dodawanie różnego rodzaju „stabilizatorów chemicznych”. Obecnie jest już zrealizowana produkcja przemysłowa alkoholu absolutnego. Podczas prób, przeprowadzonych przez „Komitet naukowy paliwa narodowego”, mieszaniny z 10, 15 i 20-tu procentami alkoholu dały wprawdzie dobre wyniki, ale trwałość ich wydała się Komitetowi niezbyt pewną; natomiast

uznano, jako zupełnie pewną, mieszaninę 50-cio procentową. 2. Moc. Liczne próby wykazały, że silnik, bez zmiany stopnia sprężania, osiąga przy zastosowaniu paliwa narodowego moc taką samą jak przy napędzie benzyną. Rozważania teoretyczne wykazują wprawdzie wyższą wartość opałową idealnej mieszanki benzynowej, ale konieczność stosowania przy tej ostatniej nadmiaru powietrza tę różnicę usuwa. W licznych próbach, przeprowadzonych przez Komitet z silnikami samochodów ciężarowych, osobowych i pławowców, osiągnięto nawet, przy użyciu mieszaniny benzyny z alkoholem, moc nieco wyższą. 3. Zużycie paliwa na jednostkę mocy. Badania laboratoryjne, przeprowadzone przez Komitet, wykazały przy zawartości alkoholu 30–50% zużycie o 5 do 15% większe, niż przy użyciu benzyny ciężkiej. Ale próby z samochodami, odbyte w warunkach normalnych na szosie, wykazały, że procent ten jest znacznie mniejszy. Tłumaczy się to tem, że bieg silnika, pędzonego benzyną, jest mniej jednostajny, z punktu widzenia regularności wybuchów, szczególnie przy zmienianiu zapotrzebowaniu mocy, co ma miejsce w samochodach. 4. Pozatem, z badań, przeprowadzonych w Anglii przez P. Ricardo, wynika, że alkohol pozwala stosować wyższy stopień sprężania, czego nie można osiągnąć z benzyną, ze względu na wymienioną wyżej nieregularność biegu silnika.

22 maja 1923

PRZEGLĄD ELEKTROTECHNICZNY

Lampa żarowa z kondensatorem

Drogą ciągłych ulepszeń lamp wypełnionych gazem rozwiązano w sposób zadawalniający problem lamp żarowych wieloświecowych. Lampy te posiadają trwałość wystarczającą przy małym zużyciu mocy. Natomiast fabrykacja lamp małoświecowych dla napięć normalnych napotyka ciągle na trudności, które nie

są dziś jeszcze w stopniu dostatecznym pokonane. Ponieważ rozwiązanie tych trudności byłoby łatwe gdyby lampy małoświecowe budowano na niskie napięcie, przeto pomysły wynalazców szły w kierunku zmniejszania w ten czy w inny sposób napięcia, przyłożonego do poszczególnych lampek. Najprostszymi rozwiązaniami, jakie tu się nasuwały, były następujące: 1) łączenie lampek po kilka szeregowo, 2) załączanie przed poszczególnymi lampkami oporów, 3) stosowanie transformatorów. Inż. L. P. Graner w artykule, pomieszczonym w holenderskim piśmie „Tydschrift voor Elektrotechniek” podaje jeszcze inny sposób obniżania napięcia na lampce, opatentowany przez fabrykę lamp elektrycznych w Holandji „Philips”. Sposób ten polega na załączeniu przed lampką kondensatora o odpowiednio dobranej pojemności. Jest jasne, że kondensator powoduje spadek napięcia, jednocześnie wywołuje tylko bardzo małe straty mocy. Na skutek całego szeregu prób zdołano zbudować kondensator, wytrzymujący wysokie napięcie przy bardzo małej objętości (10–20 cm³), oraz powodujący straty mocy, nie przekraczające 0,1 wata. Kondensator ten można umocować w oprawce żarówki.

1 maja 1923

Elektryczne wytapianie glinu
Nowy typ pieca elektrycznego systemu Bailey’a został zbudowany w „Dayton Foundries”. Zasilający go prąd transformuje się z 6 600 V na 80 V, normalne obciążenie pieca wynosi 600 f.; pirometr i liczniki służą do pomiarów i regulacji. Podobnie zastosowano 2 piece elektryczne do stapienia stopów mosiądzu o pojemności 600 f. każdy. Straty przy eksploatacji nie przekraczają 1½%; zużywana energia przy pracy ciągłej na 1 tonnę wynosi 400 kWh, przy pracy dziesięciogodzinnej – 550 kWh.

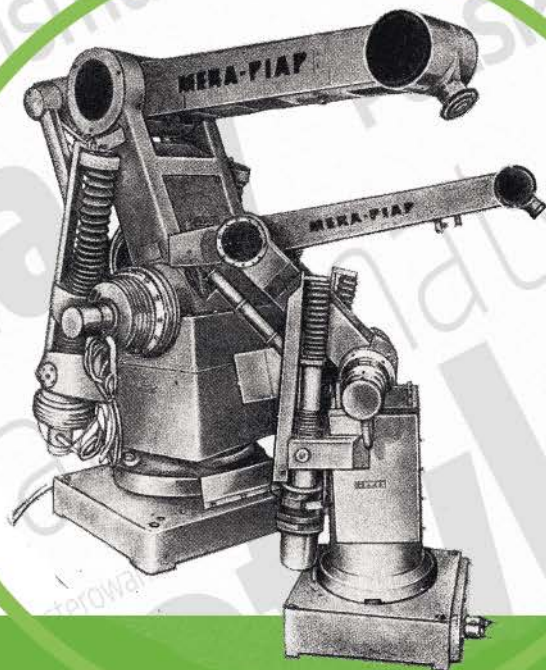
15 maja 1923



Łukasiewicz
PIAP

roborepo

roborepo.pl



REPOZYTORIUM ROBOTYKI

– CYFROWE KOMPENDIUM WIEDZY
O ROBOTYCE I NIE TYLKO

- Podsumowanie pięciu dekad rozwoju robotyki i dziedzin pokrewnych w zdigitalizowanej formie.
- Przeszłość i teraźniejszość w jednej interaktywnej bazie wiedzy.
- Bezpłatny dostęp do zgromadzonych w Łukasiewicz – PIAP prac naukowych, badawczych i rozwojowych, a także raportów z badań, opisów projektów i filmów z wdrożeń.



Fundusze
Europejskie
Polska Cyfrowa



Rzeczpospolita
Polska

Unia Europejska
Europejski Fundusz
Rozwoju Regionalnego



Projekt Repozytorium Robotyki – cyfrowe udostępnianie zasobów nauki z obszaru robotyki realizuje Sieć Badawcza Łukasiewicz – Przemysłowy Instytut Automatyki i Pomiarów PIAP w ramach Programu Operacyjnego Polska Cyfrowa. Projekt jest współfinansowany ze środków Europejskiego Funduszu Rozwoju Regionalnego.