

FIZYKA

w Szkole z Astronomią

CZASOPISMO DLA NAUCZYCIELI

360 (LXIV) indeks 35810X Nr 1 styczeń/luty 2019 CENA 27,50 zł (w tym 5% VAT)

**Paradoksy szczególnej
teorii względności**

Ile jest stanów skupienia?

Plazma

**REWOLUCJA KOPERNIKAŃSKA
OKIEM FIZYKA**

WOKÓŁ NEPTUNA
– PLANETOIDY, SEDNOIDY
I DYSK ROZPROSZONY

LEONARDO DA VINCI
CZŁOWIEK RENESANSU,
INŻYNIER RENESANSU



Odkrycia w praktyce!



- ✓ Jak i dlaczego działa GPS?
- ✓ Jak wykorzystuje się niezwykle właściwości nanocząstek?
- ✓ Jak się numerycznie prognozuje pogodę?
- ✓ Jak z kwiatów czerpać energię elektryczną?
- ✓ Jak się bada ludzki mózg?
- ✓ Jak fizycy leczą raka?

Wydanie specjalne
w wersji elektronicznej
(plik PDF)

Tylko 15 zł!

Formularz zamówienia na stronie: www.aspress.com.pl/specjalne/

Drodzy Czytelnicy!

Serdecznie witamy w nowym, 2019 roku! Mamy właśnie czas karnawału. Karnawał to oczywiście czas zabawy. Ale może pomiędzy jedną a drugą karnawałową imprezą znajdą Państwo czas na lekturę „Fizyki w Szkole”. Ten numer jest numerem zimowym. Mimo, że nie jest to najprzyjemniejsza pora roku to jednak, ze względu na długie noce jest to świetny czas na obserwacje astronomiczne. Między innymi dlatego poświęciliśmy tej tematyce dwa artykuły. Jeden autorstwa Jana Rokity dotyczy najdalszych zakątków Układu Słonecznego. Obiekty transuranowe są co prawda trudne do obserwowania z Ziemi metodami amatorskimi, są jednak na topie ze względu na misję sondy New Horizon oraz na fakt, że stanowią skarbnicę wiedzy o początkach naszego kosmicznego domu.

Innym artykułem o tematyce astronomicznej jest artykuł Pawła Wajera i Ryszarda Gabryszewskiego „Rewolucja kopernikańska okiem fizyka”, w którym autorzy w sposób szczegółowy pokazują jak z praw Newtona można wyprowadzić prawa Keplera i jak przewidywać położenie planet.

Nie tylko jednak wnioski z teorii Newtona są przedstawione w tym numerze. Na pewno miłośnicy szczególnej teorii względności znajdą coś dla siebie. Mowa oczywiście o artykule Jana Kurzyka poświęconego paradoksom wynikającym z tej teorii. Na sam koniec chciałbym polecić termodynamiczny artykuł Alfreda Zmitrowicza. Niestety jest to ostatni artykuł z tej serii, ale miejmy nadzieję, że nie ostatni artykuł tego autora.

Pozostało mi teraz tylko życzyć Państwu przyjemnej lektury i wiele szczęścia w Nowym Roku.

W imieniu redakcji
Zbigniew Wiśniewski



9 Termodynamika: główne koncepcje i etapy rozwoju. Cz. 4.
Ciepło tarcia † Alfred Zmitrowicz

Fizyka wczoraj, dziś, jutro

- 4** Od kuli plazmowej do plazmy termojądrowej
† Grzegorz Karwasz
- 16** Paradoxy szczególnej teorii względności. Część I
† Jan Kurzyk
- 22** Kobięca strona fizyki. Siła słabych oddziaływań
† Aleksandra Mielewczyk-Gryń, Marcin Zagród
- 40** Leonardo da Vinci. Inżynier renesansu
† Kazimierz Mikulski



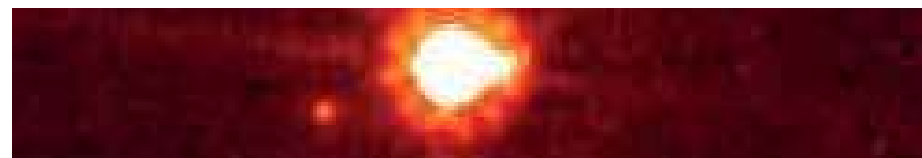
Z naszych lekcji

- 26** Rewolucja kopernikańska okiem fizyka
† Paweł Wajer, Ryszard Gabryszewski
- 32** Wykorzystanie rozumowania matematycznego do analizy równoległego połączenia oporników
† Andrzej Sokołowski
- 36** Wilhelm Eduard Weber (1804-1891)
† Tadeusz Wibig
- 38** Zasady i schematy stosowane podczas rozwiązywania zadań fizycznych
† Czesław Surowiec



Astronomia dla każdego

- 45** Blisko i daleko od Neptuna † Jan Rokita



FIZYKA

w Szkole z Astronomią

NUMER 1 LISTOPAD/LUTY 2019 Nakład 3000 egz. CENA 27,50 zł
360 (LXI) indeks 35810X ISSN 0426-3383 (w tym 5% VAT)

Komitet redakcyjny Krystyna Jabłońska-Ławiczak, Jerzy Kreiner, Andrzej Majhofer (Przewodniczący Komitetu), Zygmunt Mazur, Andrzej Szymacha, Mirosław Trociuk
Redakcja Zbigniew Wiśniewski (redaktor prowadzący – fizykas@wp.pl) **Adres redakcji** ul. Warchałowskiego 2/58, 02-776 Warszawa **Wydawnictwo** Agencja AS Józef Szewczyk, ul. Warchałowskiego 2/58, 02-776 Warszawa, e-mail: szewczyk24@gmail.com, tel. 606 201 244, www.aspress.com.pl, NIP: 951-134-91-51 **Wydawca** i **redaktor naczelny** Józef Szewczyk, szewczyk24@gmail.com **Prenumerata** www.aspress.com.pl/prenumerata-2019/, e-mail: szewczyk24@gmail.com, tel. 606 201 244 **Reklama** Jędrzej Chodakowski, jchodakowski1953@gmail.com **Skład i łamanie** Vega design **Druk i oprawa** Paper & Tinta, ul. Ceglana 34, 05-270 Nadma
Zdjęcie na okładce i wspis treści: Adobe Stock

Redakcja nie zwraca nadesłanych materiałów, zastrzega sobie prawo formalnych zmian w treści artykułów i nie odpowiada za treść płatnych reklam.

Ile jest stanów skupienia?

Od kuli plazmowej do plazmy termojądrowej

Grzegorz Karwasz

Ile jest stanów skupienia? Podręczniki szkolne odpowiadają, że trzy: ciała stałe mają kształt, ciecze nie mają kształtu a tylko objętość, a gazy starają się zająć objętość jak największą (czyli „rozprężyć”). Ale stanów skupienia, jak to pokazujemy w Toruniu na interaktywnych wykładach dla młodzieży [1] jest więcej: naliczyliśmy ich cztery, może pięć. Nie wiadomo jak zakwalifikować ciekłe kryształy, a jak kondensat Bosego-Einsteina (uporządkowany stan atomów w stanie gazowym [2]).

Kolejne typowe zdanie z podręczników to: „już starożytni Grecy”. Rzeczywiście, w odróżnieniu od Egipcjan czy mieszkańców Malty, Grecy nie budowali gigantycznych obserwatoriów astronomicznych, ale pisali książki. Arystoteles, który w jednym umyśle zawarł miłość do wiedzy (czyli filo-zofię), miał za sobą pięćset lat rozwoju nauki, w jej ogromnej różnorodności, rozsypanej w greckim archipelagu po całym Morzu Śródziemnym. Jego poprzednicy poszukiwali składników materii. Anaksymander (a może Anaksagoras) uważał, że pierwotnym składnikiem jest ziemia (albo woda, nie pamiętam). Arystoteles podsumował te poszukiwania tak we wstępie do *Fizyki* jak i *Metafizyki*. To Empedokles jako pierwszy wyróżnił cztery „elementy”, ale jak pisze Arystoteles (*Metafizyka*, 985b, 1), sam Empedokles uważał ogień za element odmienny od trzech pozostałych. I właśnie „ogniem”, czyli plazmą zajmujemy się w tym artykule.

Lampa jarzeniowa

Plazmą nazywamy „zjonizowany gaz”, gdzie zaraz wszystko wyjaśnimy. Podobno większość Wszechświata to plazma – bo z niej składają się wszystkie gwiazdy świecące na nocnym niebie (gwiazdy neutronowe to już nie plazma).

Gaz w zwykłych warunkach nie przewodzi prądu elektrycznego – gdyby było inaczej, między dwoma otworami gniazodka elektrycznego cały czas byłoby widać „sznurek”

ognia. Ale w lampie jarzeniowej nad głową prąd płynie: niektóre z atomów gazu (jest w takiej lampie głównie argon i pary rtęci) straciły elektrony (zazwyczaj po jednym elektronie) – stały się jonami (o ładunku dodatnim). W ten sposób, zarówno dodatnie jony, jak i uwolnione elektrony mogą przewodzić prąd. W którym kierunku? W rurze nad głową raz w prawo, raz w lewo – 50 razy na sekundę.

Aby zaobserwować, że rzeczywiście muszą najpierw powstać jony, aby popłynął prąd, włączamy lampę jarzeniową na ułamek sekundy – nie zapala się, a jedynie na jej końcu coś się jarzy – to żarnik uruchamiany przez układ zapłonowy (cewkę i przerywnik, tzw. starter). Żarnik wytwarza początkową „paczkę” jonów i elektronów, które są potrzebne do inicjacji wyładowania elektrycznego w całej rurze. Raz zainicjowane wyładowanie podtrzymuje się samo. (Potrzebny oczywiście zewnętrzny układ stabilizujący, ale wystarczy zwykły opornik włączony szeregowo w obwód lampy).

Wyjaśnienia wymagają dwie kwestie: ile jest jonów w gazie i dlaczego, mimo że rura świeci (czyli jest w niej jakby ogień) pozostaje zimna? Jonów w gazie jest niewiele – łatwo to policzyć. Lampa jarzeniowa „konsumuje” niewiele energii – rzędu 20 W. Oznacza to, że przepływa przez nią prąd rzędu 0,1 ampera: w ciągu 1/50 sekundy do elektrod dopływa mniej więcej 10^{16} elektronów (ładunek elektronu to $1,6 \times 10^{-19} \text{C}$). Zakładając dla uproszczenia, że są to wszystkie elektrony, które w ciągu cyklu 1/50 sekundy powstały w rurze, a objętość rury to 1 dm^3 , koncentracja elektronów (i jonów) jest więc rzędu $10^{13}/\text{cm}^3$. Ponieważ ciśnienie w rurze jest rzędu 1/10 ciśnienia atmosferycznego (a przy ciśnieniu atmosferycznym w $22,4 \text{ dm}^3$ jest $6,02 \times 10^{23}$ atomów), w rurze jarzeniowej jeden jon przypada na jakieś 100 tysięcy atomów (oczywiście, uczniowie mogą to policzyć dokładniej bez większych trudności). Taki rodzaj plazmy to w rzeczywistości *slabo zjonizowany gaz*.

Ponieważ i elektronów, i jonów jest niewiele, prawdopodobieństwo zderzeń z atomami gazu też jest niewielkie

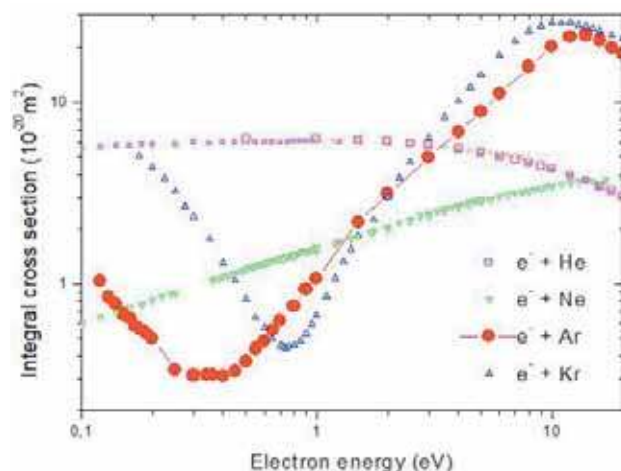
– gaz się nie grzeje. Ale pytanie, jaka jest temperatura tej plazmy, nie jest takie proste. Na rysunku 1 przedstawiam wyniki naszych badań – są to tzw. „przekroje czynne”, czyli prawdopodobieństwa zderzeń elektronów z atomami argonu. Najwyższa krzywa przedstawia *całkowite* prawdopodobieństwo zderzenia, ale nas interesuje krzywa zielona – prawdopodobieństwo, że w zderzeniu powstanie kolejny jon (bo w przeciwnym razie starter musiałby pracować cały czas).

Z wykresu wynika, że aby elektron dokonał jonizacji atomu, np. argonu, jego energia musi wynosić kilkanaście elektronowoltów – dokładniej powyżej 15,7 eV. Wynika z tego, że w rurze jarzeniowej podłączonej do napięcia 220 V jeden elektron może kilkakrotnie dokonać jonizacji.

Fizycy atomowi nie korzystają na co dzień z jednostki energii SI, tj. dżuli. Nie byłaby ona „namacalna”. W procesach atomowych fizycy (szczególnie tzw. doświadczalnicy) korzystają z jednostki eV: jest to energia, jaką nabywa elektron przyspieszony różnicą potencjału 1 V. Jest to jednostka bardzo praktyczna, bo energia niezbędna do oderwania elektronu z atomu wodoru to 13,6 eV. Zresztą, napięcie na baterii telefonu komórkowego, zazwyczaj litowej, to 3,7 V, bo to też proces przekazywania pojedynczych elektronów między elektrodami baterii. Ponieważ ładunek elektronu to $1,6 \times 10^{-19} \text{ C}$, $1 \text{ eV} = 1,6 \times 10^{-19} \text{ J}$.

Plazma niskotemperaturowa

Plazma w lampie jarzeniowej jest nie tylko słabo zjonizowana, ale jest też *niskotemperaturowa*. Aby to ocenić, potrzebujemy zamienić jednostki energii kinetycznej na temperaturę. Można tego dokonać używając stałej Boltzmanna k (albo stałej gazowej $R=8,31 \text{ J/(mol}\cdot\text{K)}$): ciepło właściwe gazu jednoatomowego przy stałym ciśnieniu wynosi $3/2R$, a więc energia termiczna jednego mola gazu (tak zwana energia wewnętrzna) wynosi $3/2 RT$; energia pojedynczego atomu wynosi $3/2 kT$ (stała Boltzmanna to stała gazowa podzielona przez liczbę Avogadro).



Ryc. 1. Prawdopodobieństwo zachodzenia określonych procesów w plazmie opisujemy za pomocą tzw. przekrojów czynnych: na zderzenia w ogólności, na wzbudzenie poziomów elektronowych (a w konsekwencji świecenie plazmy), na jonizację. Źródło: prace autora.

„Reguła” zamiany temperatury używa iloczynu kT : w jednostkach energii eV, dla temperatury pokojowej 300 K, wynosi on 25 meV.

Jak więc w gazie o temperaturze pokojowej, jaki wypełnia lampę jarzeniową może zachodzić jonizacja? Otóż temperatura gazu i energia kinetyczna elektronów nie są takie same: atomy argonu mają temperaturę 300 K, ale elektrony mają energię kinetyczną wyższą, jakieś 0,35 eV (co odpowiada minimum przekroju czynnego na rys. 1). Ale to za mało, aby dokonać jonizacji! Tak, ale pojęcie temperatury jest pojęciem statystycznym – średnio temperatura elektronów to 4000 K, ale część z tych elektronów, te „najgorętsze” mają energię wystarczającą do dokonania jonizacji atomów argonu.

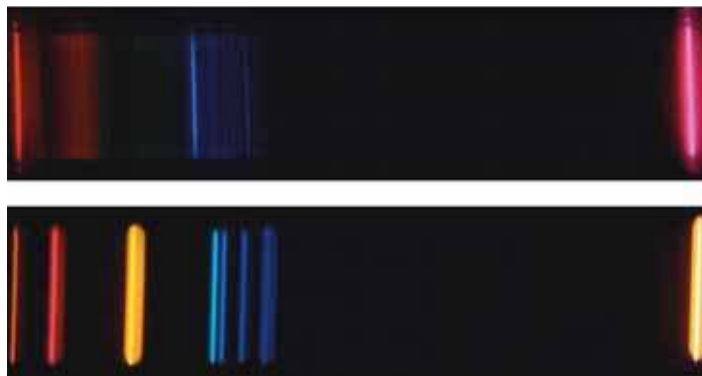
Plazma niskotemperaturowa to nie tylko lampa jarzeniowa, ale też zorza polarna (fot. 2), kula plazmowa czy wyladowania „szkolne”, w rurkach z gazem pod niskim ciśnieniem.

Ciemnia Faradaya, ciemnia Crookesa

Wyladowanie elektryczne w rurze jarzeniowej zachodzi w gazie rozrzedzonym: w tych warunkach dwie temperatury – atomów gazu i elektronów są różne. Temperatura gazu jest pokojowa: plazma, dość rozrzedzona, jest „niskotemperaturowa”. Zderzeń elektronów z atomami



Fot. 2. Zorza polarna jest przykładem plazmy niskotemperaturowej w bardzo rozrzedzonym gazie. Świecą w niej atomy azotu, tlenu i drobiny tlenku azotu (NO). Foto – Adobe Stock



Fot. 3. Nawet tania okulary dyfrakcyjne pozwalają na analizę kolorów w wyladowaniu elektrycznym (w tzw. rurkach Plückera). Plazma wodorowa (również ta na obrzeżach w reaktorze termojądrowym) jest różowa, plazma helowa – pomarańczowa. W rzeczywistości, jest to wiele linii (w wodorze, nie do końca dysocjowanym, oprócz 4 linii Balmera widać również pasma H₂).



Fot. 4. Aby zapalić lampę jarzeniową, nie trzeba nawet dotykać kuli plazmowej – wystarczy do niej zbliżyć jeden koniec (drugi jest uziemiony ręką wykładowcy). Napięcie na zewnątrz kuli może przekraczać 1 kV.

gazu jest mało, więc dwie temperatury – gazu i elektronów pozostają różne: zapalona lampę jarzeniową można spokojnie wziąć do ręki.

Dalsze obniżanie ciśnienia gazu prowadzi do coraz słabszego świecenia – widoczne stają się pewne obszary ciemne: zderzeń jest zbyt mało, aby spowodować wzbudzenie atomów gazu. Zresztą, atomy gazu ulegają wzbudzeniu i jonizacji nie tylko w zderzeniach z elektronami, ale również z jonami. Do jakiego poziomu energetycznego zostanie wzbudzony atom, zależy to od energii zderzenia i rodzaju „pocisku”.

Tak więc, zarówno w lampie jarzeniowej, jak w rurkach wyładowczych do demonstracji linii widmowych, występują rejony o świeceniu w innym kolorze lub obszary ciemne. Rurka helowa świeci intensywnie pomarańczowo, ale w okolicy katody plazma ma inny kolor – niebieskawy. Atom helu ma tylko dwa elektrony, ale poziomów energetycznych mnóstwo: linii widzialnych „gołym” okiem jest w zakresie optycznym kilkanaście (foto 3). Jaki jest kolor wyładowania, zależy od energii elektronów i koncentracji elektronów i jonów. Dziś to potrafimy modelować, ale w czasach Faradaya (około 1830 roku) i Crookesa (około 1860 roku) była to zagadka.

Prawie, że prawie potrafimy już zrozumieć, jak działa kula plazmowa. Tylko że w kuli plazmowej nie ma żadnych metalowych elektrod, jak w lampie jarzeniowej czy rurce spektralnej. Otóż kula plazmowa, teoretycznie opatentowana przez Nikołę Teslę jeszcze w XIX wieku, trafiła do sklepów dopiero jakieś 25 lat temu. Przy okazji wyprodukowania nowego rodzaju telewizorów – plazmowych, czyli płaskich, bez „rury katodowej” z tyłu.

Jak działa kula plazmowa

Zdjęcia kuli plazmowej są na stronach co drugiego centrum nauki i prawie w każdym podręczniku fizyki. Wodzenie palcem po szklanej bańce i przyglądanie się językom plazmy jest podobnie fascynujące jak oglądanie rybek w akwarium. Niestety, dużo bardziej niebezpieczne.

Niezwykłość kuli plazmowej polega na tym, że wyładowanie elektryczne zachodzi między dwoma szklanymi bańkami, a szkło, jak wiadomo, nie przewodzi prądu elektrycznego. Ale tak samo działają telewizory plazmo-

we, a także wyświetlacze dzisiejszych smartfonów: prąd przepływa, wydaje się, między dwoma płytkami szkła. W rzeczywistości, pod dolną szklaną płytką wyświetlacza smartfonu jest cienka warstwa metalu a przednia „szybka” jest z wierzchu pokryta przezroczystym a przewodzącym prąd elektryczny tlenkiem indy i cyny (tzw. ITO). Dwie przewodzące warstwy, rozdzielone dwoma szklanymi płytkami (w wyświetlaczach telefonów są to materiały bardziej wytrzymałe niż szkło, ale ciągle kruche) to jakby dwie okładki kondensatora – a przez kondensator prąd „płyń”, tylko że zmienny.

I w kuli plazmowej, i w ekranie telewizora plazmowego, płynący prąd jest to zamienne ładowanie się jednej lub drugiej szklanej powierzchni – dodatnio lub ujemnie. Ale częstotliwość takiej zamiany znaku elektrycznego okładek musi być bardzo szybka – w kuli plazmowej jest to jakieś 20-30 kHz. Rzeczywiście, wewnątrz mniejszej bańki szklanej umieszczona jest cewka dająca zmienne napięcie. Jak duże? Potrzebne jest doświadczenie.

Dlaczego lepiej nie bawić się kulą plazmową?

Najbardziej widowiskowym (a przy tym dydaktycznym) doświadczeniem z kulą plazmową jest zapalanie się lampy jarzeniowej już po zbliżeniu do kuli (wcale nie trzeba jej dotykać, foto 4). Ponieważ jarzeniówka zapala się przy napięciu sieciowym, na zewnętrznej powierzchni kuli panuje napięcie co najmniej 200 V. Z prawa Gaussa (uczniowie włoskich liceów mają je w programie) można oszacować, że wewnątrz kuli musi być to napięcie co najmniej 2 kV. Zresztą, nie jest to napięcie łatwe do pomiaru – dla zapewnienia wyładowania elektrycznego między dwoma szklanymi powierzchniami (tego rodzaju wyładowanie elektryczne nazywamy wyładowaniem z barierą dielektryczną „Dielectric barrier discharge”), nie tylko znak napięcia musi się zmieniać szybko, ale też zbrocza tych zmian powinny być prawie pionowe. Można próbować zmierzyć napięcie na zewnątrz kuli, ale potrzeba jest miernik o dużej oporności wewnętrznej (prądy, jakie generowane są przez pole elektryczne na zewnątrz kuli są niewielkie).

Dla sprawdzenia, jak wielkie są napięcia na zewnątrz kuli, jeden z autorów filmów internetowych proponuje położyć małą monetkę na kuli a następnie zbliżyć do monetki palec: przeskakuje iskra, a sam palec jest okopcony. Niezbyt mądra zabawa! Jeżeli w suchym powietrzu przeskakuje iskra o długości 1 cm, między dwoma przewodnikami panuje napięcie 30 kV. Iskra w internetowym doświadczeniu jest krótsza, ale nawet napięcie 1 kV może być śmiertelne, jeśli np. ktoś ma rozrusznik serca. Napięcie wewnątrz kuli dochodzi do 10 kV a dodatkowo są to impulsy napięcia: ot! dlaczego nie polecam kuli plazmowej jako zabawki. Natomiast jest ona nieoceniona do lekcji o czwartym stanie materii.

Kolory kuli plazmowej

Skład gazu w kuli plazmowej jest tajemnicą producenta. Studenci z Kalifornii [3], używając skomplikowanych technik stwierdzili, że jest tam neon z małym dodatkiem ksenonu. Ksenon jest najrzadszym z gazów szlachetnych, jest więc drogi. Po co ksenon w kuli? Otóż, już z roz-

ważać o lampie jarzeniowej wiemy, że jej działanie jest skomplikowane. Lampa jest wypełniona argonem, pod niskim ciśnieniem i w argonie płynie prąd, ale po rozgrzaniu się lampy, świecą w niej pary rtęci (i luminofor na wewnętrznej powierzchni, który zamienia światło nadfioletowe na widzialne). W kuli plazmowej ksenonu jest kilka procent, ale to on głównie świeci, i to jonów ksenonu jest więcej niż jonów neonu (energia jonizacji ksenonu to 12,1 eV a neonu 21,6 eV). Tak przynajmniej sugerują modele [4] stworzone w czasach, kiedy telewizory plazmowe dopiero były projektowane.

A neon? Neon pomaga w jonizacji ksenonu: wzbudzone atomy neonu zderzają się z atomami ksenonu. W tego rodzaju zderzeniu, wzbudzony elektron w atomie neonu wraca do poziomu podstawowego, a jeden elektron w atomie ksenonu zostaje oderwany. Taki proces nazywamy jonizacją Penninga.

Jasne staje się więc, że kolory kul plazmowych, a także wygląd samego wyładowania mogą być różnorodne (a modele fizyczne bardzo skomplikowane, gdyż muszą uwzględniać wiele możliwych procesów). Pokazujemy kilka przykładów na fot. 5.

Iskra, piorun, łuk spawalniczy

Niskotemperaturowe, niskociśnieniowe wyładowanie elektryczne w lampie jarzeniowej to nie jedyny rodzaj plazmy, z jaką mamy do czynienia na co dzień. W kuli plazmowej tworzą się sznury prądu elektrycznego, bo ciśnienie jest większe niż w lampie jarzeniowej – nieco mniejsze tylko od ciśnienia atmosferycznego (i znów zależy to od producenta). Aby płynął prąd w powietrzu pod ciśnieniem atmosferycznym, trzeba gaz zjonizować. Spawarka elektryczna to transformator dający napięcie kilkunastu tylko woltów, ale prąd kilkudziesięciu amperów. Spawacz najpierw dotyka elektrodą spawanego elementu, zaczyna płynąć spory prąd, elektroda paruje i wówczas trzeba sprawnie odsunąć elektrodę, na kilka milimetrów, tak aby ustabilizować prąd. Spawacze są najlepiej zarabiającymi pracownikami – niestety, nadfioletowe promieniowanie emitowane przez plazmę niszczy wzrok. A temperatura plazmy to tysiące stopni Celsjusza.

O ile łuk elektryczny jest wyładowaniem o niskim napięciu i dużym prądzie (a temperatura plazmy wystarcza do odparowania metalu), to iskra jest wyładowaniem

zachodzącym przy dużym napięciu (i małym prądzie). Maszyna elektrostatyczna daje piękne iskry, ale pamiętajmy, napięcia na jej elektrodach (czyli też w pobliżu i na kablach) są rzędu setek (!) tysięcy woltów [5]. Powstawanie iskry (i pioruna, który jest gigantyczną iskrą) jest na tyle skomplikowane, że nawet Richard Feynmann poświęcił mu oddzielny wykład.

Od wełny do układu scalonego

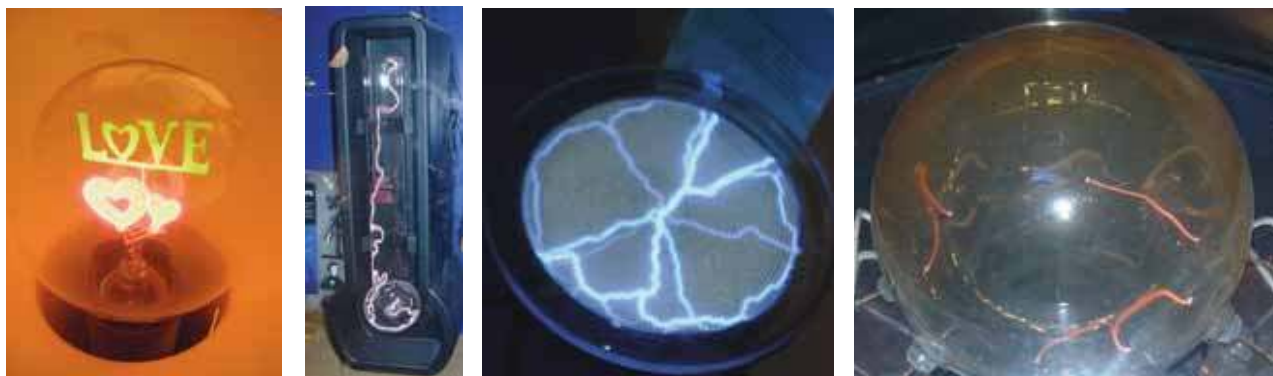
Zastosowań technologii plazmowych jest nieskończenie wiele – od czynienia higroskopijnymi ubrań, poprzez osadzenie tlenków ITO na ekranach telewizorów i smartfonów, do utwardzania wrzecion przedzarek bawełny. „Złote” warstwa na plastikowych elementach to TiN, nakładane są plazmowo. Wszystkie układy scalone muszą przejść kilka etapów plazmowych – od osadzania tlenków do ich trawienia. Każdy proces jest inny, a postęp technologiczny polega na znalezieniu takiego sposobu napyłania, trawienia, domieszkowania, jakiego konkurencja jeszcze nie zna. Kupując nowy, lepszy smartfon pamiętaj, że za nowymi możliwościami stoi praca nie tylko informatyków, ale i fizyków od plazmy.

Słońce w magnetycznym koszyku

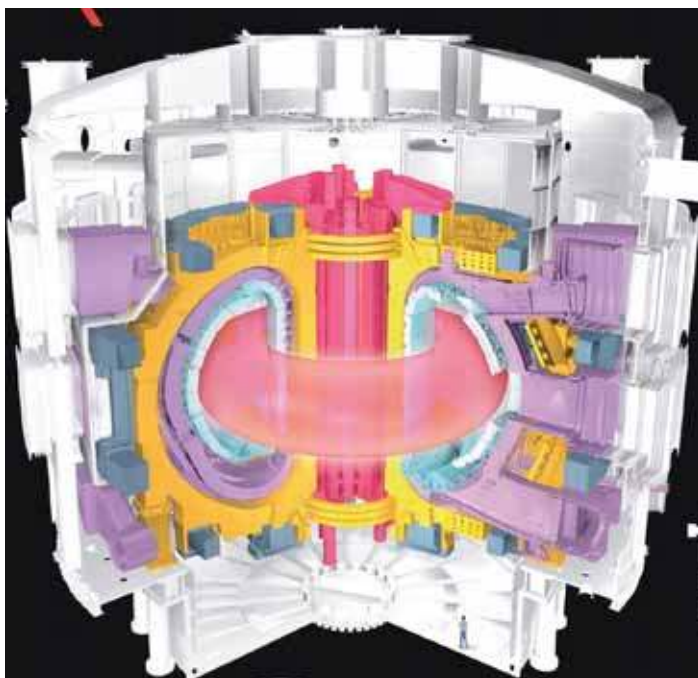
90% widzialnego wszechświata, a może i więcej, to plazma: wszystkie gwiazdy na nieboskłonie to kule gazu rozgrzanego do milionów stopni Celsjusza: wodoru, helu, ale i par węgla i żelaza. Trudno tę plazmę porównywać z lampą jarzeniową, bo nie tylko wszystkie atomy tracą elektrony, ale też tracą ich większość (jeśli nie wszystkie). To badając linie widmowe Słońca sprawdzamy, jakie poziomy energetyczne ma atom (a właściwie jon) żelaza bez dziesięciu lub dwudziestu elektronów.

I temperatury w gwiazdach są ogromne: na powierzchni co prawda jedynie kilka tysięcy stopni (5500°C Słońca, więcej np. niebieskich gigantów), ale wewnątrz Słońca – 15 milionów. Takie energie są potrzebne, aby protony (albo jądra helu) przezwyciężyły odpychanie elektrostatyczne i zbliżyły się na odległości (rzędu 10^{-15} m), gdzie działają ogromne, przyciągające siły jądrowe. Tak mogą powstać jądra węgla, tlenu, siarki – jednym słowem wszystkiego, co jest potrzebna dla życia.

Ale Słońce to nie tylko reaktor jądrowy – przede wszystkim niewyczerpane (na jeszcze 4 miliardy lat)



Fot. 5. Przykłady plazmy z witryn sklepowych w Paryżu: (a) mała żarówka wypełniona neonem pod ciśnieniem atmosferycznym, plazma widoczna w pobliżu włókna żarówki, (b) sznur plazmy (napisano „nie dotykać szkła”) i (c) płaska lampa plazmowa (chyba z domieszką tlenu) w sklepie Science Center La Villette, (d) bańka plazmowa, chyba wypełniona neonem (S.C. La Villette) – widać sznury plazmy. Foto Maria Karwasz



Fot. 6. Pierwszy reaktor termojądrowy, który dostarczy netto 500 MW energii powstaje w Cadarache we Francji, jako wspólny światowy projekt badawczy. Ma wielkość 10-piętrowego budynku, plazma zajmuje 900 m³. Wersja przemysłowa będzie gotowa w 2050 roku. Źródło: ITER.

źródło energii docierającej również do Ziemi. Jak pisze w książce astronomicznej dla dzieci [6], w ciągu sekundy Słońce „spala” tyle wodoru, ile wydobywa się węgla w Polsce przez rok. Warto byłoby mieć podobny reaktor na ziemi. I właśnie nad tym pracujemy [7].

W Cadarache, na południu Francji, powstaje reaktor, który będzie zamieniał wodór (dokładniej jądra „ciężkiego” wodoru, czyli deuteru) w hel. Podobnie jak dzieje się to w Słońcu. Ale we wnętrzu Słońca panuje ogromne ciśnienie – średnia gęstość Słońca, zbudowanego z plazmy, czyli gazu, równa jest gęstości wody. Na ziemi nie ma ogromnego pola grawitacyjnego, aby tak ścisnąć plazmę – temperatura plazmy reaktora termojądrowego musi więc być wyższa – nie 15, ale 150 milionów stopni. Nie ma materiału na zbiornik tak gorącej plazmy. No chyba że użyjemy ogromnego pola magnetycznego, które ściśnie sznur plazmy. Ale i tak na ścianki reaktora potrzebny jest wolfram – dlatego brakuje go na żarówkę.

W następnym numerze:

Hemodynamika obliczeniowa gałęzią medycyny przyszłości

Marcin Majka

Wśród najprężniej rozwijających się dziedzin nauk niewątpliwie miejsce w czołówce zajmuje medycyna. Dzięki rozwojowi techniki powstają nowe metody diagnostyczne, sprzęty chirurgiczne ale przede wszystkim lekarstwa i urządzenia umożliwiające lepsze ich poddawanie. Niestety nic z tych rzeczy nie miałyby miejsca gdyby nie współpraca lekarzy z naukowcami – fizykami, chemikami, biologami oraz inżynierami. Od dawna rozwija się mało znana gałąź

Reaktor termojądrowy ma rozmiary 10-cio piętrowego budynku a objętość plazmy, w Cadarache, to 900 m³, rys. 6. Będzie gotowy za kilka lat, a reaktor przemysłowy, dla wytwarzania energii elektrycznej już jest projektowany w Republice Korei: będzie gotowy w 2050 roku.

P.S. Naukowcy pilnie potrzebni

Dlaczego to wszystko o plazmie i tak szczegółowo? Reaktor termojądrowy w Cadarache to największy, wspólny projekt badawczy całego świata, o wartości kilkunastu miliardów euro. Tylko lot na Marsa pochłonie więcej pieniędzy. Ale gra w Cadarache jest warta świeczki: paliwa kopalne – ropa, węgiel, gaz wyczerpują się. Można dyskutować, na ile lat ich starczy, ale na pewno nie na sto. A energie alternatywne też nie zaspokoją wszystkich potrzeb świata. Z dachu domku jednorodzinnego starczy energii fotowoltaicznej tylko na własny użytek (o ile, oczywiście, świeci słońce). A energii termojądrowej, z deuteru w oceanach winno starczyć na trzy tysiące lat. Warto, aby i polscy, młodzi uczeni w tym pasjonującym wyścigu z czasem uczestniczyli!

Grzegorz Karwasz

Zakład Dydaktyki Fizyki

Uniwersytet Mikołaja Kopernika w Toruniu

Profesor G. Karwasz zajmuje się plazmą zawodowo, badając procesy elementarne, mierząc jej właściwości i proponując modele reaktorów, w tym termojądrowych: we Włoszech, Czechach, Korei, Australii. Artykuł jest poświęcony pamięci prof. Zenona Zakrzewskiego, od którego autor uczył się o plazmie.

LITERATURA

- [1] G. Karwasz i wsp. Cztery i pół stany skupienia, Pokazy interaktywne, Toruń 2017, http://dydaktyka.fizyka.umk.pl/4_i_pol_stany_skupienia/
- [2] G. Karwasz, M. Sadowska, K. Rochowicz, Toruński podręcznik do fizyki. I Mechanika. Wyd. Naukowe UMK 2009, <http://dydaktyka.fizyka.umk.pl/Ksiazki/Porecznik/1.5.pdf>
- [3] [3] M. J. Burin et al., Princeton Plasma Physics Laboratory, On filament structure and propagation within a commercial plasma globe, Physics of Plasmas 22, 053509 (2015)
- [4] <https://aip.scitation.org/doi/am-pdf/10.1063/1.4919939>
- [5] J. Meunier, Ph. Belenger, and J. P. Boeuf, Numerical model of an ac plasma display panel cell in neon-xenon mixtures, J. Appl. Phys. 78 (1995) 731
- [6] M. Sadowska, G. Karwasz, Stara, pocziwa maszyna elektromagnetyczna, Fizyka w Szkole, nr 5/2011, str. 40, http://dydaktyka.fizyka.umk.pl/Publikacje_2011/Maszyna_el_2011.pdf
- [7] G. Karwasz, Mały astronom. Przewodnik dla dzieci, Publicat, Poznań, 2016.
- [8] G. Karwasz, Słońce w(magnetycznym) koszyku, Głos Uczelni, UMK, 3/2017, str. 24, http://dydaktyka.fizyka.umk.pl/Publikacje_2017/GK_Slonce_w_koszyku_2017.pdf

łącącą wiedzę medyczną opartą na pomiarach klinicznych oraz wiedzę i umiejętności fizyków i informatyków. Jest to hemodynamika obliczeniowa, która obecną wiedzę medyczną wykorzystuje do opisu matematycznego oraz fizycznego zjawisk zachodzących w układzie krążenia człowieka. To dzięki hemodynamice obliczeniowej możemy z wyprzedzeniem przewidywać nagłe incydenty sercowo-naczyniowe i co ważniejsze, określać ich konsekwencje. Hemodynamika obliczeniowa oprócz ściśle aplikacyjnego zastosowania w medycynie posiada w sobie ogromny potencjał naukowy, pozwalający odkrywać dotąd nieznanne mechanizmy i reakcje zachodzące w układzie krążenia człowieka. Nic dziwnego, że coraz większa liczba naukowców z całego świata zaczyna się interesować łąčeniem nauk ścisłych i medycyny.

Termodynamika: główne koncepcje i etapy rozwoju. Część 4

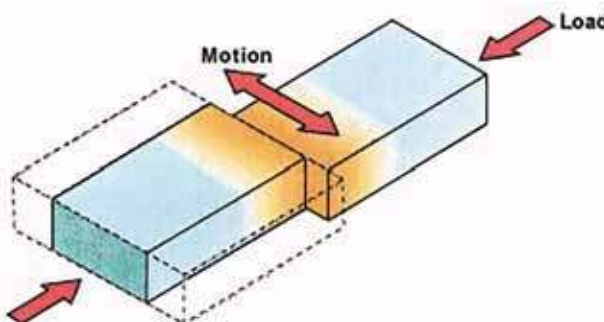
Ciepło tarcia

Alfred Zmitrowicz

Tarcie jest procesem, w którym dochodzi do zamiany pracy sił tarcia głównie na ciepło w sposób nieodwracalny. Od bardzo dawna znane były proste sposoby rozniecania ognia za pomocą tarcia. Ciepło tarcia powstaje w hamulcach ciernych pojazdów (samochodów, pociągów, samolotów), w sprzęgłach ciernych, w procesach zgrzewania tarcowego metali, w procesach cięcia (metali, drewna, itp.), w procesach wiercenia (w metalach, drewnie, skałach, itp.), w procesach szlifowania (metali, kamieni, szkła, itp.) (Foto 14, Rys. 19).

Ciepło generowane jest w smarowanych łożyskach maszyn. Ważnym zjawiskiem jest ciepło tarcia pojazdów kosmicznych o powietrze przy przechodzeniu przez atmosferę podczas powrotu z przestrzeni kosmicznej na Ziemię. W Wikipedii hasło pt. „konwersja energii” jako przykład podaje proces hamowania samochodu, gdzie energia kinetyczna samochodu zamieniana jest na energię cieplną w wyniku tarcia między tarczami hamulcowymi lub między bębniami a klockami i okładzinami hamulcowymi. Wytworzone ciepło jest

odprowadzane do otoczenia. Ilość wytworzonego ciepła zależy od ilości włożonej pracy mechanicznej sił tarcia. Pracy zużytej na pokonanie oporów tarcia nie potrafimy odzyskać za pomocą odwrócenia kierunku procesu, tracimy ją bezpowrotnie.



Rys. 19. Rysunek generacji ciepła tarcia podczas zgrzewania tarcowego, źródło: S.W. Kallee et. al., *Friction welding of aero engine components*, 2003 r.

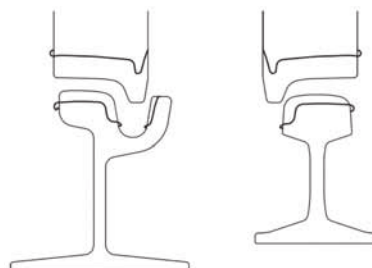


Foto 15. Zużycie bieżnika opony samochodowej, ubytki materiału spowodowane są ścieraniem gumy



Foto 14. Generacja ciepła tarcia podczas szlifowania powierzchni obrabianego materiału, źródło: Wikipedia, Schleifen

Rys. 20. Rysunek zużycia szyny i obręczy koła kolejowego, ubytki materiału spowodowane są głównie ścieraniem stali (wielkość ubytków przedstawiono w powiększonej skali), źródło: Jacek Makuch, ćwiczenie 2, profilomierz



Termodynamiczne modele tarcia, ciepła tarcia i zużycia materiałów

Zjawisku tarcia towarzyszy nie tylko wydzielanie się ciepła, ale również propagacja fali akustycznej, elektryzowanie się materiałów, uszkodzenia powierzchni zewnętrznej brył materiałów tj. zużycie materiałów wskutek ścierania (Foto 15, Rys. 20). Termodynamiczne modele tarcia, zużycia i adhezji materiałów [13, 14], ułatwiają opis sprzężeń termomechanicznych i wprowadzają podstawowe ograniczenia na modele zjawisk. W ramach termodynamiki ośrodków ciągłych, w pierwszym kroku należy określić zbiory zmiennych niezależnych i zależnych modeli tarcia, ciepła tarcia i zużycia materiałów. Zmienne te przedstawiono w Tabeli 3.

Wskaźniki dolne dotyczą dwuwymiarowego układu współrzędnych na powierzchni zewnętrznej bryły materiału i przyjmują wartości $i = 1, 2$.

Następnie bada się ograniczenia wynikające z procedur termodynamiki ośrodków ciągłych. Po pierwsze, aksjomat obiektywności materialnej ogranicza zbiór wielkości fizycznych, które mogą pełnić rolę zmiennych w równaniach konstytutywnych. W równaniach siły tarcia, strumienia ciepła tarcia i prędkości zużycia wielkościami obiektywnymi materialnie są: wektor siły tarcia, wielkość docisku, wektor prędkości poślizgu (gdyż jest różnicą prędkości cząstek dwóch stykających się elementów, inaczej mówiąc prędkością względną w punkcie), wielkość prędkości poślizgu, wielkość prędkości zużycia. Spełnienie aksjomatu obiektywności materialnej oznacza, że dwaj obserwatorzy, wyposażeni w dwa różne układy odniesienia, zaobserwują taką samą siłę tarcia, taką samą prędkość poślizgu, taką samą prędkość zużycia, itd.

Zgodnie z zasadą zachowania energii, wszystkie formy energii mechanicznej i cieplnej są addytywne i muszą być zbilansowane. Po zapisaniu bilansu energii dla układu złożonego ze stykających i trących się ciał stałych otrzymujemy warunek, który mówi, że moc siły tarcia M jest zamieniana na ciepło tarcia Q^* i energię wydatkowaną w procesie zużycia E^* . Moc siły tarcia M jest iloczynem skalarnym wektora siły tarcia o składowych F_i i wektora prędkości poślizgu o składowych P_i , tj.

$$M = F_1 P_1 + F_2 P_2 = Q^* + E^*$$

Dobrym oszacowaniem wielkości energii zużycia E^* dla niektórych materiałów (metali) jest przyrównanie

jej do energii deformacji, gdy naprężenia osiągają granicę plastyczności a odkształcenia są odkształceniami plastycznymi. Ponadto, badania eksperymentalne pokazują, że około 90% mocy sił tarcia zamienia się w sposób nieodwracalny w ciepło tarcia wnikające do trących się elementów. Innymi formami przemiany mocy sił tarcia są: energia fal dźwiękowych, energia elektryzowania, itp. Po podstawieniu modeli siły tarcia, strumienia ciepła tarcia i prędkości zużycia do równania mocy sił tarcia otrzymuje się ograniczenie na wartość współczynnika równania prędkości zużycia.

Nierówność produkcji entropii mówi, że przemiana energii mechanicznej może następować tylko w jednym kierunku. W przypadku tarcia, moc siły tarcia jest zamieniana na ciepło i energię zużycia, lecz ciepła i energii wydatkowanej w procesie zużycia nie można zamienić na tarcie. Po zapisaniu nierówności produkcji entropii dla układu materialnego złożonego ze stykających i trących się ciał stałych otrzymujemy kilka nierówności, które nakładają ograniczenia na modele siły tarcia, strumienia ciepła tarcia i prędkości zużycia. W przypadku siły tarcia nierówność ta ma następującą postać

$$M = F_1 P_1 + F_2 P_2 \leq 0$$

Oznacza to, że moc siły tarcia M jest niedodatnia dla każdej prędkości poślizgu. Innymi słowami, tarcie zawsze rozprasza energię mechaniczną. Z nierówności powyższej wynikają ograniczenia na współczynniki modeli siły tarcia (w modelu Coulomba współczynnik tarcia musi być liczbą nieujemną). Podobne warunki zachodzą w przypadku modeli prędkości zużycia i modeli strumienia ciepła tarcia. Na przykład, nierówność produkcji entropii nakłada ograniczenia na współczynniki modeli prędkości zużycia (w modelu Archarda współczynnik intensywności zużycia musi być liczbą nieujemną).

Podstawowymi modelami zjawisk tarcia, ciepła tarcia i zużycia są: model tarcia Coulomba, liniowy model strumienia ciepła tarcia i model zużycia Archarda, tj.

$$F_i = \mu N \frac{P_i}{\sqrt{(P_1)^2 + (P_2)^2}}$$

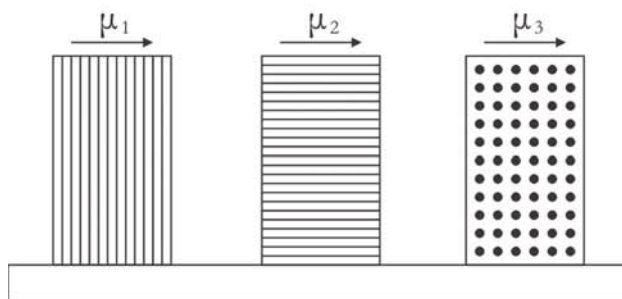
$$q^* = -wN \sqrt{(P_1)^2 + (P_2)^2}$$

$$v^* = -aN \sqrt{(P_1)^2 + (P_2)^2}$$

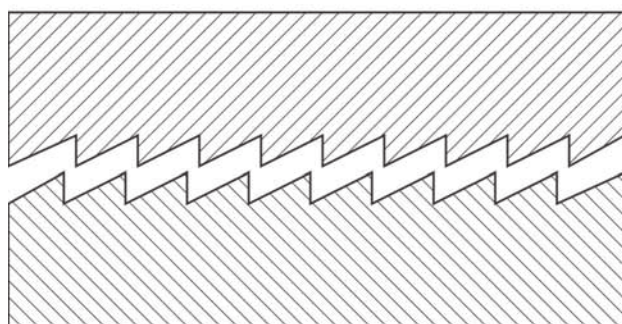
Tabela 3. Podstawowe zmienne termodynamicznych modeli tarcia, ciepła tarcia i zużycia materiałów

x_i	wektor położenia cząstki na powierzchni zewnętrznej bryły materiału	są to zmienne modeli niezależne i określają one proces termodynamiczny dowolnej cząstki materiału X_i w dowolnej chwili czasu t na powierzchni zewnętrznej bryły materiału
θ	temperatura absolutna na powierzchni zewnętrznej bryły materiału	
F_i	wektor siły tarcia na powierzchni zewnętrznej bryły materiału	są to zmienne modeli zależne od wektora położenia cząstki x_i na powierzchni zewnętrznej bryły materiału i temperatury absolutnej θ na tej powierzchni oraz ich gradientów względem współrzędnych i czasu
q^+	strumień ciepła tarcia wnikającego do wnętrza bryły materiału	
v^+	prędkość zużycia tj. prędkość ścierania powierzchni materiału	

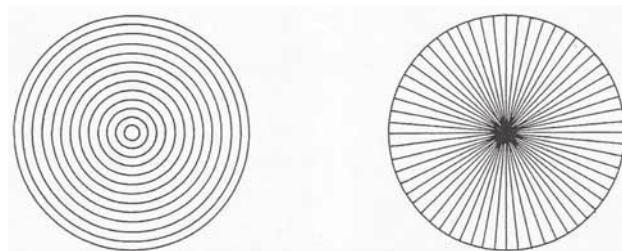
gdzie, μ współczynnik tarcia, współczynnik strumienia ciepła tarcia, współczynnik intensywności zużycia, N docisk, P_i składowe wektora prędkości poślizgu.



Rys. 21. Wielkość współczynnika tarcia kompozytu zależy od kierunku poślizgu względem zbrojenia; orientacja zbrojenia: prostopadła, podłużna i poprzeczna (względem kierunku poślizgu)



Rys. 22. Niesymetryczne tarcie dla poślizgu w prawo (małe tarcie), dla poślizgu w lewo (duże tarcie), gdy struktura chropowatości powierzchni jest niesymetryczna



Rys. 23. Przykłady struktury chropowatości powierzchni po obróbce mechanicznej: współśrodkowa i promieniowa (względem środka powierzchni)

Rozwój nowych modeli tarcia i zużycia jest stymulowany przez współczesną technikę, która stawia wymagania dokładniejszego opisu własności tarcia i zużycia materiałów inżynierskich w szczególności materiałów ze złożoną mikrostrukturą. Materiały te posiadają korzystne właściwości wytrzymałościowe i tarcia i zużycia pożądane w licznych zastosowaniach w technice.

W pojedynczych kryształach materiałów, w niektórych polimerach, ceramikach, materiałach o strukturze warstwowej, zbrojonych i laminowanych kompozytach, biomateriałach, warstwach monocząsteczkowych, opór tarcia i stopień zużycia zależą od kierunku poślizgu (Rys. 21). Zagadnienia te stanowią motywację dla badań anizotropii tarcia i zużycia materiałów. Źródłem anizotropii tarcia może być również ukierunkowana chropowatość powierzchni zewnętrznych brył materiałów (Rys. 22, Rys. 23). Poślizg może być łatwiejszy w jednym kierunku i trudniejszy w drugim. Mogą istnieć kierunki poślizgu o szczególnych własnościach, tzw. kierunki

główne tarcia. Ponadto, poślizg może inicjować ewolucję i reorientację mikrostruktury na powierzchniach ślizgowych i w warstwach wierzchnich w niektórych polimerach i materiałach warstwowych takich jak grafit i dwusiarczki molibdenu. Ewolucja mikrostruktury powoduje zmiany w procesach tarcia i zużycia. W tym przypadku, opór ruchu i stopień zużycia zależą od kierunku ruchu i dodatkowo od kształtu toru ruchu, w szczególności od krzywizny toru ruchu (tarcie niejednorodne).

W ramach termodynamiki ośrodków ciągłych można sformułować liniowe i nieliniowe modele tarcia anizotropowego. Liniowe modele tarcia anizotropowego definiują co najwyżej dwa kierunki główne tarcia oraz krzywe graniczne sił tarcia (tzw. przekroje stożka tarcia) o eliptycznych i kołowych kształtach. Nieliniowe modele opisują tarcie o dowolnej skończonej liczbie kierunków głównych i dowolnych kształtach przekrojów stożka tarcia. Modele z tensorami tarcia zależnymi od kierunku poślizgu określają tarcie niesymetryczne (tj. o różnych wartościach w zależności od znaku kierunku poślizgu). Korzystając z własności symetrii przeprowadza się klasyfikację różnych typów tarcia: izotropowe, ortotropowe, anizotropowe, osiowo-symetryczne, jednostronne, trigonalne, tetragonalne. Anizotropowe zużycie modeluje się przez rozwinięcie modelu zużycia Archarda. Zakłada się, że intensywność zużycia jest funkcją parametru kierunku poślizgu i jest pokrewna (homogeniczna) z funkcją współczynnika tarcia anizotropowego.

Można sformułować modele, które opisują ewolucję anizotropii i niejednorodności tarcia i zużycia wywołanych krzywizną toru poślizgu. Są to modele rzędu pierwszego, drugiego i wyższych ze względu na potęgi krzywizny toru poślizgu. Zmiennymi niezależnymi równań konstytutywnych są: wektor prędkości poślizgu i jego pochodna oraz tensory rzędu nieparzystego zbudowane z tych wektorów. Krzywizna toru ruchu jest źródłem: (a) dodatkowego oporu ruchu (siły typu dyssypatywnego), (b) siły reakcji więzów normalnych do kierunku poślizgu (siły typu żyroskopowego). Tym sposobem przyczynia się do powstania dodatkowego tarcia (dodatniego lub ujemnego), a trajektoria poślizgu może w sposób istotny zmieniać swój kształt.

Termodynamika w różnych obszarach wiedzy

Zjawiska cieplne występują nie tylko w maszynach cieplnych lecz również w innych zastosowaniach technicznych, kilka przykładów pokazano na fotografiach (Foto 16-17). Zjawiska cieplne występują w przyrodzie (Foto 19-21). Poniżej omówiono wybrane przykłady współczesnych zastosowań termodynamiki i nauki o cieple.

(a) *Termodynamika w chemii.* Efekty cieplne występują w różnych procesach chemicznych, np. rozpuszczanie substancji w rozpuszczalniku, przemiany fazowe (parowanie i skraplanie płynów, wiązanie i twardnienie klejów, krzepnięcie roztworów), itp. W procesach spalania paliw stałych i gazowych powstaje energia chemiczna. W reakcjach chemicznych następuje uwalnianie lub pochłanianie



Foto 16. Zabytkowa gazownia miejska w Górowie Ławieckim na Warmii z 1908 r., gaz miejski (gaz koksowniczy) wytwarzano za pomocą cieplnej obróbki węgla kamiennego



Foto 19. Globalne ocieplenie wywiera niekorzystny wpływ na obszary Ziemi pokryte lodem, na fotografii południowa część Grenlandii, na Grenlandii znajduje się 8% zasobów lodu na Ziemi, źródło: fotografie wykonano z wahadłowca



Foto 17. Schładzanie ogromnych ilości wody z elektrowni w chłodniach kominowych, woda oddaje ciepło do powietrza, źródło: galeria fotografii „Kom Spec”, chłodnie elektrowni Bełchatów



Foto 20. Wybuch wulkanu, chmury gorących popiołów i gazów, źródło: Photovolcanica by Dr Richard Roscoe, wulkan Colima (Meksyk)



Foto 18. Zanieczyszczenie powietrza przez niewłaściwe spalanie w piecach domowych, źródło: Czyste ogrzewanie, ekonomiczne spalanie drewna i węgla



Foto 21. Erupcja gejzera, kolumna wody i pary wodnej w temperaturze wrzenia, źródło: Wikipedia, gejzer Old Faithful (Yellowstone)

energii (ciepła). Reakcjom chemicznym może towarzyszyć wzrost lub obniżenie temperatury. Termodynamika chemiczna bada efekty energetyczne reakcji chemicznych. W badaniach wykorzystuje się procedury termodynamiki: zachowanie masy i energii, wzrost entropii. Ważną rolę pełnią pojęcia z zakresu chemii, np. potencjał chemiczny, gdyż różnica potencjału chemicznego jest siłą napędowej dyfuzji, rozpuszczania, przemian fazowych i reakcji chemicznych.

(b) *Termodynamika w biologii.* Organizmy biologiczne mogą istnieć wyłącznie dzięki nieustannej wymianie materii i energii z otoczeniem. W organizmach żywych ma miejsce ciągłe dostarczanie i wydalenie materii. Procesy termodynamiczne związane z życiem są przede wszystkim procesami biochemicznymi, należą do nich: metabolizm, czyli energia chemiczna pochodząca ze spalania (utleniania) pożywienia podczas trawienia; fotosynteza czyli synteza nowych składników przez rośliny,

gdzie energia potrzebna do reakcji chemicznych uzyskiwana jest ze światła słonecznego (za pomocą fotonów); itp. Organizmy żywe są otwartymi układami termodynamicznymi. Zachodzą w nich nieodwracalne procesy termodynamiczne, takie jak przepływ materii i energii, zmiany w otoczeniu, ale również zjawiska typowe dla biologii jak: adaptacja, ewolucja, samo-organizacja, dobór naturalny, itp. W biologii coraz to wyższe formy życia pojawiają się i organizują się w sposób naturalny.

(c) *Termodynamika w geofizyce*. Zjawiska w środowisku naturalnym Ziemi mają ścisły związek z ciepłem. Głównymi źródłami ciepła w procesach geofizycznych są: promieniowanie słoneczne, energia termiczna wnętrza Ziemi (ciepło geotermalne), ciepło tarcia generowane podczas ruchu obiektów geofizycznych (np. uskoków tektonicznych), itp. Termodynamika jest stosowana w badaniach zjawisk w skorupie ziemskiej i we wnętrzu Ziemi (trzęsienia ziemi, erupcje wulkanów i gejzerów, mechanika skał i gruntów) na powierzchni Ziemi (dynamika lodowców i lawin, hydromechanika akwenów wodnych [4]) i w atmosferze ziemskiej (opady atmosferyczne, naturalne ruchy mas powietrza, temperatura powietrza). Do zadań termodynamiki należą: analiza wymiany masy i ciepła (w skorupie ziemskiej, w akwenach wodnych, w atmosferze), analiza bilansów cieplnych Ziemi, poszukiwanie związków między cieplnymi i mechanicznymi procesami geofizycznymi, analiza efektu dziury ozonowej, globalnego ocieplenia i zmian klimatu Ziemi, itp. Na przykład termika morza (specjalność w dyscyplinie oceanologia) analizuje bilans cieplny mórz i oceanów w powiązaniu z atmosferą i skorupą ziemską, tj. bada czynniki kształtujące klimat Ziemi.

(d) *Termodynamika w astrofizyce i kosmologii*. W XIX w. sformułowano teorię Śmierci Ciepłej Wszechświata (Herman Helmholtz, Rudolf Clausius). Model ten zakłada, że **Wszechświat** jest układem zamkniętym, entropia jest miarą porządku oraz powstający lokalnie porządek odbywa się kosztem jego znacznego pogorszenia w pozostałej części układu. Jest to model ewolucji Wszechświata w stronę nieporządku (dezorganizacji). W XX w. sformułowano teorię Wielkiego Wybuchu (Stephen Hawking). Zgodnie z tym modelem, na początku była materia o dużej gęstości i ogromnej temperaturze skupiona w jednym punkcie. Wszechświat powstał podczas Wielkiego Wybuchu (ang. *Big Bang*) i od tamtej chwili rozszerza się i stygnie. W astrofizyce występują obiekty o bardzo wysokich temperaturach (np. jądro Słońca), generalnie obecnie obserwowany Kosmos jest zimną przestrzenią. Ekstremalnie wysokie temperatury bardzo gorących gazów występujące w niektórych obiektach w astrofizyce sprawiły, że należało uwzględnić poprawki relatywistyczne w równaniach termodynamiki. Na tej podstawie wyrosła termodynamika relatywistyczna.

(d) *Termodynamika w ekonomii*. W ekonomii zakłada się analogię między zasadami termodynamiki klasycznej a zachowaniem się systemów ekonomicznych i finansowych. Dla pojęć ekonomicznych przypisuje się odpowiednie pojęcia z termodynamiki. Należą do nich: układ lub system ekonomiczny (np. osoba, gospodarstwo

domowe, gmina, kraj), stan ekonomiczny, równowaga systemu ekonomicznego (tj. równowaga rynków kapitałowych), energia ekonomiczna (tj. kapitał lub wartość ekonomiczna), praca (jako sposób pomnażania kapitału), entropia ekonomiczna (np. jako miara ryzyka i niepewności), potencjał ekonomiczny (miara zdolności do wykonania pracy), przepływy kapitałowe (gradient kapitału jest przyczyną przepływu kapitału), naprężenia kapitałowe, itp. Następnie korzysta się z zasad termodynamiki w celu określenia relacji analitycznych między wprowadzonymi analogiami pojęć ekonomicznych.

Podsumowanie

1. Czy klasyczna termodynamika płynów (gazów i cieczy) jest lepsza od termodynamiki ośrodków ciągłych (materiałów odkształcalnych)? Odpowiadając na to pytanie trzeba uwzględnić następujące fakty.

- (a) Po pierwsze, gazy i ciecz mają inne niż materiały odkształcalne obszary praktycznych zastosowań. W technice gazy i ciecz są substancjami, które służą do transportu (przenoszenia) ciepła, masy i pędu. Płyny mogą swobodnie przemieszczać się i łatwo transportują energię (cieplną, kinetyczną), dlatego w maszynach cieplnych pełnią rolę substancji roboczej. Możemy powiedzieć, że gazy i ciecz mają za zadanie przepłynąć przez określone kanały maszyny (turbiny, sprężarki, silnika spalinowego), urządzenia (chłodziarki, wymiennika ciepła, strumienicy) i instalacji przemysłowej (kotła energetycznego, chłodni kominowej, rurociągu) i mają wykonać użyteczną pracę (silniki cieplne). Natomiast materiały odkształcalne (ciała stałe) są wykorzystywane praktycznie do przenoszenia obciążeń statycznych, dynamicznych i cieplnych działających zwykle na powierzchniach zewnętrznych ciał stałych. Materiały odkształcalne przenoszą energię i ciepło, lecz w przeciwieństwie do płynów nie ma w nich transportu masy. Potocznie możemy powiedzieć, że materiały odkształcalne mają za zadanie utrzymać (wytrzymać) nałożone na nieobciążenia. Dlatego z materiałów odkształcalnych wznoszone są budynki, drogi i mosty, konstruowane są maszyny, samochody, pociągi, statki, samoloty, narzędzia, instalacje przemysłowe, meble, telewizory, komputery, telefony, itp.
- (b) Druga podstawowa różnica wynika z budowy wewnętrznej substancji materialnych. Z punktu widzenia mechaniki klasycznej gazy i ciecz są materiałami o nieskomplikowanej strukturze wewnętrznej i pod wpływem bodźców zewnętrznych łatwo przechodzą z fazy ciekłej w gazową i odwrotnie. Natomiast materiały odkształcalne najczęściej mają bardzo bogatą budowę wewnętrzną i pod wpływem zwykle stosowanych obciążeń zewnętrznych nie zmieniają fazy tj. pozostają ciałami stałymi. Oczywiście istnieją procesy technologiczne w których metale (ciała stałe) pod wpływem ciepła zmieniają fazę ze stałej na ciekłą np. przy spawaniu i lutowaniu, w procesach hutniczych i odlewniczych

(topnienie, krzepnięcie, krystalizacja), itp. Ponadto w procesach cięcia laserowego metali ciepło sprawia, że metale zmieniają fazę stałą na gazową. Podobne zjawiska występują w materiałach zmiennofazowych (stopach metalicznych) stosowanych do magazynowania ciepła, gdzie mają miejsce przemiany fazowe typu: ciało stałe – ciecz, ciecz – ciało stałe, itp. Są to jednak szczególne przypadki zastosowań ciał stałych.

- (c) Inne są miary odkształcenia i naprężenia w płynach i materiałach odkształcalnych, wynika to z różnic w budowie wewnętrznej tych substancji. Kinematykę opisuje się w płynie wektorem prędkości, a w materiale odkształcalnym wektorem przemieszczenia. W płynach ma miejsce tylko deformacja objętościowa (mierzy się zmianę objętości), a w materiałach odkształcalnych deformacja następuje we wszystkich kierunkach (objętościowa i postaciowa) i opisuje się ją tensorem odkształcenia. W płynach siły wewnętrzne (naprężenia) mierzy się ciśnieniem i opisuje jedną składową sferycznego tensora ciśnienia, w materiałach odkształcalnych naprężenia opisuje sześć niezależnych składowych symetrycznego tensora naprężenia (drugiego rzędu). Wielkości fizyczne opisujące przepływ ciepła w przypadku płynów i ciał stałych są jednakowe, należą do nich: temperatura, strumienie ciepła, źródła ciepła, itp.
- (d) Czy można utworzyć taką termodynamikę, która nie zawiera równań ruchu substancji materialnej? Za przykład podaje się obiegi termodynamiczne (obieg Carnota i inne), gdzie nie analizuje się ruchu. Należy pamiętać, że obiegi termodynamiczne w maszynach cieplnych polegają na wykonanych kolejno kilku przemianach fazowych przez substancje robocze którymi są płyny. Ten wyjątkowy przypadek nie można uogólnić na materiały odkształcalne, gdyż w powszechnych zastosowaniach technicznych ciała stałe nie pełnią roli substancji roboczej (nie ma w nich transportu masy i najczęściej nie podlegają one przemianom fazowym). Przypomnijmy, że głównym zadaniem mechaniki jest opis ruchu. Wobec tego, nie ma mechaniki bez bilansów pędu i krętu (tj. równań ruchu), z tego powodu termodynamika materiałów odkształcalnych uwzględniła bilanse pędu i krętu.
- (e) Trudność prostego podziału substancji materialnych na płyny i ciała stałe powstaje wtedy, gdy mamy do czynienia z cieczami nienewtonowskimi o własnościach zbliżonych do ciał stałych np. cieczami lepko-sprężystymi spotykanymi w przetwórstwie polimerów, przemyśle chemicznym i spożywczym, medycynie i biologii. Termodynamika cieczy lepko-sprężystych może być tworzona w ramach termodynamiki materiałów odkształcalnych. Podobne trudności spotykamy w przypadkach: ciekłych kryształów traktowanych jako płyny, płynów w polu elektromagnetycznym i mieszanin. Kolejną trudność sprawia analiza tzw. zagadnień sprzężonych, ma to miejsce wtedy gdy analizujemy dynamikę (przepływ) plyn

nów i jednocześnie sprężyste deformacje konstrukcji stykającej się z płynem. Mówimy, że w tych przypadkach dochodzi do wzajemnego oddziaływania (interakcji) gazu lub cieczy z konstrukcją odkształcalną, np. w aero-sprężystości, w hydro-sprężystości, w zjawisku flatteru, itp. Wtedy możemy zastosować jednocześnie różne termodynamiki dla różnych substancji w celu analizy sprzężeń termomechanicznych. Istnieją szczególne przypadki ruchu ciał stałych, które opisuje się tak jakby ciała stałe były płynami, np. grunty, materiały granulowane i sypkie, lodowce i lawiny [4].

2. Czy w naszych czasach termodynamikę da się zastosować? W odpowiedzi na to pytania należy stwierdzić co następuje.

- (a) Współcześnie różne dyscypliny naukowe korzystają z termodynamiki jako atrakcyjnego wzorca metodologicznego. Generalnie, po to korzysta się z procedur termodynamiki, aby poprawnie opisać zjawiska dysypacji i we właściwy sposób uwzględnić sprzężenia termomechaniczne, aby tworzone opisy zjawisk były kompletne i miały jednoznaczne rozwiązania. Po to korzysta się z podstawowych ograniczeń, które nakłada termodynamika (np. Druga Zasada Termodynamiki) aby tworzone opisy zjawisk były zgodne z wiedzą wynikającą z eksperymentów, aby pozwalały wyjaśnić i przewidzieć zachowanie się substancji. Klasyczna termodynamika płynów i termodynamika ośrodków ciągłych mają w naszych czasach różne zakresy zastosowań i służą do modelowania i analizy symulacyjnej zjawisk ważnych z technicznego punktu widzenia.
- (b) Przykłady zastosowań termodynamiki płynów są następujące: modelowanie i symulacja procesów przepływowych i cieplnych w płynach; modelowanie zjawisk wymiany ciepła, masy i pędu w gazach i cieczach stosowanych w maszynach, urządzeniach i instalacjach przemysłowych; modelowanie obiegów cieplnych oraz obiegów chłodniczych gazów i cieczy wykorzystywanych w maszynach i innych urządzeniach; modelowanie wymiany ciepła i konwersji energii w przepływach gazów i cieczy przy przemianach fazowych (wrzenie, parowanie, skraplanie); modelowanie przemian fazowych w gazach i cieczach będących czynnikami roboczymi lub płynami chłodniczymi (podgrzewacze, parowniki, skraplacze, kondensatory); modelowanie przepływów wielofazowych z wymianą ciepła (np. ciecz-para cieczy); modelowanie przepływów turbulentnych z wymianą ciepła; badanie możliwości sterowania przepływami płynów i intensywnością wymiany ciepła (np. cieczy magnetoreologiczne); badanie procesów spalania różnych paliw, itp. Współczesna termodynamika płynów bada procesy wymiany ciepła w mikro- i nanoskali w celu zastosowania w mikro- i nano-skalowych urządzeniach ciepl-

nych. Z tego powodu przedmiotem zainteresowania badaczy są ciecze o złożonej strukturze np. nano-ciecze używane w mikro-kanalowych wymiennikach ciepła. W nano-cieczach małe cząstki ciał stałych rozmieszczone w cieczy sprawiają, że wzrasta przewodność cieplna cieczy.

- (c) Przykładami zastosowań termodynamiki ośrodków ciągłych są: modelowanie procesów dyssypacyjnych w materiałach odkształcalnych którym często towarzyszy wydzielanie się ciepła (sprężysto-plastyczność, plastyczność, lepko-sprężystość, lepko-plastyczność, pełzanie, relaksacja, pękanie, zmęczenie, tarcie, zużycie, propagacja defektów, korozja, przemiany fazowe, itp.), dotyczy to materiałów odkształcalnych oraz tych cieczy które mają własności zbliżone do ciał stałych; określanie reguł rządzących dyssypacją energii podczas procesów mechanicznych i fizycznych w materiałach odkształcalnych (termosprężystość, termo-lepko-sprężystość, termo-lepko-plastyczność, termo-sprężysto-lepko-plastyczność, magneto-elektro-sprężystość, termo-elektro-lepko-sprężystość); wyznaczanie ograniczeń nakładanych na opisy (prawa) materiałów odkształcalnych wynikających z Pierwszej i Drugiej Zasady Termodynamiki; ustalanie opisów (praw) procesów mechanicznych i fizycznych w materiałach odkształcalnych o złożonej budowie wewnętrznej z uwzględnieniem struktur mikroskopowych i wieloskalowych (np. kompozyty, ciekłe kryształy, piezoceramiki, mieszaniny wieloskładnikowe, materiały funkcyjne i wielofunkcyjne); modelowanie materiałów odkształcalnych z wadami (defektami) mikro-struktury (dyslokacje, dysklinacje, mikro-pęknięcia, delaminacje, defekty punktowe, wtrącenia); przewidywanie i kontrolowanie przemian fazowych w ciałach stałych (stopy metali i polimery z pamięcią kształtu i inne materiały inteligentne); itp. Współczesna termodynamika ośrodków ciągłych odnotowuje intensywny rozwój nowych modeli materiałów odkształcalnych i ma zastosowania praktyczne, od inżynierii po medycynę.

3. Czy do dalszego rozwoju termodynamiki może przyczynić się postęp w innych obszarach wiedzy np. energetyce lub fizyce?

- (a) W czasach współczesnych uważa się, że jednym z filarów dalszego rozwoju cywilizacji jest postęp w energetyce. Energetyka jest dyscypliną naukową w dziedzinie nauk technicznych i zajmuje się technikami pozyskiwania, przetwarzania, przesyłania i użytkowania energii (elektroenergetyka, energetyka cieplna, energetyka jądrowa, hydroenergetyka, energetyka wiatrowa, energetyka słoneczna, energetyka geotermalna, energetyka elektrochemiczna, itp.). Do obszarów zainteresowania energetyki należą między innymi takie zagadnienia jak: zrównoważony rozwój w ener-

getyce (zmniejszenie zużycia energii, zwiększenie czystości i efektywności technik wytwarzania energii), poszukiwania nowych sposobów pozyskiwania energii elektrycznej, stosowanie „obiegów zamkniętych” wykorzystania materiałów energetycznych (węgiel, węglowodory, biomasy), wychwytywanie dwutlenku węgla pochodzącego ze spalania materiałów energetycznych, energia odnawialna (pozyskiwanie, gromadzenie, transport), bezpieczeństwo elektrowni jądrowych, nowe materiały dla magazynowania i przesyłania energii, poszukiwania złóż gazu łupkowego i rozwój technik wydobywania, itp. Zjawiska cieplne są ważnym elementem energetyki tradycyjnej, mimo to nie jest pewne, że przyszłe wynalazki i odkrycia dokonywane w energetyce wniosą wkład w rozwój termodynamiki. Obie dyscypliny (termodynamika i energetyka) mogą rozwijać się w przyszłości niezależnie.

- (b) Rozwój fizyki XX wieku pokazał, że termodynamika, która w XIX wieku dała podstawy do sformułowania teorii maszyn cieplnych, nie jest teorią uniwersalną. Perspektywy dalszego rozwoju termodynamiki mogą być stworzone przez nowe odkrycia w fizyce (dotyczące ciepła, energii i entropii), nowe metody fizyki teoretycznej i doświadczalnej, nowe symulacje komputerowe zjawisk fizycznych. Nowe możliwości przyniosła teoria kwantowa jako nauka o zjawiskach w skali atomowej i sub-atomowej, zapoczątkowana przez Alberta Einsteina w 1905 r. Współcześnie rozwijana termodynamika kwantowa pozwala na interpretowanie praw mechaniki kwantowej. Na przykład informatyka kwantowa dostrzega rolę termodynamiki w wyjaśnianiu zachowania się układów kwantowych.

Alfred Zmitrowicz

Instytut Maszyn Przepływowych PAN, Gdańsk

LITERATURA

- [1] A.K. Wróblewski, *Historia fizyki*, PWN, Warszawa 2006.
- [2] S.C. Brown, *Rumford fizyk niezwykle*, Wiedza Powszechna, Warszawa 1966.
- [3] I. Müller, *A history of thermodynamics: The doctrine of energy and entropy*, Springer, Berlin Heidelberg 2007.
- [4] K. Hutter, Y. Wang, *Fluid and thermodynamics*, vol. 1,2,3, Springer International Publishing Switzerland, 2016, 2018.
- [5] K. Wilmański, *Thermomechanics of continua*, Springer, Berlin Heidelberg 1998.
- [6] C. Truesdell, *Sześć wykładów nowoczesnej filozofii przyrody*, PWN, Warszawa 1969.
- [7] C. Truesdell, *Rational thermodynamics*, Springer, New York 1984.
- [8] J. Kestin, Local-equilibrium formalism applied to mechanics of solids, *International Journal of Solids and Structures*, vol. 29, no. 14-15, 1992, p. 1827-1836.
- [9] J. Taler, P. Duda, *Rozwiązywanie prostych i odwrotnych zagadnień przewodnictwa ciepła*, WNT, Warszawa 2003.
- [10] J. Badur, *Pięć wykładów ze współczesnej termomechaniki płynów*, Wydawnictwo IMP PAN, Gdańsk 2005.
- [11] W. Nowacki, *Termosprężystość*, PWN, Warszawa 1960; *Thermoelasticity*, Pergamon Press, Oxford 1962.
- [12] P. Perzyna, *Termodynamika materiałów niesprężystych*, PWN, Warszawa 1978.
- [13] M. Fremond, *Non-smooth thermomechanics*, Springer, Berlin 2002.
- [14] M. Amiri, M.M. Khonsari, On the thermodynamics of friction and wear – a review, *Entropy*, vol. 12, no. 5, 2010, p. 1021-1049.

Paradoksy szczególnej teorii względności

Część I

Nieintuicyjność szczególnej teorii względności sprawia, że jak rzadko która teoria fizyczna, obfituje ona w eksperymenty myślowe prowadzące do zaskakujących wniosków przeczących tak zwanemu „zdrawemu rozsądkowi”. Z tego powodu eksperymenty te nazywamy paradoksami.

Jan Kurzyk

Dla wielu osób są one dowodem na nieprawdziwość szczególnej teorii względności. Jednak sprzeczności związane z tymi paradoksami pojawiają się jedynie wtedy, gdy patrzymy na nie posługując się mechaniką newtonowską. Analiza tych eksperymentów na gruncie mechaniki relatywistycznej sprawia, że sprzeczności, jakie podpowiada nam nasza intuicja znikają i paradoksy przestają być paradoksami.

Względność jednoczesności

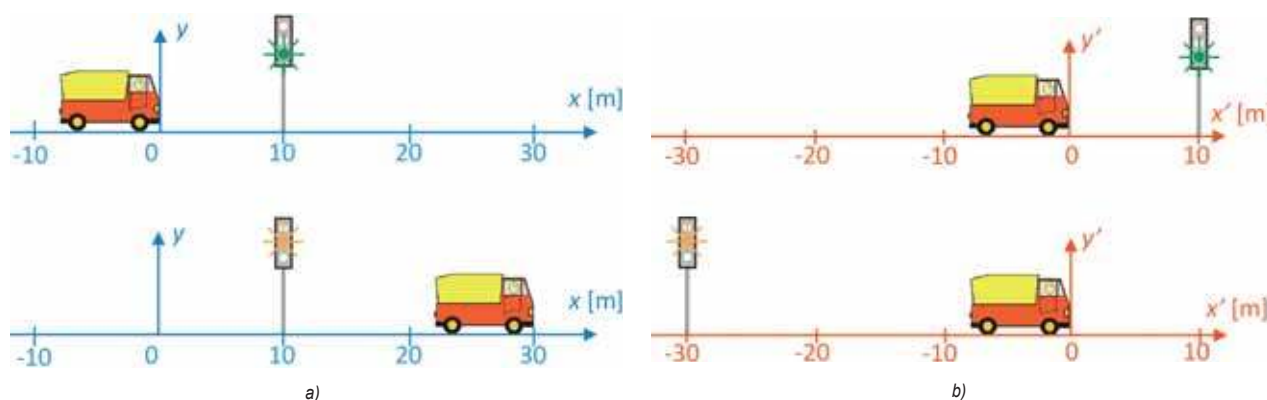
W większość przypadków pozorna sprzeczność jaką dostrzegamy w zjawiskach związanych z paradoksami szczególnej teorii względności wynika z naszego głębokiego przekonania o bezwzględnym charakterze jednoczesności zdarzeń. W pełni akceptujemy względność położenia zdarzeń [1], ale nie potrafimy zaakceptować względności jednoczesności zdarzeń. Przyjrzyjmy się obu sytuacjom.

Rozważmy dwa zdarzenia, które według nas zachodzą w dwóch różnych momentach czasu, ale w tym samym miejscu. Dla przykładu niech pierwszym zdarzeniem będzie zaświecenie się światła zielonego, a drugim zaświecenie się po pewnym czasie światła żółtego na tym samym słupie sygnalizacji świetlnej. Dla obserwatorów z układu nieruchomego względem słupa oba zdarzenia zaszły w tym samym miejscu (rysunek 1 (a)). Jednak dla obserwatorów znajdujących się w samochodzie zbliża-

jącym się do sygnalizacji świetlnej pierwsze zdarzenie zajdzie w punkcie o współrzędnej (mierzonej w kierunku ruchu) np. $x' = 10$ m (np. 10 m przed kierowcą), a drugie w punkcie o współrzędnej np. $x' = -30$ m (30 m za kierowcą, rysunek 1 (b)). Dwa stwierdzenia: pierwsze, że obserwator nieruchomy względem słupa odnotuje zaświecenie się zielonego światła, a następnie zaświecenie się żółtego światła w tym samym miejscu oraz drugie, że dla kierowcy oba te zdarzenia zachodzą w różnych punktach jego układu są dla nas oczywiste. Nie dziwi nas to, bo coś takiego obserwujemy na co dzień.

A teraz rozważmy inną sytuację. Będziemy obserwować dwa zdarzenia zachodzące według nas w dwóch różnych miejscach przestrzeni (dokładniej chodzi o punkty, których współrzędne mierzone w kierunku ruchu jednego układu względem drugiego są różne), ale w tym samym momencie (zdarzenia jednoczesne). Okazuje się, że dla obserwatorów będących względem nas w ruchu oba zdarzenia nie zajdą jednocześnie. Dla nich jedno z nich zajdzie wcześniej niż drugie. W przeciwieństwie do poprzedniej sytuacji, z takim twierdzeniem trudno nam się pogodzić. Codzienne doświadczenia pokazują coś innego. Tymczasem jest to podstawowy wniosek wypływający ze szczególnej teorii względności. Nazywamy go *względnością jednoczesności*. Przyjrzyjmy się temu bliżej.

Wyobraźmy sobie dwie latarnie: pierwszą (1) na początku i drugą (2) na końcu peronu patrząc w kierunku ruchu pociągu. Załóżmy, że w układzie związanym z peronem obie latarnie zaświeciły się w tym



Rysunek 1. (a) Obserwacja dwóch zdarzeń z układu nieruchomego względem słupa sygnalizacji świetlnej: zaświecenie się światła zielonego, a po pewnym czasie zaświecenie się światła żółtego sygnalizacji świetlnej. Zdarzenia zachodzą w tym samym miejscu układu. (b) Te same zdarzenia obserwowane z układu związanego z pojazdem poruszającym się względem słupa. Zdarzenia zachodzą w różnych miejscach układu

samym momencie. Te dwa zdarzenia jednocześnie dla obserwatorów na peronie nie będą jednak jednocześnie dla pasażerów poruszającego się pociągu. Pasażerowie pociągu stwierdzą, że najpierw zaświeciła się latarnia numer 2, a dopiero po jakimś czasie latarnia numer 1. Od razu podkreślam, że nie chodzi tu o to, że latarnie są w różnej odległości od pasażerów i z tego powodu światło z tych latarni dojdzie do obserwatorów w różnych momentach. Jest oczywiste, że obserwator oddalony od miejsca zajścia jakiegoś zdarzenia obserwuje to zdarzenie z opóźnieniem. Ale jeśli zna on swoją odległość do miejsca zajścia zdarzenia, to może uwzględnić czas lotu sygnału do niego i obliczyć w jakim faktycznie momencie zaszło zdarzenie. Byłoby to jednak trudne w przypadku obserwatora będącego w ruchu względem punktu, w którym zachodzi zdarzenie. Wymagałoby to znajomości odległości pomiędzy obserwatorem a tym punktem w momencie, w którym zachodzi zdarzenie. Znacznie wygodniej będzie przeprowadzić nasz eksperyment inaczej. Zamiast jednego obserwatora wykorzystajmy wielu obserwatorów (*nota bene* obserwatorem niekoniecznie musi być człowiek, to może być odpowiednia aparatura pomiarowa) rozstawionych zarówno wzdłuż peronu, jak i w poszczególnych przedziałach pociągu. Oczywiście warunkiem koniecznym do przeprowadzenia tego eksperymentu jest precyzyjna synchronizacja zegarów, jakimi będą dysponować obserwatorzy w obu układach.

- Synchronizacja zegarów jest kluczowym zagadnieniem szczególnej teorii względności. Według przepisu na synchronizację zegarów podanego przez Einsteina, aby zsynchronizować zegary z zegarem bazowym należałoby:
- zmierzyć odległości l_i między zegarem bazowym a pozostałymi zegarami
 - zatrzymać wszystkie zegary
 - ustawić bazowy zegar na daną godzinę t , a pozostałe na godzinę $t + l_i/c$ czyli późniejszą od bazowego o czas, jaki światło potrzebuje na pokonanie drogi między zegarem bazowym a danym zegarem
 - uruchomić zegar bazowy wysyłając jednocześnie sygnały radiowe do pozostałych zegarów

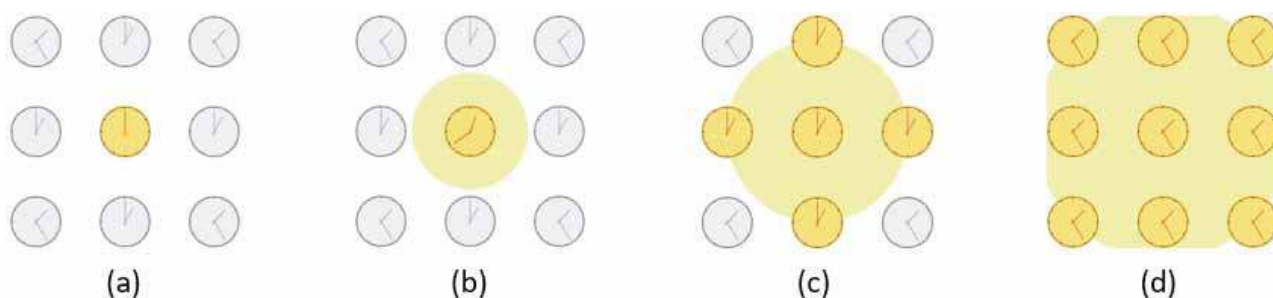
- uruchomić dany zegar w momencie dotarcia do niego sygnału z zegara bazowego.

Przepis ten został zaprezentowany na rysunku 2.

Założmy, że w obu układach (na peronie i w pociągu) przeprowadzono synchronizację zegarów. W ten sposób w układzie związanym z peronem możemy stwierdzić, czy latarnie zaświeciły się jednocześnie pytając obserwatorów, tego stojącego przy pierwszej latarni i tego stojącego przy drugiej latarni, o której godzinie zaobserwowali zaświecenie się latarni, przy których stali. Podobnie postąpimy z pasażerami pociągu. Założmy, że pociąg jest na tyle długi, że znajdzie się zarówno pasażer, który mijając pierwszą latarnię zauważy, że w tym momencie latarnia zaświeciła się, jak i pasażer, który zaobserwuje to samo w przypadku drugiej latarni. Wystarczy teraz zapytać obu pasażerów, o jakiej godzinie zaobserwowali zaświecenie się latarni, które mijali i w ten sposób rozstrzygnąć, czy oba zdarzenia były jednocześnie czy też nie.

Żeby odpowiedzieć na nasze pytanie na gruncie teoretycznym musimy przetransformować współrzędne przestrzenne i współrzędną czasową z jednego układu do drugiego. Gdybyśmy skorzystali z transformacji Galileusza [2], to otrzymalibyśmy odpowiedź taką, jakiej oczekuje nasza intuicja oparta na codziennym doświadczeniu czasu i przestrzeni – pasażerowie pociągu powiedzą, że obie latarnie zaświeciły się jednocześnie. Wynika, to z tego, że w transformacji Galileusza czas ma charakter absolutny i biegnie jednakowo w obu układach odniesienia. Transformacja Galileusza sprawdza się bardzo dobrze przy prędkościach małych w porównaniu z prędkością światła, ale nawet dla małych prędkości jest tylko przybliżeniem transformacji lepiej opisującej własności czasoprzestrzeni – transformacji Lorentza [3]. Przypomnijmy tę transformację. Transformacja współrzędnych z układu inercyjnego S do współrzędnych z układu inercyjnego S' poruszającego się z prędkością v względem układu S w kierunku osi X ma postać

$$x' = (x - vt)\gamma, \quad y' = y, \quad z' = z, \quad t' = \left(t - \frac{xv}{c^2}\right)\gamma,$$



Rysunek 2. Kolejne fazy synchronizacji zegarów.

gdzie c jest szybkością światła w próżni (w układzie SI szybkość ta jest dokładnie równa 299 792 458 m/s) oraz

$$\gamma \equiv \frac{1}{\sqrt{1-\beta^2}} \geq 1, \quad \beta \equiv \frac{v}{c} \leq 1.$$

Zaś odwrotna transformacja współrzędnych – z układu S' do układu S ma postać

$$x = (x' + vt')\gamma, \quad y = y', \quad z = z', \quad t = \left(t' + \frac{x'v}{c^2}\right)\gamma.$$

Trzeba jeszcze dodać, że powyższe wzory transformacyjne odpowiadają układom, których kierunki osi X i X' pokrywają się, o pozostałe pary osi, Y i Y' oraz Z i Z' , są do siebie równoległe. Dodatkowo, gdy zegary umieszczone w początkach obu układów współrzędnych miały się, to oba pokazywały godzinę $t = 0$, i $t' = 0$.

Wróćmy do naszego eksperymentu. Załóżmy, że w układzie S związanym z peronem pierwsza latarnia znajduje się w punkcie o współrzędnej x_1 , a druga latarnia w punkcie o współrzędnej $x_2 > x_1$. Zaświećmy obie latarnie w chwili t_0 układu S . Korzystając z transformacji Lorentza znajdujemy, że w momencie zaświecenia się pierwszej latarni będzie ją mijał pasażer pociągu znajdujący się w punkcie o współrzędnej (w układzie pociągu)

$$x'_1 = (x_1 - vt_0)\gamma$$

i stanie się to o godzinie (w układzie pociągu)

$$t'_1 = \left(t_0 - \frac{x_1 v}{c^2}\right)\gamma.$$

Zaś w momencie zaświecenia się drugiej latarni będzie ją mijał pasażer pociągu znajdujący się w punkcie o współrzędnej (w układzie pociągu)

$$x'_2 = (x_2 - vt_0)\gamma$$

i stanie się to o godzinie (w układzie pociągu)

$$t'_2 = \left(t_0 - \frac{x_2 v}{c^2}\right)\gamma.$$

Jak widzimy $t'_2 < t'_1$, czyli w układzie pociągu (S') najpierw zaświeci się latarnia numer 2, a pierwsza dopiero po czasie

$$\Delta t' = (x_2 - x_1) \frac{v\gamma}{c^2}.$$

Jeśli odstęp między latarniami jest rzędu 100 m, a prędkość pociągu rzędu 100 km/h odstęp czasu między

tymi zdarzeniami będzie rzędu 10^{-14} s (10 femtosekund). Nic więc dziwnego, że w normalnych warunkach nie obserwujemy tego efektu. Zresztą przy tak niewielkiej odległości między oboma punktami, nawet gdybyśmy dysponowali „relatywistycznym pociągiem” o szybkości rzędu 0,5 prędkości światła efekt byłby do zaobserwowania jedynie przy użyciu zegarów atomowych. Odstęp czasu między oboma zdarzeniami byłby wówczas rzędu 10^{-7} s (0,1 mikrosekundy). Aby przyjrzeć się dokładniej problemowi względności jednoczesności rozważmy prędkość będącą dużą w porównaniu z prędkością światła oraz dwa zdarzenia zachodzące w dużo większej odległości od siebie. Wówczas przedział czasu będzie odpowiednio większy.

Dla uproszczenia rachunków zarówno czas jak i odległość będziemy mierzyć w tych samych jednostkach np. w godzinach (h). Dla rozróżnienia obu jednostek, godziny, w których mierzymy odległość nazywamy godzinami świetlnymi (lh – ang. light hour). Jedna godzina świetlna jest odległością jaką światło w próżni pokonuje w ciągu jednej godziny, $1 \text{ lh} \approx 1,08 \cdot 10^{12} \text{ m} = 1,08 \text{ Tm}$ (terametrow). W tych jednostkach prędkość jest bezwymiarowa, a prędkość światła ma wartość 1 ($c = 1 \text{ lh/h}$). Transformacja Lorentza przyjmuje wtedy postać

$$x' = (x - vt)\gamma, \quad y' = y, \quad z' = z, \quad t' = (t - vx)\gamma,$$

$$x = (x' + vt')\gamma, \quad y = y', \quad z = z', \quad t = (t' + vx')\gamma.$$

Zwróćmy przy okazji uwagę na to, że używając takich jednostek dostajemy pełną symetrię wzorów transformacyjnych na współrzędną przestrzenną x i współrzędną czasową t . Obie współrzędne czasoprzestrzeni stają się równoważne!

Założmy, że nasz relatywistyczny pociąg ma długość spoczynkową $l_0 = 1 \text{ lh}$ i szybkość 0,6 lh/h (0,6 c), a odległość między latarniami wynosi 0,8 lh. Czynniki Lorentza jest równy $\gamma = 1,25$. Zwiążmy z latarnią numer 1 początek układu współrzędnych układu S (związanego z peronem), a z końcem pociągu zwiążmy początek układu współrzędnych układu S' (związanego z pociągiem). Obie latarnie zaświećmy w chwili $t = 0$ układu S . Gdzie wtedy znajduje się koniec pociągu i jaką godzinę wskazuje zegar pasażera w ostatnim przedziale? Z transformacji Lorentza znajdujemy

$$0 \text{ lh} = (x \text{ lh} - (0,6 \text{ lh/h}) \times 0 \text{ lh}) \times 1,25 = x \times 1,25 \text{ lh},$$

stąd

$$x = 0 \text{ lh}, \\ t' = (0 \text{ h} - (0,6 \text{ lh/h}) \times 0 \text{ h}) \times 1,25 = 0 \text{ h},$$

Czyli koniec pociągu znajduje się przy latarni numer 1, a zegar na końcu pociągu wskazuje godzinę 00:00:00. A gdzie w chwili $t = 0$ układu S znajduje się początek pociągu i jaką godzinę wskazuje zegar maszynisty? Z transformacji Lorentza znajdujemy

$$1 \text{ lh} = (x \text{ lh} - (0,6 \text{ lh/h}) \times 0 \text{ h}) \times 1,25 = x \times 1,25 \text{ lh},$$

stąd

$$x = 0,8 \text{ lh},$$

$$t' = (0 \text{ h} - (0,6 \text{ lh/h}) \times 0,8 \text{ lh}) \times 1,25 = -0,6 \text{ h},$$

Czyli początek pociągu znajduje się przy latarni numer 2, a zegar maszynisty pociągu wskazuje godzinę 23:24:00. Jak widzimy jest to godzina wcześniejsza niż na zegarze pasażera z ostatniego przedziału. Może się to wydawać dziwne. Przecież umówiliśmy się, że obserwatorzy w obu układach przeprowadzili synchronizację swoich zegarów. Dodatkowo ustaliliśmy, że w momencie, gdy zegary umieszczone w początkach obu układów współrzędnych mijają się, to oba wskazują godzinę 00:00:00. To ostatnie się zgadza, więc dlaczego zegar na początku pociągu pokazuje inną godzinę niż zegar na końcu.

Czyżby zegary w pociągu nie były zsynchronizowane? Ależ nie. One są nadal zsynchronizowane, ale w układzie pociągu, a my sprawdzaliśmy co widzą obserwatorzy na peronie. Zwróćmy tu uwagę na ważny fakt wynikający wprost z transformacji Lorentza. Chociaż w obu układach S i S' obserwatorzy mają układy zsynchronizowanych zegarów i ustaliliśmy, że zegary ustawione w początkach układów współrzędnych obu układów odniesienia w momencie mijania się początków układów wskazują chwilę $t = 0$ i $t' = 0$, to oba układy zegarów nie są (i nie mogą być!) zsynchronizowane ze sobą [4]. Patrząc z układu S na zegary układu S' oraz odwrotnie patrząc z układu S' na zegary układu S zobaczymy coś podobnego, do tego to co pokazuje rysunek 3.

Ta własność czasoprzestrzeni, nieobserwowana na co dzień ze względu na małe prędkości, z jakimi mamy do czynienia, jest źródłem większości nieporozumień. To dlatego według obserwatorów na peronie obie latarnie zaświecą się w tym samym momencie (patrz rysunek 4(a)), a według pasażerów pociągu w różnych momentach (patrz rysunek 4(b) i 4(c)). Zobaczmy, jak to wygląda z punktu widzenia pasażerów pociągu. Według nich peron wraz z latarniami zbliża się do nich od strony początku pociągu z szybkością 0,6 lh/h. Jak wyliczyliśmy wcześniej w momencie $t' = -0,6 \text{ h}$, czyli o godzinie 23:24:00 latarnia numer 2 dojeżdża do początku pociągu, a zegar na niej wskazuje godzinę 00:00:00. W tym momencie latarnia zaświeca się. Jak na razie wszystko jest tak jak widział to obserwator na peronie stojący obok latarni numer 2. Ale gdzie w tym momencie (mierzonym w układzie pociągu) znajduje się latarnia numer 1 i jaką godzinę wskazuje zegar związany z tą latarnią? Z transformacji Lorentza znajdujemy

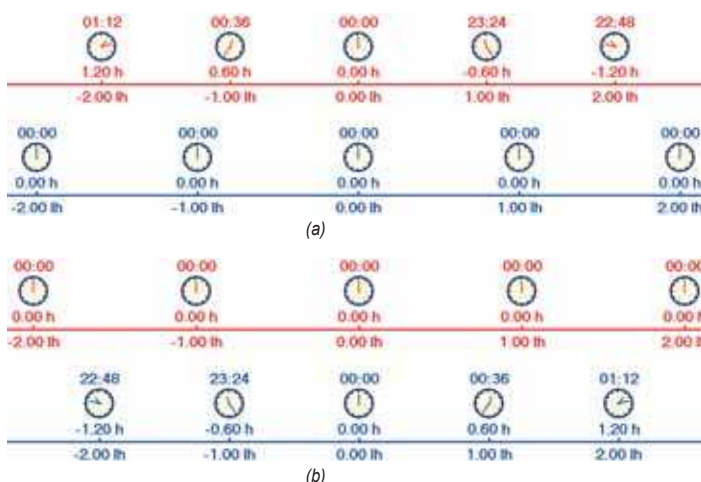
$$0 \text{ lh} = (x' \text{ lh} + (0,6 \text{ lh/h}) \times (-0,6 \text{ h})) \times 1,25,$$

stąd

$$x' = 0,36 \text{ lh},$$

$$t = (-0,6 \text{ h} + (0,6 \text{ lh/h}) \times (0,36 \text{ lh})) \times 1,25 = -0,48 \text{ h},$$

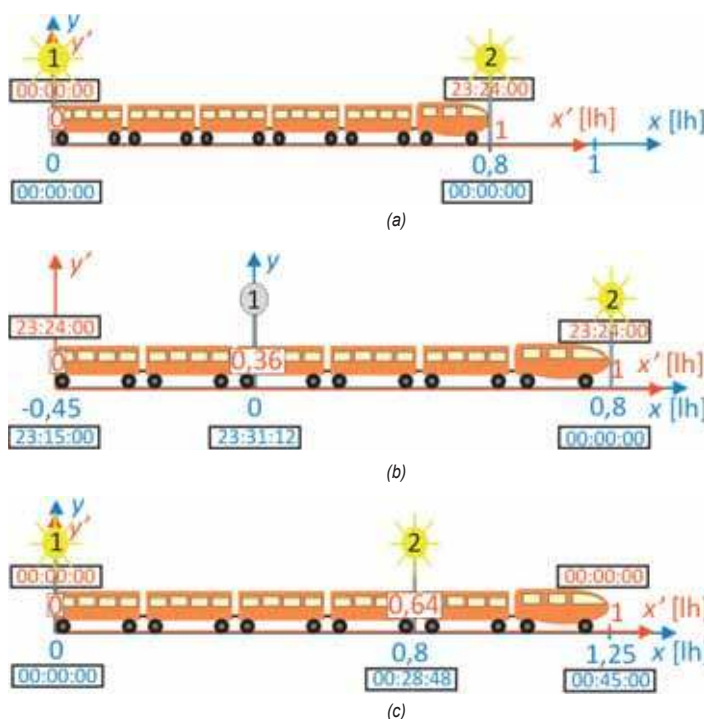
czyli pierwsza latarnia znajduje się przed końcem pociągu, w odległości 0,36 lh od niego, a zegar obok tej latarni wskazuje godzinę 23:31:12. Latarnia jeszcze nie świeci (patrz rysunek 4 (b)). Zaświeci się dopiero o godzinie 00:00:00, czyli za 0,48 h (28 minut i 48 sekund) mierzone w układzie S. Wtedy, (pozostawiam to do sprawdzenia czytelnikowi) latarnia „dojedzie” do końca pociągu, a zegar z nią związany oraz zegar na końcu pociągu wskaza godzinę 00:00:00 (patrz rysunek 4 (c)).



Rysunek 3. Dwa układy odniesienia: układ S (kolor niebieski) i układ S' (kolor czerwony) wraz z ich układami zsynchronizowanych zegarów. Układ S' porusza się względem układu S z szybkością $v = 0,6 c$ w prawo. (a) Obserwacja układu zegarów S' z układu S w chwili $t = 0$. (b) Obserwacja układu zegarów S z układu S' w chwili $t' = 0$.

Zauważmy, że obserwacje wykonywane przez poszczególnych obserwatorów z obu układów znajdujących się w momencie zachodzenia zdarzenia w miejscu zajścia zdarzenia są w pełni zgodne ze sobą. Na przykład obserwator stojący na peronie przy latarni numer 2 powie, że o godzinie 00:00:00 zaświeciła się latarnia i jednocześnie minął go początek pociągu, a zegar umieszczony w tym miejscu pociągu wskazywał godzinę 23:24:00. Maszynista pociągu powie, że o godzinie 23:24:00 mijają go latarnia numer 2 i w tym momencie zaświeciła się, a zegar obok niej wskazywał godzinę 00:00:00. Ktoś może powiedzieć, że tu też mamy do czynienia ze zdarzeniami równoczesnymi – zaświecenie latarni i jednocześnie mijanie się latarni i początku pociągu. Dlaczego w takim razie te zdarzenia są jednoczesne w obu układach? Otóż dlatego, że zaszły w tym samym miejscu. Względność jednoczesności dotyczy jedynie zdarzeń zachodzących w punktach o różnych współrzędnych X (przypomnijmy wyprowadzony wcześniej wzór $\Delta t' = (x_2 - x_1) v/c^2$).

Względność jednoczesności jest ważną własnością czasoprzestrzeni, bez zrozumienia której nie zrozumimy wielu innych efektów relatywistycznych. Proponuję dokładne przestudiowanie rysunku 4, aby oswoić się z tym pojęciem.



Rysunek 4. Pociąg o długości spoczynkowej 1 lh przejeżdża z szybkością 0,6 lh/h wzdłuż peronu, na którym w odległości 0,8 lh od siebie, mierzonej w układzie peronu, stoją dwie latarnie. Rysunek (a) pokazuje sytuację widzianą przez pasażerów pociągu w chwili o godzinie 00:00:00 ich układu. Rysunek (b) pokazuje sytuację widzianą przez pasażerów pociągu w chwili o godzinie 23:24:00, a rysunek (c) o godzinie 00:00:00 ich układu. Rysunki pokazują względność jednoczesności. Dla obserwatorów na peronie obie latarnie zaświeciły się jednocześnie. Dla pasażerów pociągu najpierw zapaliła się latarnia numer 2, a 36 minut później latarnia numer 1. Na rysunkach można dostrzec również względność kontrakcji długości. Dla obserwatorów na peronie pociąg ma długość 0,8 lh, a dla pasażerów pociągu odległość między latarniami wynosi 0,76 lh. Uważny czytelnik zauważy również względność dylatacji czasu, ale tym zajmę się w kolejnym artykule.

I jeszcze jedna uwaga. Zapisując transformację Lorentza w jednostkach, w których prędkość światła jest równa 1 zauważyliśmy, że współrzędna przestrzenna x i czasowa t stają się równoważne. W świetle tego spostrzeżenia warto zastanowić się jeszcze raz nad oboma opisanymi tu doświadczeniami. Przypomnę, że pierwsze pokazuje zdarzenia, które zachodzą w tym samym miejscu jednego układu, a w różnych miejscach układu będącego w ruchu względem tego pierwszego. Drugie pokazuje dwa zdarzenia, które zachodzą w pierwszym układzie w tym samym momencie, a w układzie będącym w ruchu względem niego w różnych momentach. Pierwsze z tych doświadczeń jest dla nas zrozumiałe i nie budzi naszych wątpliwości, a drugie trudno nam zaakceptować.

Skrócenie długości

Spróbujmy zmierzyć długość pociągu z poprzedniego przykładu. Zrealizujemy dwa eksperymenty myślowe pozwalające na zmierzenie długości poruszającego się obiektu. Pierwszy eksperyment będzie polegał na pomiarze czasu, jaki potrzebuje nasz pociąg na przejazd obok jednego z obserwatorów stojących na peronie. Jeden z tych obserwatorów odczyta moment czasu t_p , w którym mija go początek pociągu oraz moment czasu t_k , w którym minie go koniec pociągu. Długość tego odcinka czasu pomnożony przez szybkość pociągu jest długością pociągu zmierzona w układzie peronu S. Załóżmy, że

w układzie pociągu S' początek pociągu ma współrzędną x'_p , a koniec pociągu współrzędną $x'_k < x'_p$, przy czym długość pociągu zmierzona przez pasażerów wynosi l_0 , czyli $x'_p - x'_k = l_0$. Niech obserwator związany z peronem znajduje się w punkcie o współrzędnej x_0 . W punkcie tym w chwili t_p ma znaleźć się początek pociągu, czyli zgodnie z transformacją Lorentza

$$x'_p = (x_0 - vt_p)\gamma$$

W tym samym punkcie w chwili t_k ma znaleźć się koniec pociągu, czyli zgodnie z transformacją Lorentza

$$x'_k = (x_0 - vt_k)\gamma.$$

Po odjęciu tych równań stronami dostajemy

$$x'_p - x'_k = (t_k - t_p)v\gamma.$$

Ale zgodnie z naszymi założeniami $(t_k - t_p)v$ jest wynikiem pomiaru długości l pociągu poruszającego się względem peronu, a $x'_p - x'_k = l_0$ jest wynikiem pomiaru długości pociągu wykonanym przez pasażerów pociągu, czyli

$$l_0 = l\gamma$$

lub

$$l = l_0/\gamma < l_0$$

Wynika z tego, że długość poruszającego się obiektu mierzona w kierunku ruchu jest mniejsza niż tego samego obiektu spoczywającego względem obserwatorów. Nazywamy to kontrakcją lub skróceniem długości.

Sprawdźmy czy inna metoda pomiaru da ten sam wynik. Poprzedni pomiar polegał na rejestracji czasów dwóch zdarzeń zachodzących w tym samym punkcie układu związanego z peronem. Do przeprowadzenia pomiaru wystarczył jeden obserwator. Teraz skorzystajmy z tego, że mamy jeszcze innych obserwatorów ustawionych wzdłuż peronu. Dokonamy rejestracji dwóch zdarzeń jednoczesnych zachodzących w dwóch punktach toru. Będą to punkty, w których w tym momencie będą początek i koniec pociągu. Umówmy się z obserwatorami, że o godzinie t_0 każdy z nich sprawdzi, czy nie mija go właśnie początek lub koniec pociągu. Niech współrzędna obserwatora, który o godzinie t_0 zauważył, że mija go początek pociągu wynosi x_p , a współrzędna obserwatora, którego w tym momencie mija koniec pociągu wynosi x_k . Naszym wynikiem pomiaru długości poruszającego się pociągu jest różnica tych współrzędnych

$$l = x_p - x_k.$$

Zgodnie z transformacją Lorentza mamy

$$x'_p = (x_p - vt_0)\gamma \quad i \quad x'_k = (x_k - vt_0)\gamma.$$

Stąd dostajemy

$$x'_p - x'_k = (x_p - x_k)\gamma,$$

czyli ostatecznie

$$l = l_0/\gamma.$$

Jak widzimy otrzymaliśmy ten sam wynik pomiaru co poprzednią metodą.

Zwróćmy uwagę, że kontrakcja długości nie polega na ściskaniu poruszających się obiektów. Nie towarzyszy

temu żadne naprężenie. Efekt jest własnością czasoprzestrzeni. Można by go porównać do efektu, jaki obserwujemy w zwykłej przestrzeni – pręt obrócony pod pewnym kątem względem linii wzroku jest dla nas krótszy.

Skrócenie długości było postulowane jeszcze przed powstaniem szczególnej teorii względności. W roku 1889 hipotezę taką wysunął irlandzki fizyk George Fitzgerald (1851-1901) [5], a w 1891 r. holenderski fizyk Hendrik Lorentz (1853-1928) [6] zaproponował wzór określający to skrócenie (chodzi o tego samego uczonego, który jako pierwszy podał równania transformacji nazywanej dziś jego nazwiskiem).

Przez pamięć dla obu fizyków efekt ten nazywamy czasem skróceniem Lorentza lub skróceniem Lorentza-Fitzgeralda. Zarówno dla Fitzgeralda jak i dla Lorentza hipoteza o skróceniu długości była próbą wyjaśnienia wyniku doświadczenia Michelsona-Morleya [7] (Albert Michelson (1852-1931) – amerykański fizyk polsko-żydowskiego pochodzenia [8], Edward Morley (1838-1923) – amerykański fizyk i chemik [9]). Według ich hipotezy ciała ulegały skróceniu w kierunku ruchu względem hipotetycznego eteru, którego poszukiwali ówczesni fizy-

cy. Szczególna teoria względności odrzuciła pojęcie eteru. Skrócenie długości występuje zawsze, gdy obserwujemy obiekt będący względem nas w ruchu.

W następnej części artykułu przedstawię kilka paradoksów związanych ze zjawiskiem skrócenia długości. Pokażę, że sprzeczność w nich zawarta znika, jeśli uwzględnimy własności czasoprzestrzeni opisane przez szczególną teorię względności. Tym samym będę chciał wykazać, że skrócenie Lorentza nie jest efektem pozornym, czy jakąś matematyczną sztuczką. Jest to zjawisko jak najbardziej rzeczywiste.

Jan Kurzyk

Instytut Fizyki Politechniki Krakowskiej

LITERATURA

- [1] <https://www.youtube.com/watch?v=GZ6tdsFJHhk> [dostęp 28.10.2018].
- [2] https://pl.wikipedia.org/wiki/Transformacja_Galileusza [dostęp 28.10.2018].
- [3] https://pl.wikipedia.org/wiki/Transformacja_Lorentza [dostęp 28.10.2018].
- [4] J. Kurzyk, *Dlaczego działa GPS?* Fizyka w szkole, nr 2, 2017.
- [5] https://en.wikipedia.org/wiki/George_Francis_FitzGerald [dostęp 28.10.2018].
- [6] https://en.wikipedia.org/wiki/Hendrik_Lorentz [dostęp 28.10.2018].
- [7] https://pl.wikipedia.org/wiki/Do%C5%9Bwiadczenie_Michelsona-Morleya [dostęp 28.10.2018].
- [8] https://pl.wikipedia.org/wiki/Albert_Michelson [dostęp 28.10.2018].
- [9] https://en.wikipedia.org/wiki/Edward_W._Morley [dostęp 28.10.2018].

Co w fizyce piszczy

Tajemnicza piąta siła

– Przesunięcie Lamba zainspirowało nas do stworzenia takiej pięknej teorii zwanej elektrodynamiką kwantową. Ja, poprzez moje obliczenia, weryfikowałem, na ile ta teoria jest dokładna – mówi prof. Krzysztof Pachucki, laureat Nagrody Fundacji Nauki Polskiej 2018, którego doceniono za precyzyjne kwantowo-elektrodynamiczne obliczenia spektroskopowych parametrów lekkich atomów i cząsteczek.

Obecnie prof. Pachucki wciąż stara się zbadać, jak dokładne są prawa natury, które znamy. Wraz ze współpracownikami weryfikuje prawa fizyki poprzez porównanie bardzo precyzyjnych obliczeń z eksperymentami. „Jeśli się nie zgadzają, jest to dla nas sygnał, że być może o czymś nie wiemy. Bardzo możliwe, że istnieją jakieś inne, nieznanne jeszcze siły. Bardzo chcielibyśmy je poznać” – mówi.

Przyznaje, że najczęściej okazuje się, że to wyniki eksperymentu są błędne. Ale jeśli przy ponownym sprawdzeniu wciąż pojawia się różnica, to jest to podstawa do odkrycia nowych praw fizyki, nowych sił. Czym jest „piąta siła”?

Jak wyjaśnia fizyk, znane są cztery rodzaje sił. Istnieją siły elektromagnetyczne i oddziaływanie grawitacyjne. Oddziaływania silne utrzymują nukleony w jądrze. Czwarta siła to oddziaływanie słabe. Łamią one symetrię parzystości, czyli powodują, że fizyka jest inna przed lustrem, a inna za lustrem. Okazuje się na przykład, że molekuły prawo i lewoskrętne mogą mieć różne energie wibracji.

„Ale co, jeśli istnieją inne? Właśnie po to, żeby je odkryć, mierzymy przejścia w atomie wodoru oraz

innych atomach i porównujemy je z obliczeniami” – tłumaczy naukowiec. Dodaje, że choć nauka nie potrafi jeszcze opisać piątej siły, to znajduje mocne podstawy, żeby podejrzewać jej istnienie.

„Na podstawie badań oddziaływań grawitacyjnych wiemy, że istnieje ciemna materia. Co więcej, jest jej znacznie więcej niż materii widzialnej. Nikt nie wie, co to jest, wszyscy jej szukają, ale żeby można było ją zaobserwować, musi ona oddziaływać z widzialną materią. Takie oddziaływanie może być właśnie tą nieznaną nam jeszcze piątą siłą. My szukamy jakichkolwiek objawów tej siły na poziomie stołu laboratoryjnego, poprzez porównywanie bardzo precyzyjnych eksperymentów i obliczeń” – wyjaśnia.

Według laureata Nagrody FNP, każdy fizyk teoretyczny marzy o tym, że to właśnie on odnajdzie dowód na istnienie piątej siły. Ale „odkrycia nie przychodzą na zamówienie”. Badacze starają się stworzyć takie warunki eksperymentu, żeby uzyskać jak największą czułość i dokładność i żeby móc najmniejsze objawy owej piątej siły wychwycić. Zwiększają precyzję poprzez nowy dokładniejszy pomiar, np. rzędu 12 cyfr. Obliczenia najczęściej dotyczą widma atomu wodoru, atomu helu, cząsteczki wodoru.

Grupa prof. Pachuckiego ma charakter wirtualny. Naukowiec realizuje projekty we współpracy z naukowcami pracującymi w czterech różnych miejscach (Vojtech Patkos w Pradze, Władimir Yerokhin w Sankt-Petersburgu, Jacek Komasa i Mariusz Puchalski – w Poznaniu) oraz z doktorantami i magistrantami.

PAP – Nauka w Polsce, Karolina Duszczyk



Kobięca strona fizyki – dokończenie z poprzedniego wydania

Siła słabych oddziaływań

Od Marii Curie-Skłodowskiej oraz Lise Meiter do dnia dzisiejszego wiele się zmieniło, wiele nazwisk kobiet pojawiło się w społeczności akademickiej oraz na kartach historii. Można by tu przytoczyć szereg życiorysów, ale przytoczmy jeszcze dwa. Pierwszą postacią, obowiązkową przy omawianiu historii fizyczek jest druga kobieta w historii – laureatka nagrody Nobla – Maria Goeppert-Meyer.

Aleksandra Mielewczyk-Gryn, Marcin Zaród

Maria urodziła się w 1906 roku w dzisiejszych Katowicach, jednak już w 1910 roku jej rodzina przeniosła się do Getyngi, gdzie jej ojciec przyjął pozycję profesora na uniwersytecie. To również na tamtejszy uniwersytet wstąpiła w 1924 r., co rzecz jasna nie było standardem w ówczesnych Niemczech. Na początku studiów matematykę a po trzech latach wybiera fizykę. I to z fizyki, pod promotorstwem Maxa Borna (Nagroda Nobla 1954) w 1930 r. broni doktorat na temat absorpcji dwufotonowej – zjawiska eksperymentalnie udowodnionego w latach 60. i będącego podstawą działania laserów. Jednostka przekroju czynnego absorpcji dwufotonowej jest nazwana jednostką Goeppert-Mayer (GM).

Zajmuje się fizyką kwantową, najpierw w Getyndze a po zamążpójściu w Stanach Zjednoczonych. W czasach Wielkiego Kryzysu bezpłatnie wykłada fizykę na Uniwersytecie Johna Hopkinsa. W czasach późniejszych m.in. w S.A.M. Laboratory, pracowała nad izolowaniem izotopów uranu (znowu ten uran!). Od 1946 r. pracowała w Chicago jako profesor na Wydziale Fizyki w Instytucie Badań Jądrowych oraz w Laboratorium Narodowym

w Argonne (pomimo bardzo małej wiedzy na temat fizyki jądrowej!). Pracowała wtedy razem z Edwardem Tellerem oraz Enrico Fermim.

W 1948 roku rozpoczęła pracę nad liczbami magicznymi (tak nazywamy liczby protonów i neutronów, dla których wypełnione są poszczególne powłoki), do takich samych jak ona wniosków doszła grupa badaczy Haxel, Jensen i Suess, co potwierdzało jej wnioski.³⁶ Opracowuje model powłokowy jądra, za który wraz z J. Hansem Jensenem otrzymuje w 1963 r. nagrodę Nobla z fizyki³⁵. Tak sama opisuje opracowany przez siebie model: „Think of a room full of waltzers. Suppose they go round the room in circles, each circle enclosed within another. Then imagine that in each circle, you can fit twice as many dancers by having one pair go clockwise and another pair go counterclockwise. Then add one more variation; all the dancers are spinning twirling round and round like tops as they circle the room, each pair both twirling and circling. But only some of those that go counterclockwise are twirling counterclockwise. The others are twirling clockwise while circling counterclockwise. The same is true of those that are dancing around clockwise: some twirl clockwise, others twirl counterclockwise”.^{1 37} To właśnie ten model przyniósł jej sławę i pozwolił na zapisanie się na kartach historii.

¹ Tłumaczenie: „Wyobraź sobie salę pełną tańczących walców. Tancerze okrążają salę w koncentrycznych kołach. Dalej pomyśl, że w każdym kole możesz zmieścić dwa razy więcej tancerzy, jeśli jedna para wiruje w kierunku zgodnym z kierunkiem wskazówek zegara, a druga w przeciwnym. Następnie wprowadźmy dodatkową zmienną: pomyśl, że tancerze wirują w porywach i wykonują obroty. Niektóre z tych par, które wirują w kierunku zgodnym z kierunkiem wskazówek zegara i obracają się w tym samym kierunku. Obroty pozostałych par są w kierunku przeciwnym. Tak samo z parami wirującymi w kierunku przeciwnym do kierunku wskazówek zegara – niektóre wykonują obroty w tym samym kierunku, inne w przeciwnym.”

Tak dochodzimy do lat 60., kiedy to urodziła się nasza ostatnia już bohaterka – Fabiola Gianotti, jak ją określił tygodnik „Times” – *Fabulous Fabiola*³⁸, obecna dyrektorka generalna najśłynniejszego chyba instytutu badawczego – Conseil européen pour la recherche nucléaire – CERN-u. To właśnie ona była liderką zespołu, który znalazł potwierdzenie eksperymentalne istnienia „boskiej cząstki” czy tzw. bozonu Higgsa – cząstki do tej pory jedynie postulowanej w modelu standardowym.



Foto - Fabiola Gianotti, źródło: ATLAS Experiment © 2011 CERN (licencja: Creative Commons Attribution-Share Alike 4.0 International)

Większość społeczności naukowej jest zgodna, że naukowa wiedza Gianotti oraz sposób zarządzania zespołem przyczynił się do tego wielkiego dla fizyki cząstek elementarnych odkrycia (część osób postuluje nawet, że godnego nagrody Nobla).

Jest ona absolwentką Uniwersytetu w Mediolanie, gdzie otrzymała stopień doktora (1989). W 1996 roku rozpoczęła pracę na stanowisku postdoc w CERN i od tego czasu jest z związana z tą instytucją, dodatkowo od 2013 r. jest honorowym profesorem Uniwersytetu w Edynburgu, jest również członkinią wielu rad oraz komitetów najwybitniejszych jednostek naukowych takich jak m.in. FermiLab w USA czy DESY w Niemczech.³⁹

Fabiola jest autorką, bądź współautorką ponad 500 artykułów naukowych a jej tzw. indeks Hirsha to 116 (znaczy to, że 116 z jej publikacji było cytowanych przez innych badaczy minimum 116 razy). W 2009 r. przejęła kierowanie projektem ATLAS (A Toroidal LHC Apparatus)⁴⁰, który to właśnie miał na celu m.in. eksperymentalne udowodnienie istnienia bozonu Higgsa. Zespół badawczy składał się 3000 fizyków z 180 instytucji z 38 krajów, z czego tylko 20% stanowiły kobiety. Sama Fabiola twierdzi, że nigdy nie była ofiarą dyskryminacji, ale przyznaje, że była tak zdeterminowana w swojej pracy, że mogła tego po prostu nie zauważyć. Wspiera jednak wiele rozwiązań równościowych, zwłaszcza wspierających powroty do nauki młodych matek. Jedno jest pewne sam fakt, że dyrektorką tak istotnej jak CERN instytucji jest kobieta stanowi olbrzymi krok do pełnego równouprawnienia kobiet w nauce, a w fizyce w szczególności.⁴¹

Tak doszliśmy w naszej opowieści do czasów współczesnych, mamy nadzieję, że niedługo rzesze kobiet dołączą do omawianych badaczek i kobiety na stałe zagospodzą na kartach historii fizyki, nauki i świata!

Odwrócona piramida.

Przegląd sytuacji w fizyce w Polsce

Jeśli przegląd karier sławnych badaczek z historii dziedziny nam nie pomógł, to może spójrzmy na dane statystyczne o populacji osób studiujących i pracujących w naukach fizycznych, wg. raportu GUS⁴². W 2016 roku fizykę w Polsce studiowało ponad 4300 kobiet, z tego ok. 2000 w ramach licencjatu lub studiów jednolitych, a 2300 w ramach studiów magisterskich drugiego stopnia. Studentki stanowiły ponad 65% wszystkich osób studiujących na tych dwóch poziomach.

Pomimo niżu demograficznego, udział kobiet zwiększa się również w naukach matematycznych, technicznych i informatycznych. Innymi słowy: choć osób studiujących kierunki ścisłe i techniczne jest w Polsce coraz mniej, to coraz większy odsetek stanowią kobiety. Sytuacja zmienia się w momencie obrony doktoratu. W grudniu 2016 osób, które uzyskały stopień doktora nauk fizycznych w 2016 roku było 69 kobiet, czyli nieco ponad 30%. Na poziomie habilitacji, w tym samym roku to 55

osób, z czego 16 kobiet (niewiele poniżej 30%). Ukoronowaniem kariery naukowej w Polsce jest tytuł profesora. W 2016 roku, w dziedzinie nauk fizycznych, tytuł ten otrzymały 3 kobiety spośród 13 osób (ok. 23%).

Półzartem można te procesy przyrównać do cyklu Carnota. Odsetkowi kobiet w fizyce odpowiada objętość gazu doskonałego, a zamiast otoczenia układu fizycznego przyjmujemy otoczenie społeczne. W tej metaforze studia licencjackie to odpowiednik izotermicznego rozprężania, gdy liczba kobiet rośnie, za sprawą rekrutacji ze źródeł zewnętrznych. Studia magisterskie to odpowiednik rozprężania adiabatycznego, gdy nie ma wymiany z otoczeniem a procesy dzieją się wewnątrz układu. Doktorat to moment izotermicznego sprężania, gdy stosunkowo duża grupa studentek nagle gwałtownie maleje. Profesura i habilitacja to adiabatyczne sprężanie, czyli drugi wypadek – procesy wewnątrz środowiska zawodowego są silniejsze niż wymiana z otoczeniem.

Użycie metafory fizycznej do wyjaśnienia zjawisk społecznych rzecz jasna niczego nie wyjaśnia, bo kobiety w fizyce nie są gazami doskonałymi, a wyznaczenie stopni swobody w systemie społecznym jest dużo bardziej problematyczne. Udało nam się jednak dostrzec pierwsze prawidłowości. Po pierwsze: udział kobiet w grupie studiujących i badających fizykę stopniowo rośnie. Po drugie: zmiany są nierównomierne, liczba studiujących rośnie szybciej niż liczba badaczek na dalszych etapach kariery naukowej.

Te prawidłowości kierują nas w stronę pytań szczegółowych: jakie czynniki sprawiają, że licealistki wybierają studia fizyczne? Co sprawia, że spośród dużego grona absolwentek studiów magisterskich stosunkowo niewiele kończy doktoraty? Wreszcie: dlaczego na kolejnych etapach kariery w fizyce kobiet jest mniej niż na początku?

Natura czy kultura?

Wybrane wątki debaty o źródłach różnic

Zadawaliśmy te pytania znajomym naukowcom i badaczkom. Odpowiedzi były zazwyczaj zdecydowane i przekonujące, problem w tym, że wzajemnie sprzeczne.

Jeden zbiór odpowiedzi można podsumować jako „różnice w pracy naukowej między kobietami a mężczyznami są naturalne i wynikają z budowy mózgu”. Starsze pokolenie może przywołać książkę popularnonaukową „Mężczyźni są z Marsa, a kobiety z Wenus”⁴³, młodsze badania Simona Barona-Cohena, których popularny przegląd można znaleźć w książce „The Essential Diffe-

rence³⁴⁴. W dobie postępów w technologii obrazowania mózgu, ta pierwsza pozycja jest uznawana za przestarzałą, stąd omówimy nowszą pozycję. Badacz zaproponował pomiar aktywności mózgu w dwóch wymiarach: systematyzowania i współodczuwania. Mężczyźni uzyskiwali średnio lepsze wyniki w systematyzowaniu, kobiety w empatii. Wyniki w tych dwóch skalach miały być też predyktorami wyboru kariery w naukach ścisłych. Propozycja ta nie została przyjęta bez krytyki. Pominie debaty publicystyczne i polityczne, zamiast tego wskażemy, że badania zespołu Alyssi Kersey z 2018 roku wskazywały na brak istotnych różnic w zdolnościach matematycznych dzieci obu płci⁴⁵. Krytyka ta jest godna uwagi, ponieważ zespół nie ograniczył się do jednorazowego eksperymentu, ale przeprowadził badanie na kilku grupach dzieci, w wieku od 6 miesięcy do 8 lat.

Drugi zbiór odpowiedzi można podsumować jako „różnice w pracy naukowej między kobietami a mężczyznami wynikają z mechanizmów społecznych”. Tutaj można wskazać badania socjolożki Harriet Zuckerman i historyczki Margaret Rossiter⁴⁶, o systematycznym pomijaniu wkładu kobiet w autorstwie prac naukowych (tzw. efekt Matyldy), potwierdzone w badaniach eksperymentalnych. Przykładowo: zespół Silvii Knobloch-Westerwick proponował badanie eksperymentalne⁴⁷, w którym propozycje referatów konferencyjnych (tzw. abstrakty) były oceniane przez dwie losowo wybrane grupy. Pierwsza grupa oceniała abstrakty podpisane imionami i nazwiskami wyraźnie żeńskimi, druga – wyłącznie męskimi. Treść propozycji referatów była ta sama, a jednak te podpisane imionami męskimi zbierały lepsze oceny, zarówno od kobiet jak i od mężczyzn.

Jakkolwiek bliżej nam do drugiej grupy wyjaśnień, nie negujemy znaczenia badań obrazowania mózgu i innych neuronauk. Wprost przeciwnie: widzimy w nich fascynujące pole badań fizyki medycznej i kognitywistyki i trzymamy kciuki ze eksperymenty zespołu prof. Ducha (grupa toruńska, kognitywistyka) i prof. Durkę (grupa warszawska, fizyka medyczna). Nasze stanowisko jest inne. Nawet gdyby istniały statystycznie istotne różnice w budowie lub działaniu mózgu między mężczyznami a kobietami, nie musi z tego wynikać, że różnice te będą równie istotne zarówno wśród wszystkich ludzi co wśród fizyków. Godna rozważenia wydaje nam się hipoteza o tym, że zdolności manualne, talent do rozumowania abstrakcyjnego lub inne cechy powiązane z pracą naukową w tej podgrupie mogą być inny niż wśród reszty ludzi.

Gdyby hipoteza Barona-Cohana⁴⁷ była prawdziwa, to licealistki w Polsce nie powinny wybierać fizyki chętniej niż ich koledzy. Jak wiadomo z danych GUS cytowanych przez nas wcześniej, takie zjawisko nie zachodzi. Albo



Foto – Maria Goeppert-Meyer źródło: domena publiczna

więc uznamy, że mózgi studentek z Polski są biologicznie inne niż ich koleżanek ze Zjednoczonego Królestwa, albo musimy dodać hipotezy dodatkowe (np. o tym, że mężczyźni rezygnowali z fizyki na rzecz nauk informatycznych, które dają duże lepsze perspektywy kariery lub też o tym, że licealistki kierowane propagandą feministyczną kończą magisteria, bez predyspozycji do pracy). Pierwsze wyjaśnienie wydaje nam się zbyt ryzykowne, drugie osłabia znaczenie różnic biologicznych na rzecz zmiennych społecznych, trzecie wydaje nam się mało uzasadnione, bo studentki i doktorantki fizyki, z którymi rozmawialiśmy nie

sprawiały wrażenia ofiar propagandy.

Po drugie: jeśli spojrzymy na rozmaitość badań we współczesnej fizyce, wydaje nam się wątpliwe, aby dało się zidentyfikować pojedyncze cechy umysłu, które dają przewagę w całości fizyki. Fizyka współczesna bywa uprawiana przez pojedynczych teoretyków jak również przez ogromne zespoły badawcze (patrz np. opis socjologiczny CERN w badaniach Karin Knorr-Cetiny⁴⁸). Fizyka komputerowa, inżynieria materiałowa, analiza sygnałów z detektorów, kalibracja urządzeń w astrofizyce czy realizacja eksperymentów w niskich temperaturach – każda z tych specjalizacji wymaga innych predyspozycji. Różnorodność predyspozycji w różnych odmianach fizyki ściśle wiąże się z różnorodnością problemów, przed którymi stoją różne specjalizacje i metody badawcze.

Problemy, rozwiązania i problemy z rozwiązaniami

Zespół Mercedes Lorenzo przeanalizował wyniki eksperymentu dydaktycznego prowadzonego na Harvardzie w latach 1990-1997⁴⁹. Na potrzeby publikacji dokonano przeglądu literatury związanej z wyrównywaniem szans między kobietami a mężczyznami w naukach ścisłych. Wśród rozwiązań wyróżniono różne podejścia²:

Część podejść koncentrowała się na takim konstruowaniu zadań i przykładów, aby włączać doświadczenia i skojarzenia wspólne dla wszystkich płci. Nie chodzi tu o to, aby w zadaniach związanych z prędkością jazdy samochodem zastąpić pracą w kuchni, ale choćby o to, aby bohaterkami zadań czy przykładów częściej były kobiety. Innymi słowy: nie chodzi o to, aby powielać stereotypy związane z podziałem pracy domowej między płciami (dziewczynki gotują, chłopcy naprawiają), ale żeby fizyka nie kojarzyła się z czymś nie-dziewęcym.

Inne badania proponowały nacisk na spójność programu. Nacisk na zrozumienie uprzedniego materiału i powolne wprowadzanie nowych treści miało zapobiec „odpadaniu” uczniów z zajęć z nową treścią. Ważnym nurtem były eksperymenty związane z mieszaniami dyskusji i pracy w grupach z bardziej ustrukturyzowanymi formami

² Zdajemy sobie sprawę z tego, że konstrukcja podstaw programowych i liczba godzin fizyki w polskich szkołach nie zawsze pomagają w realizacji tych postulatów. Omawiamy stan literatury nie po to, aby nałożyć na nauczycieli dodatkowe obowiązki, ale aby kolejne rozporządzenia MEN były lepiej oparte na znanych nam badaniach społecznych.

pracy. Redukowanie współzawodnictwa miało redukować stereotyp wiążący nauki ściśle z rywalizacją. Uważni Czytelnicy zauważą, że „bocznymi drzwiami” wracamy tutaj do prac Barona-Cohena i innych badaczy, którzy różnice płci wiązali z różnicami między zdolnościami komunikacji lub empatii.

Nieuczciwe byłoby jednak zredukowanie całego zjawiska do okresu szkolnego, zwłaszcza, gdy przypomnimy sobie dane z GUS⁴² cytowane powyżej. Zmniejszanie się liczby kobiet na kolejnych etapach kariery naukowej jest powiązane z edukacją powszechną, ale nie ogranicza się do niej. Równie ważne są zjawiska zachodzące w trakcie edukacji wyższej i w pracy zawodowej. Ani szkoła, ani studia, ani środowisko pracy nie mają decydującego znaczenia, bo nie mamy do czynienia ze zjawiskiem gwałtownym, ale ze stopniowym. Stąd badania społeczne zjawiska mówią często o zjawisku dziurawego rurociągu, czyli nakładających się na siebie drobniejszych wykluczeniach.

Ta metafora jest użyteczna o tyle, że kieruje naszą uwagę również w stronę późniejszych etapów kariery w fizyce. Przykładowo, z badań Izabeli Wagner na temat mobilności międzynarodowej polskich elit naukowych wiemy, że wyjazd naukowy może być dla kobiet bardziej obciążający niż dla mężczyzn, ponieważ to na kobiety zazwyczaj jest wywierana presja związana z opieką nad dziećmi. Dotyczy to zarówno kilkudniowego wyjazdu na konferencję jak i kilku lat spędzonych na stażu podoktorskim. Innym momentem, w którym polityka publiczna wpływa na karierę jest kwestia emerytur. Zgodnie z nową ustawą o nauce, kobiety będą kończyły pracę naukową o 5 lat krócej niż koledzy. Emerytka może zostać na uczelni i prowadzić część zajęć, ale to nadal redukuje jej szansę na ukoronowanie kariery tytułem profesora.⁴²

Jednak metafora rurociągu też jest problematyczna. Po pierwsze praca akademicka nie jest jedyną formą pracy naukowej, ani nawet pracy w zawodzie fizyka. Badaczka w sektorze prywatnym lub konstruktorka urządzeń fizyki medycznej też korzystają z wiedzy fizycznej, rozwiązując problemy i kształtując cywilizację. Na dodatek kariera naukowa jest skonstruowana w taki sposób, że jest dużo mniej miejsc pracy po doktoracie niż absolwentów i absolwentek osób broniących doktorat. Yi Xue i Richard Carson, z amerykańskiej publicznej agencji statystycznej (U.S. Bureau of Labor) pokazali w 2015 roku⁵⁰, że w gospodarce amerykańskiej nie brakuje rąk, ani głów do pracy naukowej. Nie znamy podobnych danych dla Polski, ale wydaje nam się wątpliwe, aby zapotrzebowanie na pracowników związanych z fizyką było istotnie większe. Nie znaczący to oczywiście, że fizyka jest niepotrzebna poza obszarem badań naukowych lub technicznych, naszą tezę jest raczej, że postulat wzrostu liczby osób z wykształceniem w tym kierunku nie powinien być przyjmowany bezkrytycznie.

Podsumowanie

Czy jest coś co łączy badaczki znane z historii fizyki, studentki i badaczki współczesne? W jaki sposób prawidłowości rozpoznane w skali społeczeństw i gatunków przekładają się na pracę codzienną? Docieramy w ten sposób do jednej z fundamentalnych różnic między naukami

społecznymi a fizyką. W fizyce elektrony nie przejmują się swoją sytuacją, stąd ich opis naukowy jest dość stabilny w czasie. Jeśli coś udało się ustalić, to zazwyczaj można założyć, że to zjawisko się powtórzy w przyszłości.

W naukach społecznych uczestniczki badań reagują na treść i koncepcje socjologiczne, nie są przedmiotami badania, ale właśnie uczestniczkami, wchodzącymi w dialog z osobami badającymi. Nie czekają aż socjologia zdiagnozuje problem, próbują radzić sobie z nim na własną rękę i przy pomocy własnych przemyśleń. Próbują na różne sposoby poskładać swoją tożsamość, czasami walcząc ze stereotypami płci i nauki, a czasami je przepracowując.

Konieczność rekonstruowania własnej tożsamości, godzenia oczekiwań przypisywanych kobiecie z oczekiwaniami przypisywanymi fizykowi wymaga wysiłku. Z badań zespołu Louise Archer wynika, że ta praca emocjonalna i wysiłek związany z tworzeniem własnej tożsamości nakłada się na zwykłe problemy emocjonalne, które ma każda kobieta i mężczyzna. Według naszych intuicji, podobnego nakładu pracy emocjonalnej wymaga kontynuacja kariery naukowej. Co z tego wynika? Jak to bywa w naukach społecznych kończymy z paradoksem. Żeby zachęcić kobiety do pracy w fizyce musimy zarówno wypracować specjalne mechanizmy zachęcania dziewczynek i analizować sytuację dorosłych badaczek. Jednocześnie musimy ich pracę i naukę traktować jako coś codziennego, niewymagającego osobnych wyjaśnień.

Trudne? Być może, ale fizyczki i fizyka radziły już sobie z trudniejszymi paradoksami.

Aleksandra Mielewczyk-Gryń

Politechnika Gdańska Wydział Fizyki Technicznej i Matematyki Stosowanej

Marcin Zaród

Akademia Leona Koźmińskiego Katedra Zarządzania w Społeczeństwie Sieciowym

LITERATURA

- [36] Maria Goeppert Mayer - Biographical. Available at: <https://www.nobelprize.org/prizes/physics/1963/mayer/auto-biography/>. (Accessed: 23rd August 2018)
- [37] A Life of One's Own: Three Gifted Women and the Men They Married - Joan Dash - Google Książki. Available at: https://books.google.pl/books/about/A_Life_of_One's_Own.html?id=FQ8aAAAIAMAJ&redir_esc=y. (Accessed: 23rd August 2018)
- [38] Fabulous Fabiola: the first lady of CERN | The Times Magazine | The Times. Available at: <https://www.thetimes.co.uk/article/the-first-lady-of-cern-8f52gm50g>. (Accessed: 23rd August 2018)
- [39] Dr Fabiola Gianotti. Available at: http://www.iop.org/about/awards/hon_fellowship/hon_fellows/page_68416.html. (Accessed: 23rd August 2018)
- [40] Aad, G. et al. The ATLAS Experiment at the CERN Large Hadron Collider. *J. Instrum.* **3**, S08003–S08003 (2008).
- [41] A Celebrated Physicist With a Passion for Music - The New York Times. Available at: <https://www.nytimes.com/2018/03/07/science/fabiola-gianotti-physics-cern.html>. (Accessed: 23rd August 2018)
- [42] Szkoły wyższe i ich finanse w 2016.
- [43] Gray, J. *Mężczyźni są z Marsa, kobiety są z Wenus*. (Rebis, 2008).
- [44] Baron-Cohen, S. *The essential difference: Male and female brains and the truth about autism*. (Basic Books, 2004).
- [45] Kersey, A. J., Braham, E. J., Csumitta, K. D., Libertus, M. E. & Cantlon, J. F. No intrinsic gender differences in children's earliest numerical abilities. *npj Sci. Learn.* **3**, 221 (2018).
- [46] Rossiter, M. W. The Matthew Matilda Effect in Science. *Soc. Stud. Sci.* **23**, 325–341 (1993).
- [47] Knobloch-Westerwick, S., Glynn, C. J. & Huge, M. The Matilda Effect in Science Communication. *Sci. Commun.* **35**, 603–625 (2013).
- [48] Cetina, K. K. *Epistemic cultures: how the sciences make knowledge*. *Journal of Chemical Information and Modeling* **53**, (Harvard University Press, 1999).
- [49] Lorenzo, M., Crouch, C. H. & Mazur, E. Reducing the gender gap in the physics classroom. *Am. J. Phys.* **74**, 118–122 (2006).
- [50] Xue, Y. & Larson, R. C. STEM crisis or STEM surplus? Yes and yes. *Mon. labor Rev.* **2015**, (2015).



Rzeczony nauki bardzo często utożsamiany jest z genialnymi odkryciami, które mają miejsce raz na kilkadziesiąt czy kilkaset lat. Przykładem może tu być dzieło Kopernika „De revolutionibus orbium coelestium”, gdzie nasz wielki rodak uważany jest za inicjatora rewolucji zmian postrzegania budowy Układu Słonecznego. To niestety bardzo uproszczone spojrzenie.

Równanie Keplera, granica Laplace’a i rekurencyjne szeregi potęgowe

Rewolucja kopernikańska okiem fizyka

Paweł Wajer, Ryszard Gabryszewski

Heliocentryzm Układu Słonecznego był podnoszony już w czasach antycznych. Przywoływany najczęściej Arystarch z Samos nie był pierwszym, przed nim byli Heraklides, Platon, pitagorejczycy i wielu innych, filozofów, którzy intuicyjnie utożsamiali Słońce ze środkiem Wszechświata. Dlaczego więc idee te nie przebiły się do świadomości ówczesnych społeczeństw? Zdecydowanie większa część filozofów, łącznie z najznamienitszymi przedstawicielami świata nauki, hołdowała „naturalnemu” obrazowi świata, w którym – jak wszyscy widzieli – to Słońce porusza się wokół Ziemi. Ponadto niezbyt wysoka precyzja obserwacji i przywiązanie do artystotelejskiej wizji świata (m.in. zakładanie orbit kołowych) powodowały, że heliocentryści nie mogli udowodnić swoich racji. Błędy pomiędzy obliczonymi a obserwowanymi pozycjami planet były na podobnym poziomie jak w systemie geocentrycznym.

Kopernik nie był więc pierwszym postulatorem heliocentryzmu w nauce¹. Tak jak poprzednicy zakładał, że ciała poruszają się po okręgach. Nie zerwał też z epicyklami i deferentami, co pozwalało uzyskać lepszą zgodność obliczeń orbitalnych i obserwacji. Kopernik nie podał także żadnego dowodu na prawdziwość teorii heliocentrycznej, dowody te były mozolnie zbierane przez kolejne pokolenia. Tu warto wspomnieć trzech innych, wybitnych astronomów: Galileusza, Keplera i Newtona.

W 1608 roku holenderski optyk, Hans Lippershey wynalazł lunetę. Rok później Galileusz, opierając się

na pracach Lippersheya zbudował własny przyrząd obserwacyjny i skierował go na sferę niebieską śledząc m.in. ruch 4 największych księżyców (nazwanych znacznie później galileuszowymi) wokół Jowisza oraz fazy Wenus. Obserwacje te pokazały, że nasz układ planetarny jest zbudowany inaczej niż zakładał system geocentryczny. W szczególności istnienie faz Wenus jasno dowodziło, że planeta ta okrąża Słońce.

W tym samym roku co pierwsze obserwacje Galileusza, pojawia się *Astronomia Nova*, wielkie dzieło Johannes Keplera i jednocześnie jedna z najważniejszych pozycji w historii astronomii. Prezentuje ona wyniki ponad 10-letnich badań ruchu Marsa dając jednocześnie silne argumenty za heliocentryzmem: lepszą zgodność przewidywań pozycji planet w stosunku do obserwacji w modelu kopernikańskim dzięki wprowadzeniu idei orbit eliptycznych. Ostateczny cios geocentryzmowi zadał jednak Newton. Bazując na prawach Keplera wytłumaczył ruch planet wokół Słońca i nadał nazwę sile, która spaja Układ Słoneczny – grawitacja.

Rzeczony nauki to nieodłączny element każdej cywilizacji, który nie jest mierzony wielkimi, rzadko wykonywanymi krokami. To setki większych i mniejszych kroków, z których każdy jest równie istotny. Jednym z istotniejszych są prace Keplera.

Kepler obserwował ruch Marsa po sferze niebieskiej, dysponował także bardzo precyzyjnymi obserwacjami pozycyjnymi planety wykonanymi przez Tyhona de Brahe. Wiedział, że planety w pewnej części swojej orbity poruszają się szybciej, co wskazywało, że okrążane ciało

¹ Często wspomina się, że tzw. rewolucja kopernikańska zmieniła sposób patrzenia na budowę Wszechświata. Patrząc na historię heliocentryzmu trudno oprzeć się wrażeniu, że jest to dalece niepełna interpretacja. Powrót do teorii heliocentrycznej przez Kopernika był rewolucją głównie w sensie filozoficznym: to jeden z pierwszych w czasach nowożytnych istotnych impulsów, które doprowadziły do zmian w sposobie myślenia: o budowie świata, sposobie poznawania praw nim rządzących a przez to także do zmian postrzegania miejsca i funkcji religii w życiu społecznym. Jeszcze na początku XVII wieku opis świata zawarty w Biblii był powszechnie rozumiany w sposób dosłowny, a każde sprzeczne twierdzenie groziło oskarżeniem o herezję i skazaniem przez sąd na karę więzienia lub nawet śmierci na stosie. Renesans odsunął metafizykę od nauki opierając tę ostatnią na doświadczeniu i obserwacjach, doprowadził także do zmian społecznych: sprawił, że ludzie przestali być karani za „nieprawomyślne” poglądy (w krajach, w których później rozwinął się system demokratyczny).

(Słońce) nie może znajdować się w środku orbity kołowej lub też orbita ... nie jest kołowa. Kepler nie miał problemu ze wskazaniem ciała, wokół którego krąży Mars, był zwolennikiem teorii heliocentrycznej. W swoich pracach zastanawiał się nad powodem, który sprawiał, że planety krążą w przestrzeni. Powód ten upatrywał w Słońcu, które oddziaływało na planety siłą lub siłami (w sensie nienewtonowskim, Kepler zmarł przed urodzeniem Newtona). Jednak największy problem sprawiło wyjście poza wielowiekowy schemat orbit kołowych.

Do czasów Keplera astronomia zakładała kołowość orbit planet. Założenie to miało swój początek w metafizycznym postrzeganiu Wszechświata począwszy od arystotelesowskiej idei sfer, po których miały poruszać się ciała niebieskie. Okrąg zaś był zawsze traktowany jako ucieleśnienie doskonałości formy oraz boskiego porządku Wszechświata. Także i Kepler rozpoczął swoje poszukiwania kształtu orbit od okręgu ze Słońcem odsuniętym od środka. Metodą prób i błędów, testując także inne kształty orbit, doszedł do wniosku, że planety poruszają się po elipsach. Kilkuletnie poszukiwania właściwego kształtu orbit spowodowane były zarówno niezbyt dokładnymi obserwacjami pozycyjnymi – pamiętajmy, że do 1609 roku obserwacje były wykonywane za pomocą prostych przyrządów pozwalających określać odległości kątowe na sferze względem gwiazd – oraz małą eliptycznością orbit planet (ekscentryczność orbity Marsa wynosi zaledwie ok. 0.1). Nie mniej, wyznaczając błąd efemerydy Marsa przy założeniu eliptycznej orbity planety widać wyraźnie, że jest on zdecydowanie mniejszy, niemal zerowy w porównaniu do wyznaczeń w oparciu o teorię zakładającą orbitę kołową (Rysunek 1).

Wydana w 1609 roku *Astronomia Nova* zawierała także pierwsze dwa prawa Keplera:

- I planety poruszają się wokół Słońca po orbitach eliptycznych, a Słońce znajduje się w jednym z ich ognisk;
- I w równych odstępach czasu promień wodzący planety, poprowadzony od Słońca, zakreśla równe pola.

Trzecie prawo:

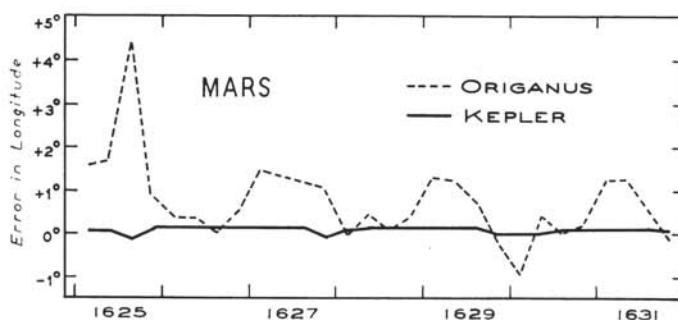
- I stosunek kwadratu okresu obiegu planety wokół Słońca do sześcienu wielkiej półosi jej orbity jest stały dla wszystkich planet (w układzie),

Kepler sformułował znacznie później i opublikował w 1619 roku w książce *Harmonices Mundi*.

Niejednostajny ruch obiegowy ciała po orbicie eliptycznej sprawia pewną trudność w obliczeniu pozycji planety na dowolny moment czasu – wymaga rozwiązania równania, które dziś nazywamy równaniem Keplera:

$$E = M + e \sin E \quad (1)$$

gdzie M oznacza kąt anomalii średniej, E – kąt anomalii mimośrodowej, zaś e – mimośród orbity. W równaniu tym M oraz e jest dane, natomiast zmienną poszukiwaną jest E . Niestety, z uwagi na to, iż występuje w powyższym równaniu bezpośrednio oraz jako $\sin E$, nie da się znaleźć prostej formuły, z której można wyliczyć E mając dane M oraz e . W praktycznych zastosowaniach równanie Keplera najczęściej rozwiązywane jest numerycznie.



Rys 1. Błąd efemerydy planety przy założeniu kołowej (linia przerywana) oraz eliptycznej (linia ciągła) orbity Marsa, wykres z publikacji [1].

Niniejszy artykuł ma za zadanie przybliżyć najciekawsze metody rozwiązania tego równania, przy czym skoncentrowaliśmy się na metodach wykorzystujących rozwinięcia anomalii mimośrodowej w różnego rodzaju szeregi oraz przedstawiliśmy analizę tych szeregów, m.in. zakres ich stosowności.

Metoda iteracji

Historycznie najstarszą, używaną już przez Keplera metodą rozwiązania tego równania jest metoda iteracji. Przekształcamy równanie Keplera do formy, w której wartość E wyznaczamy iteracyjnie w kolejnych krokach:

$$E_{n+1} = M + e \sin E_n \quad (2)$$

gdzie n jest kolejnym krokiem iteracji. W pierwszym kroku zazwyczaj zakładamy $E_0 = M$, przy czym należy pamiętać, iż wielkości E oraz M muszą być wyrażone w radianach. Rachunki przerywamy, gdy różnica wartości E w kolejnych krokach jest mniejsza niż zakładany błąd: $|E_{n+1} - E_n| < \varepsilon$, np.: $\varepsilon = 10^{-13}$.

Metoda ta jest zawsze zbieżna. Jednak uzyskanie dokładnego rozwiązania – w czasach przed wynalezieniem procesora – wymagało bardzo długich i żmudnych obliczeń. Dlatego też szukano innych, czasami nawet bardzo skomplikowanych metod, które gwarantowałyby zarówno wysoką dokładność i pozwoliłyby istotnie skrócić czas obliczeń.

Powyżej przedstawiona metoda jest metodą rzędu pierwszego, tj. w każdym kolejnym kroku różnica między rozwiązaniem dokładnym a przybliżonym zmniejsza się w przybliżeniu e razy, gdzie e jest mimośrodem orbity. Oznacza to, że uzyskanie dużej dokładności dla orbit o dużej ekscentryczności będzie wymagało znacznie więcej obliczeń niż dla orbit quasi-kołowych, tj. takich, które mają mimośród bliski zeru.

Inne metody to te bazujące na metodzie Newtona rozwiązywania równań nieliniowych. Najprostszy sposób zastosowania metody Newtona, będącej metodą rzędu drugiego, daje następującą iteracyjną procedurę znajdowania E :

$$E_{n+1} = E_n - \frac{E_n - e \sin E_n - M}{1 - e \cos E_n} \quad (3)$$

gdzie jako wartość początkową, dla małych wartości e , $E_0 = M$, dla $e > 0,8$ lepiej jest przyjąć $E_0 = \pi$. Bardziej wyrafinowane rozwiązania rzędu czwartego bazujące na meto-

dzie Newtona podał Danby. Do rozwiązywania równania Keplera można również zastosować metodę bisekcji lub siecznych.

Metody wykorzystujące szeregi nieskończone

W rozważaniach teoretycznych używa się jednak często rozwinięcia szukanej anomalii mimośrodowej w szereg. Najbardziej znany to ten pochodzący od Lagrange'a [2]:

$$E = M + \sum_{n=1}^{\infty} a_n(M) e^n \quad (4)$$

gdzie współczynniki a_n są funkcjami M , w szczególności $a_1(M) = \sin M$, $a_2(M) = 1/2 \sin 2M$, $a_3(M) = 1/8(-\sin M + 3 \sin 3M)$ itp. Wrócimy do tego równania w dalszej części artykułu. Inny przykład rozwinięcia anomalii mimośrodowej w szereg, który również rozważał Lagrange, ma postać [2]:

$$E = M + \sum_{n=1}^{\infty} b_n(e) \sin nM \quad (5)$$

Lagrange podał również wartości współczynników $b_n(e)$ dla kilku początkowych wartości n . Około pół wieku później Bessel podał ogólny wzór na obliczanie $b_n(e)$ korzystając z funkcji nazywanych obecnie jego nazwiskiem, tj. funkcji J_n Bessela. Mniej znane jest natomiast rozwinięcie E w szereg podane przez Stumpffa [3]:

$$E = \sum_{n=1}^{\infty} c_n(e) M^n \quad (6)$$

gdzie wyrazy o numerach parzystych mają wartość zero, natomiast $c_1(e) = 1/(1-e)$, $c_3(e) = -e/[6(1-e)^4]$, itp. Wartość anomalii średniej podana jest tutaj w radianach. Mimo iż powyżej przedstawione szeregi zawierają nieskończoną liczbę wyrazów, w praktyce, aby uzyskać dobre przybliżenie szukanej wartości anomalii mimośrodowej, wystarczy ograniczyć się do kilku pierwszych wyrazów szeregów. Liczba wyrazów, które należy uwzględnić zależy od dokładności z jaką chcemy uzyskać wartość E , czy jak, np. w przypadku szeregu (4) – wartości mimośrodu e . Im mniejsza wartość mimośrodu, tym wartości potęg e^n szybciej zbiegają do zera wraz ze wzrostem n . Aby to lepiej wyjaśnić rozpatrzmy 2 przykłady.

Założmy, że w obu przypadkach $M = \pi/4$ radianów. Biorąc kolejno 1, 2, 3 i 4 wyrazy rozwinięcia (4) znajdujemy kolejne przybliżenia wartości szukanej anomalii mimośrodowej. Dla Ziemi ($e=0,0167$) wynoszą one kolejno: 0,707107, 0,797207, 0,797346 i 0,797347. Widzimy zatem że wyniki dla 3 i 4 wyrazów rozwinięcia różnią się na 6 cyfrze znaczącej. Zatem z dokładnością do 5 cyfr znaczących możemy założyć, że w tym przypadku $E=0,79735$. Wykonując podobne obliczenia dla Marsa ($e = 0,093$) uzyskujemy jako kolejne przybliżenia wartości anomalii mimośrodowej następujące liczby: 0,707107, 0,851159, 0,855484, 0,855626. Otrzymane wyniki pokazują, że w przypadku orbity Marsa (większy mimośród niż dla Ziemi) aby uzyskać wartość anomalii mimośrodowej z dokładnością do 5 miejsc znaczących trzeba wykorzystać więcej wyrazów rozwinięcia (4). W szczególności dla orbit quasi-kołowych można zastosować

następujące przybliżenie służące do obliczania E : $E \approx M + e \sin M$ (wzięliśmy tylko dwa wyrazy z szeregu 4), lub wzór trochę bardziej dokładny: $E \approx M + e \sin M + 1/2 e^2 \sin 2M$ (wzięliśmy trzy wyrazy szeregu (4)).

Ponieważ artykuł dotyczy rozwiązywania równania Keplera za pomocą szeregów potęgowych, zatem szeregiem (5), który jest trygonometryczny, nie będziemy się poniżej zajmowali. Warto jednak wspomnieć, że szereg (5) jest zbieżny dla wszystkich wartości M oraz $0 \leq e < 1$.

Inaczej jest w przypadku szeregu (4). Laplace chyba jako pierwszy pytał o przedział zbieżności tego szeregu oraz jako pierwszy uzyskał na to pytanie odpowiedź: szereg ten na pewno jest zbieżny dla $e < 0,66274\dots$. Liczba ta zwana jest granicą Laplace'a², aczkolwiek Laplace granicę zbieżności szeregu oszacował jako $e < 0,66195\dots$

Zauważmy jednak, że współczynniki rozwinięcia według potęg mimośrodu są funkcją anomalii średniej. Powstaje zatem pytanie czy promień zbieżności tego szeregu jest stały, czy zależy od M i czy dla konkretnych wartości M szereg Lagrange'a może być zbieżny dla mimośródów większych od granicy Laplace'a? Odpowiedzi na te pytania są już oczywiście znane. W tym artykule chcielibyśmy nie tylko podać odpowiedź na powyższe pytanie (oraz analogiczne dotyczące szeregu Stumpffa) ale również przedstawić sposób uzyskania tej odpowiedzi.

Przez sposób uzyskania rozumiemy nie ścisłą analizę matematyczną tego zagadnienia a jedynie numeryczne oszacowania. Żeby móc numerycznie dokonać analizy tych szeregów dobrze jest znać jawnie wartości współczynników $a_n(M)$ oraz $c_n(e)$. Jawna postać tych współczynników jest oczywiście znana, aczkolwiek sposób jej uzyskania jest dość złożony. W tym artykule przedstawiony został sposób na uzyskanie tych współczynników metodą rekurencyjnych szeregów potęgowych. Metoda ta jest na tyle uniwersalna, że może być stosowana również do innych zagadnień oraz (choćby rozwiązywania równania Keplera dla orbit hiperbolicznych, czy rozwiązywania równań ruchu ciał w Układzie Słonecznym), co równie istotne, łatwo jest ją zaprogramować [4].

Metoda rekurencyjnych szeregów potęgowych

Rozpocznijmy od przykładu. Założmy, że chcemy rozwinąć w szereg potęgowy funkcję $f(x) = 1/(1-x)$. Założmy ponadto, że da ona się przedstawić w postaci szeregu potęgowego:

$$f(x) = \sum_{n=0}^{\infty} f_n x^n \quad (7)$$

Jeżeli zapiszemy równanie wyrażającą funkcję $f(x)$ w postaci $(1-x)f(x) = 1$ i podstawimy powyższy szereg w miejsce $f(x)$ (grupując uprzednio według potęg x):

$$f_0 + \sum_{n=1}^{\infty} (f_n - f_{n-1}) x^n = 1 \quad (8)$$

Widzimy zatem, że aby powyższa równość była spełniona niezależnie od wartości x musi zachodzić $f_0 = 1$ oraz

² Jeżeli x_0 jest jedynym dodatnim rozwiązaniem równania $\coth x_0 = x_0$ to granica Laplace'a jest równa wartości wyrażenia $1/\sinh x_0$, gdzie $\coth x_0$ oraz $\sinh x_0$ są to funkcje kotangens oraz sinus hiperboliczny argumentu x_0 .

$f_n = f_{n-1}$, dla $n = 1, 2, \dots$, czyli $f_1 = f_0 = 1, f_2 = f_1 = 1$ itd. Przykład może wydać się dość banalny, gdyż wiemy, że rozwinięciem w szereg potęgowy według potęg x funkcji $1/(1-x)$ jest $1 + x + x^2 + x^3 + \dots$. Spróbujmy jednak znaleźć rozwinięcie według potęg x takiej funkcji: $1/(1 - \sin x)$. W tym przypadku możemy oczywiście zastosować znany wzór Taylora³ i różniczkować wielokrotnie tę funkcję, jednak jak łatwo zauważyć, obliczanie drugiej czy kolejnych pochodnych staje się dość uciążliwe. Stosując analogiczne rozważania jak w powyższym przykładzie (znając przy tym rozwinięcie funkcji $\sin x$ w szereg potęgowy⁴) uzyskujemy następujący związek rekurencyjny:

$$f_n = \begin{cases} 1 & \text{gdy } n=0 \\ -\sum_{k=1}^n f_{n-k} b_k & \text{gdy } n>0 \end{cases} \quad (9)$$

Gdzie b_k to współczynniki rozwinięcia w szereg funkcji $1 - \sin x$ ($b_0 = 1, b_1 = -1, b_2 = 0, b_3 = 1/6$ itd.). Po rozpatrzeniu powyższych przykładów możemy wrócić do równania Keplera. Ponieważ chcemy rozwinąć anomalię mimośrodową E w funkcję czy to mimośrodu czy anomalii prawdziwej, z równania Keplera widzimy, że potrzebny będzie nam związek między E rozwiniętym w szereg oraz $\sin E$. Innymi słowy, mając dany szereg potęgowy szukamy sinusa tego szeregu.

Pokażemy teraz, w jaki sposób można uzyskać taki szereg. Załóżmy, że znamy rozwinięcie funkcji $f(x)$ w postaci szeregu potęgowego $f_0 + f_1 x + f_2 x^2 + f_3 x^3 + \dots$. Mając takie rozwinięcie, szukamy szeregu przedstawiającego funkcję $\sin f(x) = s_0 + s_1 x + s_2 x^2 + s_3 x^3 + \dots$

Załóżmy, że wartościami tej funkcji są liczby rzeczywiste. Wtedy proces znajdowania liczb s_n (i przy okazji współczynników c_n – rozwinięcia $\cos f(x)$ w szereg według potęg x) można uprościć rozpatrując funkcje:

$$e^{if(x)} = \cos f(x) + i \sin f(x) = \sum_{n=0}^{\infty} (c_n + i s_n) x^n \quad (10)$$

Różniczkując powyższą równość po x otrzymujemy:

$$if'(x) e^{if(x)} = \sum_{n=0}^{\infty} (n+1) (c_n + i s_n) x^n \quad (11)$$

Prawą stronę mamy już w postaci szeregu, teraz czas na lewą. Wyrażenie $e^{if(x)}$ dane jest przez (10), natomiast różniczkując wyraz po wyrazie szereg przedstawiający $f(x)$ uzyskujemy szereg przedstawiający $f'(x)$. Ostatecznie, część rzeczywista powyższego wyrażenia da nam szereg na współczynniki $\cos f(x)$, urojona – $\sin f(x)$:

$$c_{n+1} = -\frac{1}{n+1} \sum_{k=0}^n (k+1) f_{k+1} s_{n-k} \quad (12)$$

$$s_{n+1} = \frac{1}{n+1} \sum_{k=0}^n (k+1) f_{k+1} c_{n-k} \quad (13)$$

³ Jeżeli funkcja $f(x)$ jest n -krotnie różniczkowalną w otoczeniu jakiegoś punktu $x = a$, to wzór Taylora pozwala przybliżyć tę funkcję w otoczeniu tego punktu za pomocą wielomianu rzędu n oraz dostatecznej małej reszty. Wzór ten ma następującą postać:

$f(x) = f(a) + \frac{x-a}{1!} f^{(1)}(a) + \frac{(x-a)^2}{2!} f^{(2)}(a) + \frac{(x-a)^3}{3!} f^{(3)}(a) + \dots + \frac{(x-a)^n}{n!} f^{(n)}(a) + R(x)$, gdzie $f^{(n)}(a)$ to n -ta pochodna funkcji $f(x)$ obliczona w punkcie $x = a$, natomiast $R(x)$ jest resztą.

⁴ Rozwinięcie to ma postać: $\sin x = x - \frac{x^3}{3!} + \frac{x^5}{5!} - \frac{x^7}{7!} + \dots$

gdzie oczywiście $s_0 = 0$ oraz $c_0 = 1$.

Uzbrojeni w powyższe wyrażenia rekurencyjne możemy przejść do rozwiązywania równania Keplera za pomocą szeregów. Na początek rozwinięcie według potęg mimośrodu, czyli znany szereg Lagrange'a.

Szereg Lagrange'a

W przypadku szeregu Lagrange'a współczynniki $a_n(M)$ występujące w rozwinięciu (4) można wyrazić w następujący, jawny, tzn. nierekurencyjny sposób:

$$a_n(M) = \frac{1}{2^n n!} \sum_{k=0}^{\lfloor n/2 \rfloor} (-1)^k \binom{n}{k} (n-2k)^{n-1} \sin[(n-2k)M] \quad (14)$$

Sposób uzyskania tego wzoru nie jest jednak zbyt atrakcyjny, ponadto, do obliczeń numerycznych nie potrzebujemy znać tych współczynników jawnie a wystarczy nam jak są dane w sposób rekurencyjny. Aby uzyskać zależności rekurencyjne dla $a_n(M)$ podstawiamy szereg (4) do równania Keplera:

$$\sum_{n=0}^{\infty} a_n e^n - e \sin \left(\sum_{n=0}^{\infty} a_n e^n \right) = M \quad (15)$$

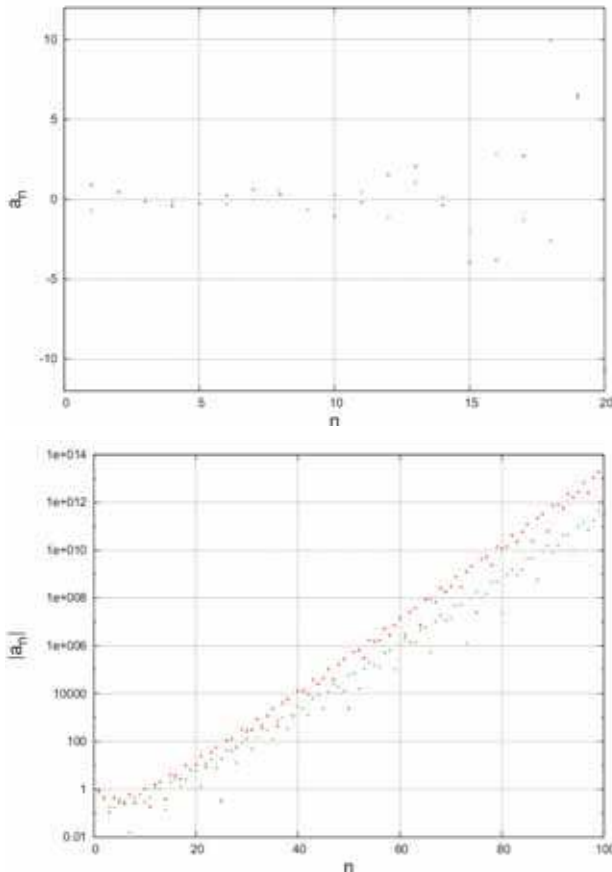
gdzie współczynniki $a_n(M)$ zapisano w skróconej postaci a_n . Powyższe równanie można zapisać zatem w postaci

$$\sum_{n=0}^{\infty} a_n e^n - \sum_{n=1}^{\infty} s_n e^n = M \quad (16)$$

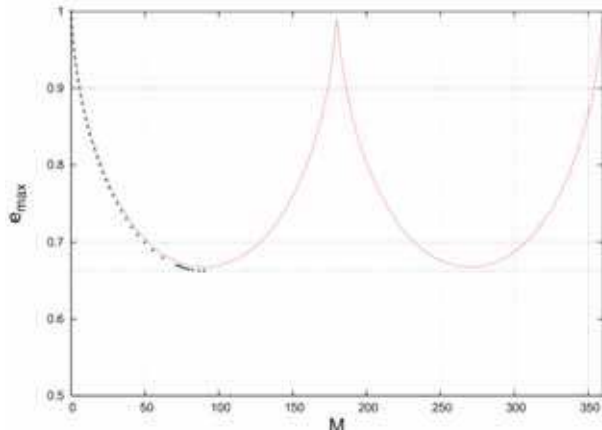
Wykorzystując zależności (11) i (12) (gdzie jako f_k przyjmujemy a_k uzyskujemy ostatecznie: $a_0 = M, s_0 = \sin M, c_0 = \cos M$ oraz $a_n = s_{n-1}$. Procedura znajdowania współczynników a_n wygląda zatem następująco. Wartości s_0 oraz c_0 podstawiamy do równań (12) i (13) z których uzyskujemy kolejno s_1, c_1, s_2, c_2 itd. Mając wartości współczynników s_0, s_1, s_2 , znajdujemy kolejno a_1, a_2 itd.

Zobaczmy zatem jak zachowują się te współczynniki dla różnych wartości M oraz n . Rysunek 2 przedstawia wykres tych współczynników dla dwóch wartości anomalii średniej $M = 60^\circ$ i $M = 225^\circ$. O ile dla małych wartości n nie szczególnie interesującego nie możemy powiedzieć o tych współczynnikach, to gdy wykreślimy ich wartości dla większych wartości n możemy już poczynić interesujące obserwacje.

Po pierwsze, jak widać z dolnego wykresu współczynniki te, a dokładniej ich moduł, rosną bardzo szybko osiągając wartości większe od miliona już dla n większego niż 60, większe od miliarda dla $n > 80$, natomiast gdy $n > 100$ to a_n jest już rzędu co najmniej biliona. Zauważamy również, że szybkość wzrostu tych współczynników zależy od M . Dla $M = 225^\circ$ rosną one wolniej niż dla $M = 60^\circ$. Widzimy również w przybliżeniu liniowy wzrost (w skali logarytmicznej) wartości $|a_n|$ dla dużych wartości n . Oznacza to, że gdy n jest duże zachodzi $|a_n| \cong C A^n$, gdzie A i C są pewnymi stałymi (w ogólności zależnymi od M).



Rys. 2. Wartości współczynników a_n dla $n \leq 20$ (rysunek górny) oraz $|a_n|$ dla $n \leq 100$ w skali logarytmicznej (rysunek dolny) policzone dla dwóch wartości M : kolor czerwony $M = 60^\circ$, zielony $M = 225^\circ$.



Rys. 3. Promień zbieżności szeregu Lagrange'a w funkcji anomalii średniej. Linia czerwona zaznaczona jest maksymalna wartość mimośrodowa, dla którego rozważany szereg jest zbieżny, linia zielona to teoretycznie wyznaczona granica Laplace'a, punkty niebieskie to promień zbieżności policzony teoretycznie.

Możemy teraz odpowiedzieć na pytanie dotyczące promienia zbieżności szeregu Lagrange'a. Dla przypomnienia: promień zbieżności szeregu potęgowego (danego w postaci (7)) to taka największa liczba r , że dla wszystkich liczb rzeczywistych x takich, że $|x| < r$ szereg ten jest zbieżny. Promień zbieżności szeregu można aproksymować w następujący sposób: $r \approx 1/\sqrt[n]{|a_n|}$ przy czym im większe n użyjemy tym dokładniejszą wartość przybliżenia wartości promienia zbieżności uzyskamy.

Zauważmy, że stała A w powyższych rozważaniach dotyczących współczynników $|a_n|$ to po prostu odwrotność promienia zbieżności szeregu. Wyznamy zatem

promień zbieżności szeregu Lagrange'a numerycznie. W obliczeniach przyjęto $n = 1700$. Wyniki tych obliczeń przedstawione są na rysunku 3. Linia czerwona zaznaczona jest maksymalna wartość mimośrodowa, dla którego rozważany szereg jest zbieżny, linia zielona to teoretycznie wyznaczona granica Laplace'a, punkty niebieskie to promień zbieżności policzony teoretycznie. Widać dość dużą zgodność rozważań numerycznych z wynikami teoretycznymi. Możemy teraz udzielić odpowiedzi na pytania zawarte we wstępie. Widzimy, iż promień zbieżności szeregu zależy od anomalii prawdziwej. Dla $M = 90^\circ$ oraz $M = 270^\circ$ jest najmniejszy i równy jest po prostu granicy Laplace'a, dla innych wartości anomalii promień zbieżności jest większy, dla M bliskich 0° , 180° oraz 360° szereg Lagrange'a jest zbieżny nawet dla e bliskiego jedności. Mamy zatem następującą interpretację granicy Laplace'a: granica ta to po prostu maksymalna wartość mimośrodowa gwarantująca, że szereg Lagrange'a jest zbieżny dla wszystkich wartości anomalii średniej.

Szereg Stumpffa

Użyteczność metody rekurencyjnych szeregów potęgowych widać szczególnie w przypadku znajdowania współczynników $c_n(e)$ występujących w szeregu (6). Stumpff wyraził te współczynniki w następujący sposób [3]:

$$c_n(e) = \frac{1}{n!} D^n x \Big|_{x=0} \quad (17)$$

gdzie

$$D^1 x = \frac{1}{1 - e \cos x} \quad (18)$$

$$D^{n+1} x = \frac{1}{1 - e \cos x} \frac{d}{dx} D^n x \quad (19)$$

W szczególności otrzymujemy

$$c_1(e) = \frac{1}{1 - e \cos x} \Big|_{x=0} = \frac{1}{1 - e \cos 0} = \frac{1}{1 - e} \quad (20)$$

$$\begin{aligned} c_2(e) &= \frac{1}{2} \frac{1}{1 - e \cos x} \frac{d}{dx} D^1 x \Big|_{x=0} = \\ &= \frac{1}{1 - e \cos x} \frac{d}{dx} \frac{1}{1 - e \cos x} \Big|_{x=0} = \\ &= -\frac{1}{1 - e \cos x} \frac{e \sin x}{(1 - e \cos x)^2} \Big|_{x=0} = 0 \end{aligned} \quad (21)$$

$$\begin{aligned} c_3(e) &= \frac{1}{6} \frac{1}{1 - e \cos x} \frac{d}{dx} D^2 x \Big|_{x=0} = \\ &= -\frac{1}{6} \frac{1}{1 - e \cos x} \frac{d}{dx} \left(\frac{1}{1 - e \cos x} \frac{d}{dx} \frac{1}{1 - e \cos x} \right) \Big|_{x=0} = -\frac{e}{6(1 - e)^4} \end{aligned} \quad (22)$$

Zauważmy, że współczynniki te obliczane są w sposób rekurencyjny. Niestety, z uwagi na konieczność obliczania w każdym kolejnym kroku pochodnej z coraz bardziej złożonej funkcji, sposobu tego nie da się prosto zaimplementować numerycznie. Wróćmy do metody szeregów

potęgowych. Podstawmy zatem prawą stronę równania (6) do równania Keplera, z zależności (12) i (13) otrzymujemy ostatecznie dla $n = 1$

$$c_1(e) = \frac{1}{1-e} \quad (23)$$

oraz dla $n > 1$

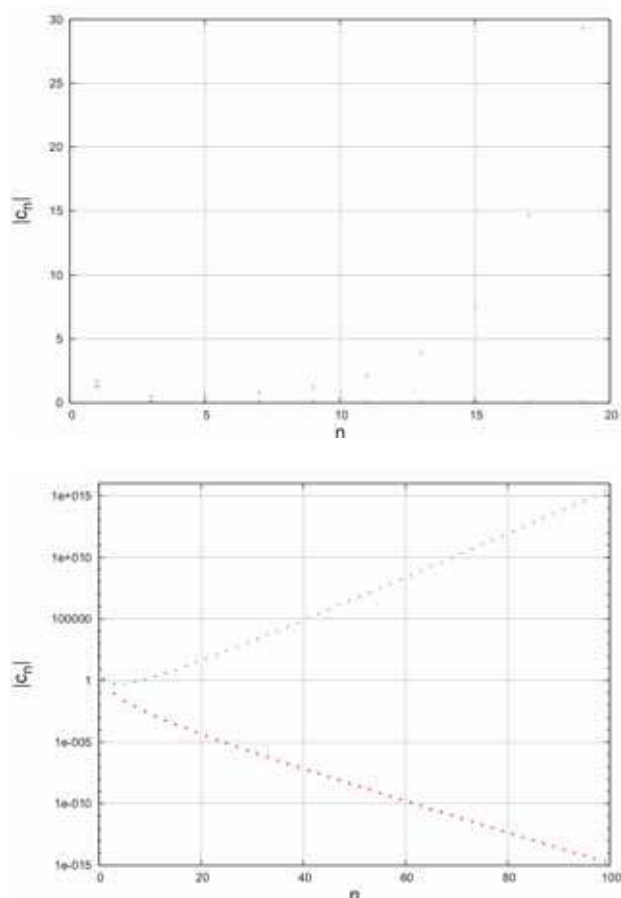
$$c_n(e) = \frac{\sum_{k=0}^{n-2} (k+1)c_{k+1}(e)c_{n-k-1}(e) + \delta(n)}{1-e} \quad (24)$$

gdzie $\delta(n) = 1$ dla $n = 1$, w pozostałych przypadkach $\delta(n) = 0$.

Współczynniki c_0, c_1, c_2 itd. obliczane są z zależności (12) i (13) gdzie jako f_k przyjmujemy $c_k(e)$ oraz oczywiście $c_0(e) = 0$.

Możemy teraz zająć się analizą współczynników występujących w szeregu Stumpffa, które zostały przedstawione na rysunku 4 dla $n \leq 20$, $n \leq 100$ oraz dwóch wartości mimośrodu. Rysunek ten pokazuje, że dla $e = 0,2$ liczby $|c_n|$ bardzo szybko maleją, natomiast dla $e = 0,4$ bardzo szybko rosną. W skali logarytmicznej wzrost ten, jak również spadek (dla $e = 0,2$) jest w przybliżeniu liniowy. Mamy zatem w przybliżeniu $|c_n(e)| \cong CA^n$, gdzie stałe A oraz C zależą od e . W szczególności dla $e = 0,2$ zachodzi $A < 1$, natomiast dla $e = 0,4$ jest $A > 1$.

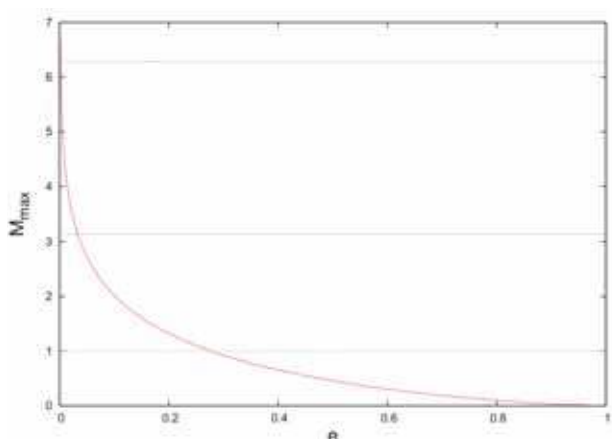
Na koniec oszacujmy promień zbieżności szeregu Stumpffa w funkcji mimośrodu. Wyniki obliczeń poka-



Rys. 4. Wartości współczynników $|c_n(e)|$ dla $n \leq 20$ (rysunek górny) oraz $c_n(e)$ dla $n \leq 100$ skali logarytmicznej (rysunek dolny) policzone dla dwóch wartości e : kolor czerwony $e = 0,2$ zielony $e = 0,4$. Wykreślone zostały współczynniki o indeksach nieparzystych.

zuje rysunek 5. Z rysunku tego widać, że tylko dla bardzo małych wartości e (rzędu 10^{-3} i mniejszych) szereg Stumpffa jest zbieżny dla wszystkich wartości M mniejszych od 2π , co pokazuje pozioma linia fioletowa. Niewiele lepiej jest, gdy $M \leq \pi$, wtedy szereg ten jest zbieżny tylko dla $e \lesssim 0,03$ (pozioma linia niebieska).

Z tego rysunku możemy również odczytać (pozioma linia zielona) dla jakich wartości mimośrodu współczynniki $|c_n(e)|$ maleją, gdy n jest dostatecznie duże. Widzimy zatem, że gdy $e \lesssim 0,27$ to wraz ze wzrostem n mamy $|c_n(e)| \rightarrow 0$. Na koniec warto podkreślić, iż wszystkie przedstawione w tej części obliczenia są przybliżone. Nie udało nam się znaleźć w literaturze informacji na temat analizy zbieżności szeregu Stumpffa.



Rys. 5. Maksymalna wartość anomalii średniej w funkcji mimośrodu (linia czerwona) gwarantująca zbieżność szeregu Stumpffa. Opis znaczenia pozostałych linii znajduje się w tekście.

Zakończenie

Równanie Keplera wiążące ze sobą mimośród orbity, anomalie średnią oraz mimośród orbity eliptycznej jest podstawowym, ale zarazem jednym z najważniejszych równań w mechanice nieba, jak również w astrodynamice. Rozwiązanie tego równania, z uwagi na to, iż jest przestępne, nie może być wyrażone w skończonej postaci. Zazwyczaj rozwiązuje się je korzystając z różnych metod iteracyjnych. W niniejszym artykule pokazaliśmy w jaki sposób można to równanie rozwiązać wykorzystując szeregi potęgowe. Przedstawiliśmy kilka rodzajów tego typu rozwiązań oraz przedstawiliśmy ich analizę, w szczególności zakres ich stosowalności. Aby znaleźć rozwiązania równania Keplera wykorzystaliśmy metodę rekurencyjnych szeregów potęgowych. Metoda ta jest wykorzystywana m.in. w znajdowaniu rozwiązań ruchu n-ciał.

dr Paweł Wajer i dr Ryszard Gabryszewski
Centrum Badań Kosmicznych PAN

LITERATURA

- [1] Owen Gingerich, „Johannes Kepler and the Rudolphine Tables” Sky and Telescope, grudzień 1971, str. 328.
- [2] Carl D. Murray, Stanley F. Dermott, „Solar System Dynamics”, Cambridge, UK: Cambridge University Press, 1999.
- [3] Karl Stumpff „On the application of Lie-series to the problems of celestial mechanics.”, Rep. NASA Tech. Note, NASA-TN-D-4460, 29 p. styczeń 1970.
- [4] Grzegorz Sitarski „Recurrent power series integration of the equations of comet's motion”, Acta Astronautica, vol. 29, no. 3, 1979, str. 401-411.

Wykorzystanie rozumowania matematycznego do analizy równoległego połączenia oporników

Propozycja dydaktyczna

Andrzej **Sokołowski**

W poniższym artykule podjęta jest próba głębszego zrozumienia właściwości obwodów elektrycznych poprzez zastosowanie interpretacji funkcji algebraicznych i wirtualnego eksperymentu. Lekcja ta jest on niejako kontynuacją dydaktycznej innowacji zapoczątkowanej wcześniej na łamach „Fizyki w Szkole” [1]. Jako fizyczny kontekst tej jednostki lekcyjnej, zostało użyte równoległe połączenie oporników. Idea ta wychodzi naprzeciw wnioskowi z badań dydaktyki fizyki, które zachęcają do częstszego i bardziej wnikliwego stosowania aparatu matematycznego w nauczaniu praw przyrody [2, 3]. Aby rozszerzyć znaczenia aparatu matematycznego, zaproponowane zostało przekształcenie znanych wzorów fizycznych na funkcje matematyczne i wykorzystanie atrybutów otrzymanych funkcji do poznania zachowania się elektrodów w połączeniach równoległych.

Tak przygotowany aparat matematyczny pozwala na wyciągnięcie wniosków, które by nie były możliwe do sformułowania korzystając ze wzorów. Jako środek eksperymentalnej weryfikacji wniosków z analizy matematycznej, proponuję wykorzystać symulacje fizyczne⁵, które są dostępne również w języku polskim. Lekcja ta może być wykorzystana w klasie gimnazjalnej lub licealnej przy założeniu, że uczniowie posiadają podstawową wiedzę o szkicowaniu i analizie funkcji wymiennych. Rys. 1 ilustruje ogólny schemat dydaktyczny wykorzystany przy opracowaniu treści tej lekcji.

Przed realizacją tej lekcji zalecane jest by uczniowie byli zapoznani z ogólnymi własnościami obwodów elektrycznych. Lekcja ta może pełnić rolę tematu rozszerzającego i jednoczenia uogólniającego matematyczne metody poznawania zjawisk przyrody. Sugerowane jest też by nauczyciel przed prezentowaniem tej lekcji, zbudował i zapoznał się jak pracuje wirtualny obwód elektryczny (według rys. 2) i wiedział np. jak użyć wirtualny amperomierz czy włączyć wyłącznik. Te umiejętności pozwolą na efektywne zastosowanie tych symulacji i skupieniu uwagi uczniów na weryfikacji ich hipotez.

Lekcja jest zbudowana w formie otwartego problemu. Uczniowie mają podany problem i próbują najpierw sami

odpowiedzieć na postawione pytania. Po sformułowaniu własnych odpowiedzi (poprawnych lub nie) nauczyciel sugeruje rozwiązanie tego problemu, demonstruje symulowany eksperyment i daje uczniom możliwość weryfikacji ich rozwiązań.

Cele lekcji

- I *Analiza obwodu elektrycznego z dwoma opornikami połączonymi równolegle, z których jeden posiada zmienną oporność, która zmienia się od wartości zerowej do nieskończenie dużej.*
- I *Zastosowanie wykresów i granic funkcji do głębszego zrozumienia oporności zastępczej i natężenia prądu w obwodzie równoległym.*

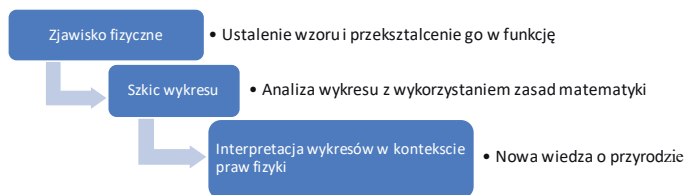
Przebieg lekcji

Rys. 2 przedstawia wirtualny obwód elektryczny. Instruktor za pomocą projektora szkolnego, demonstruje go uczniom i omawia krótko jego elementy. Celem przedstawienia wirtualnej (dynamicznej) reprezentacji tego obwodu, na początku lekcji, jest ułatwienie uczniom zrozumienia istoty problemu poprzez zobrazowanie jego opisu. Wstępnie nauczyciel nie będzie demonstrować jak ten obwód funkcjonuje. Symulacja będzie wykorzystywana bardziej szczegółowo w późniejszym etapie lekcji, aby zweryfikować hipotezy uczniów i wnioski osiągnięte z analizy sformułowanych funkcji algebraicznych. Jest konieczne by nauczyciel wyjaśnił, że przełącznik znajdujący się w wewnętrznym równoległym rozgałęzieniu obwodu (oznaczony jako R_2) służy jako opornik o zmiennym oporze; gdy jest otwarty, jego opór elektryczny będzie traktowany jako nieskończenie duży; gdy jest zamknięty jego opór elektryczny będzie traktowany jako zerowy (zobacz rys. 5 i rys. 6). Niestety symulacja ta nie posiada opcji użycia typowego opornika (podobnego do R_2) z podobnymi własnościami.

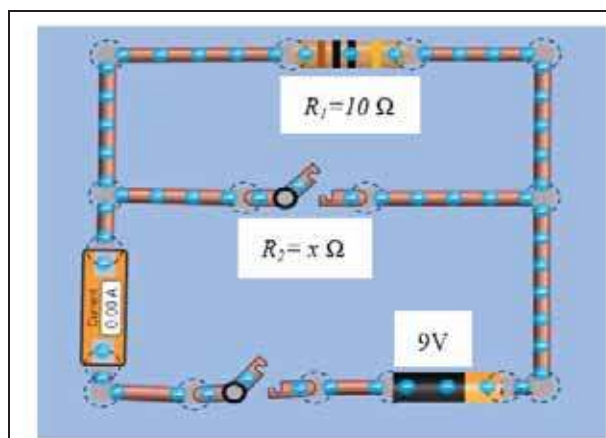
Problem

Podany jest obwód elektryczny z dwoma opornikami połączonymi równolegle. Opor elektryczny jednego z oporników oznaczony R_1 jest stały i wynosi $10\ \Omega$. Opor elektryczny drugiego opornika, R_2 , jest zmienny i może on przyjąć wartości od $0\ \Omega$ lub nieskończenie dużej wartości. Potencjał elektryczny baterii jest równy $9V$.

Na sformułowanie odpowiedzi dajemy uczniom do dyspozycji 5-10 minut. Nauczyciel może spytać uczniów o ich odpowiedzi jednak nie koryguje ich na tym etapie. Po krótkiej dyskusji nauczyciel wyjaśnia, że jedna z metod znalezienia odpowiedzi na te pytania jest przekształcenie wzorów na oporność zastępczą (lub natężenie prądu) i próba oszacowania wartości tych funkcji dla podanych wartości oporności R_2 .



Rys. 1. Proponowany proces integracji matematycznych struktur w nauczaniu fizyki [4]

Rys. 2. Obwód elektryczny. Źródło: <https://phet.colorado.edu>⁵.

Faza przekształcenia wzoru $\frac{1}{R_T} = \frac{1}{R_1} + \frac{1}{R_2}$
w funkcję i jej analiza.

Zacznijmy od odpowiedzi na pytanie 1. Uczniowie wiedzą, że

$$\frac{1}{R_T} = \frac{1}{R_1} + \frac{1}{R_2}.$$

Rozwiązując to równanie na R_T otrzymujemy

$$R_T = \frac{R_1 R_2}{R_1 + R_2}.$$

Podstawiając $R_1 = 10 \Omega$ i wyrażając przy pomocy zmiennej x oporność R_2 otrzymujemy następującą postać funkcji

$$R_T(x) = \frac{10x}{x+10}. \quad (2)$$

Funkcja (2) spełnia warunki funkcji wymiernej, która może być naszkicowana (zobacz rys. 3).

Ponieważ opór zmienny nie może być ujemny, dziedziną funkcji $R_T(x)$ jest $x \geq 0 \Omega$, która w konsekwencji zawęża analizę wykresu $R_T(x)$ do pierwszej ćwiartki układu współrzędnych. Poziome i pionowe przerywane czerwone linie reprezentują poziomą i pionową asymptotę tej funkcji i będą one wykorzystane do interpretowania odpowiedzi w kontekście postawionego problemu. Wykres (Rys. 3) posłuży do zweryfikowania odpowiedzi uczniów na pierwszy problem.

1a) Wartość funkcji $x = 0 \Omega$, dla jest zerowa, tak więc oporność zastępcza tego obwodu w tym przypadku jest zerowa $R_T(0) = 0 \Omega$.

1b) Aby znaleźć odpowiedź na to pytanie, musimy znaleźć lub oszacować wartość tej funkcji, kiedy oporność opornika R_2 jest nieskończenie duża. Wielkość nieskończenie duża jest często oznaczona przez symbol ∞ . Bezpośrednie podstawienie wartości ∞ , za x w $R_T(x)$ nie jest poprawne matematycznie. By poprawnie oszacować tę oporność musimy skorzystać z pojęcia granicy funkcji i policzyć tę granicę, kiedy $x \rightarrow \infty$. Tak więc, jeśli $x \rightarrow \infty$;

$$R_T(x \rightarrow \infty) = \lim_{x \rightarrow \infty} \frac{10x}{x+10} = \lim_{x \rightarrow \infty} \frac{10x}{x} = 10 \Omega$$

1. Jaka jest oporność zastępcza tego obwodu, jeśli opornik R_2 posiada:

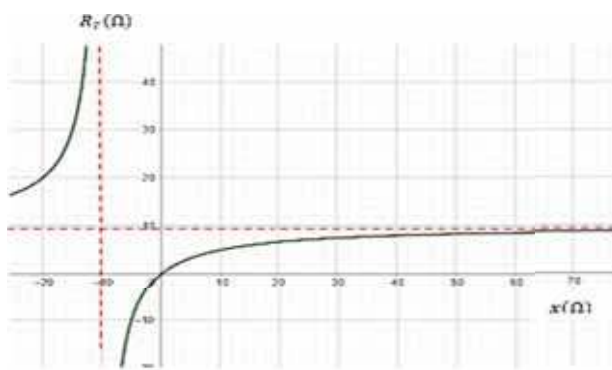
- a) Zerową oporność?
- b) Nieskończenie dużą oporność?

Policz te oporności lub sformułuj hipotezy, które będą reprezentować twoje przypuszczenia.

2. Jakie jest natężenie prądu w tym obwodzie, jeśli R_2 posiada:

- a) Zerową oporność?
- b) Nieskończenie dużą oporność?

Policz te natężenia lub sformułuj hipotezy, które będą reprezentować twoje przypuszczenia.

Rys. 3. Wykres funkcji $R_T(x)$.

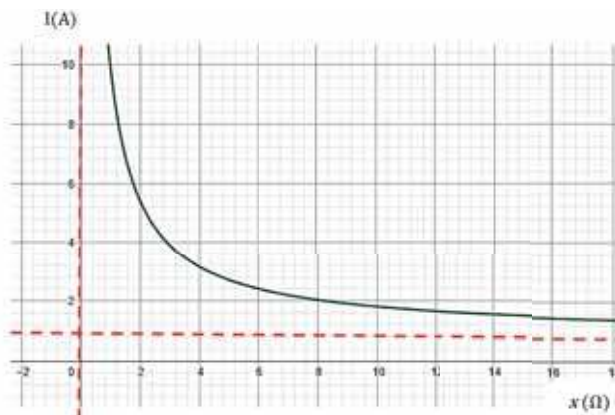
Tak więc jeśli $x \rightarrow \infty$; oporność zastępcza tego obwodu wynosi 10Ω . Warto dodać, że wartość ta odpowiada wartości równania horyzontalnej asymptoty tej funkcji (zaznaczonej czerwoną przerywaną linią na rys. 3).

Więcej o właściwościach tej asymptoty: oporność tego obwodu nigdy nie osiągnie wartości 10 ponieważ w istocie nie wiemy jak duża może być oporność x . Wartość granicy $R_T(x \rightarrow \infty)$ jest więc wielkością szacunkową. Uczniowie zwykle (mylnie) przewidują, że ta oporność będzie nieskończenie duża dla dużych wartości x . Interpretacja oporności zastępczej tego obwodu dla dużej wartości x , może być również umotywowana zachowaniem elektronów (prądu) w tym obwodzie. W przypadku dużej oporności opornika R_2 elektrony wybiorą inną bardziej im wygodną drogę, która stwarza mniejszy opór dla ich ruchu i która w tym wypadku jest częścią obwodu z opornością 10Ω . Nauczyciel może dodać, że jeśli oporność x będzie przybierać wszystkie wartości od zera do nieskończoności, to oporność zastępcza tego obwodu może być bezpośrednio odczytana z wykresu (Rys. 3).

Faza przekształcenia wzoru $I = \frac{V}{R_T}$
w funkcję i jej analiza

Uczniowie wiedzą, że natężenie prądu elektrycznego zależy od zastępczej oporności obwodu i napięcia źródła prądu zgodnie ze wzorem

$$I = \frac{V}{R_T} = \frac{9V}{R_T}.$$



Rys. 4. Wykres natężenia prądu w funkcji oporności x .

Z wcześniejszych rozważań wiemy że jeśli $x = 0 \Omega$, $R_T = 0 \Omega$. Po podstawieniu, mamy więc przypadek dzielenia przez zero, które to działanie nie jest dozwolone w matematyce. Jak więc oszacować prąd w tym przypadku? Lub jak określić działanie tego obwodu elektrycznego w tym przypadku?

Postępując podobnie, jak przy znalezieniu oporności zastępczej jedną z możliwości jest przekształcenie wzoru

$$I = \frac{V}{R_T}$$

na funkcję i oszacowanie wartości tej funkcji dla $x = 0 \Omega$. Można ten proces przeprowadzić na kilka sposobów. Jednym z nich jest znalezienie funkcji złożonej $I(R_T(x))$. Podstawiając (rów. 2) do

$$I = \frac{V}{R_T},$$

wzór ten przyjmuje następującą formę

$$I(R_T(x)) = I(x) = \frac{9}{\frac{10x}{x+10}} = \frac{9x+90}{10x}. \quad (3)$$

Funkcja ta reprezentuje również funkcję wymierną, której wykres przedstawia rys. 4.

Wykres tej funkcji i jej algebraiczna postać pozwala na znalezienie odpowiedzi na postawione pytania.

2a) Dla $x = 0 \Omega$ funkcja ta posiada pionową asymptotę.

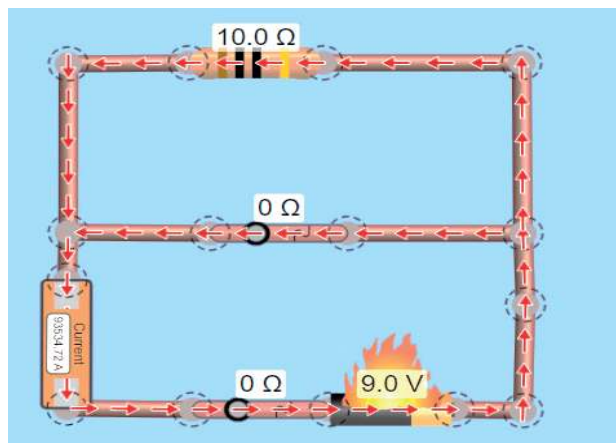
By oszacować wartość natężenia prądu możemy zastosować technikę znalezienia granicy prawostronnej funkcji $I(x)$. Tak więc

$$I(x \rightarrow 0^+) = \lim_{x \rightarrow 0^+} I(x) = \lim_{x \rightarrow 0^+} \frac{9(0.001)+90}{10(0.001)} = \infty \text{ A}.$$

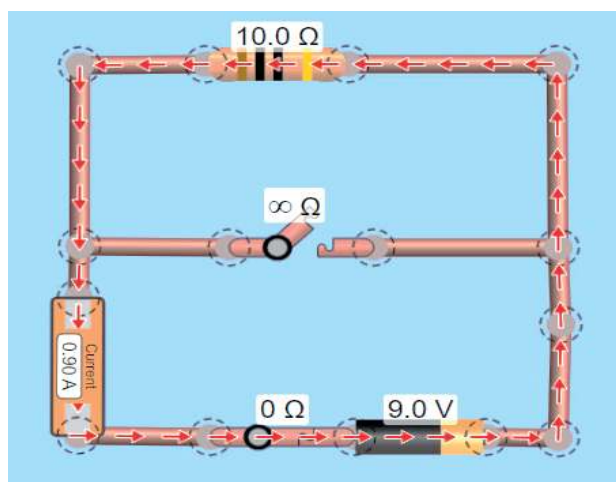
Wynik ten pokazuje, że wartość prądu osiąga wartość nieskończenie dużą dla oporności $x \rightarrow 0 \Omega$. Wynik ten ma również swoje uzasadnienie w zachowaniu elektronów. Liczba przepływających w obwodzie elektronów będzie nieskończenie duża, jeśli znajdą one drogę, która nie stawia im żadnego oporu.

2b) Odpowiedź na to pytanie możemy znaleźć poprzez oszacowanie granicy funkcji $I(x)$ kiedy oporność $x \rightarrow \infty \Omega$. Krzystając z (3) otrzymujemy

$$I(x \rightarrow \infty) = \lim_{x \rightarrow \infty} \frac{9x+90}{10x} = \lim_{x \rightarrow \infty} \frac{9x}{10x} = \frac{9}{10} \text{ A},$$



Rys. 5. Źródło: <https://phet.colorado.edu>.



Rys. 6. Źródło: <https://phet.colorado.edu>.

co oznacza, że natężenie prądu nie przekroczy 0.9 A nawet jeśli oporność x będzie bardzo duża. W tym przypadku prawie wszystkie elektrony przepłyną przez opornik R_1 .

Zastosowanie symulacji do zweryfikowania wniosków z operacji matematycznych

Wyniki operacji matematycznych, zwykle przekonujące dla uczniów oraz wykorzystanie symulacji stawia niejako kropkę nad i . Jakkolwiek oporność elektryczna jest właściwością obwodu elektrycznego i może być ona tylko policzona, natężenie prądu może być pomierzone i zobrażowane.

Weryfikacja zachowania obwodu, kiedy $x = 0 \Omega$

Nauczyciel demonstruje wcześniej zbudowany obwód (rys. 5). Poprzez włączenie przełącznika, który symbolizuje opornik R_2 , oporność tej części obwodu jest zera. Włączenie głównego przełącznika powoduje zwarcie, które jest spowodowane dużym natężeniem prądu (Przypomnijmy, że $I(x \rightarrow 0^+) = \infty \text{ A}$). Ten przypadek jest niebezpieczny, ponieważ powoduje przegrzanie baterii i w rezultacie jej samozapalanie. Uczniowie zauważą, że w pracowni fizycznej takiego doświadczenia nie przeprowadziliby.

Weryfikacja zachowania obwodu w przypadku, kiedy $x \rightarrow \infty \Omega$.

Duża oporność opornika R_2 jest przedstawiona za pomocą otwartego wyłącznika. W tym przypadku elektrony nie będą przepływać przez tą część obwodu a wybiorą drogę, która stawia im mniejszy opór, a więc będą przepływać przez R_1 . Amperomierz pokazuje natężenie prądu rzędu 0.9 A co pokrywa się z wcześniej znalezionej wartości horyzontalnej asymptoty funkcji natężenia prądu

$$(I(x \rightarrow \infty)) = \lim_{x \rightarrow \infty} \frac{9x + 90}{10x} = \lim_{x \rightarrow \infty} \frac{9x}{10x} = \frac{9}{10} A).$$

Uwagi końcowe

Uczniowie, którzy posiadają pewne przygotowanie matematyczne są bardzo zaintrygowani możliwościami zastosowania narzędzi tego przedmiotu na fizyce. Szczególnie ciekawa jest dla nich dyskusja o znaczeniu asymptot, których interpretacja w kontekście jest raczej rzadkością na lekcjach matematyki. Inną ciekawą obserwacją jest ta, że wszystkie trzy modele: algebraiczny, graficzny i wirtualny prowadzą do tych samych wniosków.

Sumując, jakkolwiek wzory fizyczne pozwalają policzyć szczególne wartości wielkości fizycznych, funkcje

algebraiczne wzbogacają i pogłębiają tą analizę i dają możliwości dla zastosowania bardziej wnikliwej interpretacji. W efekcie może to rozbudzić naukową wyobraźnię uczniów dając im bodźce do rozwijania własnego matematyczno-przyrodniczego zainteresowania związkami pomiędzy zjawiskami przyrodniczymi.

Zaproponowana lekcja jest przykładem na zastosowanie funkcji w dziale prąd elektryczny. W konsekwencji ciekawym wydawałoby się wprowadzenie na szerszą skalę takiego pojmowania fizyki w praktyce szkolnej.

Dr Andrzej Sokółowski

Lone Star College, Houston, USA

LITERATURA

- [1] Sokółowski, A., Modelowanie Zasad Newtona z Zastosowaniem Granic Funkcji-Propozycja Dydaktyczna, *Fizyka w Szkole*, 3249, LXII, indeks 35810X, Nr 2 marzec/kwiecień, 2017.
- [2] Hestenes D., „Oersted Medal Lecture 2002: Reforming the mathematical language of physics”, (2003) pp. 104-121.
- [3] Sherin L. B., „How students understand physics equations”, *Cognition and Instruction* 19, 4 (2001) pp. 479-541.
- [4] Sokółowski, A., *Scientific Inquiry in Mathematics-Theory and Practice: A STEM Perspective* (Springer, New York, 2018).
- [5] Physics Simulations, <https://phet.colorado.edu/en/simulation/circuit-construction-kit-dc-virtual-lab>, udostępnione 5 lipca, 2018.

Co w fizyce piszczy

Stabilne pary mikroobiektów

Oddziaływania elektrostatyczne mogą sprawić, że mikroobiekty w cieczy zaczną tworzyć stabilne pary – pokazały teoretyczne analizy w IPPT PAN. Jeśli to nowe zjawisko uda się potwierdzić w doświadczeniach, być może przyda się ono np. przy oczyszczaniu wody czy produkcji leków.

Oddziaływania elektrostatyczne to te same oddziaływania, które sprawiają, że włosy elektryzują się nam przy zdejmowaniu czapki. Wywołują też to „kopnięcie”, które odczuwamy, kiedy czasem po przejściu przez wykładzinę dotykamy klamki. Nie mówiąc już o piorunach podczas burzy... Elektrostatyka to nasza stara znajoma, która wydawałoby się, że już od dawna nie ma przed nami tajemnic. A jednak.

Teoretyczne analizy zespołu prof. Marii Ekiel-Jezewskiej przeprowadzone w Pracowni Fizyki Płynów Złożonych Instytutu Podstawowych Problemów Techniki PAN pokazały, że mogą istnieć konfiguracje mikroobiektów, które będą ze sobą stabilnie połączone wcale nie wiązaniami chemicznymi, ale właśnie siłami elektrostatyki. Pod warunkiem, że ma to miejsce w cieczy, a mikroobiekty te opadają pod wpływem grawitacji.

To zaskakująca wiadomość! „Zajmuję się fizyką układów w skali mikro od wielu lat, ale dotąd nie słyszałam o istnieniu mikroobiektów, które tworzyłyby takie stabilne pary” – mówi PAP prof. Maria Ekiel-Jezewska.

„Obiektami o wielkości mikrometrów są np. bakterie, glony czy krwinki czerwone” – mówi prof. Maria Ekiel-Jezewska. Okazuje się, że w środowisku wodnym można dobrać do takich opadających grawitacyjnie obiektów inne

mikroobiekty tak, że połączą się one w pary siłami elektrostatycznymi. I pary te będą funkcjonować tak, jakby tworzyły nową całość. Kiedy jeden z obiektów spróbujemy odciągnąć od drugiego, on spontanicznie będzie wracał na swoje miejsce.

Kto wie, może siły elektrostatyki – kiedy już takie stabilne pary zaobserwujemy w eksperymentach – uda się kiedyś zaprząć do produkcji nowoczesnych nośników leków. Nowe zjawisko, o którym pisze zespół z IPPT PAN może się też przydać w produkcji nowoczesnych sposobów na oczyszczanie wody. Być może wyniki tych badań przydadzą się w rozwoju mikrofluidyki, układów Lab-on-Chip, medycznej diagnostyki czy ochrony środowiska, lub w jakiś inny sposób, którego dziś nie sposób przewidzieć. Wynikami tej pracy mogą być też zainteresowani badacze pracujący przy eksperymentach z bakteriami czy glonami – wymienia autorka.

Wyniki pracy prof. Ekiel-Jezewskiej i jej doktoranta, Chrisa Trombleya amerykańskiego matematyka i fizyka, ukazały się w grudniu w prestiżowym czasopiśmie „Physical Review Letters”. To badania teoretyczne. Teraz czas dla eksperymentatorów, aby takie układy mikroobiektów znaleźć w doświadczeniach.

PAP – Nauka w Polsce, Ludwika Tomala



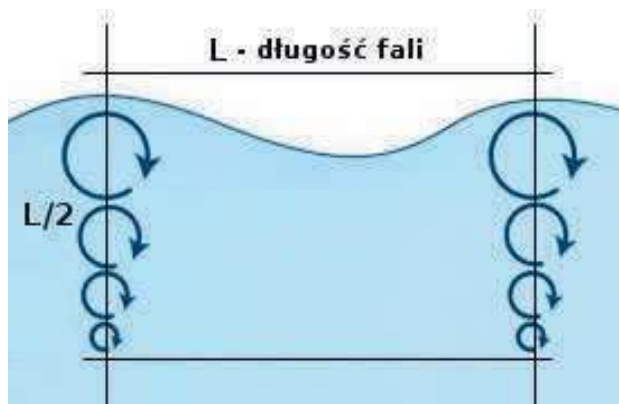
Wilhelm Eduard Weber

(1804-1891)

Tadeusz Wibig

Niemiecki fizyk XIX stulecia znany jest dziś praktycznie tylko jako wynalazca telegrafu, co jest w zasadzie niesłuszne i niesprawiedliwe. Na początku swej kariery wspólnie ze starszym bratem, anatomem, Ernstem Heinrichem, zajmował się mechaniką płynów. W książce, którą wspólnie napisali, poza fizyką przepływu cieczy przez elastyczne rurki, opisali także mechanizm falowania wody, w którym cząsteczki opisują dość złożone kołowe ruchy w płaszczyźnie pionowej równoległej do ruchu fal.

Wilhelm wspólnie ze swoim drugim, młodszym bratem, Eduardem Friedrichem, też naukowcem, anatomem i psychologiem, zajmował się krótko mechanizmem chodzenia i na ten temat też opublikowali książkę.



Doktorat, a potem i habilitację obronił z teorii organów (chodzi tu o instrument muzyczny, nie anatomicznego tym razem). Wykład Webera o organach przedstawiony na zjeździe fizyków niemieckich w Berlinie w roku 1828 zainteresował samego Gaussa, co zaowocowało przeprowadzką do Getyngi, gdzie za sprawą Gaussa dostał profesurę.

Obaj naukowcy rozpoczęli owocną współpracę w dziedzinie elektryczności i magnetyzmu. Wynikiem tej pracy było stworzenie teorii, która miała opisywać wszystkie zjawiska elektryczne i magnetyczne używając do tego celu pojęcia siły typu newtonowskiego (trzy prawa dynamiki itd.) działającej momentalnie na dowolne odległości. Dodali oni do siły opisującej oddziaływanie nieruchomych ładunków (czyli prawo Culomba) oddziaływanie ładunków poruszających się ruchem jednostajnym, czyli stałych prądów elektrycznych (prawo Ampera) i jeszcze oddziaływanie ładunków przyspieszanych (jakby prawo Faradaya).

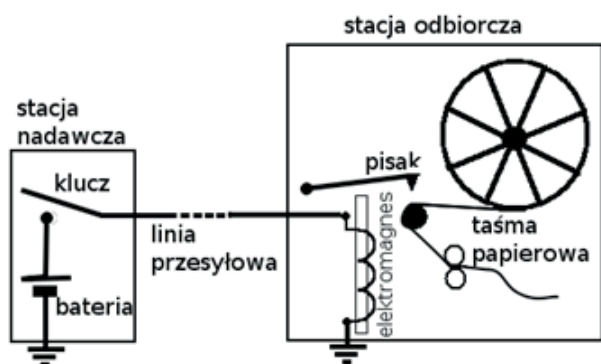
$$F = \frac{ee'}{r^2} \left(1 - \frac{1}{c^2} \frac{dr^2}{dt^2} + \frac{2}{c^2} r \frac{d^2r}{dt^2} \right)$$

Siła Webera wyrażała się wzorem, którego dziś już nikt nie pamięta, bo właściwie nie ma powodów, żeby go pamiętać. Ta bardzo elegancka i matematycznie prosta w sumie teoria okazała się błędna. Nie wytrzymała weryfikacji z rzeczywistością. Dziś oddziaływania elektromagnetyczne opisywane są znacznie bardziej skomplikowanymi, niektórzy mówią, że tylko pozornie, równaniami Maxwella. Warto jednak zauważyć, że we wzorze Wenera pojawia się litera c oznaczająca prędkość światła. To Weber po raz pierwszy użył tego symbolu, czyli, w pewnym sensie, umożliwiając niejako Einsteinowi napisanie słynnego $E=mc^2$.

W międzyczasie Gauss z Weberem wymyślili telegraf, choć nie tak do końca oni właśnie byli pierwszymi, którzy na to wpadli. Inni przed nimi, już od początku XIX wieku, próbowali użyć prądu do przesyłania wiadomości. Zbudowano nawet w skali laboratoryjnej kilka działających modeli. Dopiero jednak Gauss i Weber w roku 1833 okazali się na tyle konsekwentni, że korzystając z doświadczeń poprzedników zbudowali nadajnik i odbiornik i poprowadzili pierwszą praktycznie działającą linię telegraficzną o długości ponad kilometra. Wiodła ona z domu Webera do jego laboratorium.

Zasada działania była prosta: na jednym końcu linii było źródło prądu i przełącznik, którym można było zmieniać kierunek prądu, a po drugiej stronie bardzo czuły galwanometr – urządzenie do pomiaru przepływającego prądu, który reagował wychyleniem wskazówki w jedną (0), lub w drugą stronę (1) w zależności od tego, w którą stronę prąd płynął. Najważniejsze w wynalazku Gaussa i Webera była nie sama maszyna, a idea binarnego kodowania przekazu. Dziś jest to podstawa całej naszej cyfrowo-komputerowej cywilizacji informatycznej.

Telegraf widzimy dziś głównie w westernach. Ciągące się po horyzont rzędy telegraficznych słupów z rozpiętych na nich drutem (jednym!) łączą miasteczka pełne niedobrych rewolwerowców i dzielnych szeryfów przekazując bardzo ważne informacje w okamgnieniu. W kabinach telegrafistów znajdują się urządzenie odbiorcze z elektromagnesem przyciągającym pisak do taśmy papierowej rozwijającej się z wielkiego, wolno obracającego się koła. Rysuje on na niej kreski i kropki, które potem pracownik poczty rozszyfrowuje zgodnie ze schematem, jaki wymyślił niejaki Samuel Morse.



Doświadczenie domowe:

Telegraf – transmisja sygnałów na odległość (po jednym drucie!)

A. Potrzebne materiały:

1. Dwie diody LED świecące w różnych kolorach (albo jedna dwukolorowa, np. OSRG3131).
2. Przewód w izolacji tak długi, jak długą chcemy mieć linię telegraficzną.
3. Bateria 9V.
4. Przełącznik hebelkowy, dźwigniowy, podwójny, chwilowy (monostabilny), dwubiegunowy: ON – OFF – ON (np. KN3B-223).
5. Dwie drewniane deseczki o rozmiarach jakieś 10 cm x 10 cm na stację odbiorczą i nadawczą.
6. Taśma klejąca do mocowania, choć lepszy byłby na przykład klej na gorąco.
7. Dwa narzędzia ogrodnicze dające się wbić w ziemię (szpadle, łopaty itp.), ewentualnie dwa długie noże kuchenne, albo wielkie gwoździe itp.

1. B. Narzędzia – nożyczki lub inny przyrząd do cięcia drutu, nóż, wkrętak, lutownica, młotek.

2. C. Kolejność czynności:

1. Zlutować diody zgodnie ze schematem (no chyba, że ma się gotową podwójną).
2. Do jednej końcówki przylutować przewód transmisyjny, a do drugiej kawałek drutu, który da się elektrycznie połączyć z **łopatką uziemiającą**.
3. Zamocować diody do deseczki stacji odbiorczej.
4. Drugi koniec drutu transmisyjnego przylutować do przełącznika na stacji nadawczej wg schematu.
5. Przylutować do przełącznika krótki przewód do podłączenia szpadla uziemiającego w stacji nadawczej.

Schemat wszystkich połączeń przedstawiony jest na rysunku:



Telegrafista może też naciskając z wprawą na klucz telegraficzny wysłać wiadomość do innego telegrafisty podłączonego do tej samej linii drutów wiele mil na zachód (lub wschód, albo w innym, w zasadzie dowolnym kierunku).

W uznaniu zasług Webera, choć może nie chodziło tu akurat o telegraf, jego nazwiskiem postanowiono nazywać w obowiązującym nas układzie SI jednostkę strumienia indukcji magnetycznej. Jeśli pole magnetyczne o indukcji 1 Tesli (T) przecina prostopadle 1 m² powierzchni, to jest to właśnie strumień indukcji równy 1 Weberowi (W). Nie jest to jednostka często używana w życiu codziennym.

6. Przylutować do przełącznika baterijkę i w odpowiedni sposób połączyć wszystkie końcówki wychodzące z przełącznika (zgodnie ze schematem).
7. Zmontować elementy stacji nadawczej na deseczce.
8. Umieścić stację nadawczą i odbiorczą we właściwych miejscach – w ogródku.
9. Wbić w ziemię łopaty (szpadle, noże, gwoździe itp.) i połączyć je do odpowiednich kabelków w obu stacjach.

I telegraf gotowy jest już do pracy!

Przy włączeniu przełącznika w jedno z ustawień świeci się dioda jednego koloru, przy drugim położeniu – drugiego. Można umówić się, który kolor oznacza kreskę, a który kropkę i wystarczy teraz nauczyć się alfabetu Morse'a, aby można było zapomnieć o comiesięcznych doładowaniach telefonu, albo o innych opłatach za przesyłanie krótkich wiadomości na nieduże odległości.

Modyfikacje eksperymentu:

- A) Gdyby ktoś nie dysponował przełącznikiem, o jakim mowa powyżej, może poradzić sobie zupełnie dobrze i bez niego. Służy on przecież tylko przełączaniu biegunów baterii. Można to robić „ręcznie” dotykając do biegunów drutami uziemia do plusa i transmisyjnym do minusa, lub odwrotnie.
- B) Gdyby ktoś nie chciał używać diod (za czasów Webera i Gaussa nie było ich przecież!) może użyć, jak oni galwanometru własnej konstrukcji. Robi się go z kompasu, który owija się wielokrotnie drutem na osi N-S, tworząc elektromagnes mogący wychylić igłę raz w jedną (W), raz w drugą (E) stronę. O orientacji stacji odbiorczej należy pamiętać podczas pracy!



Zasady i schematy stosowane podczas rozwiązywania zadań fizycznych

Czesław Surowiec

Ucząc rozwiązywania zadań fizycznych w szkole podstawowej należy przywiązywać dużą wagę do przestrzegania pewnych zasad, które niekiedy można ująć w postaci schematu typowego dla danego rodzaju zadań. Korzyści wynikające ze stosowania zasad i schematów podczas uczenia rozwiązywania zadań to:

- Zachęcenie ucznia do podjęcia próby rozwiązania zadań przez zapoznanie z zasadami jakimi powinien się kierować podczas rozwiązywania i czynności, które powinien wykonywać.
- Nauczycielowi ułatwia ocenianie rozwiązania przez przypisanie odpowiedniej wagi poszczególnym etapom rozwiązania oraz ustalenie, które z nich sprawiają uczniom największe trudności, aby im pomóc w ich pokonywaniu.
- Pozwala na ujednolicenie i doskonalenie umiejętności rozwiązywania zadań począwszy od szkoły podstawowej przez szkołę średnią aż po wyższą uczelnię sprawdzoną praktycznie.
- Ułatwieniem dla ucznia i nauczyciela będzie, jeśli takie schematy i zasady będą w postaci tablicy poglądowej w pracowni fizycznej lub w postaci kserokopii w zeszytach w wersji pełnej lub skróconej zawierające tylko tekst podkreślony w schematach. Proponuję schematy rozwiązań zadań jakościowych i obliczeniowych najczęściej rozwiązywanych na lekcjach fizyki. Dla zadań doświadczalnych i graficznych proponuję zasady jakimi należy się kierować przy ich rozwiązywaniu, chociaż możliwe jest zastosowanie podobnego schematu do zadań doświadczalnych uwzględniając różnice między rodzajami tych zadań. Zainteresowanym podjęcia takiej próby polecam wykorzystanie niektórych elementów tych schematów i dostosowania ich do wymagań zadań doświadczalnych, oraz wykorzystania omówionych zasad.

Schemat rozwiązywania zadań jakościowych.

- Zapoznanie z warunkami zadania
 - uważne przeczytanie tekstu zadania,
 - wyjaśnienie znaczenia słów budzących wątpliwości lub nieznanych oraz dotyczących przyrządów, urządzeń itp.
 - powtórne przeczytanie tekstu z pełnym lub skróconym zapisem warunków,
 - sformułowanie głównego pytania zadania.
- Analiza treści zadania.
 - badanie wyjściowych danych,
 - wyjaśnienie sensu fizycznego zadania przez odpowiedź na pytanie „o jakich zjawiskach, faktach, właściwościach ciał, stanach układu itp.” mowa w danych wyjściowych oraz „jakie związki występują między nimi”,
 - dokładnie należy oglądać i analizować: rysunki, grafiki, zdjęcia, schematy itp. zawarte w zadaniu lub wykonane w trakcie jego rozwiązywania,

d) uściślić dodatkowe warunki dla otrzymania jednoznacznej odpowiedzi.

- Sporządzenie planu rozwiązania.
Opracowanie analitycznego ciągu wnioskowania myślowego zaczynającego się od pytania postawionego w zadaniu a kończącego się na warunkach opisanych w zadaniu lub wykonanym doświadczeniu sformułowania praw i zdefiniowania wielkości fizycznych.
- Zrealizowanie planu rozwiązania.
Zbudowanie syntetycznego ciągu wniosków wynikających ze sformułowań, odpowiednich praw, definicji wielkości fizycznych, opisanych właściwości, stanów ciała a kończących się na odpowiedzi na pytanie.
- Sprawdzenie odpowiedzi, które można zrealizować przez:
 - wykonanie odpowiedniego doświadczenia,
 - rozwiązanie tego zadania innym sposobem,
 - zgodność otrzymanej odpowiedzi z punktu widzenia jej sensu fizycznego i realności takiego rozwiązania.

Schemat rozwiązywania zadań obliczeniowych

- Zapoznanie się z treścią zadania zwracając uwagę na:
 - uważne przeczytaniu tekstu,
 - wyjaśnienie znaczenia słów występujących w nim, jeśli mogą się pojawić wątpliwości co do właściwego rozumienia ich znaczenia.
- Analiza treści zadania powinna nam dać odpowiedzi na pytania:
 - jakiego zjawiska dotyczy zadanie,
 - w jakich warunkach zachodzi zjawisko i jakie warunki powinniśmy uwzględnić, a jakie pominąć,
 - jakie stałe fizyczne, dane tablicowe itp. możemy wykorzystać w rozwiązywaniu zadania.
- Wypisanie danych, stałych fizycznych, danych tablicowych niezbędnych do rozwiązania zadania:
 - wypisując dane powinniśmy je wyrazić w podstawowych jednostkach układu SI wypisując miano (wymiar) również w tym układzie,
 - jeśli w treści zadania nie są podane oznaczenia literowe wielkości danych stosujemy oznaczenia literowe ogólnie przyjęte w układzie SI,
 - wartości liczbowe małe i duże zapisujemy w postaci $a \cdot 10^n$ co znacznie ułatwia obliczenia,
 - aby uniknąć wielokrotnego czytania treści zadania, zwłaszcza jeśli czas rozwiązywania zadań jest ograniczony (sprawdzian, egzamin) zapisujemy w sposób umowny (objaśniając przyjętą umowę) warunki w jakich zachodzi zjawisko.
- Wykonanie rysunku, schematu itp. ma sens, jeśli to ułatwia nam rozwiązanie zadania:
 - na rysunku wielkości dane oznaczamy za pomocą poprzednio przyjętych oznaczeń literowych pomijając ich wartości liczbowe,
 - właściwie wykonany rysunek niejednokrotnie bardzo ułatwia rozwiązanie, sprowadzając je do zapisu sytu-

acji przedstawionej na rysunku za pomocą zależności między wielkościami przedstawionymi na nim.

5. Wypisanie zależności (praw, zasad itp.) między wielkościami występującymi w zadaniu:
 - a) na podstawie analizy treści zadania, rysunku ustalamy jakie zależności stosują się do opisanego zjawiska i wypisujemy je,
 - b) jeśli zapisane zależności dotyczą wielkości wektorowych, zastępujemy je równoważnymi zapisami skalarowymi.
6. Wykonanie działań matematycznych polegających na:
 - a) doprowadzenie zapisanej pojedynczej zależności do takiej postaci, aby po jednej stronie była wielkość szukana a po drugiej dane,
 - b) jeśli mamy do czynienia z kilkoma zależnościami rozwiązujemy układ równań jakie tworzą zależności, wyrażając wielkości szukane poprzez dane, stałe fizyczne i tablicowe,
 - c) w przypadku trzech lub więcej zależności polecam metodę wyznaczników.
7. Sprawdzenie wyniku rozwiązania działania na mianach:
 - a) pozwala uniknąć zbędnych obliczeń w przypadku błędu popełnionego przy przekształcaniu zależności,
 - b) jest formą sprawdzenia szczególnie jeśli nie mamy danych liczbowych,
 - c) jeśli w rozwiązaniu zastosujemy niewłaściwe zależności zgodność wymiarów obydwu stron zależności nie potwierdza właściwego rozwiązania.
8. Analiza otrzymanego rozwiązania powinna dać odpowiedzi na następujące pytania:
 - a) kiedy rozwiązanie ma sens fizyczny,
 - b) czy rozwiązanie jest jednoznaczne,
 - c) jakie są możliwe przypadki (jeśli nie jest jednoznaczne).
9. Wykonanie obliczeń, jeśli są dane liczbowe.
 - a) obliczenia wykonujemy zazwyczaj kalkulatorem,
 - b) wyniki obliczeń zapisujemy przestrzegając zasad działania na liczbach przybliżonych.
10. Zapisanie wyniku (wyników) rozwiązania w postaci wartości liczbowej i miana w układzie SI.

Jeśli otrzymany wzór końcowy ma prostą postać matematyczną i działania na mianach są nieskomplikowane celowe jest jednocześnie działanie na liczbach i mianach a otrzymanie wyniku liczbowego łącznie z mianem.

Zasady obowiązujące podczas rozwiązywania zadań doświadczalnych.

Podstawowe etapy rozwiązywania zadań doświadczalnych są podobne do rozwiązywania zadań obliczeniowych i jakościowych, jednak posiadają swoją specyfikę:

- a) przygotowując zadanie doświadczalne należy dobrać odpowiednie przyrządy oraz należy je wstępnie sprawdzić,
- b) cechą charakterystyczną tych zadań jest praca jaką należy wykonać nad znalezieniem potrzebnych informacji dotyczących użytych przyrządów, metod pomiaru itp.
- c) w przypadku zadań obliczeniowych można założyć, że zjawisko jest idealne, w zadaniach doświadczalnych nie jest zawsze możliwe i należy liczyć się z wpływem czynników zewnętrznych, które należy wykryć i w miarę możliwości usunąć je lub ich wpływ znacząco zmniejszyć.

Podczas analizy treści zadania i planowania jego rozwiązania należy znaleźć odpowiedzi na pytania:

- a) jakie dane są niezbędne do rozwiązania,
- b) jak je otrzymać wykonując doświadczenie,
- c) w jakich jednostkach je wyrazić,
- d) jakie pomiary wykonać i z jaką dokładnością.

Poprawność rozwiązania zadań doświadczalnych można sprawdzić:

- a) metodą bezpośredniego pomiaru szukanej wielkości za pomocą odpowiednich przyrządów o wymaganej dokładności,
- b) za pomocą innego doświadczenia wyznaczając szukaną wielkość innym sposobem za pomocą innych przyrządów,
- c) porównując otrzymany wynik z danymi tablicowymi lub katalogowymi.

Niektórych rozwiązań zadań doświadczalnych np. wyznaczanie sprawności urządzenia, strat ciepła itp. nie da się sprawdzić za pomocą doświadczenia, należy wówczas omówić wpływ różnych czynników na wynik doświadczenia.

Zasady obowiązujące przy rozwiązywaniu zadań graficznych.

Rozwiązywanie zadań graficznych należy rozpocząć od:

- a) uczenia analizowania rysunków, schematów, ilustracji, zdjęć itp.
- b) odczytywania i wykreślania łatwych wykresów w celu ustalenia jakie zjawiska, właściwości ciał, stan układu itp. one przedstawiają,
- c) odkrywania na podstawie wykresów zależności ilościowych między wielkościami fizycznymi i praw z nich wynikających.

Rozwiązując zadania graficzne należy:

1. Jeśli podany jest wykres zależności między wielkościami fizycznymi należy:
 - a) odpowiedzieć na pytanie jakiego zjawiska dotyczy na każdym odcinku wykresu,
 - b) wyjaśnić sens fizyczny charakterystycznych punktów wykresu,
 - c) korzystając z liniiki z podziałką należy odczytać szukane wielkości tzn. rzędne, odcięte, pole powierzchni między wykresem a osiami współrzędnych itp.
2. Jeśli wykres nie jest dany należy na podstawie treści zadania lub tabelki sporządzić go
 - a) kreślimy osie współrzędnych zaznaczając na nich obrane wielkości i ich miana,
 - b) wybieramy określoną podziałkę tak, aby wielkości dane lub z tabelki mieściły się na osiach,
 - c) nanosimy wielkości dane otrzymując zbiór punktów, na podstawie których sporządzamy wykres,
 - d) prowadzimy linię, dla której otrzymane punkty są położone symetrycznie względem niej (lub leżą na niej),
 - e) na podstawie wykresu ustalamy zależność między wielkościami występującymi na osiach współrzędnych (prawa).

Proponowane schematy i zasady dotyczące rozwiązywania zadań mają charakter ogólny i wymagają dostosowania do zakresu programu fizyki (szkoły: podstawowa, zawodowa, technika, licea). W szkole podstawowej możemy je uprościć sprowadzając do tych, które są istotne do rozwiązania zadania, pomijając te etapy, które na tym poziomie nauczania fizyki nie występują. W szkole średniej propozycja ta pozwala doskonalić umiejętność rozwiązywania zadań i uzyskiwanie przy ocenianiu rozwiązania maksymalnej ilości punktów i najwyższych ocen.

Czesław Surowiec
Dębica

500 rocznica śmierci Leonarda da Vinci

Inżynier renesansu

Kazimierz Mikulski

Leonardo da Vinci, właściwie **Leonardo di ser Piero da Vinci**¹ (1452-1519) – włoski renesansowy malarz, architekt, filozof, muzyk, pisarz, odkrywca, matematyk, mechanik, anatom, wynalazca, geolog, rzeźbiarz.

W opracowaniach biograficznych czytamy, że urodził się i wychował niedaleko Vinci, będąc nieślubnym synem notariusza ser Piera da Vinci i chłopki Cateriny. W rozumieniu współczesnym nie miał nazwiska, człon „da Vinci” oznacza bowiem „z [miasta] Vinci”. Jego pełne imię, nadane mu przy narodzinach, to „Leonardo di ser Piero da Vinci”, czyli „Leonardo, [syn] ser Piera z [miasta] Vinci”.²

Powyższy autoportret został namalowany w 1512 roku za pomocą czerwonej kredy, kiedy Leonardo da Vinci miał 50 lat i mieszkał we Francji. Oryginalny obraz ma wymiary 33,3x 21,3 cm (13 1/8 x 8 3/8 cala). Obecnie znajduje się we wspaniałym zbiorze Biblioteca Reale, Turyn.

Często w literaturze przedmiotu zauważono, że na tym obrazie mistrz Leonardo wygląda na starszego niż jego wiek – mógł mieć około sześćdziesięciu lat, kiedy wykonał ten rysunek – i dlatego niektórzy krytycy wątpili, czy jest to podobieństwo samego siebie. Oprócz tego inne ważne elementy wskazują, że ten portret doskonale pasuje do postaci jaką przypisuje się Leonardo. Przedstawiony na obrazie czcigodny starzec z długą białą brodą, surowe oczy ocienione pod krzaczastymi brwiami, był tradycyjnym typem reprezentującym filozofów, proroków, a także Boga. Oczywiście nikt nie sugerowałby poważnie, że Leonardo świadomie wciągnął się w pozór Wszechmogącego, ale musimy pamiętać jego twierdzenie, że malarz walczy z naturą i że malarstwo jest związane z Bogiem. Ten imponujący mędrzec, który wydaje się pochodzić z innego świata, ma coś z niedefiniowalnego z magów, tego, który poprzez odkrywanie praw wszechświata wie, jak nimi manipulować.⁵

Leonardo często był opisywany jako archetyp „człowieka renesansu”, którego, wydawałoby się, niespożytej ciekawości dorównywała tylko siła jego kreatywno-



Rysunek 1. Leonardo da Vinci (1452-1519) Autoportret (ok. 1510-1515)³ oraz Portret wykonany przez Francesco Melzi⁴

Źródło: https://upload.wikimedia.org/wikipedia/commons/b/ba/Leonardo_self.jpg
https://upload.wikimedia.org/wikipedia/commons/f/f7/Francesco_Melzi_-_Portrait_of_Leonardo_-_WGA14795.jpg

ści. Uważa się go za jednego z największych malarzy wszech czasów i prawdopodobnie najwszechstronniej utalentowaną osobą w historii⁶. To właśnie talent malarski przysporzył Leonardowi największej popularności. Dwie z jego prac, *Mona Lisa* i *Ostatnia Wieczerza*, zajmują czołowe miejsca na listach najslawniejszych, najczęściej imitowanych i wspominanych portretów i dzieł malarstwa. Równie wielkie znaczenie w historii sztuki ma szkic Leonarda *Człowiek witruwiański*, przedstawiony poniżej.

Ten obraz stanowi doskonały przykład zainteresowania Leonardo proporcją. Ponadto stanowi on kamień węgielny prób Leonardo, aby odnieść człowieka do natu-



Rysunek 2. Człowiek witruwiański (wł. *Le proporzioni del corpo umano secondo Vitruvio*, „Proporcje ciała ludzkiego wg Witruwiusza”) – rysunek autorstwa Leonarda da Vinci upowszechniony przez niego około roku 1490. Stworzony ołówkiem i atramentem na papierze, przedstawia figurę nagiego mężczyzny w dwóch nałożonych na siebie pozycjach, wpisanej w okrąg i kwadrat (łac. *Homo Quadratus*, „człowiek kwadratowy”⁷). Rysunkowi towarzyszy tekst sporządzony tzw. pismem lustrzanym.

Źródło: https://upload.wikimedia.org/wikipedia/commons/2/22/Da_Vinci_Vitruve_Luc_Viatour.jpg

¹ We Włoszech nadawano wówczas nazwiska na podstawie pochodzenia, miejsca urodzenia bądź profesji.

² Źródło: https://pl.wikipedia.org/wiki/Leonardo_da_Vinci [2018-08-03]

³ **Autoportret** (wł. *Autoritratto*) – rysunek stworzony w 1513 r. przez włoskiego artystę renesansowego,

⁴ Leonarda da Vinci znajdujący się w zbiorach Biblioteca Reale w Turynie. Nie ma całkowitej pewności co do autorstwa Leonarda. Niewykluczone, że jest to falsyfikat z XIX wieku.

⁵ [https://pl.wikipedia.org/wiki/Autoportret_\(rysunek_Leonarda_da_Vinci\)](https://pl.wikipedia.org/wiki/Autoportret_(rysunek_Leonarda_da_Vinci))

⁶ **Francesco Melzi**, lub **Francesco de Melzi**, (ok. 1491-1568/70) był włoskim malarzem urodzonym w rodzinie mediolańskiej szlachty w Lombardii. Był uczniem Leonarda da Vinci. Wykonawcą testamentu. https://en.wikipedia.org/wiki/Francesco_Melzi

⁷ Źródło: <https://www.leonardodavinci.net/self-portrait.jsp>



Rysunek 3. Portret damy z gronostajem (zwyczajowo: Dama z łasiczką lub rzadziej Dama z gronostajem) – obraz olejny autorstwa Leonarda da Vinci, namalowany około 1489 roku, wykonany w technice olejnej, z użyciem tempery, na desce orzechowej o wymiarach 54,7 na 40,3 cm. Przedstawia Cecylię Gallerani, kochankę księcia Ludwika Sforzy.
Źródło: https://upload.wikimedia.org/wikipedia/commons/e/e1/The_Lady_with_an_Ermine.jpg

ry. Encyklopedia Britannica online stwierdza: „*Leonardo przewidywał świetną grafikę ludzkiego ciała, którą wyprodukował dzięki swoim anatomicznym rysunkom, a człowiekiem witrawiańskim jako kosmologię*

del minor mondo (kosmografię mikrokosmosu), który uważał działanie ludzkiego ciała za analogię do działania wszechświata.”⁸

W Polsce znane jest także dzieło *Dama z gronostajem*, ze względu na to, że jest jedyną pracą artysty, jaka znajduje się w polskich zbiorach. Do czasów dzisiejszych przetrwało najprawdopodobniej 15 jego obrazów.

Leonardo da Vinci jako inżynier

W wielu opracowaniach opisuje się Leonarda jako inżyniera, który tworzył projekty wyprzedzające czasy jego życia. Jest autorem opracowania przedstawiającego koncepcje śmigłowca, czołgu, a także wykorzystania podstaw tektoniki płyt oraz podwójnego poszycia burt statku i wielu innych innowacji. Niestety względnie mała liczba jego pomysłów została wcielona w życie za jego życia, ale niektóre z nich, takie jak automatyczną nawijkarkę do szpul czy maszynę do sprawdzania wytrzymałości drutu na rozciąganie, zaadoptowano w świecie techniki bez większego rozgłosu⁹.

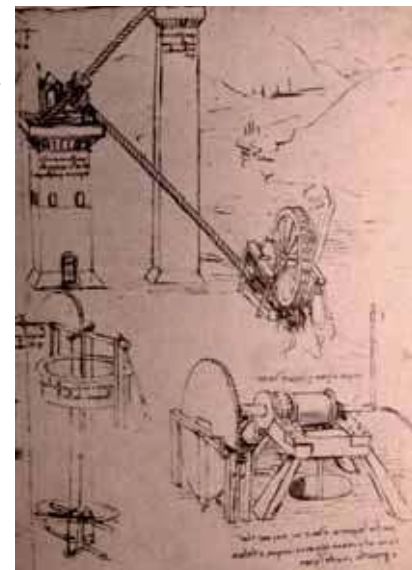
Praktyczne odkrycia i projekty

Leonardo da Vinci opracowywał wynalazki, gdy nie było jeszcze patentów, nie można zatem z całą pewnością stwierdzić, jak wiele z nich weszło do użytku, wywierając wpływ na życie wielu ludzi.

Z ostatnio przeprowadzonych badań wynika, iż Leonardo jest także wynalazcą zamka kołowego, nad którym zaczął pracować w 1493 r. Historyk Claude Blair odkrył, iż w prowincji Friuli, w północnych Włoszech, zaczęto wytwarzać zamki kołowe najpóźniej w 1510 r., a Leonardo pracował tam przez pewien czas, wykonując pomiary związane z fortyfikacjami na miejscowym zamku.

Rysunek 4. Grafika przedstawiająca różne urządzenia hydrauliczne proponowane przez Leonarda

Źródło: https://upload.wikimedia.org/wikipedia/commons/f/fa/Leonardo_machines.JPG



W obszarze hydrauliki, studia Leonarda nad ruchem wody zainspirowały go do zaprojektowania maszyn, które wykorzystywały jej siłę. Większość prac z dziedziny hydrauliki wykonał dla Ludwika Sforzy. Leonardo, przebywając we Francji, rozważając możliwości skanalizowania błotnistych terenów koło zamku Romorantin wpadł na pomysł wykonania aparatu do czerpania wód podziemnych. Urządzenie pompowało wodę za pomocą śruby obracającej się wewnątrz cylindra¹⁰. Śruba Archimedesowa jest maszyną prostą, to duża spirala, umieszczona w drewnianym cylindrze. W czasie pracy dolny koniec śruby zanurzony jest w wodzie, a obrót śruby wymusza jej ruch do góry¹¹.

Maszyny wojenne

W zachowanych notatkach Leonarda znajduje się szereg maszyn wojennych, m.in. czołg miotający małe kamienie, wprawiany w ruch przez dwóch mężczyzn napędzających wał korbowy. Mimo iż rysunek wyglądał na całkiem skończony, mechanika najwyraźniej nie była do końca dopracowana. Model zbudowany według projektu z 1484 r. mógł zaledwie z dużym wysiłkiem obracać się w miejscu, ale nigdy nie poruszał się prosto. Kolejna maszyna miała z przodu cztery kosy przymocowane do mechanizmu obracającego. Siłę napędową stanowiły konie. W notatkach Leonarda można też znaleźć działą, które były opisywane jako: *miotające małe skały niczym nawałnica powodując wielką panikę u wroga, straty oraz chaos*.¹²

Według Michaela Mosleya, producenta filmu dokumentalnego *Leonardo* dla stacji BBC, mogą być dwa powody, dla których projekty maszyn wojennych zrekonstruowane przez współczesnych inżynierów okazały się wadliwe.

Pierwszym z nich mogła być ochrona w pełni funkcjonalnych wynalazków przed skopiowaniem.

Drugi mógł odnosić się do światopoglądu Leonarda da Vinci, który w gruncie rzeczy był pacyfistą i celowo wpro-

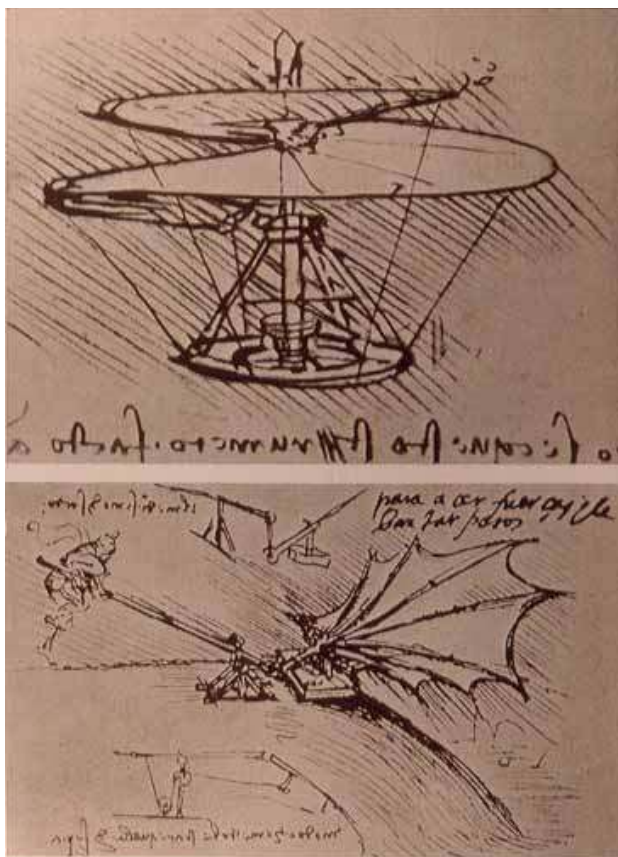
⁸ Są to opinie osób takich jak: Giorgio Vasari, Giovanni Antonio Boltraffio, Baldassare Castiglione, „Anonimo” Gaddiano, Berensen, Hipolit Taine, Henry Fuseli, Rio, Bortolon. [za] https://pl.wikipedia.org/wiki/Leonardo_da_Vinci#cite_note-2

⁹ Więcej szczegółów na stronie <https://archive.is/20120713164636/sites.google.com/site/leonardodavincivitruius/> oraz na stronie o adresie <http://leonardodavinci.stanford.edu/submissions/clabaugh/history/leonardo.html>

¹⁰ Źródło: <http://leonardodavinci.stanford.edu/submissions/clabaugh/history/leonardo.html>

¹¹ https://upload.wikimedia.org/wikipedia/commons/5/54/Szkic_%C5%9Bmig%C5%82owca.jpg

¹² https://upload.wikimedia.org/wikipedia/commons/b/b9/Leonardo_tank.JPG Więcej na ten temat w Archimedesowy gwint. „Fizyka w Szkole” nr 5, wrzesień/październik 2012, s. 19-22



Rysunek 5. Projekt ornitoptera Leonarda da Vinci, który w 1485 roku zaczął badać lot ptaków. Uznał, że ludzie są zbyt ciężcy i niewystarczająco silni, aby latać przy użyciu skrzydeł po prostu przymocowanych do ramion. Dlatego naszkicował urządzenie, w którym lotnik kładzie się na desce i wykonuje dwa duże, membranowe skrzydła za pomocą dźwigni ręcznych, pedałów nożnych i układu kół pasowych.

Źródło: https://upload.wikimedia.org/wikipedia/commons/5/50/Leonardo_da_Vinci_helicopter_and_lifting_wing.jpg

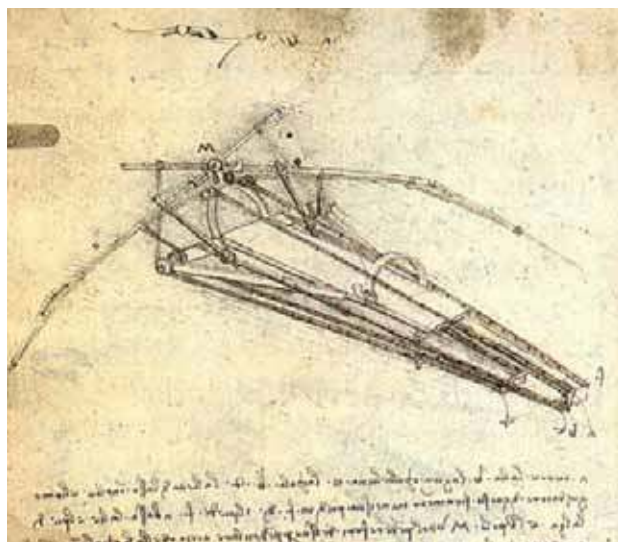
wadził błędy do swoich projektów, aby przywódcy wojskowi, dla których pracował nie mogli użyć ich w walce.¹³

W dalszej części wypowiedzi czytamy: „Niektóre z najbardziej rewolucyjnych projektów renesansowego geniuszu, takie jak jego wojskowy czołg, szybowiec i kombinizon do nurkowania, zawierały prostą, ale możliwą do skorygowania wadę, która stała się widoczna tylko wtedy, gdy zostały zbudowane”.¹⁴

Maszyny do latania

Jak głosi legenda, w niemowlęctwie nad kołyską Leonarda przeleciał jastrząb. Wspominając to wydarzenie, da Vinci uważał je za prorocze. Leonardo był przez wiele lat zafascynowany zjawiskiem lotu, tworząc wiele opracowań, w tym *Codex na temat lotu ptaków* (ok. 1505), a także wykonaniem planów kilku maszyn latających, takich jak trzepotanie skrzydeł (ornitoptera) i maszyna ze spiralą.

Niektóre z tych projektów okazały się skuteczne, podczas gdy inne testowały się gorzej.¹⁵ Jeden z jego projek-

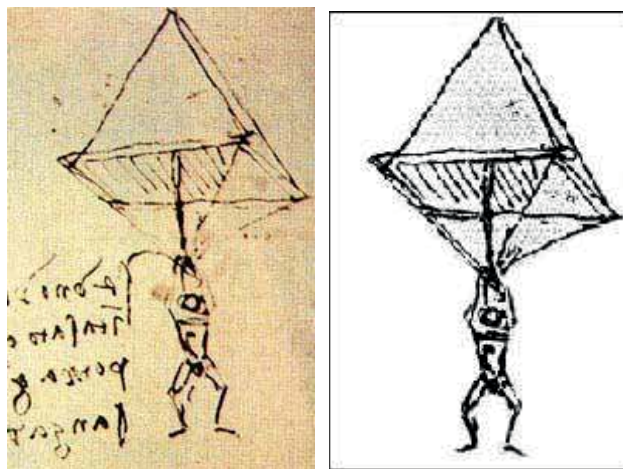


Rysunek 6. Projekt latającej maszyny (1488), Institut de France, Paris

Źródło: https://upload.wikimedia.org/wikipedia/commons/d/d4/Design_for_a_Flying_Machine.jpg

tów ukazuje helikopter podnoszony przez wirnik, wprawiany w ruch przez czterech mężczyzn. Leonardo nie ograniczał się tylko do projektowania latających maszyn z mechanicznymi skrzydłami, zaprojektował także spadochron. Lotnik miał się go trzymać za pomocą rąk. Projekt został pomyślnie przetestowany 26 czerwca 2000 przez brytyjskiego skoczka spadochronowego, Adriana Nicholasa.¹⁶

Wpadł również na pomysł lekkiej lotni. Projekt powstał w latach 1478–1480. Skrzydła były rozpięte na szkieletie siatkowym. Jej ogon był rozpostarty w wachlarz jak u ptaka.



Rysunek 7. Szkic Da Vinci był na marginesie zeszytu. Oryginalny projekt został napisany przez da Vinci w zeszycie w 1483 roku. W notatce towarzyszącej rysunkowi czytamy: „Jeśli mężczyzna ma długość gumowanego płótna o długości 12 jardów z każdej strony i 12 metrów wysokości, może skakać z jakiegokolwiek wielką wysokość bez obrażeń.”

Źródło: <http://news.bbc.co.uk/2/hi/science/nature/808246.stm>

Na rysunku stopy pilota są umieszczone na „m”, a ciało na „a, b”. Konstruktor wyraził nie myśl o tym, jak pilot ma kontrolować lot za pomocą sznurków. Niestety z rysunku nie wynika, gdzie jest przód a gdzie ogon szybowca.

¹³ <http://www.euro-met.com.pl/historia-srubby,693.html>

¹⁴ https://pl.wikipedia.org/wiki/Leonardo_da_Vinci#cite_note-61

¹⁵ Maszyny bojowe Da Vinci „zaprojektowane, by zawodzić” <https://www.theage.com.au/world/da-vinci-war-machines-designed-to-fail-20021214-gdunixhv.html>

¹⁶ Tamże

Przypominająca nowoczesną lotnię szybowiec Leonardo z elementami sterującymi opiera się na czystym ślizganiu się bez trzepotu. Jednak Leonardo miał obsesję na punkcie możliwości latania ludźmi za pomocą trzepoczących skrzydeł jak ptaki. Idea ta zainspirowała też innych. Pionier lotnictwa Otto Lilienthal (1848-1896) zbudował ponad 10 samolotów, głównie szybowców, rozciągając tkaniny nad wierzbowymi łaskami. Choć niedokładnie potwierdzili wiele pomysłów Leonarda. Podobnie jak on, nie mógł całkowicie uwolnić się od pojęcia trzepotania jako środka napędowego.¹⁷

Konstrukcja Leonarda dla podwodnego aparatu oddechowego składa się z rur trzcinowych połączonych skórą ze stalowymi pierścieniami, aby zapobiec ich zgnieceniu pod wpływem ciśnienia wody. Probówki są przymocowane do maski na twarz, a na drugim końcu do pływaka w kształcie dzwonu, aby zachować otwory nad wodą.

Inne propozycje innowacyjnego podejścia do ciekawych zagadnień natury

Leonardo poświęcił większość swojego życia na zrozumienie natury i jej opis. Realizował eksperymenty i dokonywał szczegółowej obserwacji. Opanował rysunek i malarstwo a jego uważne oko i twórczy umysł, służyły mu, aby dokonać naukowych obserwacji. Jego notatki łączą szczegółową obserwację z rysunkami z eksperymentów. Nawet jeśli nie podjął się eksperymentów, opisał, co można by osiągnąć i aktualnie wypróbować. Wiele z jego sugestii zapowiadało badania naukowe jeszcze przed wiekami. Na przykład:

- I Leonardo odrzucił wieczny ruch, zrozumiał zasadę względnego ruchu i zapowiedział trzecie prawo Newtona na dwa stulecia jego sformulowaniem: „*Dla każdego działania występuje przeciwna i równa reakcja*”.
- I Odrzucił pogląd, że powódź biblijna była odpowiedzialna za odkładanie skamieniałości wiele mil od ich źródła i pozwoliła wywnioskować istnienie bardzo długich rozpiętości czasu geologicznego.
- I Dokonując analizy ludzi i zwierząt, Leonardo dokonał wielu odkryć anatomicznych i fizjologicznych.
- I Badał optykę i percepcję za pomocą subtelnych eksperymentów. Wyjaśniał, dlaczego niebo jest niebieskie, argumentując, że światło ma skończoną prędkość i podróżuje w linii prostej. Wnioskował o istnieniu powierzchni w oku, która odbiera światło z szerokiego pola widzenia.
- I Leonardo sformułował prawo przepływu prądów: „*Cały ruch wody o jednolitej szerokości i powierzchni jest silniejszy w jednym miejscu niż przy innym, ponieważ woda jest tam płytsza niż na drugiej*”.¹⁸

Nowy sposób myślenia i działania

Artyści i rzemieślnicy w czasach Leonarda wiedzieli, jak budować i naprawiać znane rodzaje maszyn. Jednak pomysł stworzenia nowych rodzajów maszyn nie przyszedłby im do głowy. Leonardo wypracował nowe podejście do maszyn. Rozumiał, że dzięki zrozumieniu, w jaki sposób działa każda oddzielna część maszyny, może je modyfikować i łączyć na różne sposoby w celu

ulepszenia istniejących maszyn lub tworzenia wynalazków, których nikt wcześniej nie widział¹⁹.

Wynalazki Da Vinci

Fascynacja Leonardo maszynami prawdopodobnie rozpoczęła się w dzieciństwie. Niektóre z jego najwcześniejszych szkiców wyraźnie pokazują, jak działają różne części maszyn. Jako uczeń w pracowni artysty Verrocchio, Leonardo zaobserwował i używał różnych maszyn. Studiując je zdobył praktyczną wiedzę na temat ich konstrukcji, struktury i działania. Rozpoczął wówczas pisanie pierwszych systematycznych wyjaśnień, jak działają maszyny i jak można łączyć elementy maszyn. Jego talenty jako ilustratora pozwoliły mu narysować swoje mechaniczne pomysły z wyjątkową jasnością.

Pięćset lat po tym, jak zostały naniesione na papier, wiele z jego szkiców można z łatwością wykorzystać jako plany do stworzenia doskonałych modeli roboczych. Da Vinci opisał i naszkicował pomysły na wiele wynalazków. Ale nie wydaje się, żeby zostały kiedykolwiek zbudowane i przetestowane podczas jego życia. Choć jego notatki sugerują, że chciał zorganizować i opublikować swoje pomysły, zmarł, zanim mógł osiągnąć ten cel. Po jego śmierci jego notesy były ukryte, rozproszone lub zagubione, a jego wspaniałe pomysły zostały zapomniane. Minęły stulecia, zanim inni wynalazcy wpadli na podobne pomysły i doprowadzili je do praktycznego zastosowania.²⁰

Inne pomysły, szkice i idee

W historii techniki Leonardo da Vinci zapisał się jako fascynacja maszynami. Prawdopodobnie to zainteresowanie rozpoczęło się w dzieciństwie, a niektóre z jego najwcześniejszych szkiców wyraźnie pokazują, jak działają różne części maszyn. Już jako uczeń w pracowni artysty Verrocchio²¹, Leonardo zaobserwował i używał różnych maszyn. Studiując je zdobył praktyczną wiedzę na temat ich konstrukcji i struktury.

Wraz z bogatymi w szczegóły rysunkami, Leonardo rozpoczął pisanie pierwszych bardzo systematycznych wyjaśnień dotyczących działania maszyny wraz z możliwością łączynie poszczególnych elementów proponowanych maszyn. Posiadany przez niego talent ilustratora pozwolił mu narysować swoje mechaniczne pomysły z wyjątkową jasnością i precyzją. Pięćset lat po tym, jak zostały zaprezentowane na papierze, jak zauważają autorzy opracowań o Leonardzie, wiele z jego szkiców można z łatwością wykorzystać wręcz jako plany do stworzenia doskonałych modeli roboczych. Ale wydaje się, że te bardzo nieliczne z nich zostały kiedykolwiek zbudowane i przetestowane

¹⁷ https://en.wikipedia.org/wiki/Leonardo_da_Vinci#Observation_and_invention

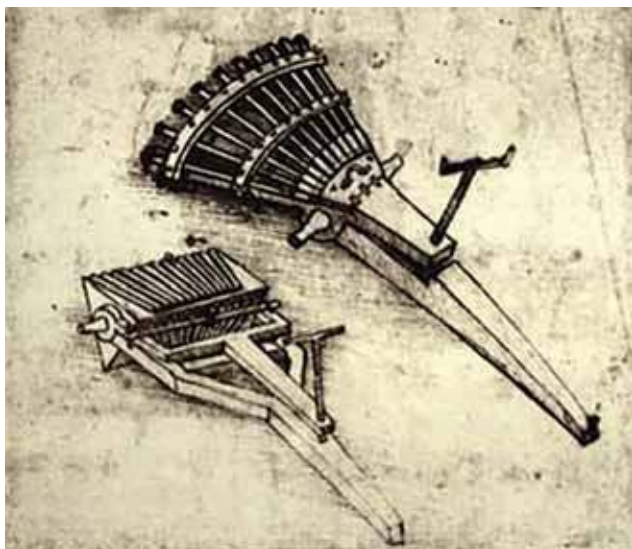
¹⁸ Spadochron Da Vinci leci, autor: Dr Damian Carrington z BBC News Online

¹⁹ <http://news.bbc.co.uk/2/hi/science/nature/808246.stm>

²⁰ <http://www.bl.uk/onlinegallery/features/leonardo/glider.html>

²¹ Źródło: <http://www.bl.uk/onlinegallery/features/leonardo/insi-ghts.html>

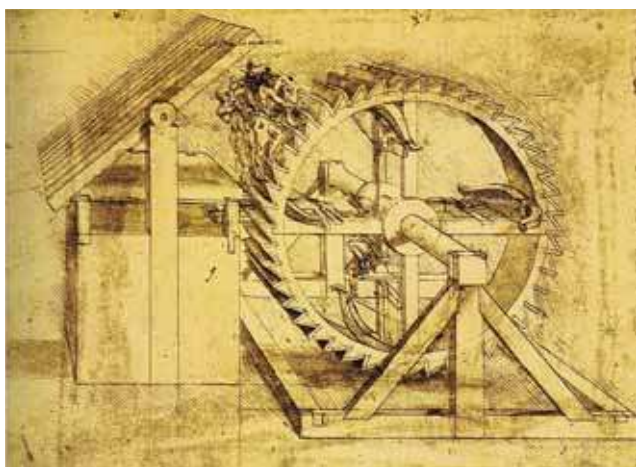
Szkice i modele maszyn Leonarda. Foto – Adobe Stock, wikimedia commons, <https://www.leonardodavinci.net/leonardo-da-vinci-biography.jsp>



Maszyna, będąca propozycją karabinu maszynowego, składała się z trzech zestawów dział, osadzonych na obracającym się trójkątnym bębnie. Rząd pistoletów umieszczono po każdej stronie bębna. Kiedy pierwszy zestaw broni wystrzelił, siła eksplozji wyrzuciła by te pistolety z powrotem, przynosząc następny zestaw dział na szczyt, gotowy do strzału.



Model przedstawia maszynę, którą jest czołgiem napędzanym przez człowieka. W czołgu ustawiono wiele armat, z możliwością oddania strzału we wszystkich kierunkach. Dolny obraz pokazuje dno zbiornika. Cztery osoby pracowałyby na kołach, aby poruszyć czołg.



Kusza to maszyna składająca się z koła, czterech kuszy, tarczy i stojaka na tarczę. Koło jest obracane przez ludzi idących na zębatym kole. Wewnątrz koła znajduje się łucznik strzelający z kuszy, gdy się zbliżają do niego. Strzały są wyrzucane przez stojak, a tarcza chroni mężczyzn chodzących po kole.



Ten kamienny miotacz składa się z podstawy, zakrzywionej belki, kół zębatych i mechanizmu zwalniającego.

Kamień ułożony byłby na łukowatej belce drewnianej, a mechanizm zwalniający byłby dokręcany za pomocą przekładni.

Kiedy promień zostanie uwolniony, siła uwolnienia wyrzuci kamień w powietrze.

podczas jego życia. Chociaż jego notatki sugerują, że chciał zorganizować i opublikować swoje pomysły, zmarł, zanim mógł osiągnąć ten ważny cel, a po jego śmierci notatki były ukryte, rozproszone lub zagubione, a jego wspaniałe pomysły zostały zapomniane. W efekcie minęły stulecia, zanim inni wynalazcy wpadli na podobne pomysły i doprowadzili je do praktycznego zastosowania.

Jakie były wynalazki Leonarda da Vinci?

Leonardo da Vinci opisał i nakreślił pomysły na wiele wynalazków, które były setki lat przed ich czasem, ale bardzo niewiele zostało kiedykolwiek zbudowanych lub przetestowanych podczas jego życia. Zebrany materiał

i umieszczony na stronach internetowych pozwala przeglądać szkice maszyn Leonarda i jednocześnie sprawdzić swoją wiedzę na temat tego oraz ustalić jakimi urządzeniami ostatecznie się stały i funkcjonują współcześnie.

dr Kazimierz Mikulski
Maksymilianowo

²² Źródło: <https://www.mos.org/leonardo/inventor>

²³ Źródło: <https://www.mos.org/leonardo/activities/inventions-quiz>

²⁴ Właściwie Andrea di Michele di Francesco Cione, ur. 1435, Florencja, zm. 30 VI 1488, Wenecja(?), <https://encyklopedia.pwn.pl/haslo/Verrocchio-Andrea;3992588.html>

Poddasze Układu Słonecznego cz. 3.

Blisko i daleko od Neptuna

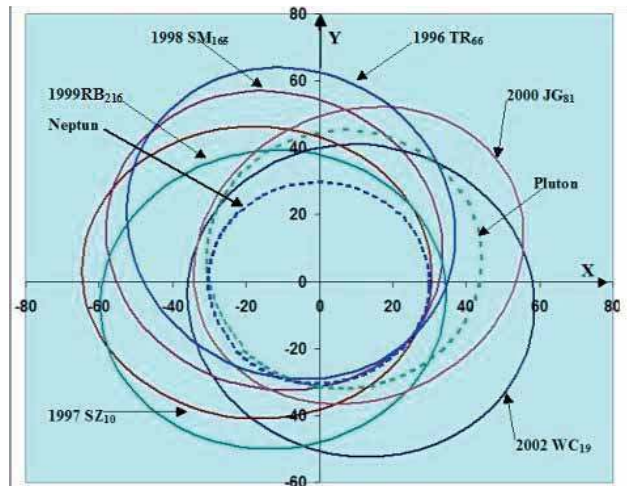
Jan Rokita

Uwaga: wszystkie dane, dotyczące liczebności poszczególnych klas obiektów są podane wg stanu na dzień 15 lutego 2015 roku i do czasu ukazania się tego tekstu mogły się oczywiście zmienić.

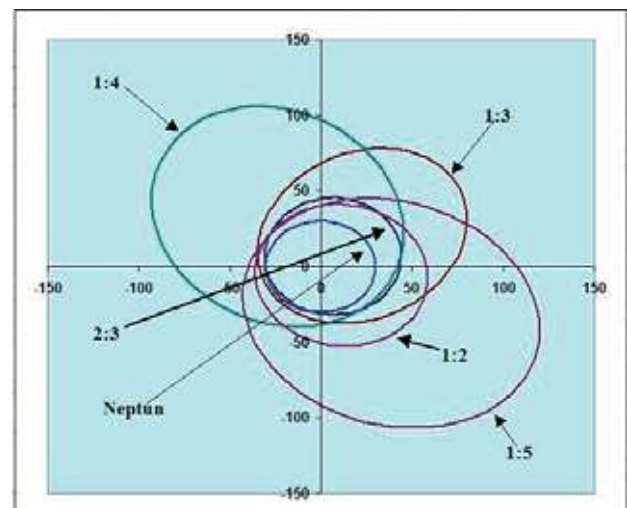
Wiemy już, że za granice występowania obiektów typu cubewano przyjęło się z jednej strony plutonki, pozostające w rezonansie orbitalnym 2:3 z Neptunem, a z drugiej twotina, dla których ten rezonans jest równy 1:2. Plutonki mamy 132 pewne i 120 prawdopodobnych, twotin – tylko 26. I jedno, i drugie zostały „zgarnięte” przez migrującego na zewnątrz Układu Słonecznego Neptuna i – jak pokazują symulacje komputerowe – tuż po zakończeniu migracji mogły być równie liczne, a nawet z przewagą twotin, jednak rezonans 1:2 okazuje się być mniej stabilny, niż 2:3 i przez 4 miliardy lat większość z nich po prostu zostało wyrzuconych ze swych orbit. Jak istotny jest dla nich ten rezonans, pokazuje rysunek obok, na którym widać, jak ładnie peryhelia twotin układają się wzdłuż orbity tego gazowego olbrzyma. Zainteresowanych odsyłam do bardzo interesującej (naprawdę!) pracy E. I. Chiang & A. B. Jordan **On the Plutinos and Twotinos of the Kuiper Belt**, dostępnej w internecie w wersji pdf.

Jest wśród nich dość duża, szacowana na około 450 km średnicy planetoida (**119979**) **2002 WC₁₉**, posiadająca prawie 240 kilometrowego satelitę, prawie 300-kilometrowa (**26308**) **1998 SM₁₆₅**, też z dużym, prawie 100-kilometrowym satelitą, które ze średnią gęstością, szacowaną na 0,5 g/cm³ wydają się równie porowate (porowaty lód?), jak plutonek **1999 TC₃₆**. Ich orbity mają mimośrodki między mniej więcej 0,25, a 0,4 i nachylenia do ekliptyki na ogół mniejsze od 20°, natomiast o ich właściwościach fizycznych wiadomo niewiele.

Z innych rezonansów orbitalnych z Neptunem, wyrażających się całkowitą proporcją, mamy 5 obiektów w rezonansie 1:3 (z angielskiego nazywanych **threetino**) i jeden 1:4 (**fourtino**). Threetina obiegają Słońce raz na prawie pół tysiąclecia, po orbitach o wielkich półosiach rzędu 63 AU, fourtino (jest na razie tylko jedno: **2003 LA₇**) – raz na 672,66 roku w średniej odległości 76,8 AU. Cechą orbit ciał w tych „dużych” rezonansach są znaczne mimośrodki (rzędu 0,5 lub więcej – patrz zresztą rysunek obok). Również i te obiekty mają właściwości fizyczne owiane, póki co, mgłą tajemnicy. Największą średnicę wydaje się mieć fourtino (231 km), ale jest to tylko oszacowanie na podstawie jasności, z przyjętym standardowym albedo 0,09. O bycie drugim fourtino podejrzewana jest również **2011 UP₄₁₁**. O rezonans w stosunku 1:5 podejrzewana była



Orbity kilku twotin porównane w płaszczyźnie ekliptyki z orbitami Neptuna i Plutona (linie przerywane).



Porównanie orbity twotina (2002 WC19), threetina (2003 LG7), fourtina (2003 LA7) i 2003 YQ179 (hipotetycznie w rezonansie 1:5) z orbitami Neptuna i Plutona (rezonans 2:3).

asteroida **2003 YQ₁₇₉**. Gdyby okazało się to prawdą, byłby to najbardziej odległy od Słońca – ze znanych dzisiaj – obiekt Układu Słonecznego w rezonansie z Neptunem. Pół jego wydłużonej orbity (mimośród 0,585) ma wartość 88,5 AU, a okres orbitalny to 810,5 roku. Odległość aphelium – to budzące szacunek 142,3 AU, natomiast peryhelium jest niewiele dalej, niż Neptun: 37,3 AU. I właśnie 23 października 2019 roku nasza planetoida minie peryhelium, a będzie to pierwszy raz od 20 kwietnia 1169 roku... Podejrzewa się jednak, że taka, a nie inna proporcja okresów orbitalnych 2003 YQ₁₇₉ i Neptuna jest efektem tylko przypadkowej koincydencji.

W przypadku ruchów orbitalnych istotne są też rezonanse, wyrażane wymiernymi ułamkami, czego dowodem mogą być **przerwy Kirkwooda** w głównym pasie planetoid, z których obiekty zostały powyrzucane przez oddziaływanie grawitacyjne potężnego Jowisza. Dla Neptuna, choć jest mniejszy i mniej masywny, takie rezonanse też są ważne. O jednym już wiemy: to 2:3 dla plutonków. Po 19 przedstawicieli mają proporcje 2:5, 3:5 i 4:7, 9 – 3:4, 5:9 i 3:7, po 5 – 4:5 i 4:9, trzech ma rezonans 2:7, po dwóch – 5:8, 5:8 i 2:11 i wreszcie po jednym 3:8, 3:10, 4:13 i 6:11.

O ten ostatni rezonans podejrzewana jest planeta karłowata Makemake, ale potwierdzenie tego faktu wymaga dalszych, obejmujących większy przedział czasowy badań. Wśród obiektów o rezonansie 2:5 jest też duża, prawie 1000-kilometrowa asteroida (84522) 2002 TC₃₀₂, mocny kandydat na planetę karłowatą. Równie mocnym wydaje się 636-kilometrowa 42301) 2001 UR₁₆₃, jeden z najbardziej czerwonych obiektów transneptunowych, pozostający z gazowym olbrzymem w rezonansie orbitalnym 4:9.

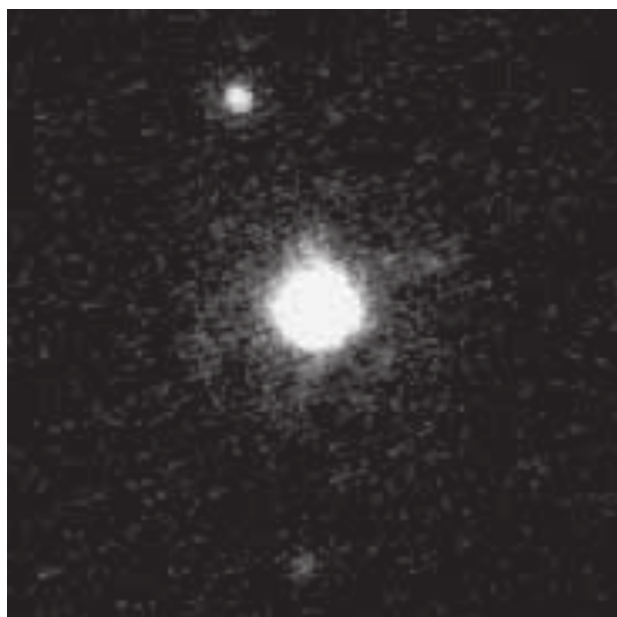
Za potwierdzony uważa się słaby rezonans 7:12, w którym znajduje się **Haumea**, kolejna z planet karłowatych i z tego powodu nie wlicza się jej do klasycznych cubewano, choć parametry jej orbity (wielka półoś 43,218 AU, mimośrodek 0,191, nachylenie płaszczyzny do płaszczyzny ekliptyki 28,19° i okres obiegu wokół Słońca 284,12 lat) są podobne do „gorących” obiektów tej klasy. Imię samej planety pochodzi od hawajskiej bogini płodności i porodu dzieci, z tego samego źródła wzięto imiona jej dwóch satelitów: **Hi’iaka** i **Namaka** – to jej córki.

Jest obiektem niezwykle ciekawym sama w sobie. Przede wszystkim, jest rekordzistką pod względem szybkości rotacji: raz na 3,9 godziny. Żaden inny obiekt Układu Słonecznego nie wiruje tak szybko. Musi to, oczywiście, odbić się na jej kształcie: Haumea jest silnie spłaszczoną elipsoidą o wymiarach 1920 × 1540 × 990 km. Dzięki satelitom można było wyznaczyć masę układu Haumea + satelity i okazała się całkiem spora: 28% masy układu Plutona, przy czym prawie cała jest skupiona w Haumei, a więc i jej średnią gęstość, która okazała się zaskakująco wysoka: mieści się w granicach 2,6-3,3 g/cm³, podczas gdy dla typowych obiektów Pasa Kuipera jest znacznie niższa, poniżej 2 g/cm³. To pierwsza wskazówka, że jakieś zderzenie pozbawiło w przeszłości ten glob większości lodu, pozostawiając tylko cienką warstwę.

Wskazówka druga: jej spektrum pokazuje, że jest to wodny lód w sporej części krystaliczny, który trwale może istnieć w temperaturach powyżej 110K, w niższych przechodzi w formę amorficzną, a proces ten jeszcze przyspiesza promieniowanie kosmiczne. Na powierzchni Haumei jest poniżej 50K, więc musiała ona zostać odnowiona stosunkowo niedawno.

Po trzecie – silne ślady krystalicznego, prawie czystego lodu wodnego są również cechą widma obu satelitów, które najprawdopodobniej są więc odłamkami dawnej lodowej powłoki planety. I to zderzeniu zawdzięcza ona tak szybką rotację. Ocenia się, że mogło ono nastąpić ponad 100 milionów lat temu. Dzisiejsza powierzchnia Haumei jest wysoce jednorodna, a oprócz wodnego lodu jest tam 8-procentowy dodatek czerwonych związków organicznych, cyjanowodoru i jego soli. Nasza karłowata planeta jest, z jasnością do +17,3^m, trzecim co do jasności, po Plutonie i Makemake, obiektem Pasa Kuipera, możliwym do zaobserwowania za pomocą większych przyrządów amatorskich.

W uzyskiwaniu informacji o powierzchni i wymiarach ciała centralnego pomocne mogą być przejścia przed nim satelitów, niestety Hi’iaka ostatni raz przesłoniła Haumeę w roku 1999 i nie zrobi tego ponownie przez kolejne 130



Haumea i jej dwa księżycy

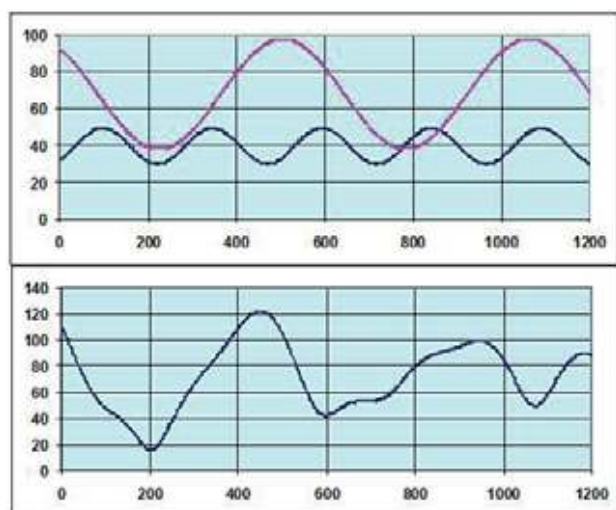
lat. Na szczęście Namaka co jakiś czas przesłania macierzysty glob, można też obserwować wyjątkową rzecz w świecie naturalnych satelitów obiektów Pasa Kuipera – wzajemne okultacje Hi’iaki i Namaki.

Wynikiem wspomnianego zderzenia z przeszłości jest powstanie całej **kolizyjnej rodziny planetoid**, które są jego odłamkami. To jedyna taka rodzina wśród obiektów transneptunowych (w głównym pasie planetoid, między Marsem i Jowiszem, jest takich więcej). Oprócz samej Haumei i jej satelitów wchodzi w jej skład 10 zidentyfikowanych obiektów, z których jeden ma średnicę, szacowaną na 332 km, a cztery – między 100, a 200 km.

Szacunki tych średnic poczyniono jednak tylko na podstawie blasku z pewnym założonym albedo, więc jeśli w rzeczywistości jest ono inne – to i wymiary też. Półosie ich orbit mają wartości około 43,5 AU, mimośrodów około 0,125, a nachylenia płaszczyzn orbit do płaszczyzny ekliptyki około 27°. I tu od razu zauważmy, że dla Haumei te wartości są nieco inne: półoś wielka – to 42,9 AU, a mimośrodek – zaledwie 0,056 i jest to zapewne efekt rezonansowego oddziaływania Neptuna. Odłamki te (które są zapewne fragmentami – podobnie, jak i satelity – pierwotnej powłoki lodowej macierzystego obiektu, tak przynajmniej sugeruje ich widmo) rozleciały się po zderzeniu z względnymi prędkościami mniejszymi od 150 m/s.

Scenariusz kolizji nie jest do końca pewny, niektórzy badacze sugerują, że były dwa zderzenia, a drugie rozproszyło materiał, który zebrał się w jedno, większe ciało po pierwszym. Nie jest pewne miejsce zderzenia (niekoniecznie tam, gdzie rodzina jest dzisiaj) i czas (od nieco ponad 100 milionów do ponad miliarda lat wstecz).

Interesujące jest, że z numerycznego modelowania migracji Neptuna i rozpraszania asteroid, opisanego we wspomnianej wcześniej pracy Changa i Jordana, wynika także – i to w proporcjach rzeczywiście obser-



U góry: zmiany odległości (w AU) od Słońca Plutona (czarna linia) i Eris (czerwona) w ciągu 1200 lat, licząc od 01.01.2016. U dołu – wzajemna odległość obydwu obiektów w tym samym okresie. Widać, że za ok. 800 lat obiekty „zamieniają się” kolejnością w Układzie Słonecznym, jednak ich wzajemna odległość nigdy nie będzie przesadnie mała.

wowanych – obecność obiektów w rezonansach wymiernych.

Planetoidy trojańskie

A skoro mówimy o obiektach w rezonansie orbitalnym z Neptunem, nie możemy pominąć tych, które są w rezonansie 1:1, czyli o jego **planetoidach trojańskich**. Na dzień dzisiejszy znanych jest 12 takich asteroid: 9 w okolicy punktu libracyjnego L4 układu Słońce-Neptun (wyprzedzają Neptuna w ruchu orbitalnym o około 60°) i 3 w okolicach punktu L5 (te się o około 60° późnią). Ich orbity, poza owymi przesunięciami, są bardzo podobne do orbity Neptuna. Mogą być one bardziej liczne, ale poszukiwania są utrudnione przez to, że punkty te są akurat na tle Drogi Mlecznej, czyli tam, gdzie jest pełno słabych gwiazd.

Jest wśród nich parę ciał o znacznych, ale nie aż tak wielkich rozmiarach. Jedyny obiekt, posiadający imię

własne, to 76-kilometrowa **(385571) Otrera** (to imię pierwszej królowej Amazonki z greckiej mitologii) w okolicach punktu L4. W tych samych okolicach jest prawie 100-kilometrowa **(385695) 2005 TO₇₄**, prawie 160-kilometrowa **2001 QR₃₂₂**, około 140-kilometrowa **2006 RJ₁₀₃** i inne, o których rozmiarach wiadomo, że są rzędu kilkudziesięciu kilometrów lub... nic. Punkt L5 – to m.in. 100-150-kilometrowa asteroida **2008 LC₁₈** i 90-180-kilometrowa **2011 HM₁₀₂**, do której była szansa skierować, po drodze do Plutona, sondę New Horizons. Zrezygnowano z tego ze względu na zbyt małą przepustowość łączności z Ziemią i skoncentrowano się na głównym celu misji. Ponieważ są to obiekty słabe, więc badania spektroskopowe są nadzwyczaj trudne i o właściwościach ich powierzchni nie wiadomo w zasadzie nic.

Niezwykle interesujący jest 200-kilometrowy obiekt **(316179) 2010 PL₆₅**, który skacze (jak koń trojański, mówią astronomowie) między punktami L4 i L5. Też jest w rezonansie 1:1 z Neptunem i efektywnie zakreśla wokół niego orbitę w kształcie podkowy – jest więc jego quasi-satelitą (pamiętaj! Czytelnicy artykułu w numerze 4/2014 „Fizyki w Szkole”: Marcin i Tomasz Majka „Podróż na drugi ziemski księżyc” o planetoidzie **3753 Cruithne?**). O podobne zachowanie podejrzewane są jeszcze trzy asteroidy.

Wyjaśnijmy od razu, że planetoidy trojańskie ma większość planet Układu Słonecznego. Oczywiście rekordy ilości bije Jowisz (ma ich 6 i ćwierć tysiąca), ale 7 ma Mars, po jednym Ziemia, Uran, a nawet Wenus, tylko o Trojańczykach Merkurego i Saturna na razie nic nie wiadomo.

Dysk rozproszony

Poza rejonem występowania cubewano, częściowo zając się z jego zewnątrz, znajdziemy stosunkowo słabo „zaludniony” (na dzień 25.02.2015 r. – 144 obiekty), za to bardzo rozległy obszar, nazywany **dyskiem rozproszonym**. To naprawdę kawał przestrzeni: należące do niego obiekty mogą niewiele różnić się parametrami orbit od obiektów Pasa Kuipera, ale też mieć półosie orbit rzędu stu, a nawet więcej jednostek astronomicznych. Z tego względu dzieli się go często na dysk **bliski** (typowe obiekty dysku rozproszonego) i **rozszerzony**, do którego należące ciała nazywane są często również **obiettami odłączonymi**. Kluczowym kryterium jest wspomniany już **parametr Tisseranda**, opisujący wpływ Neptuna: dla pierwszej grupy jest on mniejszy, a dla drugiej większy od 3.

W odróżnieniu od obiektów Pasa, które choć krążą wokół Słońca bliżej Słońca, ale nigdy się do niego, wskutek niewielkiej ekscentryczności orbit, przesadnie nie zbliżają i te orbity są dość stabilne, ciała z dysku rozproszonego poruszają się po elipsach bardziej wydłużonych (mimośrody nawet powyżej 0,8, a nawet 0,9!), o peryheliach położonych bliżej i jeśli mijają je akurat wtedy, gdy i Neptun jest niedaleko, ich ruch może zostać (nawet znacząco) zaburzony. Spory jest także rozrzut nachyleń płaszczyzn orbit do ekliptyki: od praktycznie zerowych do prawie połowy kąta prostego.

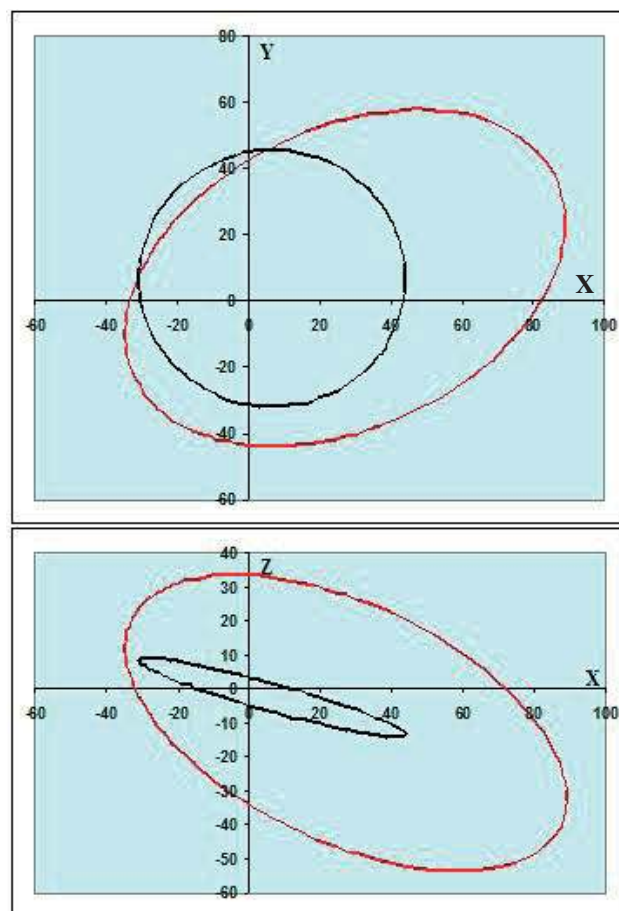
Spośród obiektów dysku rozszerzonego wyodrębnia się szczególną klasę planetoid o nazwie **sednoidy**, których peryhelia są tak daleko, że Neptun nic im nie może zrobić. Na dzień dzisiejszy zawiera ona całe dwie, plus parę prawdopodobnych. Podziały i definicje dysku rozproszonego, a nawet nazwy jego różnych rejonów, nie są jednolite i silnie zależą od tego, jaka grupa badaczy o nim mówi lub pisze (choćby to, czy dysk rozproszony to jeszcze Pas Kuipera, czy już nie, różne grupy różnie do tego podchodzą), dlatego Czytelnicy nie powinni być zdziwieni spotykając je w literaturze.

Najbardziej znana z tej grupy obiektów jest **(136199) Eris**, czwarta, najbardziej odległa z obecnie uznawanych pozaneptunowych planet karłowatych i... pośrednia przyczyna „degradacji” Plutona. Ze średnicą 2326 km jest tylko niewiele mniejsza od Plutona: raptem o 46 km (plus-minus błędy pomiarowe średnic obu obiektów), za to o $\frac{1}{4}$ (dokładnie: 1,27-krotnie) od niego masywniejsza. To ostatnie udało się tak dokładnie ustalić dzięki temu, że Eris ma 150-kilometrowego satelitę, nazwanego Dysnomia, obiegającego ją raz na 15,774 dnia. W greckiej mitologii Eris była boginią niezgody, a Dysnomia jej córką, co wydaje się trafnym wyborem ze względu na zamieszanie po degradacji Plutona, ale i... wokół jej własnego imienia, bo zespół odkrywców chciał ją nazwać Xena, a jej satelitę Gabrielle od bohaterki znanego również w Polsce serialu „Xena – wojownicza księżniczka”. Gabrielle była przyjaciółką Xeny, bardzo ogłędnie mówiąc. Nie przeszło.

Jej orbita jest typowa dla obiektów dysku rozproszonego: wielka półś ma wartość 67,78 AU, jednak mimośrd 0,44 sprawia, że aphelium znajduje się aż 97,65 AU od Słońca (Eris minęła je w ósmej dekadzie ubiegłego stulecia i teraz bardzo powoli zbliża się do niego, będąc obecnie w odległości około 92 AU), za to peryhelium – ledwie 37,91 AU, bliżej, niż aphelium Plutona! Może się więc zdarzyć, że Pluton będzie bardziej odległy od Słońca, niż Eris – i stanie się tak np. na przełomie 28-go i 29-go stulecia. Podczas najbliższego peryhelium Eris, za około 200 lat, również Pluton będzie w pobliżu swojego – i „przyrodzona” kolejność tych obiektów się nie zmieni (patrz rysunek). Ale to w tym okresie zbliżą się one do siebie na 15,3 AU, choć zderzenie obu obiektów nie jest, przynajmniej jeśli nic nie zmieni radykalnie ich orbit, możliwe dzięki aż 44-stopniowemu nachyleniu orbity Eris do płaszczyzny ekliptyki. Oczywiście i „rok” na tym obiekcie musi być odpowiednio długi: trwa niemal dokładnie 558 naszych.

Dzięki temu nachyleniu wędrówki Eris po niebie nieco się różnią od wędrówek planet, a nawet i samego Plutona. Obecnie (od roku 1929 do 2036) przemieszcza się ona na tle gwiazdozbioru Wieloryba. Wcześniej (1840-1875) była w Feniksie, a potem w Rzeźbiarzu, w roku 2036 wejdzie do gwiazdozbioru Ryb, w 2065 do Barana, w 2128 do Perseusza, a w 2173 – dotrze aż do Żyrafy.

Stosunkowo duża średnia gęstość Eris (około $2,5 \text{ g/cm}^3$) sugeruje, że w znacznie większej części, niż Pluton składa się ona z materiałów skalnych. Modele wnętrza dopuszczają, dzięki wewnętrznym rozpadom promieniotwórczym, istnienie – na granicy skalnego rdzenia i lodowego płaszcza

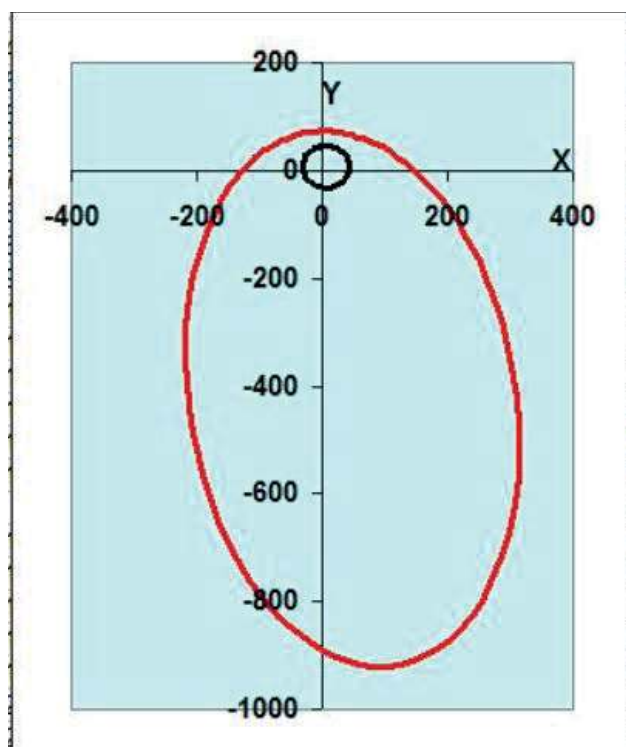


Orbity Plutona (czarna linia) i Eris (czerwona): u góry w płaszczyźnie ekliptyki, u dołu – w płaszczyźnie do niej prostopadłej. Łatwo stwierdzić, że zderzenie tych ciał jest, póki co, niemożliwe.

– oceanu ciekłej wody. I to na obiekcie tak niewiarygodnie odległym od Słońca i zimnym (obecnie tylko nieco powyżej 30K). Przecież obecnie to 3 razy dalej, niż Pluton!

Badania spektroskopowe Eris wskazują na obecność metanowego lodu, podobnie, jak na Plutonie. Jednak jest ona bardziej szara od niego. Badacze sądzą, że wskutek wielkiej odległości od Słońca zamarzający metan spokojnie i równomiernie osiada na powierzchni, pokrywając i maskując w ten sposób złoża czerwonych tholinów (czyli – to już wiemy – organicznego błota). Ten metanowy lód zaskoczył naukowców, bo albo w pobliżu peryhelium metan topi się, paruje i ulatuje w Kosmos, albo zamienia się na tholiny i osiada na powierzchni w postaci bardziej złożonych związków. Albo więc Eris stale przebywała z dala od Słońca, albo... ma jakieś wewnętrzne źródło tego gazu. Dla przykładu w widmie omawianej wcześniej Haumei jest silna obecność lodu wodnego, ale nie metanu. Jak widać, to bardzo ciekawy glob, niestety obliczono, że podróż na niego nawet z wykorzystaniem asysty grawitacyjnej Jowisza potrwałaby prawie 25 lat, więc choć zbadanie tak odległego i zimnego obiektu byłoby niezwykle interesujące, zapewne przyjdzie poczekać, aż zbliży się w okolice peryhelium. Niestety, to nie będzie za naszych czasów.

Generalnie, ten szary odcień obiektów dysku rozproszonego był zaskoczeniem. Choć jego masę ocenia



Orbity Plutona (czarna) i Sedny (czerwona) w płaszczyźnie ekliptyki.

się na zaledwie 0,01 do 0,1 masy Ziemi, jest tu parę ciekawych obiektów, nawet na tyle dużych, żeby rozważać je jako kandydatów na kolejne planety karłowate. Historycznie pierwszym ciałem, zaliczonym do tej klasy, była 340-kilometrowa asteroida (**15874**) **1996 TL₆₆**. Pół wielka jej orbity, to prawie 84 AU, co przy mimośrodku, równym 0,583 pozwala jej zbliżać się do Słońca na 35 AU w peryhelium (minęła je w 2001 roku), zaś w aphelium oddalać się na prawie 133 AU. Nic więc dziwnego, że jej okres orbitalny – to dostojne 769 lat z kawałkiem. O jej właściwościach fizycznych wiadomo tylko tyle, że jej blask zmienia się w bardzo niewielkim stopniu, co sugeruje w miarę jednorodną powierzchnię z niewielkimi, miejscowymi różnicami albedo. Później, po analizie archiwalnych zdjęć okazało się, że zarejestrowany na nich wcześniej obiekt o oznaczeniu (**48639**) **1995 TL₈** też należy do dysku rozproszonego.

To bardzo ciekawy obiekt. Okrąży Słońce raz na prawie 381 lat po orbicie o wielkiej półosi 42,5 i raczej umiarkowanym mimośrodku około 0,239, leżącej praktycznie dokładnie w płaszczyźnie ekliptyki: ich wzajemne nachylenie – to zaledwie 0,24°. Średnicę szacuje się na 350 km, ale nasza planetoida ma też dużego, 160-kilometrowego, mniej więcej 10-krotnie lżejszego satelitę, obiegającego ją raz na około pół dnia (ziemskiego). Jego orbita nie została jeszcze dokładnie określona, ale separacja obydwu ciał, to tylko... 420 km. Wiadomo też, że jest tam zimno: 38K.

Kolejna podwójna planetoida, to (**82075**) **2000 YW₁₃₄**, której składniki mają średnice 431 i 237 km. Jaj rok – to 445,5 lata ziemskie, wielka półoś orbity – 58,3 AU i nie do końca wiadomo, czy jest to obiekt dysku rozproszonego, czy w rezonansie orbitalnym 3:8 z Neptunem.



Sedna (po lewej) i trzy zdjęcia 2012 VP113, wykonane w dwugodzinnych odstępach 5 grudnia 2012 r. (kolorowe punkty otoczone owalem), ukazujące ruch asteroidy na tle gwiazd.

343-kilometrowa **2007 TG₄₂₂**, także kandydatka na planetę karłowatą, ma peryhelium 35,579 AU od Słońca (minęła je w 2005 roku), ale pół wielka jej orbity – to 501 AU, a mimośrodek o wartości 0,92779 wynosi ją w aphelium aż na odległość 967 AU. Jej „rok” – to 11200 naszych, poprzednie przejście przez peryhelium miało więc miejsce gdzieś w początkach neolitu...

Kolejna asteroida z dysku rozproszonego, **2004 XR₁₉₀**, jest godna uwagi nie tylko ze względu na rozmiar (580 km średnicy), ale i na orbitę: niezbyt ekscentryczną (mimośrodek 0,11, więc jej odległość od Słońca zmienia się od 51,4 do 64 AU, niewiele – jak na obiekty dysku), za to o nachyleniu do płaszczyzny ekliptyki aż 46,6°.

Ponad 12700 lat trwa obieg wokół Słońca niezbyt wielkiej planetoidy (38-83 km średnicy) (**87269**) **2000 OO₆₇**, która porusza się dostojnie (jej średnia prędkość orbitalna wynosi 0,88 km/s) po orbicie o wielkiej półosi prawie 544,5 AU i mimośrodku aż 0,962, nachylonej pod kątem 20° do płaszczyzny ekliptyki. Obiekt o oznaczeniu (**225088**) **2007 OR₁₀** jest z kolei największym w Układzie Słonecznym obiektem bez imienia. Ze średnicą 1500 km lub nieco mniejszą jest porównywalny z Haumeą i Makemake, co czyni go wysoce prawdopodobnym kandydatem na planetę karłowatą. Okrąży Słońce raz na 547 lat po orbicie o wielkiej półosi prawie 67 AU i mimośrodku 0,5, nachylonej do ekliptyki pod kątem prawie 31° i – co nie jest zbyt typowe dla obiektów z tamtego rejonu – jest wybitnie czerwony ze względu na obecność metanu, więc i zapewne tholinów. Nie jest wykluczone, że tak duży glob jest w stanie utrzymać jakąś atmosferę, co powinno sprzyjać ich powstawaniu. Oczywiście analiza widma asteroidy wskazuje też na obecność lodu wodnego. Jak widać – nie brakuje w dysku rozproszonym rzeczy ciekawych i godnych uwagi, a wymieniliśmy tylko niektóre.

Sednoidy

Sednoidy, czyli obiekty na tyle odległe od Słońca, że grawitacja Neptuna nie ma już na nie wpływu, wzięły swą nazwę od (**90377**) **Sedny**. Ta asteroida obiega Słońce raz na 11400 lat w średniej odległości 508 AU, po orbicie o ekscentryczności aż 0,85, a mimo to jej peryhelium znajduje się w odległości 76 AU. Za to aphelium – to prawie 940 AU. Jest duża, ma prawie 1000 km średnicy. Wiruje wokół osi raz na nieco ponad 10 godzin, na pewno jest zimna (temperatura powierzchni jest szacowana na 12-33K), ma wysokie albedo (0,41) i zapewne składa

się ze skał i lodu. Jest niewiele mniej czerwona od Marsa (pod względem czerwieni zajmuje po nim drugie miejsce w Układzie Słonecznym) i najprawdopodobniej wkrótce zostanie uznana za planetę karłowatą.

! Drugi oficjalnie uznawany sednoid, **2012 VP₁₁₃**, ma najbardziej odległe peryhelium spośród wszystkich znanych dziś obiektów Układu Słonecznego: 80,5 AU od Słońca. Obiega je w ciągu około 4300 lat w średniej odległości 263 AU, a mimośród orbity – to 0,696, czyli sporo, ale też mniej, niż w przypadku Sedny. Rozmiary obiektu nie są dokładnie znane, a szacunki mieszczą się w dość dużym przedziale od 300 do 1000 km. Podobnie jak Sedna jest czerwona i – co jest bardzo interesujące – ma bardzo podobny do niej, ale i do innych obiektów o wielkiej półosi orbity większej, niż 150 AU (na przykład **(148209) 2000 CR₁₀₅**, który obiega Słońce w średniej odległości 223 AU w ciągu 3345 lat) parametr orbity, nazywany **argumentem peryhelium** (informuje on, jaki jest kąt, patrząc ze Słońca, między kierunkiem na peryhelium, a punktem, w którym orbita ciała przebiega płaszczyznę ekliptyki, przechodząc na jej północną stronę; ten punkt nosi nazwę **węzła wstępującego**): wynosi on około 300°. Podobne są też nachylenia płaszczyzn orbit do płaszczyzny ekliptyki. Wyda-

je się to sugerować wspólne pochodzenie tych ciał, ale jakie – to na razie tajemnica. Branych jest pod uwagę parę możliwości: Pochodzą z Układu Słonecznego, ale zostały wyrwane z pierwotnych (bliższych Słońca) pozycji przez przechodzącą w pobliżu „obcą” gwiazdę;

! Z pierwotnych orbit wyrwała je masywna i obiegająca Słońce w dużej odległości jeszcze nieodkryta planeta;

! Ciało „wrywające” – to jeszcze nieodkryty gwiazdny towarzysz Słońca (obiegający je brązowy karzeł);

! We wczesnej historii Słońca – to ono wyrwało je z obcego systemu planetarnego podczas bliskiego przejścia.

Czerwonawe zabarwienie, a więc obecność metanu i – wysoce prawdopodobna – tholinów, wydają się świadczyć o tym, że sednoidy raczej nigdy zbyt nie zbliżały się ani do Słońca, ani do innych gwiazd i ich życie upłynęło w wiecznym mrozie. Na dobrą sprawę, badacze zastanawiają się, a przynajmniej niektórzy, czy nie są one przypadkiem przedstawicielami najbardziej wewnętrznej części (ciągle przecież hipotetycznego, choć większość astronomów głęboko wierzy w jego istnienie) **Obłoku Oorta**. (cdn)

Źródła zdjęć: strona misji New Horizons NASA, Wikipedia.

Co w fizyce piszczy

Zdjęcia z pasa Kuipera

Podczas kiedy my świętowaliśmy Nowy Rok słynna sonda New Horizon nie próżnowała i dokonywała nowych



Źródło: <https://www.sciencedaily.com/releases/2019/01/190102164307.htm>

odkryć. Jedno z nich było odkrycie, że obiekt Pasa Kuipera planetoida Ultima Thula ma kształt do złudzenia przypominający śniegowego bałwanika. Należy więc przypuszczać, że powstała wskutek zderzenia dwóch obiektów. Badania Pasa Kuipera wykonywane przez sondę są pierwszymi tego typu w historii i pozwalają nam wyciągać wnioski co do genezy Układu Słonecznego, gdyż materiał, z którego Pas Kuipera jest zbudowany jest pozostałością pierwotnego budulca naszego kosmicznego domu.

Misja na drugą stronę Księżyca

Druga strona Księżyca zawsze fascynowała ludzkość. Nawet bardziej niż pierwsza. Przykładem może tu być słynna niegdyś książka Jerzego Żuławskiego *Na srebrnym globie*, w której autor opisuje podróż dwójki bohaterów na niewidoczną stronę Księżyca. W książce tej niewidoczną stronę Księżyca przedstawiono jako pełną wspaniałej przyrody krainę. Co prawda pogląd ten trudno było uzasadnić naukowo, ale wielu marzycieli miało nadzieję,

że naukowcy się mylą. Ostateczny dowód na nieprawdziwość tych tezy o cywilizacji księżycowej dostarczyła misja sondy Łuna 3, która dostarczyła pierwszych zdjęć niewidocznej strony Księżyca.

Skoro wiadomo było, że na drugiej stronie Księżyca nie ma życia, to zainteresowanie nią zmalało. Nikt nie chciał tam lądować. Nikt, aż do teraz. Właśnie miało miejsce udane lądowanie Chińskiej sondy Chang'e-4. Lądowanie miało pomyślny przebieg i sonda rozpoczęła przysyłanie danych. Na miejsce lądowania wybrano dno największego księżycowego krateru uderzeniowego. Ponieważ sonda „nie widzi” Ziemi, komunikacja z ziemskim centrum dowodzenia odbywa się za pośrednictwem innego chińskiego satelity o nazwie Queqiao. Za ciekawostkę może posłużyć fakt, że sonda jest wyposażona w małą hermetyczną biosferę zawierającą nasiona traw i jajka owadów. Może więc po zakończeniu misji na Księżycu już będzie życie.

Źródło: <https://physicsworld.com/a/chinas-change-4-spacecraft-makes-historic-landing-on-far-side-of-the-moon/>

Prenumerata 2019



Nasze czasopismo ukazywało się:

- 1 rok** – gdy Maria Skłodowska-Curie w Stanach Zjednoczonych zbierała pieniądze na zakup grama radu dla Instytutu Radowego w Warszawie (obecnie Centrum Onkologii – Instytut im. Marii Skłodowskiej-Curie w Warszawie).
- 7 lat** – gdy wynaleziono wózek sklepowy
- 12 lat** – gdy wynaleziono radar
- 14 lat** – gdy uruchomiono pierwszy reaktor jądrowy
- 26 lat** – gdy wynaleziono laser
- 27 lat** – gdy wynaleziono światłowód
- 33 lata** – gdy pierwszy człowiek, Jurij Gagarin, wyleciał w Kosmos
- 62 lata** – gdy uruchomiono pierwszą stronę www
- 64 lata** – gdy wysłano pierwszego SMS
- 79 lat** – gdy opracowano iPhone'a
- 82 lata** – gdy zaprezentowano pierwszego iPada

91 lat – będzie ukazywała się „Fizyka w Szkole”
w 2019 roku, dzięki Waszej prenumeracie!

Prenumerata 2019 – formularz zamówienia znajdziecie na
www.aspress.com.pl/prenumerata-2019/

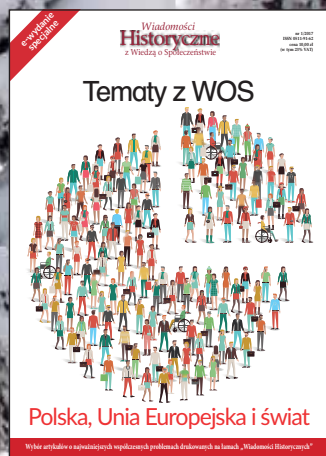
Wydania specjalne

(wersje elektroniczne – pliki PDF)

2018



2017



2016

