

FIZYKA

w Szkole z Astronomią

CZASOPISMO DLA NAUCZYCIELI

388 (LXIV) indeks 35810X Nr 5 wrzesień/październik 2023 CENA 40,00 zł (w tym 8% VAT)

KOMPUTERY I ALGORYTMY KWANTOWE

Przetwarzanie informacji
zapisanej w stanach
kwantowych

Przełomowe ogniwo termofotowoltaiczne

Żywooty fizyków Piotr Curie

Pasy Van Allena

Promieniowanie kosmiczne
widziane z kosmosu

O problemach z zasadami dynamiki Newtona

DOŚWIADCZENIA

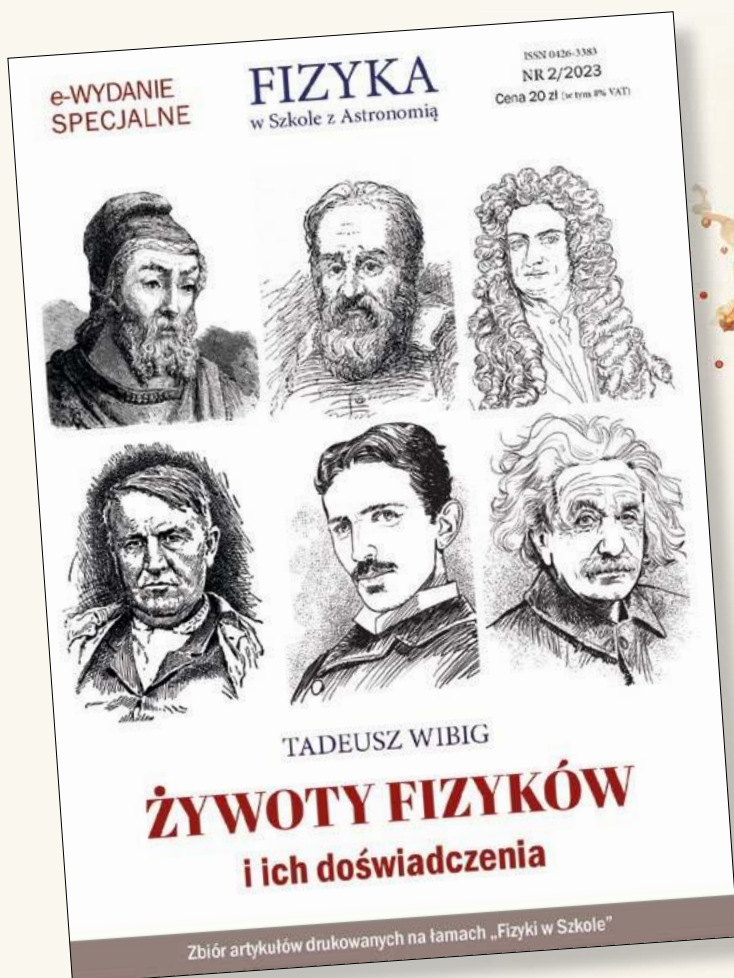
Włoskowatość, aerodynamika
i rozpraszanie światła

FIZYKA U OKULISTY

epiassa.pl c02670302



Wielcy fizycy i ich doświadczenia



- ✓ Od Archimedesesa po Einsteina
- ✓ Najważniejsze odkrycia
- ✓ 27 wybitnych postaci

**Cena
20 zł**
(w tym 8% VAT)



**WYDANIE SPECJALNE FIZYKI W SZKOLE
W WERSJI ELEKTRONICZNEJ – PLIK PDF**

Szczegóły i formularz zamówienia na www.aspress.com.pl/wydania-specjalne/

eprasa.pl c0287c7302

Szanowni Państwo

Dokładnie w momencie oddawania do druku tego numeru „Fizyki w Szkole” dowiedzieliśmy się, że Szwedzka Akademia Nauk przyznała najbardziej prestiżową nagrodę w fizyce, czyli Nagrodę Nobla. O ile tytuł zeszłorocznego Nobla jest na tyle skomplikowany, że jego objaśnianie stało się motywem przewodnim całej serii artykułów w naszym czasopiśmie, to uzasadnienie tegorocznej nagrody jest o wiele bardziej zrozumiałe dla laika. W tym roku laureatami zostali:

Pierre Agostini, Francuz pracujący na Uniwersytecie Stanu Ohio, Ferenc Krausz, Węgier pracujący Instytut Optyki Kwantowej im. Maxa Plancka w Monachium i na Uniwersytecie Technicznym w Monachium oraz Anne L'Huillier, Francuska pracująca na Uniwersytecie w Lund (Szwecja). Wspomniani naukowcy zostali nagrodzeni za metody eksperymentalne generowania attosekundowych impulsów światła do badania dynamiki elektronów w materii. Attosekunda jest to czas jaki potrzebuje światło, aby przebyć drogę rzędu średnicy atomu. Jest to też czas w jakim zachodzą elementarne zjawiska kwantowe.

Posiadając takie narzędzie jesteśmy w stanie badać takie procesy jak przeskoki elektronów między poziomami elektronowymi, powstawanie i zrywanie wiązań międzyatomowych i wiele innych, które do tej pory były domeną teoretyków od fizyki kwantowej. Posiadając takie narzędzia możemy też projektować działanie nanoukładów, jak i komputerów kwantowych. Możemy też badać procesy przeskoków elektronowych w neuronach, czyli szukać odpowiedzi na takie pytania jak czym właściwie jest myślenie.

W imieniu redakcji

Zbigniew Wiśniewski

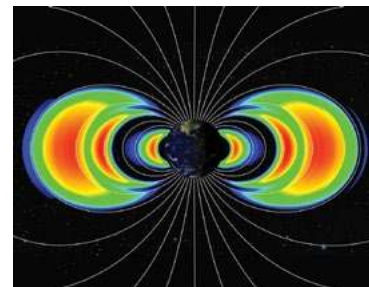
Fizyka wczoraj, dziś, jutro

4 Fizyka u okulisty i optometrysty ■ Tomasz Kubiak

Zmysł wzroku dostarcza nam około 80% informacji ze środowiska zewnętrznego. Jeśli zatem widzimy prawidłowo, możemy w pełni sprawnie funkcjonować w otaczającym nas świecie.

14 CREDO-Maze: Promieniowanie kosmiczne widziane z kosmosu – pasy Van Allena ■ Tadeusz Wibig

W badaniach pierwotnego promieniowania kosmicznego najlepiej wznieść się ponad ziemską atmosferę, a przynajmniej bardzo, bardzo wysoko.



20 Za co Nobel 2022? Część 5. Komputery i algorytmy kwantowe ■ Jan Kurzyk

27 Żywy fizyk. Piotr Curie ■ Tadeusz Wibig

Wszyscy wiedzą, że był mężem największej polskiej uczzonej Marii Skłodowskiej-Curie. Ale poza tym Piotr Curie był także naukowcem i to naukowcem dużego formatu.

30 Fizyczne elementy w supernowoczesnym wytwarzaniu energii elektrycznej ■ Kazimierz Mikulski

W Massachusetts Institute of Technology (MIT) i National Renewable Energy Laboratory w USA (NREL) zbudowano ogniwa termofotowoltaiczne (TPV) o sprawności ponad 40%.



Z naszych lekcji

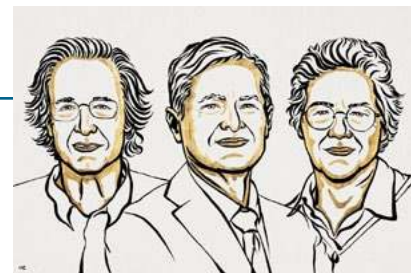
39 O problemach z zasadami i nie tylko... (Miniatura dydaktyczna) ■ Waldemar Reńda

44 „Doświadczenia skutków rzeczy pod zmysły podpadających...” – włoskatość, aerodynamika i rozpraszanie światła ■ Stanisław Bednarek



Astronomia dla każdego

47 Wybrane zagadnienia z mechaniki nieba. Cz. IV – Zadania cz. 2 ■ Marcin Wesołowski



3 października Królewska Szwedzka Akademia Nauk ogłosiła laureatów Nagrody Nobla z fizyki 2023. Laureatami zostali: Pierre Agostini, Ferenc Krausz oraz Anne L'Huillier.

Jak wskazano w komunikacie, troje laureatów zostało uhonorowanych za danie ludzkości nowych narzędzi do badania świata elektronów wewnątrz atomów i cząsteczek. Pokazali oni sposób na tworzenie ekstremalnie krótkich pulsów światła, które mogą mierzyć niezwykle szybkie procesy, w których elektrony poruszają się lub zmieniają energię.

FIZYKA
w Szkole z Astronomią

NUMER 5 WRZESIEŃ/PAŹDZIERNIK 2023
388 (LXIII) indeks 35810X ISSN 0426-3383

CENA 40,00 zł
(w tym 8% VAT)

Komitet redakcyjny Krystyna Jabłońska-Lawniczak, 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/, e-mail: szewczyk24@gmail.com, tel. 606 201 244 **Reklama** Jędrzej Chodakowski, jchodakowski1953@gmail.com **Skład i łamanie** ScanSystem.pl Ewa Szelażyńska **Druk i oprawa** Paper & Tinta, ul. Ceglana 34, 05-270 Nadma

Zdjęcie na okładce: Dreamstime

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

Fizyka u okulisty i optometrysty

Zmysł wzroku dostarcza nam około 80% informacji ze środowiska zewnętrznego. Jeśli zatem widzimy prawidłowo, możemy w pełni sprawnie funkcjonować w otaczającym nas świecie. Niestety współcześnie u coraz większej części populacji ujawniają się zaburzenia czynności narządu wzroku. W celu uzyskania pomocy udajemy się do różnych specjalistów, którzy w swojej pracy bazują nie tylko na osiągnięciach fizjologii i medycyny, ale przede wszystkim optyki oraz biofizyki. Przyjrzyjmy się zatem, jak fizyka pomaga w badaniu oraz usprawnianiu procesu widzenia.

Tomasz Kubiak

Wielu z nas w przypadku problemów ze wzrokiem w pierwszej kolejności kieruje się do lekarza okulisty. Ten jednak specjalizuje się przede wszystkim w diagnozowaniu i leczeniu chorób oczu. Badaniem i monitorowaniem wad refrakcji a także doбором odpowiedniej korekcji w postaci okularów lub soczewek kontaktowych współcześnie zajmują się natomiast coraz powszechniej optometryści. Co ciekawe, w Polsce przedstawiciele tego zawodu zdobywają fachową wiedzę oraz odpowiednie kwalifikacje na dedykowanym kierunku studiów, prowadzonych m.in. na wydziałach fizyki uniwersytetów w Poznaniu czy Warszawie.

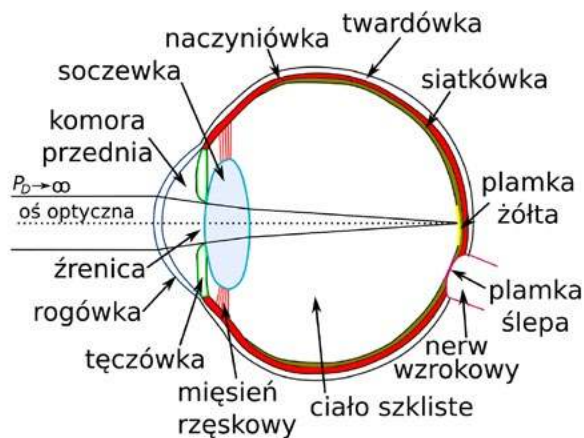
Z kolei za przygotowanie okularów korekcyjnych wg sporządzonej recepty oraz dobór odpowiednich powłok an-

tyrefleksyjnych odpowiedzialni są optycy okularowi. Nie trudno domyślić się, iż reprezentanci wszystkich wymienionych profesji muszą posiadać gruntowną wiedzę z optyki oraz biofizyki. Wiele zagadnień może okazać się fascynujących również dla szerokiego grona miłośników nauk przyrodniczych, dlatego w niniejszym artykule ich podstawy zostaną przybliżone wszystkim zainteresowanym.

Prawidłowe widzenie

Na wstępie należy przypomnieć, iż stali czytelnicy Fizyki w Szkole mieli już sposobność zapoznania się z budową oka (Rys. 1) oraz funkcjonowaniem układu wzrokowego. Zagadnienia te bowiem zostały omówione w odrębnym artykule.¹ Przypomnijmy tylko, iż światło wpada do oka przez przezroczystą rogówkę. Strukturę tę charakteryzuje współczynnik załamania $n = 1,376$ i zdolność skupiająca +43 dioptrie, co stanowi około 2/3 mocy

¹ T. Kubiak, Od biofizyki układu wzrokowego do złudzeń optycznych, Fizyka w Szkole z Astronomią, nr 6 (2019), s. 4-10.



Rys. 1. Schemat budowy gałki ocznej.

łamiącej całego oka. Promienie świetlne ulegają zatem na niej refrakcji a następnie biegną dalej przez ciecz wodnistą ($n = 1,336$), wypełniającą komorę przednią oka.

Ilość światła, która przedostaje się do wnętrza gałki ocznej, jest regulowana przez tęczówkę. Jej funkcjonowanie porównuje się często do działania przysłony w aparacie fotograficznym. Tak jak nachodzące na siebie metalowe listki regulują wielkość apertury, a tym samym ilość światła padającego na element światłoczuły w aparacie, tak dwa układy działających antagonistycznie mięśni zmieniają wymiary żrenicy, czyli okrągłego otworu w tęczówce. Nawet w warunkach domowych, przy użyciu lampki ze ściemniaczem oraz lustra, można łatwo sprawdzić doświadczalnie rozszerzanie się żrenicy przy słabym oświetleniu i wyrażanie jej zwężanie przy ekspozycji na silne światło.

Tytułem dygresji warto dodać, że natężenie oświetlenia zewnętrznego nie jest jedynym czynnikiem wpływającym na szerokość żrenic. I niekoniecznie musimy zaraz odwoływać się do działania środków psychoaktywnych. Humanisci mawiają bowiem, iż oczy są zwierciadłem duszy. Zainspirowani tym powiedzeniem fizjologów i psychologów w swoich badaniach udowodnili pozytywny odbiór wrażeń wzrokowych, przejawiający się właśnie w stopniu rozszerzenia żrenic.² Z własnych obserwacji wiemy natomiast, że powiększają się one również w sytuacji zgoła odmiennej, tj. u osób ogarniętych silnym lękiem (w rezultacie działania układu współczulnego). Wracając do fizyki widzenia, wypada jeszcze nadmienić, iż średnica żrenicy zmniejsza się także, gdy zbliżamy obserwowany obiekt do oka (podczas akomodacji).

Układ optyczny naszego oka wykazuje nastawność, czyli umożliwia nam ostre widzenie przedmiotów znajdujących się w różnych odległościach. W tym przypadku kluczowa jest oczywiście rola soczewki. Dwuwypukła, przezroczysta struktura, zbudowana z warstw charakteryzujących się niejednakowymi gęstościami optycznymi, może bowiem zmieniać swój kształt na skutek działania mięśni rzęskowych.³ Gdy patrzymy w dal, są one rozluź-

nione a soczewka spłaszczona. Jeśli natomiast obserwujemy obiekt położony blisko, mięśnie rzęskowe napinają się a soczewka przyjmuje bardziej uwypuklony kształt. Dzięki temu zwiększa się załamanie promieni świetlnych, zmienia ogniskowa a tym samym zdolność skupiająca układu.

Jeśli oko nie akomoduje (moc układu optycznego jest najmniejsza)⁴, na jego osi optycznej można zlokalizować punkt, który jest ostro odwzorowany na siatkówce. Nazywamy go punktem dali P_D , a odwrotność jego odległości od wierzchołka rogówki s_D określamy mianem refrakcji oka R (wyrażanej w dioptriach):

$$R = \frac{1}{s_D}$$

Dla oka miarowego P_D zlokalizowany jest w nieskończoności (patrz rys. 2a), zatem refrakcja wynosi 0. Nietrudno domyślić się, iż dzięki nastawności oka w sposób ostry mogą zostać odwzorowane wszystkie obiekty położone pomiędzy punktami dali i bliży. W tym ostatnim przypadku soczewka uzyskuje największą moc, czyli mamy do czynienia z okiem maksymalnie akomodującym. Warto również dodać, że zdolność nastawczą oka charakteryzuje amplituda akomodacji A_A . Wyznaczamy ją, stosując prostą formułę:

$$A_A = \frac{1}{s_D} - \frac{1}{s_B} = R - \frac{1}{s_B}$$

s_D i s_B oznaczają tu odpowiednio odległości punktu dalekiego i bliskiego oka. Ostatni z wymienionych punktów u dzieci znajduje się kilka centymetrów od oka, ale z wiekiem się od niego oddala, co oczywiście skutkuje zmniejszeniem zakresu akomodacji.

W ramach przypomnienia podstawowych informacji o funkcjonowaniu układu wzrokowego trzeba jeszcze wspomnieć, że promienie świetlne, których kierunek został zmieniony przez soczewkę, biegną dalej przez galaretowate ciało szkliste ($n = 1,336$) i finalnie trafiają na siatkówkę, gdzie prawidłowo powinno znajdować się ognisko układu optycznego oka. Na siatkówce zlokalizowane są bowiem receptory wzrokowe, czyli pręciki i czopki. To właśnie one, przy wsparciu innych wyspecjalizowanych komórek, odbierają bodźce świetlne i zamieniają je na impulsy elektryczne, które przez nerw wzrokowy przesyłane są do mózgu.

Wady refrakcji

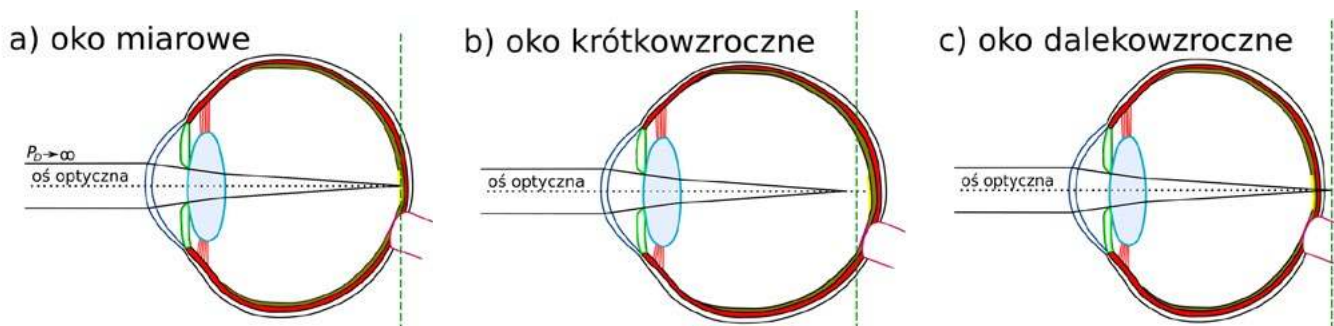
Ostre widzenie z różnych odległości zdecydowanie ułatwia codzienne funkcjonowanie. Jest ono również niezbędnym atrybutem przedstawicieli wielu zawodów, np. pilotów, kierowców, operatorów maszyn czy osób wykonujących prace precyzyjne. Niestety, badania statystyczne pokazują, iż ponad 54% Polaków ma stwierdzoną wadę lub chorobę oczu.⁵ Oczywiście najczęściej mamy do czynienia z tzw. ametropią. Niemiarkowość ta wynika

² Zainteresowanych odsyłam do artykułu: K. Kuraguchi, K. Kanari, *Enlargement of female pupils when perceiving something cute*, Scientific Reports (2021), 11:23367.

³ Współczynnik załamania w części centralnej (tzw. jądrze) soczewki przyjmuje wartość $n = 1,41$, natomiast w jej regionach zewnętrznych $n = 1,33$.

⁴ Gdy soczewka nie akomoduje, jej zdolność skupiająca wynosi ≈ 19 dpt., zatem całkowita moc układu optycznego oka wynosi wówczas ≈ 62 dpt.

⁵ Informacja na podstawie: <https://pap-mediroom.pl/zdrowie-i-styl-zycia/wielkie-badanie-wzroku-coraz-wiecej-polakow-ze-stwierdzona-wada-lub-choroba> (dostęp 15.09.2023).



Rys. 2. Bieg promieni świetlnych przez oko: a) miarowe (emmetropia); b) krótkowzroczne (miopia); c) dalekowzroczne (hiperopia).

z niedostosowania zdolności skupiającej układu optycznego oka do długości gałki ocznej danego człowieka, przez co promienie świetlne skupiane są, w zależności od rodzaju wady, przed albo za siatkówką.

W dobie pracy przy komputerze oraz nagminnego korzystania z urządzeń mobilnych najpowszechniejszą wadą refrakcji stanowi krótkowzroczność, określana przez fachowców również terminem miopia. Nietrudno domyślić się, iż osoby cierpiące na tę przypadłość nie widzą ostro obiektów znajdujących się w oddaleniu. Równoległa wiązka światła, biegnąca z założenia z nieskończoności, skupiana jest bowiem w ich oku przed siatkówką (patrz Rys. 2b). Przyczyną może być: zbyt długa gałka oczna (wydłużenie już o 1 mm w stosunku do jej standardowego osiowego wymiaru ≈ 24 mm objawia się wadą refrakcji -3 dpt); nadmiernie wypukła rogówka lub soczewka (zmniejszenie promienia krzywizny rogówki o 1 mm skutkuje wadą -6 dpt) albo zwiększony współczynnik załamania światła w rogówce lub soczewce. Co ciekawe, u cukrzyków hiperlikemia wywołuje obrzęk osmotyczny soczewki i przesuwają refrakcję w stronę krótkowzroczności.

W przypadku miopii zdefiniowany wcześniej punkt dali P_D leży w skończonej odległości przed okiem, stąd odległość s_D a tym samym refrakcja są ujemne. Wady do -3 dpt, klasyfikowane jako niskie, pojawiają się najczęściej w wieku szkolnym, gdy dzieci zmuszone są przez wiele godzin skupiać wzrok na przedmiotach położonych blisko, np. książkach, zeszytach albo tabletach, pracując dodatkowo w pomieszczeniach ze sztucznym oświetleniem. Tego typu krótkowzroczność szkolna stabilizuje się zazwyczaj około 18-20 roku życia, czyli wraz z zakończeniem wzrostu gałki ocznej.

Z kolei wysoka miopia (wada większa niż -6 dpt) rozwija się zazwyczaj już w bardzo wczesnym dzieciństwie na skutek kombinacji czynników genetycznych i środowiskowych. Miłośników historii zainteresuje na pewno fakt, że mechanizm krótkowzroczności wraz ze sposobem korekcji tej wady za pomocą soczewek wklęsłych opisał już w 1604 roku wybitny niemiecki astronom Johannes Kepler.⁶ Uczony ten nie ograniczał się jednak tylko do dyskusowania o problemach z ostrością wzroku. Interesowała go bowiem również anatomia prawidłowa gałki ocznej oraz szeroko pojęta fizyka widzenia.

Niewyraźny obraz oddalonych obiektów to oczywisty objaw krótkowzroczności. U części społeczeństwa mamy jednak do czynienia z sytuacją zgoła odmienną. Niektórzy widzą bowiem stosunkowo dobrze z daleka, ale mają wyraźny problem ze skupianiem wzroku na przedmiotach umiejscowionych blisko. Tak dzieje się m.in. w przypadku dalekowidzów. W nadwzroczności, określanej również przez fachowców jako hiperopia czy hipermetropia, równoległa wiązka światła w oku nieakomodującym ogniskowana jest za siatkówką (Rys. 2c). Przyczyna tkwi w zbyt krótkiej gałce ocznej (wówczas normalna moc układu optycznego okazuje się dla niej zbyt słaba) albo niedostatecznie zakrzywionej rogówce lub soczewce (załamanie promieni świetlnych jest niewystarczające). Bliskie obiekty są zatem nieostre, a dalej położone przedmioty widziane zdecydowanie lepiej, ale kosztem wzmożonej eksploatacji mięśni rzęskowych, zmieniających krzywiznę soczewki.

Długotrwałe nadmierne napięcie akomodacji (tym większe, im bliżej znajduje się oglądany przedmiot), będące próbą skompensowania nieskorygowanej nadwzroczności, prowadzi do szeregu dalszych dolegliwości, m.in. zmęczenia oczu czy bólów głowy. W przypadku oka nadwzrocznego punkt dali P_D jest pozorny, leży za okiem w miejscu przecięcia przedłużeń ogniskowanych promieni z osią optyczną. Odległość s_D a tym samym refrakcja są zatem dodatnie. Za niską dalekowzroczność uznaje się zazwyczaj wadę do $+2,5$ dioptrii, natomiast za wysoką tę, wynoszącą ponad $+6$ dpt.

Co ciekawe, u noworodków nadwzroczność występuje fizjologicznie ze względu na mały wymiar osiowy oka. W pierwszych latach życia, wraz z szybkim wzrostem ciała, gałka oczna powinna jednak osiągnąć już odpowiednią długość.⁷ Dodatkowo u dzieci występuje duży zakres akomodacji, który pozwala nawet do pewnego stopnia kompensować nadwzroczność. Ze względu na obciążenia środowiskowe czy genetyczne może jednak ujawnić się niemiarkowość zdecydowanie większa niż wynikałoby to z naturalnego etapu rozwoju młodego organizmu. Poglębiająca się wada wymaga oczywiście odpowiedniej korekcji, o czym będzie jeszcze mowa w dalszej części artykułu. Warto jednak jeszcze wspomnieć, iż pionierskie badania nad dalekowzrocznością prowadził już w XIX

⁶ W pracy Keplera „Astronomiae Pars Optica” znajdowały się m.in. rysunki obrazujące funkcjonowanie narządu wzroku.

⁷ Do około 8 roku życia zachodzi proces tzw. emmetropizacji, czyli stopniowego zmniejszania się nadwzroczności.

stuleciu holenderski fizjolog i oftalmolog Franciscus Cornelis Donders. Uważa się, że jako pierwszy oficjalnie rozróżnił on nadwzroczność od starczowzroczności, określanej dziś mianem presbiopii.

Każdy z czytelników zaobserwował zapewne, że część osób po 40-tym roku życia oraz niemal wszyscy seniorzy używają okularów do czytania. Mają oni bowiem kłopot ze skupianiem wzroku na małych literkach w książkach czy na wyświetlaczach urządzeń mobilnych. Początkowo próbują radzić sobie z problemem, oddalając znacznie tekst od oczu, co przez osoby związane zawodowo z optometrią nazywane jest potocznie syndromem „zbyt krótkiej ręki”. Przyczyną tych problemów, określaną właśnie jako presbiopia, jest stopniowa, naturalna utrata elastyczności przez soczewkę wewnątrzgałkową oraz osłabienie mięśni oka. Grubsza i sztywniejsza soczewka staje się mniej podatna na zmianę kształtu, dlatego dostosowanie ostrości wzroku wymaga dłuższego czasu oraz większego wysiłku akomodacyjnego.

U osób ze starczowzrocznością problemy z ostrym widzeniem pogłębia dodatkowo słabe światło, przy którym źrenica jest bardziej rozszerzona. Dlatego ważne jest zapewnienie odpowiedniego oświetlenia czytanego tekstu oraz stosowanie okularów korekcyjnych. Generalnie uznaje się, że presbiopia nie jest chorobą ani wadą refrakcji związaną z nieprawidłową budową gałki ocznej, a raczej procesem fizjologicznym, będącym następstwem starzenia się organizmu i ograniczenia zdolności akomodacyjnej oczu. Zapewne dotknie w przyszłości każdego z nas, lecz osoby nadwzroczne odczuwają jej następstwa szybciej.

Wracając do tematyki wad refrakcji, koniecznie należy wspomnieć jeszcze o niezborności, która zapewne znana jest szerokiemu gronu czytelników pod nazwą astygmatyzm. W literaturze popularnej często możemy spotkać się z obrazowym porównaniem kształtu gałki ocznej obciążonej tą niedoskonałością optyczną do piłki do rugby, podczas gdy prawidłowo powinna ona być kulista. Niesferyczność centralnej części rogówki stanowi podstawową przyczynę astygmatyzmu.

Trzeba jednak wspomnieć, że za niezborność odpowiadać może również deformacja soczewki lub jej niewłaściwe ustawienie, aczkolwiek sytuacja ta ma miejsce bardzo rzadko. Generalnie brak symetrii obrotowej rogówki sprawia, iż moce optyczne układu są różne w płaszczyznach prostopadłych do osi optycznej (poszczególnych przekrojach południkowych gałki ocznej). Obraz nie zostanie zatem dobrze odwzorowany, gdyż światło dla różnych południków zostanie skupione w innych miejscach. Zamiast jednego punktowego ogniska zlokalizowanego na siatkówce możemy więc mieć do czynienia z wieloma, położonymi w odmiennych od niej odległościach. Dwa skrajne ogniska odpowiadają najczęściej prostopadłym do siebie przekrojom głównym, dla których występuje największa i najmniejsza moc optyczna. Różnica obu tych mocy refrakcyjnych stanowi właśnie miarę astygmatyzmu.

Co ciekawe, niezborność współwystępuje z krótko- albo nadwzrocznością. Dla przykładu astygmatyzm krótkowzroczny zwykły cechuje się położeniem jednego

z ognisk na siatkówce (w tym przekroju oko jest miarowe) a drugiego przed nią (dla tego przekroju występuje krótkowzroczność). Osoby cierpiące na astygmatyzm często skarżą się na bóle głowy, mają problemy z wyraźnym widzeniem, kontury obiektów wydają się być dla nich zniekształcone a światła lamp oraz reflektorów w nocy rozmyte. Regularne badanie wzroku oraz odpowiednio dobrana korekcja optyczna są niezbędne, aby osoby z niezbornością odzyskały komfort widzenia i pozbyły się dodatkowych dolegliwości towarzyszących tej wadzie.

Inne problemy ze wzrokiem

Omówione dotychczas wady refrakcji to najczęstsze problemy, z jakimi zgłaszamy się do optometrysty. Nie są to jednak wszystkie dysfunkcje układu wzrokowego, na które warto spojrzeć z punktu widzenia biofizyki. Jedną z najbardziej kłopotliwych i uważanych przez pacjentów za wstydliwą dolegliwością jest chociażby zez, nazywany też strabizmem. Objawia się on niezdolnością do jednoczesnego skierowania spojrzenia obojga oczu na jeden punkt. Oko dominujące patrzy wówczas na wybrane miejsce, natomiast drugie w zupełnie innym kierunku. U zdrowego człowieka spoglądającego w dal osie widzenia obojga oczu muszą być praktycznie równoległe. Takie wzorcowe ustawienie oczu określa się mianem ortoforii. Z kolei w obecności bodźca do fuzji osie widzenia powinny przecinać się w punkcie fiksacji. O zbieżności osi, która wywołana jest obserwacją przedmiotów bliskich, informuje wyrażana w dioptriach pryzmatycznych konwergencja K .

$$K = \frac{PD}{s}$$

gdzie: PD – odległość między środkami gałek ocznych, s – odległość punktu fiksacji od prostej łączącej środki obojga oczu. Warto zauważyć, że przy zbliżaniu się obserwowanego obiektu konwergencja K oraz akomodacja A wzrastają. Niestety, ruch gałek ocznych, z których każda kontrolowana jest przez sześć mięśni zewnątrzgałkowych, czasami nie jest prawidłowy. Niektóre z mięśni mogą bowiem być osłabione, porażone, pracować w sposób nieskoordynowany lub posiadać uszkodzone unerwienie.

Specjaliści wyróżniają wiele typów zezów jawnych, m.in. zbieżny (czyli ezotropię, gdzie gałka oczna odchyła się w stronę nosa), rozbieżny (egzotropię, z okiem skierowanym na zewnątrz twarzy), pionowy ku górze (hipertropię) czy pionowy ku dołowi (hipotropię) a także zezy skośne. Istotę opisanych wad najlepiej obrazuje Rys. 3. Dla zainteresowanych warto wspomnieć, że czasami podczas badania okulistycznego, gdy jedno z oczu zostanie zakryte, a tym samym widzenie obuoczne wyłączone, ujawnia się zez ukryty (znany też jako foria lub heteroforia). Do jego powstania może przyczyniać się wada refrakcji, występująca w jednym z oczu, stąd zawsze należy pamiętać o właściwej korekcji.

Wśród najpoważniejszych chorób oczu, z jakimi pacjenci zgłaszają się do okulisty, wymienia się jaskrę. Nieleczona prowadzi do nieodwracalnego uszkodzenia nerwu wzrokowego. Ten ważny nerw, który przesyła

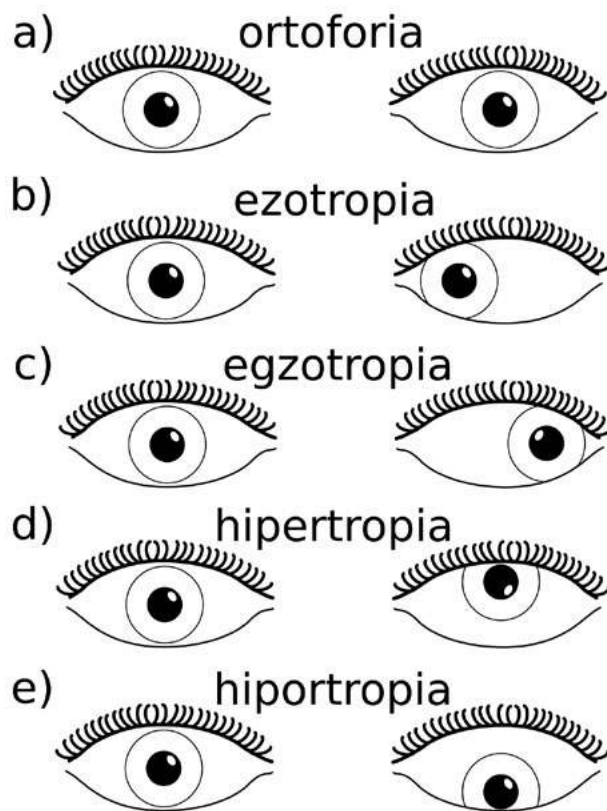
impulsy z receptorów siatkówki do mózgu, scharakteryzował po raz pierwszy w XVIII w. włoski fizyk Felice Fontana. Badania teoretyczne i eksperymentalne nad biofizycznym mechanizmem powstawania jaskry prowadzono jednak znacznie wcześniej, a ich wyniki opublikował w 1622 roku angielski oftalmolog Richard Bannister.

Zablokowanie kanalików Schlemma, będących naturalnym odpływem cieczy wodnistej, wypełniającej komorę przednią oka, powoduje zwiększenie ciśnienia płynu wewnątrz gałki ocznej. Prawidłowo ciśnienie wewnątrzgałkowe powinno zawierać się w przedziale 8-21 mmHg. Jego wzrost na skutek upośledzonego drenażu, a tym samym nadmiernego gromadzenia się cieczy wodnistej, skutkuje uciskiem i postępującym uszkodzeniem nerwu wzrokowego. Pacjent zauważa wówczas objawy w postaci pojawiających się niczym halo obwódok wokół źródeł światła a także ubytku obwodowego pola widzenia.

W tym miejscu warto przytoczyć pewną ciekawostkę. Przypuszcza się bowiem, że niezwykle sposób przedstawiania świata na obrazach Vincenta van Gogha wynikał właśnie z jego problemów z narządem wzroku. Oprócz deuteranopii⁸, o której świadczą mają stosowane przez niego kolory, wymienia się też jaskrę. Dowodem wydaje się „efekt halo” wokół ciał niebieskich na obrazach „Gwiazdzista noc” i „Droga z cyprysem i gwiazdą” czy też wokół gazowych lamp z dzieła „Nocna kawiarnia”.

Zaćma, znana również pod nazwą katarakta, stanowi kolejną powszechnie spotykaną chorobę degeneracyjną oczu. U osób po 65 roku życia może bowiem dojść do zmętnienia soczewki w jednym albo w obojgu oczu. Stopniowa utrata przejrzystości tej struktury objawia się poprzez różne symptomy, m.in. mgłę zaburzającą ostrość widzenia, zmianę zabarwienia soczewki, problemy z widzeniem barw czy zwiększoną wrażliwość na światło. Ostatecznie prowadzi do ślepoty. Wśród czynników sprzyjających rozwojowi zaćmy wymienia się m.in. wolne rodniki, czyli atomy lub cząsteczki posiadające jeden lub więcej niesparowanych elektronów a także różne przyczynki środowiskowe, np. narażenie na promieniowanie UV, palenie tytoniu, itp.⁹ Trzeba również mieć na uwadze, że oprócz zaćmy starczej istnieje także jej postać wrodzona, pojawiająca się najczęściej u dzieci już po narodzeniu.

Osoby w podeszłym wieku narażone są na jeszcze inne, poważne schorzenie oczu, t.j. zwyrodnienie plamki żółtej, oznaczane skrótem AMD od ang. *Age-related Macular Degeneration*. Jak sama nazwa wskazuje, problem ten dotyczy niedużego obszaru w centrum siatkówki (patrz ponownie Rys. 1). To tutaj znajduje się strefa najostrzejszego widzenia, czyli największe skupisko czopków. Warto przypomnieć, że to właśnie dzięki tym światłoczułym receptorom przy dobrym oświetleniu dostrzegamy kolory oraz szczegóły obrazu. AMD powoduje stopniową utratę widzenia środkowego (występuje zwiększający się z czasem mroczek centralny przy zachowanym widzeniu ob-



Rys. 3. Ortoforia (a) oraz różne rodzaje zeza (b-e).

wodowym) i barwnego (kolory przypominają te na wyblakłych fotografiach). Dodatkowo pacjenci skarżą się na zniekształcone widzenie przedmiotów i falowanie linii, jakie w normalnych warunkach powinni uznać za proste. Zwyrodnienie plamki żółtej to choroba wieloczynnikowa. Postać wysiękowa ma związek z wytwarzaniem w siatkówce nieprawidłowych, nieszczelnych naczyń krwionośnych, które są odpowiedzialne za obrzęki oraz wylewy krwi. Łagodniejsza odmiana sucha objawia się postępującymi uszkodzeniami komórek nabłonka barwnikowego siatkówki w wyniku nagromadzenia się złogów tłuszczowych. Oczywiście oba procesy są skomplikowane i wieloetapowe. Według naukowców mechanizmy rodnikowe również odgrywają w nich pewną rolę.

Diagnostyka wad wzroku i chorób oczu

Wiedząc już, jakie problemy z układem wzrokowym mogą skłonić nas do wizyty u optometrysty czy okulisty możemy przyjrzeć się fizycznym metodom stosowanym w diagnostyce wad refrakcji oraz chorób oczu.

Zapewne każdy z czytelników w swoim życiu nawet kilkukrotnie miał wykonywane badanie ostrości wzroku. Oczywiście można przeprowadzać je na kilka sposobów. Tradycyjnie wykorzystuje się tzw. tablicę Snellena¹⁰, czyli białą planszę z jedenastoma rzędami czarnych liter (rys. 4). W standardowej wersji zawiera ona 9 różnych znaków alfabetu (C, D, E, F, L, O, P, T, Z), przy czym na samej górze umieszczone jest pojedyncze duże E. Literki w kolejnych liniach (gdy patrzymy ku dołowi) stają się

⁸ Deuteranopia to odmiana ślepoty barw, w której osoba nie rozpoznaje koloru zielonego lub myli go z czerwonym.

⁹ O wykrywaniu wolnych rodników można przeczytać w artykule: T. Kubiak, Spektroskopia elektronowego rezonansu paramagnetycznego i jej przykładowe zastosowania w biofizyce i fizyce medycznej, *Fizyka w Szkole z Astronomią*, nr 2 (2016), s. 4-7.

¹⁰ Nazwa upamiętnia działającego w Utrechcie holenderskiego oftalmologa Hermanna Snellena, który zaproponował tablicę w 1862 r.

odległość w m		ostrość wzroku
50	E	0,1
25	F P	0,2
16,5	T O Z	0,3
12,5	L P E D	0,4
10	P E C F D	0,5
7,1	E D F C Z P	0,7
6,2	F E L O P Z D	0,8
5,0	D E F P O T E C	1,0
4,0	L E F O D P C T	1,3
3,3	F D P L T C E O	1,5
2,5	F E Z O L C F D T	2,0

Rys. 4. Tablica Snellena służy do badania ostrości wzroku.

coraz mniejsze. Po lewej stronie dodatkowo naniesione są wartości odległości, z jakich posiadacze dobrego wzroku lub poprawnej korekcji powinni bez problemu przeczytać znaki z danej linii.

Istnieją też alternatywne wersje tablic, np. przeznaczone dla analfabetów oraz małych dzieci. Wówczas optotypami, czyli znakami graficznymi, które trzeba rozpoznać, mogą być obrazki (np. ustandaryzowane symbole LEA: kwadrat, kółko, jabłko, domek)¹¹ albo „widełki”, czyli E ułożone w różnych orientacjach. Nietrudno domyślić się, że współcześnie zawieszane na ścianie, tekturowe plansze Snellena coraz częściej zastępuje się zamontowanym prostopadle do osi widzenia pacjenta i sterowanym pilotem monitorem optotypów. Łatwo wówczas w zależności od potrzeb zmieniać wyświetlane symbole na litery, cyfry czy obrazki. Badania wskazują, że optymalne natężenie oświetlenia podczas ustalania ostrości wzroku u osób z krótkowzrocznością zawiera się w przedziale 400 – 500 lx.¹² Pacjent ma za zadanie odczytywać każdym okiem osobno (drugie pozostaje wówczas przesłonięte) kolejne rzędy optotypów z odległości 5 m (w krajach anglosaskich 20 stóp, czyli ≈ 6 m). Wynik testu, czyli tzw. visus, oznaczany jako V , zapisuje się w postaci ułamka:

$$V = \frac{d}{D}$$

Jego licznik stanowi właśnie odległość d , z jakiej literki były czytane, natomiast mianownik wskazuje dystans D , dla którego oko miarowe widziałoby prawidłowo ostatni z właściwie rozpoznanych przez badanego symboli na tablicy testowej.

Dobra ostrość wzroku to wynik: 5/5; 20/20 albo 1,0 w przypadku stosowania ułamków dziesiętnych. Generalnie im większa liczba pod kreską ułamkową (albo mniejsza cyfra po przecinku), tym gorszy wzrok pacjenta. Przykładowo wynik 20/40 (czyli 0,5) wskazuje, iż człowiek dotknięty wadą refrakcji musi znaleźć się dwukrotnie bli-

żej obiektu, żeby widzieć go tak, jak osoba pozbawiona tej przypadłości. Rezultat poniżej 0,1 wskazuje tzw. ślepotę praktyczną, natomiast 0 oznacza ślepotę całkowitą. Dla zainteresowanych warto wspomnieć, że w przypadku niewielkiego odsetka populacji można spotkać się z sytuacją zgoła odmienną, tzn. z wynikiem na poziomie 20/10, czyli 2,0. Oznacza to nadzwyczajnie dobry wzrok. Jego posiadacz z odległości 5 m widzi tak, jak inni zdrowi ludzie z 2,5 m, ma zatem szczególne predyspozycje do wykonywania zawodów typu pilot wojskowy.

Tablice Snellena oprócz sprawdzania ostrości wzroku pomagają także dobrać optymalną korekcję. W tym przypadku potrzebny jest dodatkowo zestaw szkieł próbnych. W specjalnej kasce optometrysta posiada ich nawet kilkaset. Cechują się różnorodnymi wartościami sfery oraz cylindra. Są też soczewki pryzmatyczne. Wybrane szkła umieszcza się w specjalnych oprawkach o regulowanym rozstawie (Fot. 1) i zakłada pacjentowi. Procedurę powtarza się, aż korekcja zostanie uznana za najlepszą. Współcześnie coraz częściej oprawki próbne zastępuje profesjonalne urządzenie zwane foropterem (Fot. 2). Jego zamontowana na ruchomym ramieniu głowica zawiera zestaw sferycznych oraz torycznych soczewek a także filtrów i pryzmatów. W manualnych aparatach starszej generacji były one zmieniane i ustawiane w osi optycznej oka pacjenta ręcznie. Najnowsze foroptyery elektroniczne są natomiast sterowane przy użyciu komputera, umożliwiając zdecydowanie szybszy dobór korekcji okularowej.

Podczas badania optometrycznego, oprócz zweryfikowania zdolności rozpoznawania odległych obiektów, konieczne jest również sprawdzenie ostrości wzroku do blizy. Ma to szczególne znaczenie u osób starszych, potencjalnie dotkniętych prezbiopią. Pacjentowi pokazuje się próbnik z linijkami tekstu napisanymi czcionkami o różnej wielkości. Musi je przeczytać z odległości 30-40 cm. Wyrazy na białym tle tablicy utworzone są z czarnych, standaryzowanych liter, przy czym najdrobniejsze znaki znajdują



Fot. 1. Oprawy próbne ułatwiają właściwy dobór soczewek korygujących wadę wzroku.

¹¹ Symbole LEA do pediatrycznego badania wzroku wprowadziła fińska okulistka Lea Hyvärinen.

¹² Na podstawie: L. P. Tidbury i inni, Fiat Lux: the effect of illuminance on acuity testing, *Graefes Arch. Clin. Exp. Ophthalmol.* 254 (2016), 1091-1097.

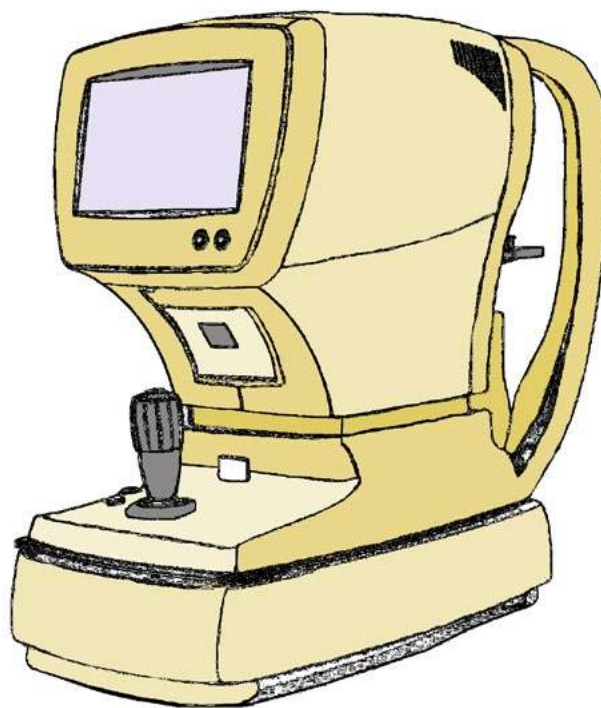


Fot. 2. Głowica foroptera zawiera zestaw soczewek, filtrów i pryzmatów.

się u góry a największe na dole. Wynik badania podaje się w postaci ułamka, którego licznik wskazuje wartość liczbową przypisaną najmniejszej poprawnie odczytanej czcionce, a mianownik oznacza dystans dzielący tablicę od oczu. Idealny rezultat to ostrość do blizy 0,5/30.

Na sztydach gabinetów okulistycznych i salonów optycznych często możemy napotkać informację, że przeprowadza się w nich komputerowe badanie wzroku. Urządzeniem wykorzystywanym podczas tej procedury jest autorefraktometr (Rys. 5). Pacjent siada, opiera brodę oraz czoło o podpórki i patrzy przez celownik na określony obrazek (np. domek), a maszyna automatycznie ustala charakterystykę jego wzroku. Przykładowy wynik dla pacjenta z astygmatyzmem przedstawia Rys. 6 (użyte symbole: sfery (sph), cylindra (cyl) i osi (ax) zostaną omówione w dalszej części artykułu).

Oczywiście tego typu aparat diagnostyczny wykorzystuje w swoim działaniu osiągnięcia takich dyscyplin jak: fizyka, informatyka czy automatyka. Autorefraktometry wysyłają wiązkę światła podczerwonego na siatkówkę osoby badanej i analizują promienie odbite. Wykorzystanie perspektywy zbieżnej sprawia, że wyświetlany obiekt wydaje się znajdować w dali, natomiast pojawiające się czasami zamglenie pomaga rozluźnić akomodację. Niekiedy nie w pełni się to udaje i u osób młodych autorefraktometry zamkniętego pola wskazują na bardziej ujem-



Rys. 5. Autorefraktometr służy do komputerowego badania wzroku.

ną wartość refrakcji niż jest ona w rzeczywistości. Warto zaznaczyć, iż wynik pomiarów automatycznych stanowi dla specjalisty jedynie wskazówkę przy dalszym doborze właściwej korekcji wzroku pacjenta.

Dzięki zintegrowaniu autorefraktometru z keratometrem możliwe jest wyznaczenie przedniej krzywizny rogówki. Odbywa się to na podstawie pomiaru wielkości obrazu odbitego, gdyż wspomniana struktura anatomiczna przypomina swoim kształtem zwierciadło wypukłe. Jeśli dwa centralne promienie krzywizny rogówki (odpowiednio najbardziej płaski i stromy) różnią się, wskazuje to na astygmatyzm rogówkowy. Co ciekawe, idea keratometrii pojawiła się już w 1619 r. za sprawą niemieckiego astronoma Christopha Scheinera, ale na jej praktyczną realizację trzeba było poczekać wiele lat. W 1880 roku Portugalczyk Antonio Plácido wprowadził bowiem jakościową keratometrię ręczną. Wykorzystuje ona specjalny krążek z naprzemiennie ułożonymi czarnymi oraz białymi okręgami, których odbicia od przedniej powierzchni rogówki obserwuje się podczas badania.

Francuski inżynier i oftalmolog Louis Émile Javal wraz ze swym uczniem Hjalmarem Augustem Schiøtzem konstruowali natomiast w 1881 r. urządzenie nazywane dziś keratometrem (lub oftalmometrem) Javala. Ten wyposażony w żarówkę aparat rzutuje na rogówkę dwie figury schodkowe a operator stara się ustawić w jednej linii i zetknąć ich obrazy.¹³ Dla sferycznej rogówki odbicia figur powinny stykać się również po obrocie ramion przyrządu o 90°. Jeśli się rozchodzą albo nakładają, mamy do czynienia odpowiednio z astygmatyzmem odwrotnym albo prostym. Warto dodać, że obecnie keratometr Javala ma

¹³ Z budową keratometru Javala zapoznać się można na stronie Muzeum Gdańskiego Uniwersytetu Medycznego: <https://gumed.edu.pl/64893.html> (dostęp z 12.09.2023)

NAME : _____

DATE : _____

[REF	DATA]		
UD:	12.00	CL:	-
<R>	SPH	CYL	AX
	-0.75	-0.75	101
	-1.00	-0.50	97
	-1.00	-0.50	96
AUE	-1.00	-0.50	98
<L>	SPH	CYL	AX
	-0.75	-0.75	74
	-0.75	-0.75	79
	-1.00	-0.50	65
AUE	-0.75	-0.75	72
PD = 64mm			

Rys. 6. Wydruk z badania autorefraktometrem pacjenta z astygmatyzmem. Podano wartości sfery (sph), cylindra (cyl) i osi (ax).

przede wszystkim znaczenie historyczne, gdyż specjaliści wykorzystują urządzenia bazujące na systemach komputerowych oraz sensorach optycznych.

Wspomniany wcześniej norweski naukowiec Hjalmar Schiøtz przyczynił się do rozwoju okulistyki poprzez jeszcze jeden wynalazek – tonometr impresyjny służący do pomiaru ciśnienia wewnątrzgałkowego. Poprzez przyłożenie siły o znanej wartości pozwalał on zmierzyć odkształcenie (wgłębienie) znieczulonej rogówki. Oczywiście opór, jaki stawiała rogówka wciskanemu w nią, obciążonemu ciężarkiem metalowemu trzpieniowi, zależał od ciśnienia w oku. Gdy było ono wysokie, rogówka nie odkształcała się przy standardowej masie ciężarka i trzeba było zwiększyć obciążenie.

Nietrudno domyślić się, że ta prymitywna metoda w dzisiejszych czasach nie ma już racji bytu w diagnostyce jaskry. Zastąpiły ją m.in. tonometria aplanacyjna Goldmanna oraz bezkontaktowa (*ang. air-puff*). Pierwsza z wymienionych wymaga zabarwienia fluoresceiną (barwnikiem fluorescencyjnym) filmu łzowego powlekającego rogówkę. Pod wpływem nacisku tworzy on fluorujący w świetle niebieskim, okrągły menisk. Jest on widoczny w mikroskopie (w postaci zielonożółtych półokręgów, które należy odpowiednio zestawić), gdy do rogówki przykładana jest końcówka tonometru z podwójnym pryzmatem. W myśl reguły zaproponowanej w oparciu o prace fizyków Armanda Imberta i Adolfa Ficka ciśnienie wywierane przez ciecz znajdującą się we wnętrzu kuli jest równe sile potrzebnej do jej odkształcenia podzielonej przez wielkość spłaszczonej powierzchni. W pomiarach zakłada się zatem, że średnie ciśnienie wywierane na rogówkę przez tonometr aplanacyjny Goldmanna odpowiada ciśnieniu wewnątrzgałkowemu. W praktyce należy jednak pamiętać o grubości rogówki centralnej.

W przypadku drugiej ze współcześnie stosowanych metod, czyli tonometrii bezkontaktowej spłaszczenie rogówki wywoływane jest skierowanym na jej powierzchnię strumieniem powietrza. Liczy się potrzebny do uzyskania odpowiedzi rogówki czas działania podmuchu. Za

pomocą odpowiedniego algorytmu przeliczany jest on na ciśnienie. Wprawdzie dokładność badania jest mniejsza, ale szybkość wykonania i nieinwazyjność sprawia, że nadaje się ono do badań przesiewowych.

Wśród przyrządów stosowanych przez optometrystów ważne miejsce zajmuje skiaskop (fot.3b), czyli retinoskop manualny. Pomaga on w sposób obiektywny, tzn. bez aktywnego udziału pacjenta, zdiagnozować wadę jego wzroku. W przypadku skiaskopii dynamicznej pozwala on dodatkowo ocenić odpowiedź oczu na zastosowany bodziec do akomodacji, np. soczewkę lub tekst umieszczony na tabliczce nałożonej na przyrząd. Generalnie niewielkie urządzenie wyposażone jest w źródło rzutujące wiązkę światła na siatkówkę, zwierciadło półprzepuszczalne pozwalające zmieniać jej położenie oraz wizjer obserwacyjny do oglądania obrazu poruszającej się plamki światła odbitego od dna oka. W przypadku oka miarowego obraz



Fot. 3. Przyrządy wykorzystywane przez optometrystów: a) oftalmoskop; b) skiaskop.

plamki powstaje w nieskończoności, przy nadwzroczności za okiem (a refleks porusza się w kierunku zgodnym z emitowaną wiązką), natomiast w krótkowzroczności przed okiem (refleks przemieszcza się w stronę przeciwną niż promień oświetlający).

Wziernikowanie dna oka odbywa się standardowo z wykorzystaniem oftalmoskopu (fot.3a). Ten przyrząd również ma prostą budowę, składa się z uchwyty oraz głowicy wyposażonej w źródło światła, zestaw soczewek powiększających oraz filtrów barwnych. Diagnosta wykorzystuje go przede wszystkim do wykrywania patologicznych zmian w obszarze gałki ocznej, m.in. występowania zwyrodnienia plamki żółtej, zapalenia nerwu wzrokowego czy odwarstwienia siatkówki. Ocena naczyń krwionośnych i struktur anatomicznych wnętrza oka możliwa jest właśnie dzięki zwrotnemu odbiciu światła emitowanego przez głowicę oftalmoskopu.

Zastosowanie specjalistycznych przyrządów oraz testów bazujących na osiągnięciach fizyki wspomaga sprawną diagnostykę rozmaitych chorób bądź umożliwia ustalenie posiadanej przez pacjenta wady refrakcji. W tym ostatnim przypadku końcowym etapem pozostaje jeszcze dobór odpowiedniej korekcji z wykorzystaniem okularów bądź soczewek kontaktowych. Przyjrzyjmy się zatem bliżej tej kwestii.

Korekcja wad refrakcji

Zapewnienie ostrego widzenia u pacjentów z wadą refrakcji możliwe jest dzięki zastosowaniu prawidłowo dobranych soczewek korekcyjnych. Wówczas ich ognisko obrazowe musi znaleźć się w punkcie dalekim oka, tzn. soczewka wraz z układem optycznym oka powinna dawać na siatkówce ostry obraz przedmiotu zlokalizowanego w nieskończoności. Moc takiej soczewki wyliczymy ze wzoru:

$$D_s = \frac{1}{s_D} = \frac{R}{1 - lR}$$

gdzie: l – odległość soczewki od płaszczyzny głównej układu optycznego oka, R – refrakcja oka.

Krótkowzroczność, w której równoległa wiązka światła skupiana jest w oku przed siatkówką, koryguje się za pomocą soczewek rozpraszających (ujemnych). Z kolei w przypadku nadwzroczności, gdzie równoległa wiązka światła w oku nieakomodującym ogniskowana jest za siatkówką, stosuje się soczewki skupiające (dodatnie). Warto nadmienić, że oddalenie soczewki od oka wymusza, w zależności od wady, zwiększenie mocy soczewki „minusowej” albo zmniejszenie mocy soczewki „plusowej”. Oczywiście krzywizna i grubość soczewek dobierane są według stopnia korekcji.

Z ciekawą sytuacją mamy do czynienia w przypadku prezbiopii. Dawniej u osób starszych stosowano przede wszystkim okulary dwuogniskowe z korekcją odpowiednio do patrzenia na dal (górna strefa) oraz do czytania (część dolna). Tego typu rozwiązanie zaproponował już

w XVIII w Benjamin Franklin, który szkła bifokalne uzyskał poprzez połączenie połówek dwóch różnych soczewek okularowych przeznaczonych do dali i bliży.

Motywację dla wybitnego Amerykanina tym razem nie stanowiła ciekawość naukowa, a raczej prozaiczna chęć ułatwienia sobie życia i możliwość zrezygnowania z irytującej go konieczności zmiany okularów. Współcześnie coraz częściej u osób ze starczowzrocznością stosowane są okulary progresywne (wielooogniskowe). W ich przypadku pomiędzy górną strefą do patrzenia w dal a dolną do bliży istnieje łagodne przejście (gradient mocy), dzięki czemu możliwe jest ostre widzenie także z odległości pośrednich.

Osoby cierpiące na astygmatyzm powinny nosić soczewki, które pozwolą połączyć przesunięte względem siebie ogniska, powstające na skutek przejścia promieni świetlnych przez niesferyczną powierzchnię rogówki. Szkła cylindryczne, bo o nich mowa, po raz pierwszy w praktyce zastosował angielski astronom George Biddell Airy już w XIX w. W jednym z przekrojów posiadają one zerową moc, natomiast w drugim moc odpowiadającą wartości astygmatyzmu. Generalnie zdolność zbierająca soczewki cylindrycznej zależy od współczynnika załamania (indeksu) materiału z jakiego została wykonana oraz promienia krzywizny walca, którego pobocznica stanowi jej powierzchnię łamiącą.

Warto jednak zauważyć, iż niezborności często towarzyszy krótko- albo dalekowzroczność. Aby skorygować obie wady, trzeba nie tylko połączyć punkty, w których skupiają się promienie świetlne, ale również to wspólne ognisko przesunąć na siatkówkę. Rozwiązaniem są w tym przypadku soczewki toryczne (fot. 4). W praktyce opisać je można jako złożenie soczewek: sferycznej (sfery) o takich samych mocach we wszystkich przekrojach oraz cylindrycznej (cylindra). Dodajmy jeszcze, że oś cylindra znajduje się w jego przekroju o zerowej mocy. Przyjmuje wartości 0° do 180° i określa ustawienie cylindra w oprawie. Korygowany cylindrycznie południk jest do niej prostopadły.

Po wizycie u optometrysty, gdy stwierdzony zostanie astygmatyzm, pacjent otrzymuje receptę na okulary. Dla każdego oka osobno widnieją na niej wartości sfery, cylindra jak również osi, aby optyk mógł właściwie przy-



Fot. 4. Soczewki sferocylindryczne wykorzystywane są do korekcji astygmatyzmu.

gotować soczewki sferocylindryczne. Przykładowy zapis to: sph -0,25, cyl -1,0, oś 45°, co oznacza: sferę -0,25 dpt, cylinder -1,0 dpt w osi 45°. Co ciekawe, można dokonać przeliczenia mocy szkieł, aby cylinder miał znak przeciwny, w naszym przypadku dodatni. Wówczas sumuje się wartości sfery oraz cylindra, zmienia znak cylindra na przeciwny (bez ingerencji w jego wartość bezwzględną) oraz modyfikuje oś o 90°. Dla powyższego przykładu będziemy zatem mieli: sph -1,25; cyl +1,0; oś 135°. Same soczewki sferocylindryczne mogą mieć różnorodną konstrukcję, najczęściej spotyka się jednak meniskową, w której powierzchnia przednia jest wypukła, a tylna wklęsła. Cylinder soczewki torycznej znajduje się wówczas na tej tylnej powierzchni.

Ponadto przy doborze okularów ważny jest pomiar rozstawu źrenic PD (*ang. pupillary distance*). W zasadzie najlepiej wyznaczyć odległości od centrum każdej z nich do środka nasady nosa (bo twarz nie musi być idealnie symetryczna), aby później właściwie ustawić osie optyczne soczewek na linii widzenia.

W przeszłości soczewki okularowe wykonywano ze szkła, obecnie z nietłukącego się plastiku. Dodatkowo współcześnie na ich powierzchnię nanosi się różnorodne powłoki, m.in. filtr UV, warstwy antyrefleksyjne, antystatyczne, hydrofobowe czy uodporniające na zarysowania. Wiele osób uważa, że okulary nadają twarzy interesujący wygląd. Są jednak i tacy, którzy dostrzegają przede wszystkim ich wady, czyli fakt, że ramki ograniczają pole widzenia, a szkła brudzą się oraz zaparowują, tzn. skrapla się na nich para wodna, gdy osoba nosząca okulary zimą wchodzi do ciepłego pomieszczenia.

Alternatywą dla okularów mogą być soczewki kontaktowe. Nie wszyscy jednak mogą je nosić, przeciwwskazaniem są chociażby alergie, stany zapalne, zespół suchego oka czy po prostu dyskomfort pacjenta. Soczewki kontaktowe muszą być zbudowane z biokompatybilnego materiału, który posiada odpowiednie właściwości optyczne (tj. wysoką przezroczystość oraz stosowny współczynnik załamania) a także mechaniczne (elastyczność i wytrzymałość). Powinny przepuszczać O₂, aby rogówka, na którą są zakładane, pozostała dotleniona.

Na rynku dostępne są soczewki miękkie jednodniowe (jedorazowe, np. silikonowo-hydrożelowe) przeznaczone do noszenia tylko w dzień. W sprzedaży mamy też soczewki wielokrotnego użytku, dwutygodniowe, miesięczne a nawet roczne. Wymagają one jednak pielęgnacji (czyszczenia i dezynfekcji), ściągania na noc i trzymania zanurzonych w specjalnym płynie a także bezwzględnego przestrzegania higieny przy każdym dotknięciu. Miękkie soczewki kontaktowe charakteryzowane są przez szereg parametrów, m.in. grubość, średnicę zewnętrzną oraz promień krzywizny tylnej powierzchni. Co ciekawe, moc plusowej soczewki kontaktowej, oznaczmy ją jako D_{sk} , jest większa od mocy D_o odpowiadającej jej soczewki okularowej:

$$D_{sk} = \frac{D_o}{1 + lD_o}$$

gdzie: l – odległość tylnej powierzchni soczewki okularów od wierzchołka rogówki ($l < 0$).

Z kolei w przypadku krótkowzroczności moc soczewek (minusowych) powinna być mniejsza niż moc szkieł okularowych. Oczywiście soczewki kontaktowe, w zależności od wady, mogą posiadać różną geometrię (sferyczne, toryczne, dwuogniskowe, wieloogniskowe). Interesująca wydaje się koncentryczna budowa tych multifokalnych. Pierścienie o wzrastającym promieniu, tworzące strefy o naprzemiennych mocach do bliży i dali, otaczają centrum przeznaczone do patrzenia w bliż albo dal i ułatwiają ostre widzenie z różnych odległości.

Warto wspomnieć jeszcze o wielorazowych, twardych soczewkach kontaktowych zrobionych z gazoprzepuszczalnego materiału RGP (od *ang. Rigid Gas Permeable*). Mają one zastosowanie w przypadku skomplikowanych wad wzroku (np. silnego, nieregularnego astygmatyzmu). Specjalne, indywidualnie dopasowywane i zakładane tylko na noc twarde soczewki kontaktowe wykorzystuje się również do tzw. ortokorekcji wad refrakcji, takich jak niewielka krótkowzroczność czy astygmatyzm. Pod wpływem nacisku wywieranego przez stabilnoksztaltną soczewkę na centralną część rogówki, stopniowo zmienia się jej krzywizna a tym samym poprawia się wzrok. Niestety efekt nie jest stały i po zaprzestaniu noszenia soczewek (w czasie do kilku tygodni) rogówka powraca do pierwotnego kształtu, a tym samym wada znowu się pojawia.

Okulista oraz optometrysta to niezwykle ciekawe zawody, których podstawę stanowi specjalistyczna wiedza z optyki oraz biofizyki widzenia. Nie bez znaczenia są też informacje z zakresu fizyki polimerów, gdyż na rynku pojawiają się nowsze materiały do produkcji soczewek kontaktowych. Dzięki praktycznemu zastosowaniu osiągnięć nauki udaje się coraz większej grupie ludzi przywrócić komfort widzenia i mam nadzieję, że po przeczytaniu tego artykułu czytelnicy będą mieli tego świadomość.

dr Tomasz Kubiak
biofizyk i fizyk medyczny
Wydział Fizyki UAM w Poznaniu

LITERATURA:

- [1] A. Przekoracka-Krawczyk, P. Jaśkowski, Biofizyka układu wzrokowego [w:] Biofizyka pod red. F. Jaroszyka, PZW, Warszawa 2008, s. 541-582.
- [2] J. Pniewski, Soczewki kontaktowe wieloogniskowe. Podstawy własności optycznych. Optyk Polski, 5(69) 2022, s.20-23.
- [3] M. H. Nizankowska, Okulistyka – podstawy kliniczne, PZW, Warszawa 2007.
- [4] A. Styszyński, Korekcja przebiopii soczewkami kontaktowymi, Optyk Polski 1 (65) 2022, s. 58-60.
- [5] A. Przekoracka-Krawczyk, R. Naskręcki, R. Dysfunkcja akomodacji i metody jej badań, Optyka 6, 2010, s. 24-30.
- [6] A. Styszyński, Powtórka z optyki (cz. III), Optyk Polski 3(61), 2021, s. 44-51.
- [7] M. Molska, R. Naskręcki, Astygmatyzm – klasyfikacja, metody badania i korekcji, Optyka 1(32) 2015, s. 36-38.
- [8] M. Czaińska, Astygmatyzm – charakterystyka wady, Optyka 5(42), 2016, s. 28-31.
- [9] D. Olkowska, Preziopia – metody korekcji, Optyka 1(68), 2021, s. 40-44.
- [10] A. Styszyński, Widzenie obuoczne, cz. 4, Optyk Polski 5 (69) 2022, s. 58-60.
- [11] S. Meredith, Jak dbać o wzrok, słuch, smak i węch, Reader's Digest, Warszawa 2009, s.18-113.
- [12] U. Ostermeier-Sitkowski, Lepiej widzieć bez i w okularach. Ćwiczenia wzroku, Warszawa 2006.
- [13] M. L. Kwitko, M. Ross, Oczy, KDC, Warszawa 2003.
- [14] J. Wosik, J. Pniewski, Dokładność i powtarzalność obiektywnego pomiaru wady refrakcji konwencjonalnym autorefraktometrem, autorefraktometrem otwartego pola oraz aberrometrem u młodych dorosłych, Optyka 5(60), 2019, 34-38.

CREDO-Maze: Promieniowanie kosmiczne widziane z kosmosu – pasy Van Allena

Foto – Dreamstime

Artykuł ten jest piątym z serii poświęconej projektowi „Kosmos widziany z Łodzi” będącym realizacją szerokiej akcji udostępniania młodzieży nowoczesnej aparatury naukowej mającej w końcowym efekcie pokazać, a może i nauczyć młodych, ciekawych świata ludzi metod, jakimi posługuje się współczesna nauka w poszukiwaniu praw rządzących Wszechświatem. Dostarczane szkołom zestawy pomiarowe stają się istotnym rozwinięciem projektu CREDO (Cosmic Ray Extremely Distributed Observatory) i wszyscy, którzy przyłączą się do nas, staną się uczestnikami niezwykle podróży w nieznane zakamarki Kosmosu¹

Tadeusz Wibig

W badaniach pierwotnego promieniowania kosmicznego najlepiej wznieść się ponad ziemską atmosferę, a przynajmniej bardzo, bardzo wysoko. Można robić to na balonach, które osiągają już prawie tę nieistniejącą granicę kosmosu, ale z natury swej i z prawa Archimidesa jest jasne, że balonem unosić się można wyłącznie w powietrzu, choćby i bardzo, bardzo rzadkim, ale zawsze jest ono niezbędne. O pomiarach balonowych zresztą już mówiliśmy, pójdźmy więc dalej.

W roku 1957 ludzkość wyszła poza atmosferę. Sputnik 1 był, jak wskazuje nawet jego nazwa (*Искусственный Спутник-1*), najprostszym satelitą, jakiego można było wymyślić. Jego zasadniczym zadaniem było nadawanie przez radio mniej więcej trzy razy na sekundę krótkiego „pip”. Nadawał tak przez trzy tygodnie. Po trzech miesiącach spadł wyhamowując na resztkach atmosfery. Pokazał, że można, a skoro można, to lawina ruszyła.

Miesiąc po Sputniku 1 wystrzelono Sputnika 2. Jego najważniejszy ładunek stanowił pies, a właściwie bezdomna suczka Łajka, pierwsze żywe stworzenie w przestrzeni kosmicznej, której tragicznej historii nie będziemy tu opisywać. Poza nią na pokładzie Sputnika 2 były

spektrometry, które mierzyły promieniowanie słoneczne w zakresie ultrafioletowym i rentgenowskim, a co szczególnie nas interesuje Sergiej Nikołajewicz Vernow namówił Korolowa, by umieścić na pokładzie dwa liczniki Geiger–Müllera pozwalające mierzyć strumień naładowanych cząstek promieniowania kosmicznego w przestrzeni wokół Ziemi. Działały one 10 dni.

Sputnik 2 krążył po wydłużonej orbicie rozciągającej się od ~200 do 1600 km nad Ziemią i tak się nieszczęśliwie dla Vernowa złożyło, że będąc na wysokiej orbicie nie był widoczny z terytorium Związku Radzieckiego. Widać go było wtedy z Australii i tamtejsi obserwatorzy przestrzeni kosmicznej odbierali wysyłane przezeń dane. Dane te były oczywiście zakodowane i bez algorytmu pozwalającego na ich odszyfrowanie, dla Australijczyków były bezużyteczne. Z dzisiejszego punktu widzenia to, że poprosili Rosjan o ten algorytm, nie wydaje się czymś dziwnym, pamiętajmy jednak, że wszystko to działo się w roku 1957, na pięć lat przed kryzysem kubańskim, kiedy świat dzieliły nieprzenikalne kurtyny i wymiana informacji z jednej strony na drugą była czymś głęboko abstrakcyjnym. Rosjanie oczywiście odmówili, a więc i Australijczycy zachowali dane dla siebie i też nie podzieliли się nimi z drugą stroną.

¹ Ponieważ mamy nadzieję, że na tym się nie skończy, prosimy i zachęcamy zaciekawione, a może nawet i zainteresowane osoby, nauczycieli o kontakt z nami. Im więcej nas będzie, tym łatwiej będzie zrobić następny krok i podjąć kolejne wyzwania

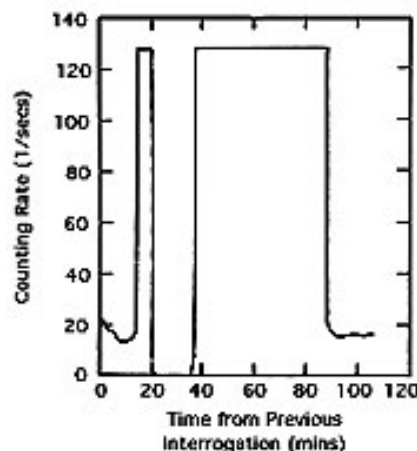
Tak, czy inaczej, wystrzelenie Sputników było niewątpliwym sukcesie Rosjan.

Ale Amerykanie byli tuż, tuż. Sam wyścig kosmiczny rozpoczął się i to po obu stronach żelaznej kurtyny oczywiście wcześniej. Po stronie amerykańskiej osobą zaangażowaną weń szczególnie był James Alfred Van Allen. Po wojnie, w której służył bohaterstwo między innymi na pancerniku USS Washington, wrócił na Uniwersytet Johna Hopkinsa, gdzie zorganizował grupę do badań wysokich warstw atmosfery z użyciem zdobycznej rakiety V-2. W 1948 roku udało mu się wystrzelić niewielką suborbitalną rakietę z serii Aerobee, która wyniosła aparaturę do badania promieniowania kosmicznego na wysokość ponad 100 km.

Van Allen z kolegami opracował też dość oryginalną technologię wielostopniowych ракет na paliwo stałe o nazwie Rockoon wynoszonych na spore wysokości przez balony, aby dopiero po osiągnięciu 20 km odpalić właściwą rakietę. Do września 1957 roku wystrzelono ponad trzydzieści Rockoonów. Była to część działań w ramach kierowanego między innymi przez Van Allena wielkiego programu naukowego nazwanego Międzynarodowym Rokiem Geofizycznym. Jego częścią był też i ogłoszony w 1955 roku przez Amerykanów Project Vanguard, który zakładał wystrzelenie na orbitę okołoziemską sztucznego satelity. Marynarka zamierzała wykorzystać do tego celu nową potężną trzystopniową rakietę o nazwie, oczywiście, Vanguard.

Próba dokonana 8 grudnia 1957 roku nie była udana. Zakończyła się wybuchem na platformie startowej. Ależ było śmiechu po drugiej stronie żelaznej kurtyny. Niedoszłego amerykańskiego satelity, parafrazując nazwę radzieckiego „Sputnika” nazywano „Flopnik”, „Kaputnik”, „Opsnik”, „Dudnick” and „Stayputnik”. Doszło do tego, że na forum ONZ radziecki delegat oficjalnie zapytał, czy Stany Zjednoczone nie są, aby zainteresowane otrzymaniem od Związku Radzieckiego pomocy przeznaczonej generalnie dla „krajów słabo rozwiniętych”.

Ale w końcu 1 lutego 1958 roku, po niecałych dwóch miesiącach od katastrofy Vanguarda, spóźnieni Amerykanie wystrzelili swojego pierwszego satelity, Explorera 1 używając innej czterostopniowej rakiety Juno. Miał



Rysunek 1: Dane o ilości promieniowania zmierzone przez Explorera 1.

on szczęśliwie, jak się później okazało, orbitę sięgającą 2500 km i wyposażony był, o co zadbał Van Allen w liczni G-M, by mierzyć promieniowanie kosmiczne w otoczeniu Ziemi. I mierzył, ale to, co zmierzył zaskoczyło wszystkich. Pokazujemy to na Rys. 1.

Zaskakujące wyniki

Gdy tak sobie leciał nie spodziewając się niczego i mierzył strumień promieniowania kosmicznego o wartościach w okolicy 20 zliczeń na sekundę (lewy fragment rysunku), zupełnie nagle po upływie kilkunastu minut (od umownego startu zegara) liczba zliczeń poszybowała gwałtownie do 128/sek! To, że zatrzymała się na tym poziomie na kilka minut nie jest dziwne, bo układy zliczające zostały skonstruowane tak, że „nasycaly się” na tym właśnie poziomie. „128” oznacza „128 albo i więcej”. Projektując układy elektroniczne nikt nie spodziewał się, że strumień może być aż tak duży.

Coś jeszcze dziwniejszego zdarzyło się po 20 minutach. Liczba zliczeń spada do zera! Licznik Geigera-Müllera nie pokazuje NIC! Czy znaczy to, że nagle promieniowania kosmicznego zabrakło? To oczywiście niemożliwe, a przynajmniej nikt nie potrafił wymyślić fizycznego mechanizmu, który wymiałałby wszystkie cząstki z przestrzeni, w której chwilę wcześniej było ich mnóstwo (>128/sek). Aby zrozumieć, co tak naprawdę pokazywała aparatura Explorera 1 trzeba dokładnie przestudiować zasadę działania licznika Geigera-Müllera.

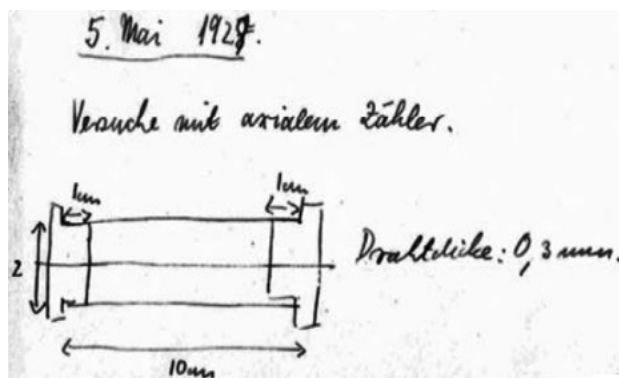
Licznik Geigera (Geigera-Müllera) jest jednym z najbardziej znanych i rozpoznawalnych przez wszystkich instrumentów naukowych. Uznaje się powszechnie, że wymyślił go Walter Müller pracując z promotorem swojego doktoratu Hansem Geigerem wiosną 1928 roku w Kilonii [1]. Szkic z jego laboratoryjnego dziennika pokazuje Rys. 2 [2].

Historia jednak zaczęła się dużo wcześniej.

Hans Geiger na początku XX wieku w roku 1906 obronił doktorat z badań wyładowań elektrycznych w gazach na uniwersytecie w Erlangen w Niemczech i wyjechał do Manchesteru w Anglii, gdzie został asystentem w laboratorium Rutherforda, który właśnie pracował nad rozpraszaniem cząstek alfa na złotej folii. Po dwóch latach wspólnie skonstruowali komorę jonizacyjną do zliczania

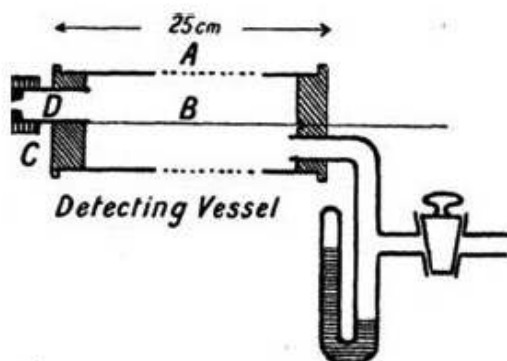


Sputnik 1 – commons.wikimedia.org



Rysunek 2: Schemat licznika GM z zeszytu Mullera (1928).

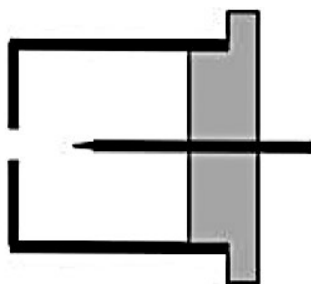
cząstek alfa. Jeśli przyjrzeć się jej konstrukcji opublikowanej w roku 1908 [3] (Rys.3) nawet pobieżnie od razu musi się zauważyć podobieństwo z młodszym o 20 lat szkicem Mullera.



Rysunek 3: Licznik cząstek alfa Rutherforda i Geigera (1906).

Ale to jeszcze nie cała historia.

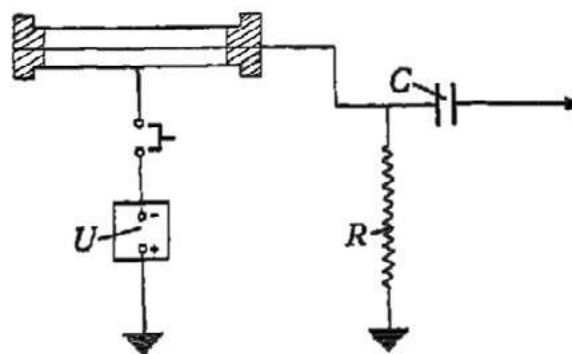
W 1912 roku Geiger wrócił do Niemiec, został dyrektorem Physikalisch-Technische Reichsanstalt w Berlinie i rok później opublikował wymyślony przez siebie „licznik punktowy” [4] (Rys.4).



Rysunek 4: Licznik punktowy Geigera (1913)

A potem zaczęła się Wojna Światowa i Geiger poszedł na front jako oficer artylerii. Trudy życia w okopach podkopały definitywnie jego zdrowie, ale udało mu się mimo wszystko przeżyć i po wojnie ponownie poświęcił się nauce trafiając najpierw do Kilonii.

Tamże zajął się wspólnie swoim doktorantem Müllerem tematem, który zaowocował udoskonaloną konstrukcją licznika nazywanego do dziś ich imionami.



Rysunek 5: Schemat elektryczny licznika GM z pracy [5] (1938).

Skoro już wyjaśniliśmy aspekt historyczny przejdźmy do opisu zasady działania licznika GM.

Jego budowa jest prosta (można ją porównać z konstrukcją cepa). Składa się z rurki, w oryginale była mosiężna, zamkniętej hermetycznie z obu końców, w oryginale zaślepki ebonitowymi, przez środek której poprowadzony jest cienki drucik. Najważniejsze w liczniku GM jest to, czego nie widać na rysunkach: jest on wypełniony gazem pod niskim ciśnieniem. Najlepiej, aby gaz był szlachetny, a ciśnienie to mniej więcej 1/10 ciśnienia atmosferycznego. Między rurką a drut przykłada się dość wysoki potencjał elektryczny. Zależnie od konstrukcji może on sięgać tysięcy woltów.

Jak działa licznik GM?

Powiedzmy, że przez licznik, a dokładniej: przez gaz w rurce, przelatuje cząstka naładowana. Cząstka ta oczywiście jonizuje cząsteczki (atomy) gazu, co prowadzi do powstania wewnątrz rurki pewnej liczby ciężkich jonów dodatnich i odpowiadającej im liczby swobodnych, lekkich elektronów. Pole elektryczne panujące między drutem a obudową oddziałuje na jedne i drugie. Elektrony, przyspieszane są w stronę drutu, a jony dodatnie w przeciwną.

Jak wiemy ze szkoły, pole elektrostatyczne w pobliżu drutu rośnie jak $1/r$, a więc w pobliżu cienkiego drutu rośnie bardzo szybko (bo bardzo małe jest r w mianowniku) i elektrony blisko drutu uzyskują energie na tyle duże, że teraz zdarzając się z innymi neutralnymi jeszcze atomami gazu mogą je same jonizować. Wybite z tych atomów nowe elektrony przyciągane są także przez drut i one też po jakimś czasie zaczynają same jonizować kolejne atomy. Tak powstaje solidna kaskada szybkich elektronów, która w końcu dociera do anody (elektrody dodatniej, do drutu) i można by uznać, że to już koniec. Wystarczy zebrać te elektrony i dać to jakiś impuls elektryczny. Tak działa z grubsza licznik proporcjonalny. Sygnał z niego wychodzący jest z grubsza proporcjonalny do liczby kaskad, jakie doszły do anody, czyli cząstek, jakie przeszły przez licznik. Ale te kilka, kilkaset, a nawet kilka milionów elektronów to jednak w sumie bardzo mało. Nie o to Geigerowi i Müllerowi chodziło.

Licznik GM wzmacnia sygnał, który wywołało przejście jednej naładowanej cząstki miliardy razy, a wszystko to przez zderzenia rozpędzonych elektronów z drutem, do którego przecież w końcu dążyły. Jeśli napięcie przyłożo-

ne do drutu jest dostatecznie duże, każde takie zderzenie powoduje gwałtowne wyhamowanie elektronu, a temu towarzyszy zawsze emisja tak zwanego promieniowania hamowania, po angielsku (z niemieckiego) *bremsstrahlung*, czyli fotonów, i są to głównie fotony z obszaru ultrafioletowego. To one właśnie rozbiegając się we wszystkie strony wewnątrz rury licznika dodatkowo jonizują wszystko, co napotkają na swojej drodze i to w wyniku tego procesu po przejściu cząstki w czasie mierzonym w mikrosekundach cały licznik wchodzi w stan przewodzenia i następuje w nim coś w rodzaju wyładowania elektrycznego. Aby ograniczyć jego skutki zasilanie jest połączone z katodą (elektrodą ujemną, rurką mosiężną) przez bardzo duży opór elektryczny rzędu milionów omów. Po przepłynięciu na drut całego ładunku zgromadzonemu początkowo na elektrodach licznika, który jest w końcu rodzajem kondensatora, wyładowanie zanika i licznik powoli zaczyna się ładować, by być gotowym na rejestrację następnej cząstki. Czas ładowania, jest zdecydowanie czasem martwym licznika. W tym czasie on po prostu nie działa.

Jeśli zrozumieliśmy działanie licznika GM, jest już teraz zupełnie jasne skąd w wynikach przesyłanych przez Explorera pojawiły się te pozornie tajemnicze zaniki sygnałów. Po prostu czas martwy liczników był dłuższy niż typowy odstęp czasu pomiędzy kolejnymi bombardującymi go cząstkami promieniowania kosmicznego. Sami członkowie zespołu Van Allana nie od razu wpadli na to rozwiązanie. Pierwszy zasugerował je zatrudniony w nim student Van Allana Carl McIlwain. Razem z kolegą Ernie Rayem sprawdzili oni tę roboczą hipotezę doświadczalnie naświetlając licznik GM potężną dawką promieniowania X. Gdy „zatkali” tak swój testowy licznik i okazało się, że mają rację, przyczepili na drzwiach pokoju Van Allana kartkę z napisem, który przeszedł do historii [6]:

„SPACE IS RADIOACTIVE”.

Okazało się, że w przestrzeni wokół Ziemi jest tajemniczy obszar, w którym liczba cząstek jonizujących gaz w liczniku jest bardzo, bardzo duża. Co więcej obszar ten ma swoje dość wyraźnie określone granice, poza którymi wszystko wygląda „normalnie”. Oczywiście nie trzeba wspominać, że nikt się czegoś takiego nie spodziewał. Zakładając, że jest to zjawisko naturalne i wytworzone siłami przyrody, powinno charakteryzować się ono symetrią obrotową, a to oznaczało, że Ziemię otacza pas intensywnego promieniowania [7].

Nic więc dziwnego, że po pierwszym Explorerze, zanim jeszcze stało się jasne, co zmierzył, wystartował Explorer-3 (Explorer 2 uległ awarii zanim jeszcze wszedł na orbitę), a potem i Explorer 4,

Rosjanie ciągle aktywni w kosmicznym wyścigu wystrzelili w tym samym czasie Sputnika-3 bogato wyposażonego w aparaturę do pomiaru cząstek promieniowania kosmicznego. Niestety nie wszystko przebiegło zgodnie z planem. Zawiódł system transmisji danych, a konkretnie magnetofon do magazynowania informacji i późniejszego wysłania ich do odbiorników na terenie Związku Radzieckiego, a wszystko to dla ominięcia problemów, na jakie napotkał Sputnik 2. Mimo to jednak właśnie Sputnik-3 odkrył, że za pierwszym pasem Van Allana zaobserwowa-

nym przez Explorera jest i drugi pas, pas zewnętrzny, a co dziwniejsze oba one są wyraźnie od siebie odseparowane i w tej przerwie promieniowanie kosmiczne zachowuje się, jak należy. Kolejna zagadka domagała się rozwiązania.

Pasy Van Allena

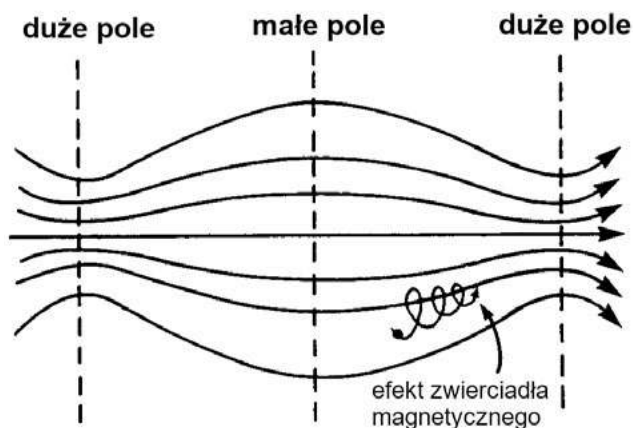
Skąd się więc wzięły, jak powstały i czym właściwie są pasy Van Allena?

Cząstki naładowane, jak wiadomo, oddziałują z polami elektrycznymi i magnetycznymi. W przestrzeni wokół Ziemi pól elektrycznych zasadniczo nie ma i nikt się ich nie spodziewa, za to pola magnetyczne są. Wiemy o tym od tysiącleci, od czasów, gdy Chińczycy wymyślili kompas. Pola, które odchylają igłę kompasu, istnieją i w przestrzeni wokółziemskiej i odchylają tam trajektorie cząstek naładowanych.

Szczegółowo opisywaliśmy ten problem w artykule w drugim tegorocznym numerze „Fizyki w Szkole”, nie będziemy więc się powtarzać. Ustalmy, że wiemy, iż w specyficznych konfiguracjach pola, cząstki naładowane mogą być w nim zamknięte. Takie „butelki magnetyczne” nie są specjalnie trudne do zrobienia, ani też nie są niczym nadzwyczajnym w przyrodzie. Wystarczy, że strumień pola początkowo równoległego, stałego w czasie i przestrzeni, „ściśniemy”, co można sobie wyobrazić odwołując się do obrazka „linii sił” i już mamy magnetyczną pułapkę (Rys.6).

Oczywiście, jeśli cząstki mają energie dużą, na tyle dużą, aby ich promień Larmora był większy od rozmiarów poprzecznych pułapki, uciekną z niej, a więc pułapka jest pułapką tylko na cząstki o energiach mniejszych niż jakaś konkretna wartość zależna od rozmiarów i pola magnetycznego pułapki.

Typowe cząstki promieniowania kosmicznego o energiach miliardów elektronowoltów (~1 GeV) nie mają szans, aby dać się złapać w ziemskich polach magnetycznych. Cząstki uciekające ze Słońca owszem mogą, ale pojawia się inny problem, o ile jakaś cząstka jest złapana w pułapkę i jako taka, nie może się z niej wydostać, o tyle cząstka o takiej samej energii, która chciałaby wejść do magnetycznej butelki i poruszać się wewnątrz po takim torze, jak ta złapana, do takiej butelki nie wejdzie. Pole magnetyczne zabrania nie tylko wychodzenia, ale i wcho-



Rysunek 6: Schemat pułapki magnetycznej.

dzenia do pułapki. Taka to pułapka! Oczywiście każda cząstka do wnętrza wejdzie, jeśli będzie się poruszała na zewnątrz na przykład wzdłuż linii sił pola. Ale taka cząstka, po ewentualnym jednym odbiciu od drugiego końca z butelki wyjdzie, jak weszła - po tej samej linii. Skąd się więc biorą cząstki w butelkach Van Allana?

Puławki magnetyczne

Jest co najmniej kilka możliwości. Najprostszą i wcale nie tak egzotyczną, jak mogłoby się wydawać, jest ta, że cząstki te w ogóle powstają już we wnętrzu pułapki.

Wiedząc, że atmosfera bombardowana jest cały czas cząstkami promieniowania kosmicznego o wysokich energiach, dla których magnetyczne pułapki Van Allana nic nie znaczą, musimy też wiedzieć, że wszystkie one prędzej, czy później zderzą się z jądrem jakiegoś atomu w ziemskiej atmosferze. Zderzenia takie mogą zapoczątkować wielkie pęki atmosferyczne, które rejestrować będziemy w Projekcie CREDO-Maze, ale z drugiej strony mogą jednocześnie rozbić to jądro, jak to ponad sto lat temu zauważyli Irena i Fryderyk Joliot-Curie [8]. W wyniku takiej reakcji z jądra wybijane mogą być nukleony.

Często zdarza się, że są to pojedyncze neutrony. Neutrony te mają energie niewielkie w porównaniu z pierwotną cząstką promieniowania kosmicznego, sięgają one jednak kilkudziesięciu milionów elektronowoltów ~ 10 MeV i z takimi energiami podróżują swobodnie w przestrzeni w kierunkach zasadniczo dowolnych. Jako nienaładowane elektrycznie nieczułe są na wszechobecne pola magnetyczne. Mogą też oczywiście wpaść do wnętrza butelki magnetycznej zupełnie nie zdając sobie z tego sprawy.

Neutrony są obecne w materii, z jakiej jesteśmy zrobieni, są też we wszystkim, co nas otacza. W każdym jądrze atomowym jest z grubsza tyle neutronów, co protonów (oczywistym wyjątkiem jest wodór) i jakoś nas to nie przeraża. Neutrony szczęśliwie są stabilne. Ale stabilne są tylko wewnątrz atomowych jader. Siły jądrowe zapewniają im stabilność. Gdy jednak neutron wyrwie się na wolność, gdy stanie się wolnym neutronem, los jego może być ciekawy, ale krótki. Średnio żyje on (w spoczynku) niecałe 15 minut, a po tym czasie niestety rozpada się. Produktami jego rozpadu są proton, elektron (ładunek elektryczny musi się zachowywać!) i neutrino elektronowe.

Z racji masy proton unosi prawie całą energię kinetyczną neutronu, a więc, jeśli neutron miał kilkanaście MeV, to i proton ma mniej więcej tyle. I jeśli tak się nieszczęśliwie, albo szczęśliwie, złożyło, że neutron rozpadł się wewnątrz pułapki magnetycznej Van Allana, to przy odpowiednim kącie, powstały proton może zostać w nią na trwałe złapany. Mechanizm ten zaproponowany został podobno przez Vernova i Lebedinskiego w lipcu 1958 roku. Na pewno zaś w tym samym czasie ukazała się w Physical Review Letters praca Singera [9] na ten sam temat.

Mechanizm ten może wyjaśnić istnienie wewnętrznego pasa Van Allana, który, jak pokazały pomiary składa się głównie z kilkudziesięciomewowych protonów. Dla odmiany na pas zewnętrzny składają się głównie elektrony o energiach 0.1-10 MeV. Kwestia ich pochodzenia jest niejasna, a co najmniej złożona.

Część z nich to na pewno okruchy wiatru słonecznego, który wysyłany jest przez naszą najbliższą gwiazdę nieustannie, acz nie zawsze w tej samej ilości. Poza wahaniem okresowymi (cykl 11-letni i inne, o których już wspominaliśmy w poprzednich tekstach w „Fizyce z Szkoły”) zmienia się też bez żadnej zapowiedzi i regularności. Słońce co jakiś czas wyrzuca miliardy ton materii rozprzeczonych do tysięcy kilometrów na sekundę. Zjawiska te nazywa się Koronalnymi Wyrzutami Masy. Materia ta jest rzecz jasna zjonizowana, czyli składa się zasadniczo z protonów i elektronów z wzmrożonym w nią polem magnetycznym. Czasem plazma ta dociera do Ziemi, a dokładnie do ziemskiej magnetosfery powodując w niej duże zamieszanie. Dokładny mechanizm tych zjawisk także nie jest znany do końca i ciągle jest tematem intensywnych badań.

Rozważając wyżej puławki magnetyczne zakładaliśmy, że pola nie zmieniają się w czasie. Jeśli jednak mamy do czynienia z polami, które zmieniają się trochę z czasem w okresach porównywalnych z okresem, z jakim krążą w pułapce uwięzione cząstki, to pułapka ta staje się niedoskonała. Niektóre cząstki mogą ją opuścić, ale jednocześnie pojawia się możliwość, że coś z zewnątrz do pułapki wpadnie. Jeśli przyjąć, że pola magnetyczne w zewnętrznym pasie Van Allana są zaburzane przez burze geomagnetyczne, przez nadlatującą ze Słońca plazmę ze swoimi polami, to coś z tej plazmy, na przykład niektóre elektrony o odpowiednich energiach i kierunkach mogą wtargnąć do pułapki i pozostać tam, gdy butelka znów zamknie się po przejściu fali kosmicznej.

Na skutek losowych, przypadkowych i niewielkich zaburzeń pól magnetycznych mogą też docierać do wnętrza pułapki elektrony przyspieszane w warstwach wewnętrznych na przykład przez bardzo długie fale radiowe emitowane w komunikacji z łodziami podwodnymi, czy fale generowane wyładowaniami elektrycznymi wysoko w atmosferze. Naukowcy rozważają wiele różnych mniej lub bardziej egzotycznych mechanizmów i każdy z nich może dodawać trochę niskoenergetycznych elektronów do zewnętrznego pasa Van Allana.

Specjalnie pominęliśmy tu kwestię szczególnie destrukcyjnego wpływu człowieka na fizykę bliskiej przestrzeni kosmicznej spowodowaną eksplozjami bomb atomowych na granicy przestrzeni kosmicznej w latach 60. ubiegłego wieku, bo nie ma się czym chwalić. Naruszyliśmy wtedy strukturę szczególnie wewnętrznego pasa Van Allana i efekty tych ekscesów trwały latami. Dostarczyło to wprawdzie danych dla testowania modeli, ale był to jedynie niewielki pozytywny efekt uboczny. Wiemy dzięki niemu, że, jeśli tylko chcemy możemy utworzyć wokół Ziemi pasy radiacyjne, które mogą być zabójcze dla satelitów, a i niestety, szkodliwe dla potencjalnych astronautów.

Czy pasy Van Allana są szkodliwe?

Geometria ziemskiego, dipolowego w zasadzie pola magnetycznego, które na szerokościach równikowych jest generalnie słabsze niż w obszarach podbiegunowych w oczywisty sposób wyznacza kształt pasów Van Allana. Oplatają one Ziemię w obszarze równikowym właś-

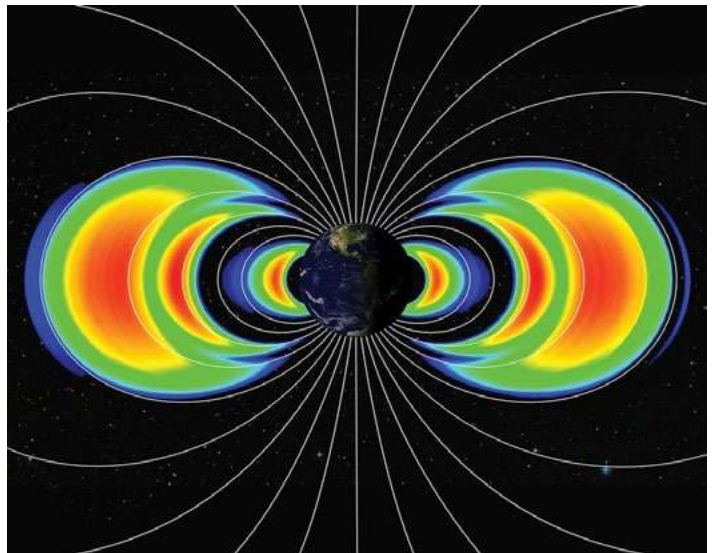
nie na wysokościach od 1000 do 12000 kilometrów (pas wewnętrzny) i od 13000 do 60000 km (pas zewnętrzny). Pas zewnętrzny z natury swojej jest zmienny bardziej. Wewnętrzny także zależy trochę od aktualnej aktywności Słońca i w niektórych miejscach może zbliżać się czasem do powierzchni Ziemi nawet na 200 km (Anomalia Południowego Atlantyku), co sprawia pewne problemy przelatującym nad tym obszarem na niskich orbitach satelitom.

Problemy te nie są jakoś szczególnie dokuczliwe. Częste w pasach Van Allena jest dużo, ale znowu nie tak dużo, jak niektórzy myślą. Są tacy, którzy uważają, że astronauta lecący na Księżyc przecinając pasy Van Allena na pewno by zginął za sprawą zabójczego promieniowania kosmicznego i wnoszą stąd, że cały program Apollo był jednym wielkim oszustwem, manipulacją i ściemą. Aby dać naukowy odpór tym sądom, wykonajmy proste rachunki.

Na rysunku 7 kolorami zaznaczono dawkę do pochłonięcia w czasie 1 sekundy. Regiony niebieskie to 0.001 mSv/s, zielone 0.01 mSv/s, żółte 0.05 mSv/s, pomarańczowe 0.1 mSv/s i wreszcie czerwone 0.5 mSv/s. Jednostka Sv (siwert) jest jednostką układu SI odnoszącą się do działania promieniowania jonizującego na organizmy żywe i odpowiada przyjęciu przez organizm energii 1 J na każdy kilogram masy ciała. Istnieje też inna jednostka pochłoniętej dawki: można ją mierzyć w „bananach” *Banana equivalent dose* (BED). Definicja jest bardzo prosta. 1 banan to dawka promieniowania przyjęta przez człowieka, gdy zje jednego banana. Banan jak wiemy zawiera potas, a ^{40}K jest radioaktywny! 1 BED to 0.1 μSv . Przeświecenie zęba to 500 BED, badanie mammograficzne to jakieś 4000 BED.

Rysunek 7 wykonany jest „w skali” i obszary żółty, pomarańczowy i czerwony mogą zajmować na drodze statku kierującego się poprzez nie ku Księżycowi odcinek o długości najwyżej 2 promieni Ziemi. Dla uproszczenia przyjmijmy, że jest to 12500 km. Statek przecina pasy Van Allena z prędkością około 25000 km/godzinę i przez trzy najbardziej „promieniotwórcze” obszary podróżował będzie mniej więcej pół godziny, to jest przez 1800 sekund. W tym czasie dawka jaką pochłonie wyniesie około 100-200 mSv.

Kwestie ochrony radiologicznej i sposobu przeliczania ilości elektronów czy protonów o energiach, z jakimi mamy do czynienia w pasach Van Allena, nie są proste. Czy 100 mSv to dużo, czy mało? Naukowo określona dawka śmiertelna (LD 50/30) dla człowieka przyjęta w krótkim czasie to 3-5 Sv (50 milionów bananów). Z drugiej jednak strony są dane, które wskazują, że promieniowanie i to w dawkach rzędu nawet 200 mSv ma pozytywny wpływ na zdrowie ludzi. Nazywa się to efektem *hormezy radiacyjnej* i podaje się tu przykład pacjentek leczonych na gruźlicę za pomocą promieniowania jądowego dawkami rzędu 100 do 200 mSv wśród których szanse zachorowania na raka okazały się znacznie mniejsze niż przeciętne. Podobnie zanotowano zmniejszoną śmiertelność i prawdopodobieństwo zachorowania na białaczkę wśród tych mieszkańców Hiroshimy i Nagasaki, którzy zostali napromieniowani dawkami w zakresie do 200 mSv [10]. Sprawa jest bardzo niejasna i badania wciąż trwają.



Rysunek 7. Pasy Van Allena widziane przez Van Allen Probes.

Trwają też ciągle badania pasów Van Allena. W 2012 roku NASA w ramach misji *Radiation Belt Storm Probe* wystrzeliła dwie sondy nazwane potem sondami im. Van Allena. Miały one zbadać dokładnie strukturę przestrzenną pasów i jej zmienność. Ku zaskoczeniu wszystkich sondy te wykryły niemal od razu trzeci pas Van Allena [11], na który składają się ultra-relatywistyczne elektrony i który znajduje się pomiędzy pasami wewnętrznym i zewnętrznym.

Tadeusz Wibig

Katedra Fizyki Teoretycznej Uniwersytetu Łódzkiego

Projekt „Kosmos widziany z Łodzi” jest dofinansowany ze środków budżetu Państwa (SONP/SN/516075/2021).

- [1] Geiger, H.; Müller, W. (1928). „Elektronenzählrohr zur Messung schwächer Aktivitäten” *Die Naturwissenschaften*, **16** (31): 617-618; doi:10.1007/BF01494093. S2CID 27274269.
- [2] Sebastian Korf (2013), „How the Geiger Counter started to crackle: Electrical counting methods in early radioactivity research”, *Ann. Phys. (Berlin)* **525**, A88–A92 (2013); DOI 10.1002/andp.201300726
- [3] Rutherford, E.; Geiger, H. (1908). „An electrical method of counting the number of particles from radioactive substances”, *Proceedings of the Royal Society Series A. London*. **81** (546): 141–161. Bibcode:1908RSPSA..81..141R. doi:10.1098/rspa.1908.0065.
- [4] Geiger H. (1913). „Über eine einfache Methode zur Zählung von α - und β -Strahlen”, *Verhandlungen der Deutschen Physikalischen Gesellschaft* **15**: 534–539.
- [5] Nunn A. (1938) *The mechanism of the geiger counter*, *Rep. Prog. Phys.* **5** 390; DOI 10.1088/0034-4885/5/1/329
- [6] James A. Van Allen, J. A., Carl E. McIlwain C. E. i George H. Ludwig G. H. (1959) „Satellite Observations Of Electrons Artificially Injected Into The Geomagnetic Field”, *Earth, Atmospheric, And Planetary Sciences* **45** 1152-1171; DOI: 10.1073/pnas.45.8.1152.
- [7] Van Allen, J. A. (1959), *The geomagnetically trapped corpuscular radiation*, *J. Geophys. Res.*, **64**(11), 1683–1689, doi:10.1029/JZ064i011p01683.
- [8] Irène Joliot-Curie – Nobel Lecture. NobelPrize.org. Nobel Prize Outreach AB 2023. Tue. 7 Feb 2023. <https://www.nobelprize.org/prizes/chemistry/1935/joliot-curie/lecture/>
- [9] Singer S. F., (1958) *Trapped Albedo Theory of the Radiation Belt*, *Phys. Rev. Lett.* **1**, 181 <https://doi.org/10.1103/PhysRevLett.1.181>
- [10] Moskal P. (2011) „Dawki promieniowania jądowego”, *Foton* **112**, 9
- [11] Cowen, R., (2013) „Ephemeral third ring of radiation makes appearance around Earth”. *Nature*; <https://doi.org/10.1038/nature.2013.12529>

Za co Nobel 2022?

Część 5. Komputery i algorytmy kwantowe

Dyskusja na temat splątania kwantowego miała początkowo charakter czysto akademicki. Jednak wraz z eksperymentalnym potwierdzeniem istnienia tej przedziwnej cechy układów kwantowych zaczęła powstawać nowa dziedzina nauki i techniki – wykorzystująca możliwość manipulowania stanami kwantowymi, a tym samym przetwarzaniem informacji zapisanej w stanach kwantowych.

Jan Kurzyk

W poprzedniej części opisałem niektóre bramki kwantowe będące podstawowymi elementami obwodów kwantowych zdolnych do realizacji kwantowych algorytmów. Rozwój tych nowoczesnych technologii może znaleźć przeróżne zastosowania. Najczęściej w tym kontekście mamy na myśli skonstruowanie w pełni funkcjonalnego komputera kwantowego. Co rusz czytamy o kolejnych krokach przybliżających nas do tego celu. W artykule opiszę w jakim punkcie rozwoju techniki komputerów kwantowych znajdujemy się. Omówię również dwa najbardziej znane algorytmy kwantowe: algorytm Grovera i algorytm Shora.

Jak zwykle sugeruję zapoznanie się z poprzednimi odcinkami cyklu [1-4], gdyż w artykule będę posługiwał się pojęciami, które omówiłem wcześniej.

Komputery kwantowe

Najbardziej oczekiwanym urządzeniem wykorzystującym informacje kwantowe jest komputer kwantowy. Idea takiego komputera powstała w latach 80. ubiegłego wieku. W 1980 roku amerykański fizyk Paul Benioff opublikował artykuł, w którym przedstawił kwantową maszynę Turringa [5]. Richard Feynman w 1982 roku zwrócił uwagę na ograniczenia związane z symulacjami dużych układów kwantowych wynikające z czasu obliczeń na

klasycznych komputerach. Czas takich obliczeń rośnie wykładniczo z liczbą cząsteczek symulowanego układu. Feynman zasugerował, że rozwiązaniem tego problemu byłby komputer wykorzystujący fizykę kwantową, dla którego czas tego typu obliczeń wzrastałby wielomianowo z liczbą cząsteczek.

W tradycyjnych komputerach procesor komputera dysponuje pewną ilością szybkiej pamięci nazywanej rejestrem. W komputerze kwantowym mamy do czynienia z rejestrem kwantowym, który zamiast bitów korzysta z kubitów. Każdy z nich może być superpozycją dwóch stanów $|0\rangle$ i $|1\rangle$ z amplitudami prawdopodobieństwa, których kwadrat modułu może być dowolną liczbą rzeczywistą z przedziału od 0 do 1. W tym sensie można powiedzieć, że komputer kwantowy jest komputerem analogowym. W przypadku rejestru klasycznego możemy przetwarzać za pomocą bramek logicznych pojedyncze bity lub ich zestawy.

Do przetwarzania informacji w rejestrze kwantowym korzystamy z bramek kwantowych, które wprowadzają nie tylko zmiany w pojedynczych kubitach, ale również w relacjach między nimi. Kubity rejestru możemy splątywać a tym samym operacje na nich przeprowadzane mogą wpływać na stan całego rejestru kwantowego. Dzięki temu komputer kwantowy może wykonywać obliczenia równolegle działając jednocześnie na wszystkie kubity rejestru kwantowego. Nazywamy to paralelizmem kwantowym.

Obliczenia wykonywane na komputerze kwantowym polegają na ustawieniu stanu początkowego rejestru kwantowego, a następnie na działaniu na rejestr odpowiednią sekwencją bramek kwantowych, co prowadzi do ewolucji stanu układu kubitów rejestru. Ostatnim etapem obliczeń jest pomiar stanu rejestru. Jak już wiemy pomiar prowadzi do kolapsu stanu będącego superpozycją różnych stanów bazowych do jednego ze stanów bazowych (chyba, że ewolucja kwantowa doprowadziła już układ do jednego z tych stanów). Oznacza to, że obliczenia kwantowe mają charakter probabilistyczny i powtórzenie obliczeń może doprowadzić do innego wyniku. Należy zatem obliczenia powtarzać kilka (lub więcej) razy, aby uzyskać statystykę pozwalającą na osiągnięcie żądanej dokładności. Temu problemowi przyjrzymy się w następnych punktach na przykładach algorytmów Grovera i Shora.

Probabilistyczność obliczeń kwantowych nie jest największym problemem w konstrukcji komputera kwantowego. Jesteśmy w stanie osiągnąć każdą założoną dokładność i pewność obliczeń kwantowych. Jest jednak inny poważny problem. Jest nim dekoherencja układu kwantowego. Stany kwantowe są niesłychanie delikatne i niestabilne. Jakakolwiek interakcja układu kwantowego z otoczeniem może negatywnie wpłynąć na stan superpozycji i zakłócić przetwarzanie informacji wprowadzając błędy niemożliwe do wykrycia. Źródłem dekoherencji mogą być zmieniające się pola elektryczne i magnetyczne czy promieniowanie termiczne pobliskich źródeł. Aby wydłużyć czas koherencji układy są maksymalnie izolowane od otoczenia. Stosuje się schładzanie do temperatur bliskich zera bezwzględnego (nawet do 20 milikelwinów [6]), ekranowanie od zewnętrznych pól elektrycznych i magnetycznych, a także nadmiarowość obliczeń i różne skomplikowane techniki korekcji błędów.

Oprócz dekoherencji wynikającej z oddziaływania układu mechanicznego z otoczeniem poważnym problemem jest tak zwany szum kwantowy. Jest on konsekwencją nieokreśloności stanów kwantowych związanych np. z fluktuacjami energii zwłaszcza w najniższych stanach energetycznych zgodnie z zasadą nieoznaczoności Heisenberga. Dlatego nawet schładzanie układu do temperatur bliskich zera bezwzględnego nie jest w stanie uchronić komputera kwantowego przed powstawaniem błędów. Jak pokazały badania z 2020 roku poważnym źródłem dekoherencji jest również promieniowanie kosmiczne.

Czym większą liczbą kubitów chcemy operować tym trudniej utrzymać je odpowiednio długo w stanie superpozycji i splątania. Najlepsze obecnie komputery kwantowe takie jak na przykład Hummingbird firmy IBM, czy googłowski komputer Sycamore operują kilkudziesięcioma kubitami¹ i są w stanie utrzymać je w stanie koherencji w czasie rzędu milisekund. Cały proces obliczeniowy musi zmieścić się w tym czasie.

Daleko nam jeszcze do konstrukcji w pełni skalowalnego komputera kwantowego, który mógłby znaleźć praktyczne zastosowanie. Po pierwsze taki komputer powinien mieć

dużo większą liczbę kubitów operacyjnych – co najmniej kilkaset². Utrzymanie w superpozycji tak dużej liczby kubitów przez odpowiednio długi czas będzie bardzo trudne.

Kwantowa korekcja błędów, jaką należałoby wprowadzić w komputerach kwantowych zdolnych do praktycznego zastosowania może znacząco osłabić przyspieszenie wynikające z przewagi algorytmów kwantowych. Ponadto algorytmy kwantowe zapewniające znaczące przyspieszenie obliczeń w porównaniu z klasycznymi algorytmami dotyczą wybranych zagadnień. W innych zagadnieniach przetwarzanie dużej ilości niekwantowych danych jest nadal dużym wyzwaniem dla komputerów kwantowych.

Osób wątpiących w osiągnięcie tzw. supremacji kwantowej jest bardzo wiele. W maju 2023 roku w renomowanym czasopiśmie *Nature* ukazał się artykuł Michaela Brooksa pod tytułem „*Quantum computers: what are they good for?*” („Komputery kwantowe: do czego służą?”). Autor artykułu bardzo sceptycznie wypowiada się na temat komputerów kwantowych. Na wstępie odpowiada na zadane w tytule pytanie twierdząc, że „na razie nie służą do niczego” („For now, absolutely nothing”). Do podobnych wniosków prowadzą komunikaty zamieszczane w ostatnich miesiącach w czasopiśmie *Association for Computing Machinery*. Angielski fizyk i prezenter telewizyjny Paul Davis argumentował, że 400-kubitowy komputer kwantowy wszedłby w konflikt z kosmologiczną informacją związaną z zasadą holograficzną będącą jednym z aksjomatów teorii strun.

To prawda, że istnieją już i działają małe komputery kwantowe realizujące konkretne zadania, ale ich znaczenie ma, póki co tylko wartość doświadczalną. Należy jednak odnotować, że postęp w tej dziedzinie jest wyjątkowo szybki i obiecujący. Współczynnik błędów, który 20 lat temu wynosił 30% obecnie jest niższy niż 1%. Prawdą jest jednak również to, co zauważył Brooks, że obecnie komputery kwantowe nie są do niczego przydatne.

Być może komputery kwantowe nigdy nie wypną klasycznych komputerów i będą używane tylko do rozwiązywania ograniczonej liczby problemów. Mimo to wysiłki wkładane w rozwój tej technologii z pewnością są opłacalne. Nawet jeśli te badania nie zakończą się stworzeniem w pełni funkcjonalnego komputera kwantowego, to jak uczy historia, mogą zaowocować powstaniem nowych zupełnie nieoczekiwanych technologii i odkryć naukowych.

Algorytm Grovera

W 1996 roku indyjsko-amerykański informatyk Lov Grover przedstawił kwantowy algorytm wyszukiwania elementu w nieuporządkowanej bazie danych. Załóżmy, że mamy nieuporządkowaną bazę zawierającą N elementów. Ponumerujemy je liczbami od 0 do $N - 1$. Będziemy potrzebować dwóch rejestrów kwantowych połączonych w jeden rejestr. Pierwszy z nich składa się z n kubitów, gdzie n jest najmniejszą liczbą naturalną taką, że $2^n \geq N$ (jeśli rozmiar bazy nie jest potęgą dwójki, to część kubitów nie będzie używana). Drugi rejestr zawiera tylko jeden pomocniczy kubit.

¹ Hummingbird zaprezentowany przez IBM w 2020 roku ma rejestr złożony z 65 kubitów.

² Julie Bobroff w książce „*Czy da się przechodzić przez ściany?*” [8] ocenia tę liczbę nawet na milion!

Na początku wszystkie kubity reprezentujące elementy bazy danych znajdują się w stanie $|0\rangle$, a kubit pomocniczy w stanie $|1\rangle$. Następnie inicjujemy stan kwantowy rejestru – przepuszczamy każdy z kubitów przez bramkę Hadamarda wprowadzając wszystkie kubity w stan superpozycji $\frac{1}{\sqrt{2}}(|0\rangle + |1\rangle)$. Po tej operacji każdemu elementowi bazy odpowiada taka sama amplituda prawdopodobieństwa $\left(\frac{1}{\sqrt{2}}\right)^n = \frac{1}{\sqrt{N}}$ (patrz rysunek 1, pierwszy od góry).

Gdybyśmy teraz wykonali pomiar stanu rejestru to dostalibyśmy z takim samym prawdopodobieństwem dowolny z elementów bazy. Przez bramkę Hadamarda przepuszczamy również kubit pomocniczy, który początkowo był w stanie $|1\rangle$. Teraz będzie on w stanie superpozycji $\frac{1}{\sqrt{2}}(|0\rangle - |1\rangle)$. Kubit pomocniczy jest wykorzystywany do implementacji bramki kwantowej realizującej kluczową dla algorytmu funkcję. Jest to funkcja odpowiadająca za wskazanie szukanego elementu. Argumentem funkcji jest numer elementu z bazy danych i , a wartością jest 1, jeśli i reprezentuje szukany element i 0 w przeciwnym wypadku.

Działanie bramki kwantowej realizującej tę funkcję sprawi, że amplituda prawdopodobieństwa stanu bazowego odpowiadającego szukanemu elementowi stanie się ujemna, a pozostałe amplitudy pozostaną dodatnie (patrz rysunek 1, lewa kolumna). Niestety pomiar nie daje nam możliwości odróżnienia amplitud dodatnich od ujemnych. Musimy zrobić coś więcej. Stosujemy następną bramkę realizującą operację obrotu wokół wartości średniej [9]. Taka operacja powoduje wzmocnienie wskazanej amplitudy i zmniejszenie pozostałych amplitud (patrz rysunek 1, prawa kolumna).

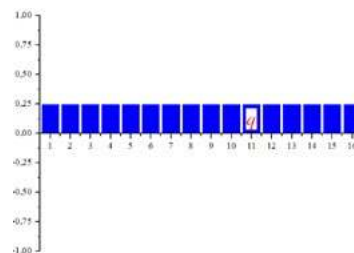
Amplituda prawdopodobieństwa szukanego elementu zachowuje się po kolejnych krokach iteracji jak sinusoida [10]

$$a_k \approx \sin\left(\frac{2k+1}{\sqrt{N}}\right).$$

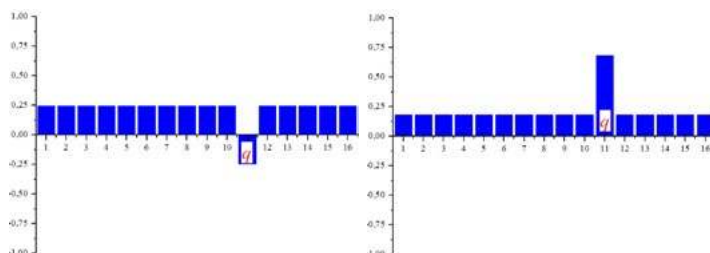
Proces przerywamy z chwilą zbliżenia się do maksimum, czyli po około $\frac{\pi}{4}\sqrt{N}$ krokach. Prawdopodobieństwo sukcesu, czyli prawidłowego odszukania elementu w bazie jest wówczas rzędu $1 - 1/N$. Dla przykładu, jeśli nasza baza zawiera milion elementów, to będziemy potrzebowali komputera kwantowego z rejestrem 20 bitowym ($10^{20} = 1\,048\,576$). Po wykonaniu tysiąca iteracji algorytmu Grovera znaleźlibyśmy szukany element z prawdopodobieństwem ok. 0,999999. Kilkukrotne powtórzenie algorytmu da nam już praktycznie absolutną pewność co do poprawności obliczeń. W przypadku przeszukiwania bazy za pomocą klasycznego komputera średnia liczba kroków wynosi $N/2$, a w skrajnym przypadku musimy wykonać N kroków. A zatem algorytm Grovera daje nam kwadratowe przyspieszenie w stosunku do algorytmów klasycznych.

Kodowanie RSA

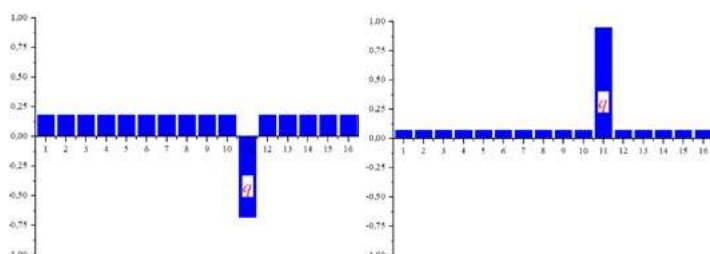
Popularność idei komputera kwantowego wzrosła gwałtownie po 1994 roku, kiedy to amerykański matematyk Peter Shor przedstawił algorytm, nazywany dziś algorytmem Shora, umożliwiający rozkład liczb naturalnych na czynniki pierwsze. Klasyczne algorytmy radzą sobie z tym problemem w czasie rosnącym wykładniczo wraz



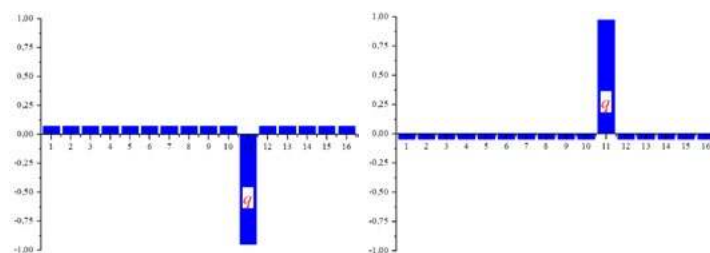
$$k = 0, a_0 = b_0 = 0,25, p_q = 1/16 = 0,0625.$$



$$k = 1, a_1 \approx 0,688, b_1 \approx 0,188, p_q \approx 0,473$$



$$k = 2, a_2 \approx 0,953, b_2 \approx 0,0781, p_q \approx 0,908$$



$$k = 3, a_3 \approx 0,980, b_3 \approx -0,0508, p_q \approx 0,961$$

Rysunek 1. Amplitudy prawdopodobieństwa a szukanego elementu i b pozostałych elementów w kolejnych trzech krokach k algorytmu Grovera podczas szukania elementu q w 16 elementowej bazie. Prawdopodobieństwo poprawnego odszukania elementu po kolejnych iteracjach oznaczono przez p_q .

z rozmiarem rozkładanej liczby, zaś algorytm Shora działający na komputerze kwantowym wykonywałby takie zadanie w czasie wielomianowym.

Dlaczego algorytm Shora dotyczący tak wydawałoby się akademickiego problemu jakim jest rozkład liczby na czynniki pierwsze rozbudził tak bardzo wyobraźnię naukowców i nie tylko? Otóż komputer kwantowy realizujący algorytm Shora zagroziłby metodzie RSA – najbardziej popularnej metodzie szyfrowania wiadomości, z której korzysta niemal każdy człowiek i prawie każda firma czy instytucja.

Algorytm RSA został opracowany w 1977 przez Rona Rivesta, Adiego Shamira oraz Leonarda Adlemana (nazwa jest akronimem utworzonym z pierwszych liter nazwisk autorów algorytmu). Jest on stosowany zarówno do szyfrowania, jak i do podpisów cyfrowych. RSA należy do algorytmów asymetrycznych. Oznacza to, że stosowane są dwa tzw. klucze – klucz publiczny i klucz prywatny. Pierwszy jest ogólnodostępny i służy do szyfrowania, a drugi jest tajny i służy do deszyfrowania wiadomości. Sam algorytm jest prosty.

- Znajdujemy dwie duże liczby pierwsze³ p i q (najlepiej o porównywalnej długości w bitach, ale równocześnie różniące się znacznie wartościami).
- Obliczamy iloczyn tych liczb $n = pq$. Tak wyliczona liczba n jest tzw. liczbą półpierwszą⁴.
- Obliczamy wartość tzw. funkcji Eulera

$$\varphi(n) = (p-1)(q-1).$$
- Wybieramy losowo liczbę e ($2 < e < \varphi(n)$) względnie pierwszą⁵ z $\varphi(n)$.
- Wyznaczamy liczbę $d < \varphi(n)$, taką, że reszta z dzielenia iloczynu $d \cdot e$ przez $\varphi(n)$ jest równa 1.

Wyznaczone liczby n i e tworzą tzw. klucz publiczny służący do szyfrowania, a liczby n i d tworzą klucz prywatny służący do deszyfrowania.

Proces szyfrowania przebiega następująco: dzielimy tekst na bloki, którym możemy w sposób jednoznaczny przypisywać liczby naturalne $m < n$. Sposób przypisywania tekstowi liczb m musi być oczywiście taki sam, jaki będzie stosował adresat podczas deszyfrowania. Ale obie osoby – nadawca i adresat stosują to samo oprogramowanie, więc z tym nie ma problemu. Następnie korzystamy z klucza publicznego (e, n) – każdą z liczb m podnosimy do potęgi e i wyliczamy resztę z dzielenia tej potęgi przez n

$$c = m^e \bmod n.$$

Powstałe w ten sposób liczby c wysyłamy adresatowi. Teraz adresat deszyfruje wiadomość – wykonuje podobną operację co nadawca tylko korzysta ze swojego prywatnego (tajnego) klucza (d, n). Każdą z odebranych liczb c podnosi do potęgi d i wylicza resztę z dzielenia tej potęgi przez n

$$m = c^d \bmod n.$$

Wyliczone w ten sposób liczby są równe liczbom m , które wysłał nadawca. Na koniec wyliczone liczby adresat zamienia na odpowiednie zbitki znaków tekstowych według tej samej reguły, którą stosował nadawca. Oczywiście klucz publiczny nie nadaje się do deszyfrowania. Nie odtworzy nam liczb m . Musimy znać klucz prywatny.

Przyjrzyjmy się temu algorytmowi na przykładzie z małymi liczbami pierwszymi.

- Weźmy $p = 5$ i $q = 7$.
- Utworzona z nich liczba półpierwsza ma wartość $n = pq = 35$.
- Wartość funkcji Eulera wynosi

$$\varphi(n) = (p-1)(q-1) = 24.$$
- Losujemy liczbę e względnie pierwszą z 24. Załóżmy, że wylosowaliśmy liczbę $e = 5$. Jest to pierwsza z liczb względnie pierwszych z 24 spełniających warunek

$$2 < e < \varphi(n).$$
- Szukaną liczbą d jest 29 ($29 \cdot 5 = 145$, $145 \bmod 24 = 1$).

- Klucz publiczny stanowią liczby (5, 35), a klucz prywatny liczby (29, 35).

Założmy, że chcemy zakodować fragment tekstu, któremu odpowiada liczba $m = 12$. Liczymy

$$c = 12^5 \bmod 35 = 17.$$

Wysyłamy odbiorcy liczbę 17. Teraz odbiorca deszyfruje wiadomość:

$$m = 17^{29} \bmod 35.$$

Tak dużej potęgi nie policzymy na komputerze. Języki programowania mają zwykle ograniczenia na wielkość liczb. Ale nas interesuje wartość tej potęgi po operacji dzielenia modulo 35. Możemy zatem wykonać takie działanie

$$\begin{aligned} m &= 17^{8+8+8+4+1} \bmod 35 \\ &= \left((17^8 \bmod 35)^3 \cdot (17^4 \bmod 35) \cdot 17 \right) \bmod 35 \\ &= (16^3 \cdot 11 \cdot 17) \bmod 35 = 765962 \bmod 35 = 12. \end{aligned}$$

Jak widać odtworzyliśmy liczbę zakodowaną przez nadawcę.

Przeglądając się algorytmowi RSA widzimy, że jedyne czego nie znamy, aby złamać szyfr to znajomość liczby d . Ale ponieważ znamy liczby n i e (klucz publiczny jest możliwy do zdobycia), to wystarczyłoby znaleźć liczby p i q , których iloczyn jest równy n , a następnie znaleźć liczbę d , co jest proste i jednoznaczne, gdyż liczby e i n znamy. A zatem cały problem sprowadza się do rozłożenia liczby n na czynniki pierwsze. Jednak w przypadku dużych liczb pierwszych jest to bardzo skomplikowane i czasochłonne zadanie i w tym tkwi siła algorytmu RSA.

W 1991 roku firma RSA Security opublikowała listę dużych liczb półpierwszych [11] i ogłosiła zawody (RSA Factoring Challenge) [12] na rozkład tych liczb na czynniki pierwsze. Za rozłożenie niektórych liczb z listy wyznaczono pieniężne nagrody. Konkurs rozpoczął się w marcu 1991 roku, a zamknięty został w maju 2007 roku. Pierwsza liczba z listy RSA-100 (100 oznacza liczbę cyfr dziesiętnych, w zapisie binarnym liczba ma 330 pozycji) została rozłożona już w kwietniu 1991 roku przez Arjena Lenstra.

RSA-100 = 152260502792253336053561837813263742
 9718068114961380688657908494580122963258952897
 654000350692006139
 RSA-100 = 379752279369436739228088727554456278
 54565536638199 × 400946909509208810306837352927
 61468389214899724061

Największą z rozłożonych na czynniki pierwsze liczb RSA była liczba RSA-250 (250 cyfr dziesiętnych, 829 bitów). Wynik ogłoszono w 2020 roku, 13 lat po zamknięciu konkursu. Biorąc jako jednostkę referencyjną procesor Intel Xeon Gold 6130 o szybkości 2,1 GHz obliczenia wymagały użycia 2700 takich jednostek przez rok. Liczby

³ Liczba pierwsza jest liczbą naturalną większą od 1 mającą tylko dwa dzielniki naturalne – jedynkę i siebie samą.

⁴ Liczba półpierwsza jest liczbą naturalną będącą iloczynem dokładnie dwóch, niekoniecznie różnych liczb pierwszych, np. liczba $15 = 3 \cdot 5$ jest liczbą półpierwszą.

⁵ Dwie liczby a i b są względnie pierwsze, jeśli największym ich wspólnym dzielnikiem jest 1, np. 6 i 35 są względnie pierwsze

powyżej RSA-250 do dzisiaj nie zostały rozłożone na czynniki pierwsze.

To właśnie trudności rozkładu dużych liczb na czynniki pierwsze stanowią siłę klasycznej kryptografii opartej na algorytmach typu RSA. Najszybszy obecnie klasyczny algorytm faktoryzujący liczby większe niż 10^{100} wymaga wykonania liczby operacji rzędu $e^{1.9 \cdot (\ln N)^{1/3} (\ln \ln N)^{2/3}}$, gdzie N jest faktoryzowaną liczbą. Dla liczby $N \sim 10^{400}$ dostajemy liczbę operacji rzędu 10^{29} . Zakładając moc obliczeniową komputera rzędu 1 TFLPOS (10^{12} operacji zmiennoprzecinkowych na sekundę) dostajemy czas obliczeń rzędu 10^9 lat. Tymczasem wspomniany wyżej algorytm Shora wykonany na komputerze kwantowym poradziłby sobie z zadaniem faktoryzacji liczby w czasie wielomianowym. Liczba operacji byłaby rzędu $(\ln N)^2 \ln N$. Dla liczby rzędu 10^{400} dostajemy liczbę operacji rzędu 10^4 , czyli 25 rzędów mniejszą niż w przypadku najszybszych klasycznych algorytmów!

Algorytm Shora

Algorytm przedstawiony przez Petera Shora w 1994 roku [13] składa się z dwóch części: klasycznej i kwantowej. Część klasyczna rozpoczyna się od wylosowania liczby naturalnej a spełniającej warunek $1 < a < N$. Następnie znajdujemy $\text{NWD}(a, N)$, czyli największy wspólny dzielnik liczb a i N . Do tej operacji możemy zastosować wydajny klasyczny algorytm Euklidesa [14]. Jeśli okazałoby się, że $\text{NWD}(a, N) \neq 1$, to a jest pierwszą z szukanych liczb pierwszych, a drugą jest $N / \text{NWD}(a, N)$. W przeciwnym wypadku uruchamiamy kwantową część algorytmu Shora.

Jeśli doszliśmy do etapu uruchomienia części kwantowej, to liczby a i N są liczbami względnie pierwszymi. Musimy teraz znaleźć najmniejszą liczbę naturalną r taką, że reszta z dzielenia potęgi a^r przez N wynosi 1, czyli

$$a^r = 1 \pmod{N}.$$

Dla znalezionej liczby r liczba $a^r - 1$ jest całkowitą wielokrotnością N . Jeżeli wyznaczona liczba r jest nieparzysta, to musimy wrócić do pierwszego kroku – wylosować inną liczbę a i powtórzyć obliczenia. W przeciwnym wypadku liczbę $a^r - 1$ możemy rozłożyć na iloczyn dwóch liczb naturalnych

$$a^r - 1 = (a^{r/2} - 1)(a^{r/2} + 1).$$

Jest oczywiste, że liczba N nie może być dzielnikiem liczby $a^{r/2} - 1$, gdyż r , a nie $r/2$ jest najmniejszą liczbą naturalną, dla której N dzieli się bez reszty przez $a^r - 1$. Szukamy zatem największego wspólnego dzielnika obu liczb

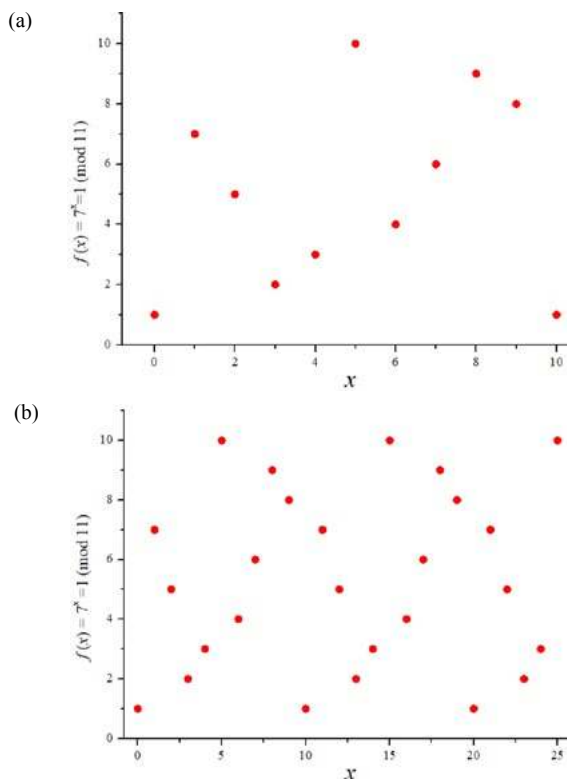
$$d = \text{NWD}(N, a^{r/2} - 1).$$

Jeśli $d = 1$, to musimy powtórzyć obliczenia dla nowego a . W przeciwnym wypadku znaleźliśmy nietrywialne dzielniki liczby N . Są nimi d i $N / \text{NWD}(N, a^{r/2} - 1)$.

Jeżeli liczba N jest także dzielnikiem liczby $a^{r/2} + 1$ (nie musi tak być), to obie liczby: $a^{r/2} + 1$ i $a^{r/2} - 1$ są szukanyymi liczbami p i q .

Zadaniem kwantowej części algorytmu jest znalezienie omawianej wyżej liczby r . Funkcja

$$f(x) = a^x = 1 \pmod{N}$$



Rysunek 2. Wykres funkcji $f(x) = 7^x = 1 \pmod{11}$ (a) w przedziale o długości równej okresowi funkcji, (b) w przedziale o długości równej dwa i pół okresu.

jest funkcją okresową z okresem równym szukanej liczbie r . Jednak w przeciwieństwie do prostych funkcji okresowych typu sinus czy cosinus ta funkcja nie cechuje się żadną prawidłowością, która na podstawie wartości z przedziału z wewnątrz okresu mogłaby wskazywać na prawdopodobną wartość okresu (patrz rysunek 2). Dlatego w przypadku szukania liczby r musielibyśmy wyliczać wartości funkcji tak długo, aż znaleźlibyśmy wartość taką, która już wystąpiła.

Komputer kwantowy zdolny do rozwiązania problemu ma dwa rejestry: wejściowy i wyjściowy (choć w tym przypadku ich nazwy mogą być mylące). Każdy z nich powinien mieć n kubitów, przy czym n jest liczbą naturalną spełniającą nierówność $2^n \geq N$. Dla osiągnięcia lepszej efektywności rejestry powinny mieć rozmiar co najmniej 2 razy większy. Podwojenie rozmiaru rejestru zapewni większą liczbę pełnych okresów funkcji. Można pokazać, że podwojenie rozmiaru rejestru daje wystarczającą dokładność znalezienia r .

Na początku kubity obu rejestrów wprowadzamy w stan superpozycji za pomocą n bramek Hadamarda. Stan rejestru jest teraz superpozycją stanów poszczególnych kubitów

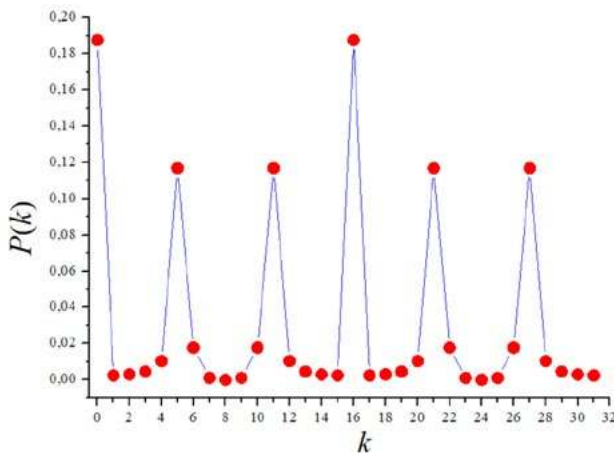
$$|x\rangle = \frac{1}{\sqrt{2^n}} \sum_{i=0}^{2^n-1} |i\rangle.$$

Dla przykładu, gdy $n = 4$ mamy

$$|0\rangle \equiv |0000\rangle, |1\rangle \equiv |0001\rangle, |2\rangle \equiv |0010\rangle, \dots, |15\rangle \equiv |1111\rangle.$$

Superpozycja stanów rejestru wejściowego $|x\rangle$ i wyjściowego $|y\rangle$ ma postać

$$|x, y\rangle = \frac{1}{\sqrt{2^n}} \sum_{i=0}^{2^n-1} |i, 0\rangle.$$



Rysunek 3. Prawdopodobieństwa stanów k po operacji IQFT dla $a = 5$, $N = 21$ i $n = 5$ ($2^5 = 32$).

Na tej superpozycji przeprowadzamy operację tak zwanego potęgowania modularnego i otrzymujemy stan

$$|x, y\rangle = \frac{1}{\sqrt{2^n}} \sum_{i=0}^{2^n-1} |i, a^i \pmod{N}\rangle.$$

Na przykład, jeśli $a = 2$ i $N = 15$, to dostajemy

$$|x, y\rangle = \frac{1}{4} (|0, 1\rangle + |1, 2\rangle + |2, 4\rangle + |3, 8\rangle + |4, 1\rangle + |5, 2\rangle + |6, 4\rangle + |7, 8\rangle + |8, 1\rangle + |9, 2\rangle + |10, 4\rangle + |11, 8\rangle + |12, 1\rangle + |13, 2\rangle + |14, 4\rangle + |15, 8\rangle).$$

Rejestr wyjściowy zawiera teraz powtarzający się wzorzec czterech elementów, a rejestr wejściowy jest jednorodną superpozycją wszystkich możliwych stanów. Informacje o okresie są przechowywane w relacjach między rejestrem wejściowym a wyjściowym, gdyż kubity tych rejestrów są ze sobą splątane. Aby znaleźć okres naszej funkcji przeprowadzamy na rejestrze wejściowym odwrotną kwantową transformatę Fouriera IQFT (ang. the Inverse Quantum Fourier Transform). Ponieważ kubity rejestru wejściowego są splątane z kubitami rejestru wyjściowego operacja IFQF wywołuje szereg interferencji kwantowych. Dla większości stanów będą to interferencje destrukcyjne i ich amplitudy spadną do zera lub prawie do zera, a dla kluczowych stanów interferencje spowodują, że ich amplitudy prawdopodobieństwa staną się bardzo wysokie w porównaniu z pozostałymi. Po tej operacji stan rejestru wejściowego przyjmie postać

$$|x\rangle = \frac{1}{\sqrt{2^n L}} \sum_{k=0}^{2^n-1} \sum_{l=0}^{L-1} e^{-2\pi i \cdot k(q_0 + lr)/2^n} |k\rangle.$$

Liczba L w powyższym wyrażeniu jest liczbą okresów r mieszczących się w 2^n , a liczba q_0 jest nieznaną przypadkową liczbą naturalną, która będzie inna po każdym uruchomieniu algorytmu. Jeśli teraz zmierzmy stan rejestru wejściowego, to superpozycja stanów dozna kolapsu do jednego z 2^n stanów bazowych $|k\rangle$ z prawdopodobieństwem

$$P(k) = \frac{1}{2^n L} \left| \sum_{l=0}^{L-1} e^{-2\pi i \cdot klr/2^n} \right|^2.$$

Zauważmy, że prawdopodobieństwo to nie zależy od nieznanego q_0 . Wynika, to z faktu, że

$$\left| e^{-2\pi i \cdot kq_0/2^n} \right|^2 = 1.$$

Możliwe są dwie sytuacje: $e^{-2\pi i \cdot kr/2^n} = 1$ lub $e^{-2\pi i \cdot kr/2^n} \neq 1$. W pierwszym przypadku dostajemy prawdopodobieństwo:

$$P(k) = \frac{L}{2^n}.$$

W drugim przypadku dostajemy

$$P(k) = \frac{1}{2^n L} \left| \frac{1 - e^{-2\pi i \cdot kLr/2^n}}{1 - e^{-2\pi i \cdot kr/2^n}} \right|^2.$$

Jak widzimy w pierwszym przypadku prawdopodobieństwo nie zależy od k . Można pokazać [15], że wówczas możliwych jest tylko r stanów z niezerowym prawdopodobieństwem równym

$$P(k) = \frac{L}{2^n} = \frac{1}{r}.$$

A zatem kilkukrotne wykonanie algorytmu Shora da nam kilka wartości k oddalonych od siebie o L lub wielokrotność L . Możemy teraz znaleźć r korzystając z algorytmu Euklidesa. Tak byłoby w rozważanym wyżej przypadku z $a = 2$, $n = 4$ i $N = 15$. Dostawalibyśmy w wyniku pomiaru liczby ze zbioru 0, 4, 8, 12. To pozwoliłoby nam na znalezienie $L = 4$, a w konsekwencji r ($r = 16/4 = 4$). Te obliczenia wykonujemy na klasycznym komputerze.

Przyjrzyjmy się drugiemu przypadkowi. Tutaj też do znalezienia r użyjemy klasycznego komputera. Teraz w dla zdecydowanej większości stanów $|k\rangle$ prawdopodobieństwo będzie niezerowe, ale dla takich, dla których liczba 2^n jest wielokrotnością kr lub jest bliska tej wielokrotności prawdopodobieństwo będzie dużo większe od pozostałych. Taką sytuację pokazuje rysunek 3, dla $a = 5$, $N = 21$ i $n = 5$ ($2^5 = 32$). Jak widzimy wyraźnie większe prawdopodobieństwa dostajemy dla stanów: 0, 5, 11, 16, 21 i 27. Tych stanów jest 6, dokładnie tyle ile wynosi okres r badanej funkcji. Spośród wymienionych stanów największe prawdopodobieństwo cechuje stany, dla których zachodzi dokładna równość

$$kr = l2^n, \quad 0 \leq l < r.$$

Dla pozostałych mamy tylko przybliżoną równość. Ich amplituda prawdopodobieństwa rozkłada się na sąsiednie stany, dlatego jest mniejsza niż w przypadku stanów, dla których zachodzi powyższa równość. (patrz rysunek 3).

Załóżmy, że w wyniku pomiaru dostaniemy liczbę k . Możemy zatem wyliczyć ułamek $k/2^n$, który jest równy (dokładnie lub w przybliżeniu) wartości ułamka l/r

$$\frac{k}{2^n} \approx \frac{l}{r}, \quad 0 \leq l < r.$$

Nie znamy co prawda liczby l , ale wiemy, że l/r jest liczbą wymierną i możemy znaleźć r stosując metodę tak zwanych ułamków łańcuchowych [16]. Zapisujemy liczbę $k/2^n$ w postaci ułamka łańcuchowego rozwijając go do momentu, gdy ułamkowa część mianownika będzie dosta-

tecnie mała (lub wręcz zerowa). Tę część pomijamy i dostajemy przybliżenie liczby l/r . W ten sposób powinniśmy dostać $r-1$ (pomijając trywialny przypadek dla $k=0$) ułamków: $1/r$, $2/r$, $(r-1)/r$. Formalny schemat postępowania łącznie z algorytmem możemy znaleźć w [17]. Tu pokażę, bez wyprowadzenia, jak to wygląda na rozważanym wyżej przykładzie dla $a=5$, $N=21$ i $n=5$ i stanami specjalnymi $k=5, 11, 16, 21$ i 27 . Liczby $k/2^n$ wynoszą odpowiednio $0,15625$, $0,34375$, $0,5$, $0,65625$ i $0,84375$. Możemy je rozłożyć na następujące ułamki łańcuchowe i ich przybliżenia:

$$0,15625 = 0 + \frac{1}{6 + \frac{1}{2 + \frac{1}{2}}} \approx \frac{1}{6} \approx 0,15384,$$

$$0,34375 = 0 + \frac{1}{2 + \frac{1}{1 + \frac{1}{10}}} \approx \frac{1}{2 + \frac{1}{1}} = \frac{1}{3} \approx 0,33333,$$

$$0,5 = 0 + \frac{1}{1 + \frac{1}{1}} = \frac{1}{2} = 0,5,$$

$$0,65625 = 0 + \frac{1}{1 + \frac{1}{1 + \frac{1}{1 + \frac{1}{10}}}} \approx \frac{1}{1 + \frac{1}{1 + \frac{1}{1}}} = \frac{2}{3} \approx 0,66667,$$

$$0,84375 = 0 + \frac{1}{1 + \frac{1}{5 + \frac{1}{2 + \frac{1}{2}}}} \approx \frac{1}{1 + \frac{1}{5}} = \frac{5}{6}.$$

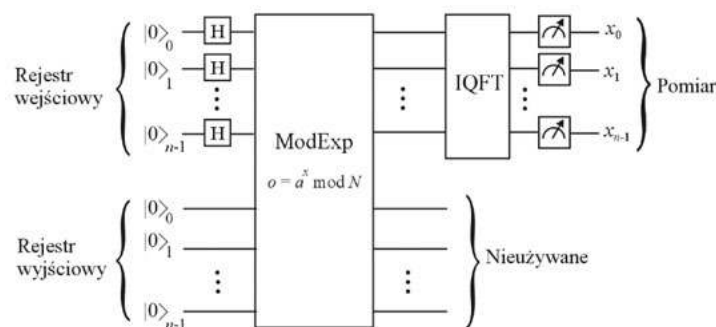
Najmniejszym wspólnym mianownikiem liczb $\frac{l}{r} = \frac{1}{6}, \frac{1}{3}, \frac{1}{2}, \frac{2}{3}, \frac{5}{6}$ jest 6. Po sprowadzeniu tych ułamków do wspólnego mianownika 6 ich liczniki będą liczbami naturalnymi $1 \leq l < 6$. Wobec tego dostajemy:

$$\frac{l}{r} = \frac{1}{6}, \frac{2}{6}, \frac{3}{6}, \frac{4}{6}, \frac{5}{6}.$$

Algorytm Shora, podobnie jak opisany wyżej algorytm Grovera i inne algorytmy kwantowe jest algorytmem probabilistycznym. Prawdopodobieństwo uzyskania którejś z r wartości k z maksymalnym prawdopodobieństwem jest dużo większe od pozostałych. Ale prawdopodobieństwa pozostałych stanów są niezerowe (poza niektórymi) i może się zdarzyć, że pomiar da nam inną wartość k niż te z maksymalnymi prawdopodobieństwami. Dlatego dla zwiększenia dokładności i pewności obliczeń algorytm musimy powtórzyć kilka razy.

Schemat obwodu kwantowego realizującego opisane operacje przedstawia rysunek 4.

Czy to działa? Czy powinniśmy się obawiać łamania naszych zaszyfrowanych wiadomości i podpisów elektronicznych? Po raz pierwszy działanie algorytmu Shora



Rysunek 4. Schemat układu kwantowego realizującego kwantową część algorytmu Shora.

na komputerze kwantowym zademonstrowano w 2001 roku. IBM za pomocą swojego komputera kwantowego z siedmio-kubitowym rejestrze rozłożył wtedy na czynniki pierwsze liczbę 15 ($15 = 3 \times 5$). Dopiero 11 lat później, bo w 2012 roku inna grupa badawcza sfaktoryzowała liczbę 21 ($21 = 3 \times 7$). W roku 2019 podjęto próbę użycia algorytmu Shora na 20-kubitowym komputerze IBM Q System One do sfaktoryzowania liczby 35. Niestety próba zakończyła się niepowodzeniem z powodu zbyt dużej liczby błędów. Jak widzimy na razie szyfrowanie metodą RSA jest niezagrożone.

W poprzednim odcinku zapowiadałem zajęcie się kwantową kryptografią, ale w trakcie pracy postanowiłem zmienić kolejność tematów. W konsekwencji kryptografią kwantową zajmę się w następnej, ostatniej już części cyklu.

dr Jan Kurzyk

Katedra Fizyki Politechniki Krakowskiej

LITERATURA

- [1] J. Kurzyk, *Za co Nobel 22. Część 1. Interpretacja kopenhaska, paradoks EPR, stany splątane*. Fizyka w Szkole 1, 2023.
- [2] J. Kurzyk, *Za co Nobel 22. Część 2. Nierówność Bella, eksperymenty ze splątanymi fotonami*. Fizyka w Szkole 2, 2023.
- [3] J. Kurzyk, *Za co Nobel 22. Część 3. Kubyty, stany Bella, teleportacja*. Fizyka w Szkole 3, 2023.
- [4] J. Kurzyk, *Za co Nobel 22. Część 4. Bramki kwantowe*. Fizyka w Szkole 4, 2023.
- [5] P. Benioff, *The computer as a physical system: A microscopic quantum mechanical Hamiltonian model of computers as represented by Turing machines*. Journal of Statistical Physics, vol 22 no 5.
- [6] https://en.wikipedia.org/wiki/Quantum_computing [Dostęp 07.07.2023].
- [7] M. Brooks, *Quantum computers: what are they good for?* Nature, 25 maja, 2023. Wersja elektroniczna: <https://www.nature.com/articles/d41586-023-01692-9> [Dostęp 07.07.2023].
- [8] J. Bobroff, *Czy da się przechodzić przez ścianę?* JK Wydawnictwo, Łódź 2022. Przekład Łukasz Musiał.
- [9] M. Sawerwain, J. Wiśniewska, *Informatyka Kwantowa. Wybrane obwody i algorytmy*. PWN, Warszawa 2015.
- [10] <https://docplayer.pl/154602183-Komputer-quantowy-zasady-funkcjonowania-dr-hab-inz-krzysztof-giario-politechnika-gdanska-wydzial-eti.html> [Dostęp 07.07.2023].
- [11] https://en.wikipedia.org/wiki/RSA_numbers. [Dostęp 11.07.2023].
- [12] https://pl.wikipedia.org/wiki/RSA_Factoring_Challenge. [Dostęp 11.07.2023].
- [13] P. W. Shor Polynomial-Time Algorithms for Prime Factorization and Discrete Logarithms on a Quantum Computer. <https://arxiv.org/pdf/quant-ph/9508027.pdf>. [Dostęp 11.07.2023].
- [14] https://pl.wikipedia.org/wiki/Algorytm_Euklidesa. [Dostęp 11.07.2023].
- [15] <https://prefetch.eu/learn/concept/shors-algorithm>. [Dostęp 15.07.2023].
- [16] https://pl.wikipedia.org/wiki/U%C5%82amek_%C5%82a%C5%84cuchowy. [Dostęp 25.08.2023].
- [17] <https://stem.mitre.org/quantum/quantum-algorithms/shors-algorithm.html>. [Dostęp 15.07.2023].

Żywoty fizyków

Piotr Curie

Tadeusz Wibig

Piotr Curie (1859-1906) w naszym kraju jest znany powszechnie. Wszyscy wiedzą, że był mężem największej polskiej uczonej Marii Skłodowskiej-Curie. Ale poza tym Piotr Curie był także naukowcem i to naukowcem dużego formatu. Nim spotkał w Paryżu Marię i raptownie zmienił swe naukowe zainteresowania, miał na koncie osiągnięcia, które trafiły do podręczników szkolnych, a ostatnio wdzierają się szeroko do naszych domów ułatwiając i umilając nam życie codzienne.

Rodzina Curie pochodziła z Miluzy w Alzacji. Jego dziadek Paul Étienne François Gustave, jak i jego ojciec Eugène byli lekarzami. Eugène studiował w Paryżu nauki przyrodnicze i medycynę, był nawet autorem kilku książek między innymi o trujących grzybach, rosiczce i zastosowaniu miedzi w medycynie. Wierzył on, że intelekt i osobowość jego syna mogą być najlepiej rozwinięte poprzez prywatne korepetycje i być może ten niekonwencjonalny charakter edukacji Piotra sprawdził się, bo w wieku 14 lat zaczął ujawniać się jego talent i zdolności matematyczne. Spowodowało to, że mając 16 lat rozpoczął studia na Wydział Nauk na Sorbonie i tamże uzyskał dyplom z fizyki w 1878 roku. Z powodów finansowych nie mógł kontynuować studiów i musiał podjąć pracę na uczelni jako kiepsko opłacany laborant i tamże dokonał swojego pierwszego wielkiego odkrycia.

W tym miejscu trzeba wspomnieć o tym, że Piotr miał brata, starszego o cztery lata Pawła Jakuba, nazywanego na ogół po prostu **Jacques'em Curie**. W cieniu innych osób w tym wielu noblistów noszących to samo nazwisko, informacje o dzieciństwie i młodych latach Jakuba zagubiły się. Czegoś dowiadujemy się dopiero, gdy w końcu lat 70 obaj bracia pracowali jako asystenci w laboratoriach na Wydziale Nauk w Paryżu, Piotr w Laboratorium Fizyki, a Jakub w Laboratorium Mineralogii.

Jakub, być może jako ten starszy i zaawansowany naukowo bardziej, nadał kierunek działań, jakie wspólnie podjęli. Na początku XIX wieku znany z osiągnięć w dziedzinie optyki **Sir David Brewster** zajął się zjawiskiem znanym już od stu lat polegającym z grubsza na tym, że zmiany temperatury niektórych ciał, kryształów prowadzą do indukowania się na ich powierzchniach ładunków elektrycznych. Nazwał to efektem **piezoelektrycznym**. W końcu XIX wieku był to temat zajmujący wielu wybitnych fizyków, wśród których wymienić należy też **Lorda Kelvina** i **Antoine'a Césara Becquerela**, a w Paryżu w Laboratorium Mineralogii **Charlesa Friedela**. To on zatrudnił Jakuba, który razem z bratem badał piezoelektryczność w roku 1880. W zasadzie bez większych sukcesów,



ale metody eksperymentalne opracowane przez nich pozwoliły im na odkrycie i zbadanie mechanizmów stojących za efektem **piezoelektrycznym**, pojawianiem się potencjałów elektrycznych na powierzchniach ciał poddawanych różnym naprężeniom.

Bracia nie docenili wagi swoich odkryć. W XIX wieku nikt nie przewidywał, jakie może mieć ono znaczenie. Nie wiedzieli, że przydarzy się ludzkości I Wojna Światowa i po morzach zaczną grasować łodzie podwodne. W 1917 roku **Paul Langevin**, doktorant Piotra, a także bliski znajomy Marii (inna historia) opracował przetwornik ultradźwiękowy do użytku na łodziach podwodnych i tak wynalazł sonar. Dziś zastosowań piezoelektryków policzyć nie sposób, każdy ma je choćby w telefonie, w komputerze, czy w zegarku. Ale wróćmy do roku 1880.

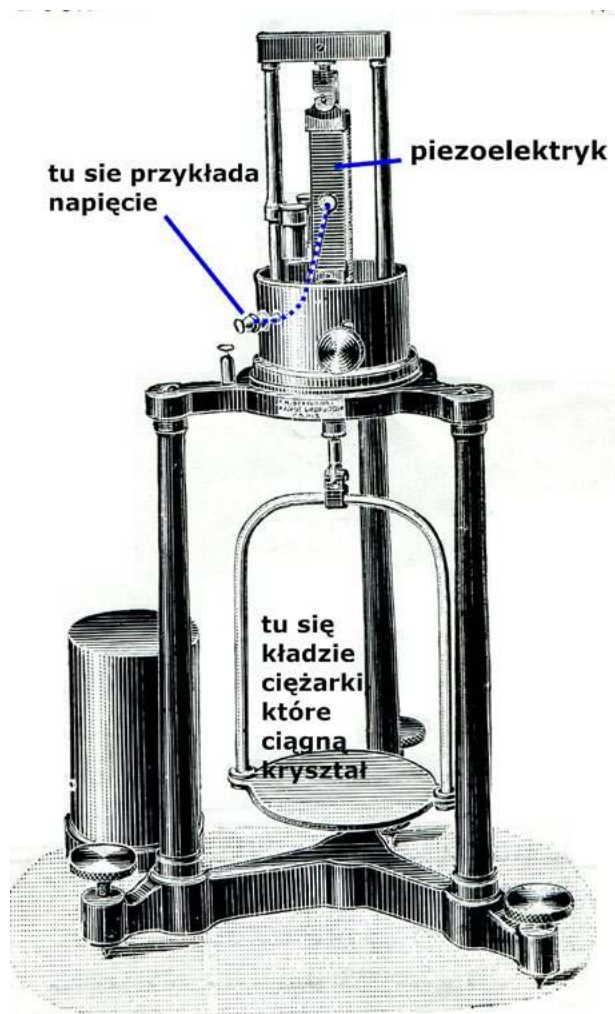
Jak powiedzieliśmy, bracia Curie nie zdali sobie sprawy z wagi swoich wyników i nie opublikowali szczegółów swoich badań w żadnym solidnym naukowym czasopiśmie. Zadowolili się notatkami w *Comptes rendus de l'Académie des sciences*. Były to krótkie doniesienia przedstawiane przez ich aktualnych kierowników: Charlesa Friedela, i Paula Desainsa, szefa Piotra. Notatek tych ukazało się w sumie siedem. W notatce z 1881 roku bracia sformułowali pięć praw dotyczących polaryzacji kryształu turmalinu pod wpływem przyłożonego ciśnienia:

1. Dwa końce turmalinu wydzielają taką samą ilość elektryczności lecz o przeciwnych znakach
2. Ilość elektryczności uwalniana przez pewien wzrost ciśnienia co do wartości jest równa, lecz ma przeciwny znak do ilości wytwarzanej przez równy jej spadek ciśnienia.

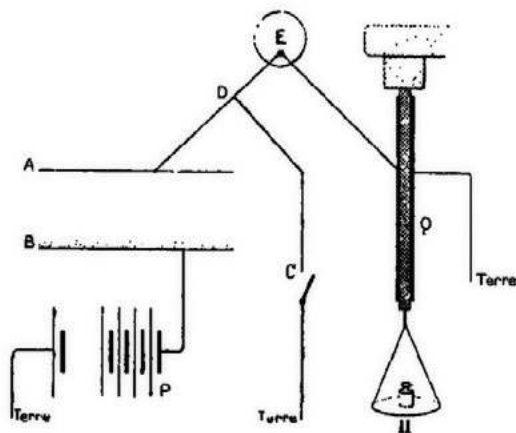
3. Wielkość ta jest proporcjonalna do zmiany ciśnienia.
4. Jest ona niezależna od długości turmalinu.
5. Dla tej samej zmiany nacisku na jednostkę powierzchni jest ona proporcjonalna do powierzchni.

Jeszcze w roku 1881 **Jonas Ferdinand Gabriel Lippmann**, jeden z bardziej znanych fizyków luksemburskich (nobel w 1908 za kolorową fotografię) i promotor doktoratu Marii Skłodowskiej Curie przewidział teoretycznie, że efekt piezoelektryczny powinien dać się odwrócić. Nie tylko osiąga się różnicę potencjałów ściskając kryształy, ale, jeśli się do kryształu przyłoży pole elektryczne, to kryształ sam się ściśnie. Przewidywania te bracia Curie potwierdzili eksperymentalnie niemal natychmiast i wyniki opublikowali także w *Comptes rendus*.

Podsumowując swoje prace w kolejnej notatce wskazali też potencjalne zastosowania wykrytego przez siebie efektu. Może on służyć jako standardowe źródło potencjału elektrycznego, można też zbudować przyrząd do pomiaru pojemności elektrycznej, a także do precyzyjnego pomiaru niewielkich ładunków elektrycznych. I tak naprawdę, jak okazało się po latach, to ostatnie zastosowanie było przynajmniej dla jednego z braci bardzo cenne. Przyczyniło się w jakimś stopniu do otrzymania nagrody Nobla przez Piotra i Marię.



Rys. 1. Przynrząd do pomiaru napięcia



Rys. 2. Schemat aparatury do badania radioaktywności.

W sumie idea była bardzo prosta. Rysunek 1 pokazuje przyrząd braci Curie do pomiaru ładunku (napięcia). Na solidnym statywie wisi podłużny kryształ kwarcu. Z boków przyklejona jest do niego metalowa folia. Do niej doprowadzany jest interesujący nas ładunek (napięcie). Kryształ ten obciąża szalka od zwykłej wagi laboratoryjnej. Na niej kładzie się zwykle odważniki i ich sumaryczny ciężar pokazuje z jakim napięciem mamy do czynienia.

Schemat całej aparatury do badania radioaktywności małżeństwa Curie pokazuje schematycznie rysunek nr 2 zaczerpnięty z pracy doktorskiej Marii:

Widać na nim komorę jonizacyjną, w której Maria umieszczała badaną próbkę materiału radioaktywnego i elektrometr piezoelektryczny Piotra.

Przeskoczyliśmy chwilowo dwadzieścia lat, wróćmy więc do roku 1880 i doświadczeń dwóch braci Curie. Ich owocna współpraca nie trwała długo. W 1883 r. Jakub wyprowadził się z Paryża, gdyż otrzymał propozycję zostania wykładowcą mineralogii na Uniwersytecie w Montpellier. Piotr także zmienił miejsce pracy. Został kierownikiem laboratorium w nowo powstałej École Municipale de Physique et de Chimie Industrielles w Paryżu, a później został tam profesorem fizyki ogólnej i teorii elektryczności. Pozostał tam przez 12 lat. W 1895 r. uzyskał stopień doktora nauk za prace o magnetyzmie i wpływie temperatury na własności magnetyczne różnych ciał (*Propriétés magnétiques des corps à diverses températures*). Przy okazji odkrył prawo znane jako **prawo Curie**, którego nie należy mylić z **prawem Curie-Weissa**, choć oba dotyczą zachowania się namagnesowania w funkcji temperatury. Prawo Curie-Weissa opisuje namagnesowanie w pobliżu **punktu Curie (temperatury Curie)** i występuje w nim stała zwana **stałą Curie**.

W tym samym roku mniej więcej czasie poznał **Marię Skłodowską** i ożenił się z nią w roku 1895. Miało to olbrzymie znaczenie dla jego dalszej naukowej kariery, ale to temat na zupełnie inną historię.

W 1900 roku powrócił na Sorbonę jako profesor na Wydziale Nauk Ścisłych, a trzy lata później otrzymał wraz z Marią nagrodę Nobla z fizyki za badania nad zjawiskiem promieniotwórczości.

19 kwietnia 1906 roku przejechał go wóz konny.

Doświadczenie domowe:**Efekt piezoelektryczny****A. Potrzebne materiały**

1. Przetwornik piezoelektryczny czasem zwany buzzerem najlepiej w wersji z przylutowanymi przewodami, co oszczędzi nam trochę pracy. Do kupienia jest on niemal za grosze. Można go też wymontować z jakiegoś urządzenia, zabawki, samochodu, której wydaje dźwięki/.
2. Dioda LED.

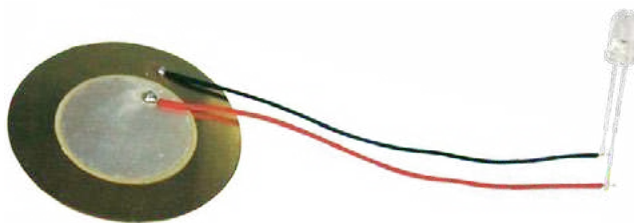
dla bardziej zaawansowanych dodatkowo:

3. Woltomierz (miernik uniwersalny).
4. Zwykła dioda prostownicza i kondensator elektrolityczny o pojemności kilku mikrofaradów, opornik o oporze rzędu kilkuset omów i prosty włącznik, bez którego oczywiście można się łatwo obejść.

B. Narzędzia: lutownica i niezbędne materiały do lutowania (konieczne dla ambitnych, ale w prostej wersji można się bez tego obyć)

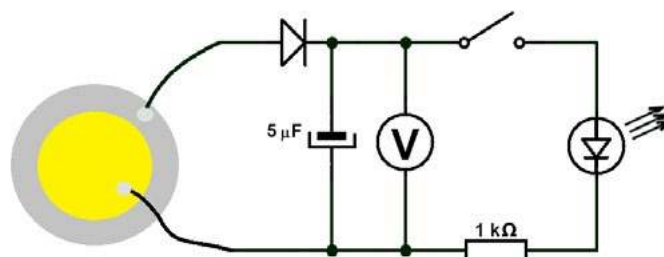
C. Kolejność czynności

1. Obie elektrody diody przylutowujemy do przewodów przetwornika i to właściwie wszystko.
2. Teraz wystarczy postukać palcem w płytkę przetwornika i dioda LED powinna zacząć mrugać.



dla ambitnych:

3. Można wmontować w nasz prosty układ kondensator i diodę prostowniczą.



4. Zginając ostrożnie przetwornik, stukając wąż ładujemy kondensator.
5. Gdy miernik pokaże napięcie 2-3 V, możemy rozładować go przez diodę LED. Będzie świeciła wyraźnie jaśniej i dłużej.

120 lat od pierwszego Nobla dla Marii Skłodowskiej-Curie i Piotra Curie

120 lat temu, 12 października 1903 r. Maria Skłodowska-Curie wraz z mężem Piotrem Curie zostali laureatami Nagrody Nobla w dziedzinie fizyki - za odkrycie zjawiska promieniotwórczości i badania nad nim. Drugą połowę tejże nagrody otrzymał Henri Becquerel, który pierwszy zaobserwował przenikliwe promieniowanie rudy uranu. Skłodowska była pierwszą kobietą, którą w ten sposób uhonorowano.

Wszystko zaczęło się od odkrycia przez Wilhelma Roentgena w roku 1895 niewidzialnego promieniowania elektromagnetycznego, które pozwoliło prześwietlać - nieprzejrzyste dla widzialnego światła - obiekty. Nowe zjawisko szybko znalazło zastosowanie w medycynie, a w roku 1901 Roentgen jako pierwszy został nagrodzony fizycznym Noblem.

7 listopada 1911 Maria Skłodowska-Curie otrzymała Nagrodę Nobla w dziedzinie chemii - za odkrycie polonu i radu, wydzielenie czystego radu i badanie właściwości

chemicznych pierwiastków promieniotwórczych. Był to pierwszy przypadek przyznania Nagrody Nobla po raz drugi, w dodatku w innej dziedzinie wiedzy. Do dziś nie udało się to żadnej innej kobiecie. W sumie tylko kilka osób otrzymało Nagrodę Nobla więcej niż raz, z czego tylko dwie w różnych dyscyplinach - drugą był Linus Pauling - z chemii i pokojową. John Bardeen miał dwa Nobla z fizyki, Frederick Sanger oraz Karl Barry Sharpless - po dwa z chemii.

Dwukrotnej noblistce udało się przekonać rząd Francji do przeznaczenia środków na budowę Instytutu Radowego (obecnie Instytut Curie). Powstał w 1914 w celu prowadzenia badań podstawowych z zakresu chemii, fizyki i medycyny dotyczących promieniotwórczości i izotopów promieniotwórczych, w tym leczenia nowotworów. Prace te przyniosły dalsze cztery nagrody Nobla - w tym dla córki Marii Skłodowskiej-Curie, Irène oraz zięcia - Frederica Joliot-Curie. (PAP - Nauka w Polsce)

W następnych wydaniach polecamy m.in.

- Kondycja nauczania fizyki w polskich szkołach, czyli kulisy 48. Zjazdu Fizyków Polskich
- Analityczne formułowanie funkcji położenia ciała w ruchu ukośnym
- Praca i energia w szkolnym programie fizyki



Fizyczne elementy w super nowoczesnym wytwarzaniu energii elektrycznej

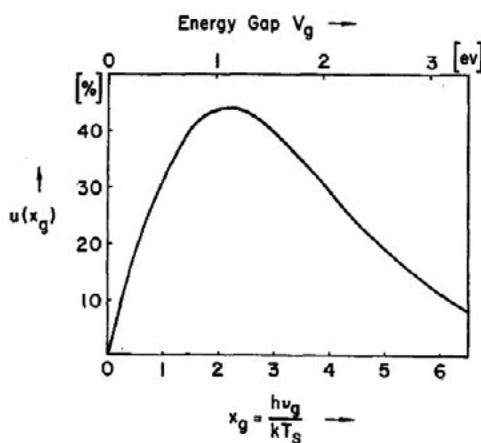
Na łamach *Physics World* poinformowano o wyróżnieniu pracowników z Massachusetts Institute of Technology (MIT) i National Renewable Energy Laboratory w USA (NREL) za zbudowanie ogniwa termofotowoltaicznego (TPV) o sprawności ponad 40%.¹ Ogniwa fotowoltaiczne mają wydajność od około 18% do 22%. MIT i NREL opracowali przełomowe ogniwo termofotowoltaiczne (TPV) o wydajności dużo lepszej.²

Kazimierz Mikulski

Już w 2022 r. I. Dumé³ pisze: „*Ogniwa termofotowoltaiczne osiągają sprawność 40 proc.*”. To przewyższa istniejące generatory ciepłne na paliwo stałe. Dokonano także znacznego przekroczenia średniej sprawności turbinyowego wytwarzania energii.

Elementy historii fotowoltaiki

W 1839 r. Alexandre Edmond Becquerel (1820-1891) eksperymentował z elektrodami metalowymi i elektroli-tem. Mając 19 lat, stworzył pierwsze na świecie ogniwo fotowoltaiczne. Uważał, że „*świecenie światłem na elektrodę zanurzoną w roztworze przewodzącym wytworzy prąd elektryczny*”. Badania wskazały, że energia fotowoltaiczna była mała. W eksperymencie umieścił chlorek



Fotografia 1. Ogniwo termofotowoltaiczne (TPV) zamontowane na radiatorze przeznaczonym do pomiaru wydajności ogniwa. Źródło: https://geekweek.interia.pl/technologie/news-termofotowoltaika-moze-generowac-prad-bez-przerwy-to-prawdzi,nld,6328657#utm_source=paste&utm_medium=paste&utm_campaign=chrome <https://physicsworld.com/a/thermophotovoltaic-cells-top-40-percent-efficiency/>

Rysunek 1. Zależność sprawności krańcowej $u(x_0)$ od przerwy energetycznej V_a półprzewodnika. Źródło: Detailed Balance Limit of Efficiency of p-n Junction Solar Cells* W. SHOCKLEY AND H. J. QUEISSER, JOURNAL OF APPLIED PHYSICS VOLJTME 32, K11MBER 3 M R ell, 1961 http://metronu.ulb.ac.be/npaully/art_2014_2015/shockley_1961.pdf s. 513

¹ Alinie LaPotin, Henry A.; <https://physicsworld.com/a/thermophotovoltaic-cells-top-40-percent-efficiency/>

² <https://swiatoze.pl/40-wydajnosci-z-ogniwa-termofotowoltaicznego-rekord-efektywnosci-tpv/> <https://android.com.pl/news/477097-ogniwo-termofotowoltaiczne/>

³ Isabelle Dumé współredaktor *Physics World*; autorka artykułów z zakresu fizyki

Fotografia 2. Alexandre Edmond Becquerel (1820-1891). Syn Antoine'a Césara i ojciec Henriego, jednego z odkrywców promieniowości. Źródło: https://en.wikipedia.org/wiki/Edmond_Becquerel

Fotografia 3. Willoughby Smith (1828-1891), angielski inżynier elektryk, odkrył fotoprzewodnictwo selenu. To doprowadziło do wynalezienia ogniw fotoelektrycznych. Źródło: https://www.wikiwand.com/en/Willoughby_Smith

Fotografia 4. William Grylls Adams (1836-1915), profesor filozofii w King's College w Londynie https://en.wikipedia.org/wiki/William_Grylls_Adams. Źródło: https://en.wikipedia.org/wiki/William_Grylls_Adams#/media/File:WilliamGryllsAdams.jpg



srebra w kwaśnym roztworze i oświetlił go, gdy był podłączony do platynowych elektrod. Powstało napięcie i popłynął prąd. Dziś określamy to zjawisko fotowoltaiczne, jako „**efekt Becquerela**”.⁴

W 1873 r. Willoughby Smith odkrył, że selen może działać jako fotoprzewodnik.⁵ W eksperymencie do powlekania elektrod platynowych użyto chlorku lub bromku srebra. Po oświetleniu elektrod powstawało napięcie.⁶

Selen jest podstawowym składnikiem cienkowarstwowych paneli słonecznych CIGS (oznacza miedź, ind, gal i selen).⁷ W 1873 r. wybrał pręty selenowe do obwodu testowego i wydawało się to dobrym rozwiązaniem w badaniach. Odkrył, że przewodnictwo tych prętów uległo zmianie pod wpływem silnego światła. Smith opisał badania w artykule przedstawionym na spotkaniu *Towarzystwa Inżynierów Telegrafów* 12 lutego 1873 r. pt. „*Właściwości elektryczne selenu i wpływ na nie światła*”⁸ i jako komunikat: „*Wpływ światła na selen podczas przejścia prądu elektrycznego*” w *Nature* – 20 lutego 1873 r. Odkrycie właściwości fotoelektrycznych selenu doprowadziło do rozwoju ogniw fotoelektrycznych.⁹

W 1876 r., W. G. Adams i R. E. Day zastosowali selen w fotowoltaice odkrytej przez Becquerela. Selenowe ogniwa były dowodem na to, że stały materiał może zmieniać światło w elektryczność bez ciepła czy ruchu.¹⁰ Elektryczność powstaje ze światła bez ruchomych części. Obaj odkryli, że „*ultra-czerwone lub ultrafioletowe promienie mają niewielki wpływ lub nie mają go wcale; a także, że intensywność działania zależy od mocy oświetlającej światła, będącej bezpośrednio pierwiastkiem kwadratowym tej mocy oświetlającej*”.¹¹

W 1883 r., Charles Fritts (1850-1903), stworzył pierwszą działającą komórkę selenową,¹² pokrywając selen cienką warstwą złota. Ogniwo osiągnęło współczynnik konwersji energii na poziomie 1-2% . Ogniwa wytworzyły prąd, jak donosi Fritts, „*który jest ciągły, stały i z znaczną siłą nie tylko pod wpływem światła słonecznego, ale także przyćmionego, rozproszonego światła dziennego*”. Fritts jest uznawany za wynalazcę ogniw solarnych, ale pierwsze ogniwa opatentowano w 1941 r.¹³

W 1887 r. został zaobserwowany efekt fotoelektryczny przez fizyka H. R. Hertza¹⁴, w którym światło jest wykorzystywane do uwalniania elektronów z powierzchni ciała stałego w celu wytworzenia energii. Stwierdził, że proces



Fotografia 5. Pierwszy na świecie panel słoneczny na dachu, zainstalowano w 1884 r. w Nowym Jorku. Wykorzystano selenowe ogniwa Frittsa. Źródło: https://en.wikipedia.org/wiki/Charles_Fritts#cite_note-1 <https://cleantechnica.com/2014/12/31/photovoltaic-dreaming-first-attempts-commercializing-pv/>

⁴ Edmond Becquerel: człowiek za panelami słonecznymi <https://solenergy.com.ph/solar-panel-philippines-edmond-becquerel/>

⁵ <https://swiatoze.pl/jasne-jak-slonce-historia-fotowoltaiki/>

⁶ https://en.wikipedia.org/wiki/Edmond_Becquerel

⁷ <https://technica.inc/cleantech/willoughby-smith/>

⁸ https://books.google.pl/books?id=5CpDAQAAMAAJ&pg=PA423&redir_esc=y#v=onepage&q&f=false

⁹ <https://technica.inc/cleantech/willoughby-smith/> W 2017 r. jego praca nad selenem uzyskała nagrodę Emmy w dziedzinie technologii i inżynierii za „Koncepcję of Opto-Electric Transduction” przyznana przez National Academy of Television Arts and Sciences. https://www.wikiwand.com/en/Willoughby_Smith

¹⁰ <https://solektro.pl/krotka-historia-o-przyszlosci-fotowoltaiki/>

¹¹ Adams, William Grylls; Dzień, Richard E (1 stycznia 1877). „V. Działanie światła na selen”. doi: 10.1098/rspl.1876.0024 <https://royalsocietypublishing.org/doi/10.1098/rspl.1876.0024> https://en.wikipedia.org/wiki/William_Grylls_Adams

¹² według Mariusa Paulescu i innych *Weather Modeling and Forecasting of PV System Operations*, Springer Verlag 2013, S. 1

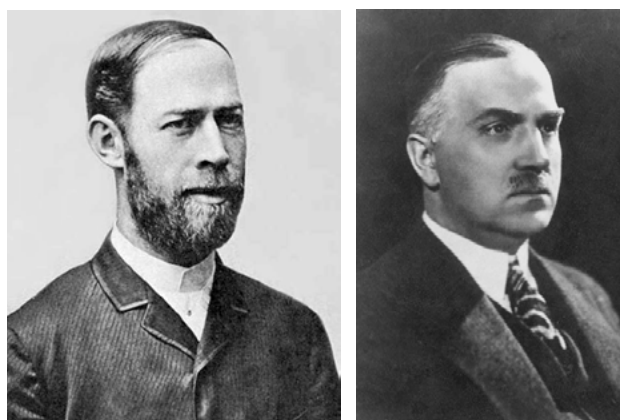
¹³ <https://swiatoze.pl/jasne-jak-slonce-historia-fotowoltaiki/>

¹⁴ <https://physics.info/photoelectric/>

ten wytwarzał więcej mocy pod wpływem światła ultrafioletowego niż intensywnego światła widzialnego. Późniejsze badania J. J. Thomsona wykazały, że ta zwiększona czułość była wynikiem nacisku światła na elektrony, cząstkę, którą on odkrył w 1897 r. Natomiast A. Einstein otrzymał Nagrodę Nobla za wyjaśnienie tego efektu.¹⁵

Wkład w rozwój technologii PV miał Polak, Jan Czochralski, który w 1916 r. opracował metodę produkcji krzemu monokrystalicznego. Ponad 20 lat później umożliwiła ona powstanie krzemowych ogniw słonecznych. **Metoda Czochralskiego** jest najstarszą i jedną z najpowszechniej wykorzystywanych na świecie do produkcji paneli słonecznych.

W latach 1904–05 G. Cove opracował „generator energii słonecznej”, który wystawił w budynku Metropole w Halifax w Nowej Szkocji w Kanadzie. Poza zdjęciem



Fotografia 6. Heinrich Rudolf Hertz (1857-1894), niemiecki fizyk; jako pierwszy udowodnił istnienie fal elektromagnetycznych przewidzianych przez J. C. Maxwella. Źródło: https://en.wikipedia.org/wiki/Heinrich_Hertz#/media/File:Heinrich_Rudolf_Hertz.jpg. Źródło: https://pl.wikipedia.org/wiki/Albert_Einstein#/media/Plik:Einstein1921_by_F_Schmutzer_2.jpg

Fotografia 7. Jan Czochralski (1885-1953) polski chemik, metaloznawca, wynalazca 19 sierpnia 1916 r. „metodę Czochralskiego”. Źródło: https://pl.wikipedia.org/wiki/Jan_Czochralski

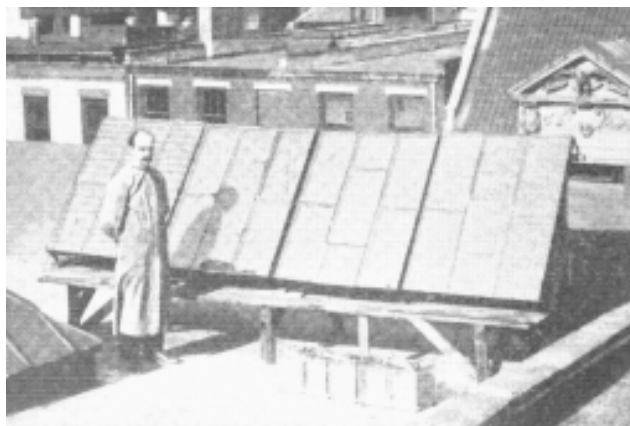
nie ma żadnych danych o tym panelu. Jego moc wyjściowa i wydajność były na tyle niezwykle, że amerykańscy inwestorzy wysłali eksperta do Halifax. Po badaniach sprowadzono Cove’a do Stanów Zjednoczonych, aby kontynuować rozwój urządzenia.¹⁷

W latach 1909–10 Cove utrzymywał warsztat w Nowym Jorku do produkcji słonecznych generatorów elektrycznych.¹⁸ Używając miedzi, **nieświadomie** dokonał przekształcenia generatora termoelektrycznego w „**generator termofotowoltaiczny**”. Działa tak samo jak fotowoltaiczne ogniwo słoneczne, ale na innej długości fali, a widmo słoneczne jest z zakresu od 0,5 do 2,9 elektronowoltów (eV), od podczerwieni do ultrafioletu. Półprzewodnik z pasmem wzbronionym od 1 do 1,7 eV wydajniej przekształca światło widzialne w energię elektryczną. Półprzewodnik z pasmem wzbronionym od 0,4 do 0,7 eV wydajniej przekształca krótkofalową energię słoneczną w podczerwieni w energię elektryczną (*generator termofotowoltaiczny*). Zauważył, że generator słoneczny przekształca zarówno ciepło, jak i światło w energię elektryczną. Generator termofotowoltaiczny dopasowuje się nie tylko do ogona podczerwieni widma słonecznego, ale także do bezpośredniego widma płonącego płomienia lub gorącej do czerwoności powierzchni emitującej, ogrzewanej przez spalanie drewna lub gazu ziemnego. Przetwarza dolną część widma widzialnego na energię elektryczną, nawet jeśli jest to bardzo nieefektywne.

Zbadał stopy cynku i antymonu i odkrył, że stop cynku o zawartości 40-42% daje najwyższe napięcie (w porównaniu do 35% cynku w $ZnSb$). Po odkryciu Zn_4Sb_3 , pasmo wzbronione tego półprzewodnika sprawiło, że przestał on działać, gdy był wystawiony na ciepło z pieca.

Działał lepiej, gdy był wystawiony na działanie energii słonecznej, przekształcając znacznie więcej widzialnego widma światła słonecznego w elektryczność.¹⁹

W 1954 r. D. Chapin, C. Fuller i G. Pearson zbudowali pierwsze krzemowe ogniwo fotowoltaiczne (PV) w la-



Fotografia 8. George Cove urodził się w Amherst w Nowej Szkocji w 1863 lub 1864 r. Źródło: https://en.wikipedia.org/wiki/George_Cove

Fotografia 9. Cove stoi przy swoim trzecim prototypie panelu słonecznego. Źródło: „Wytwarzanie energii za pomocą promieni słonecznych”, *Popular Electricity*, wyd. 2, nr 12, kwiecień 1910 r., s. 793. <https://journals.lib.unb.ca/index.php/MCR/article/view/17744/22231https://solar.lowtechmagazine.com/pl/2021/10/how-to-build-a-low-tech-solar-panel.html>

¹⁵ Mikulski K., 100 rocznica przyznania Nagrody Nobla Einsteinowi za interpretacji efektu fotoelektrycznego zewnętrznego. „Fizyka w Szkole z Astronomią” nr 4, lipiec/sierpień 2021, s. 26-29

¹⁶ Mikulski K., Jan Czochralski i jego rok 2013 „Fizyka w Szkole” nr 4, lipiec/sierpień 2013, s. 4-7

¹⁷ Czytaj: <https://solar.lowtechmagazine.com/2021/10/how-to-build-a-low-tech-solar-panel.html>

¹⁸ <https://diysolarforum.com/threads/george-cove-a-forgotten-solar-power-pioneer.29549/>

¹⁹ <https://solar.lowtechmagazine.com/2021/10/how-to-build-a-low-tech-solar-panel.html#firef:18>



Fotografia 10. Calvin Souther Fuller (1902-1994) amerykański chemik fizyczny w AT&T Bell Laboratories; pracował przez 37 lat od 1930 r. do 1967 r. Członek zespołu zajmującego się badaniami podstawowym. Wyprodukował pierwsze ogniwo słoneczne o wysokiej wydajności. Źródło: https://en.wikipedia.org/wiki/Calvin_Souther_Fuller#/media/File:Calvin_S_Fuller_Diffuse.jpeg

Fotografia 11. Daryl Muscott Chapin (1906-1995) był amerykańskim fizykiem, znanym ze współwynalezienia ogniwa słonecznych w 1954 r. podczas pracy w Bell Labs. Źródło: https://en.wikipedia.org/wiki/Daryl_Chapin#/media/File:Daryl_Chapin.jpg

Fotografia 12. Gerald L. Pearson (1905-1987) był fizykiem, którego praca nad krzemowymi prostownikami w Bell Labs doprowadziła do wynalezienia ogniwa słonecznego. Źródło: http://www.invent.org/hall_of_fame/388.html ; https://en.wikipedia.org/wiki/Gerald_Pearson

boratorium Bell Labs. **Po raz pierwszy** technologia fotowoltaiczna dostarczyła wystarczająco dużo energii, by zasilić urządzenie elektryczne przez kilka godzin. Ogniwo pracowało z wydajnością 4%. To pionier współczesnych krzemowych ogniwa fotowoltaicznych.

W 1941 r., amerykański badacz półprzewodników R. S. Ohl opatentował ogniwo słoneczne. Odkrył, „*że jeśli skieruje się fotony (światło) na pewne substancje metaliczne, ich powierzchnia emituje elektrony, natomiast gdy światło pada na inne substancje, ich powierzchnia elektrony pochłania.*” Praca nad diodami doprowadziła do opracowania pierwszych krzemowych ogniwa słonecznych, o sprawności powyżej 5%. Russell Shoemaker Ohl²⁰ (1898-1987) jest znany z opatentowania ogniwa słonecznego (patent USA 2 402 662, jako „*Urządzenie światłoczułe*”).

W latach 1953–1956 wzrosła produkcja krzemowych ogniwa słonecznych. Fizycy z Bell Laboratories odkryli, że krzem jest wydajniejszy niż selen i zbudowano ogniwo słoneczne o sprawności 6%. Od 1958 r. ogniwa słoneczne zyskały na znaczeniu po umieszczeniu na satelicie Vanguard I.²¹ Pod koniec lat 80. XX w. produkowano ogniwa krzemowe, do których wytworzenia użyto arsenku galowego, a ich wydajność przekraczała 20%. W 1989 r. ogniwo słoneczne z koncentratorem, skupiającym za pomocą soczewek światło słoneczne, osiągnęło wydajność 37%.²² Do 2006 r. najlepsze krzemowe ogniwa słoneczne miały sprawność ponad 40%, przy średniej przemysłowej ponad 17%.

W 2016 r. odkryto możliwość pozyskania energii słonecznej bez słońca. Zespół badawczy z University of California, Berkeley i Australian National University odkrył nowe właściwości nanomateriału. Jedną z tych właściwości nazwano *magnetyczną dyspersją hiperboliczną*, co oznacza, że materiał świeci po podgrzaniu. W połączeniu z ogniwami *termofotowoltaicznymi* mogłyby zamieniać ciepło w energię elektryczną bez potrzeby używania światła słonecznego.²³

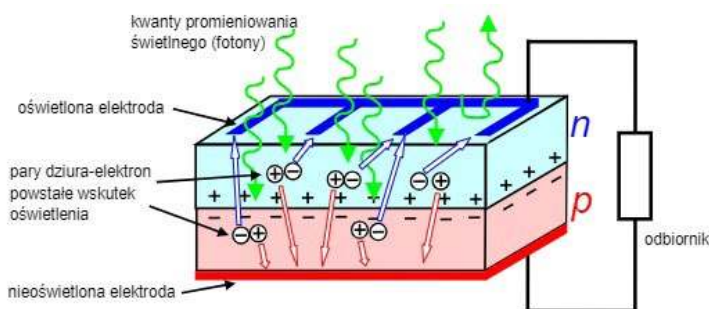
Ogniwa termofotowoltaiczne wykorzystują ciepło i światło

Naukowcy z MIT zaprezentowali hybrydowe ogniwo, które przekształca światło słoneczne – promieniowanie

podczerwone, na ciepło, absorbowane przez komórki do wytwarzania energii elektrycznej. Ogniwo fotowoltaiczne zbudowane jest z elementów półprzewodnikowych (*krzemu, germanu, selenu*), w którym następuje konwersja fotowoltaiczna energii promieniowania słonecznego w energię elektryczną. Wykorzystano półprzewodnikowe złącza typu <p-n> (z ang. *positive-negative*). Pod wpływem fotonów o energii większej, niż szerokość przerwy energetycznej półprzewodnika, elektrony przemieszczają się do obszaru <n>, a dziury do obszaru <p>. Pojawia się napięcie elektryczne. Przewodnictwo elektryczne i modele materiałów w ciałach stałych opisuje pasmowa teoria przewodnictwa elektrycznego.

Zwykle panele słoneczne pochłaniają część widma promieniowania, o odpowiedniej długości fali. Do ogniwa docierają fotony o energii potrzebnej do pokonania pasma wzbronionego (z ang. *band gap*). Panele zbudowane z zastosowaniem półprzewodnikowego krzemu reagują w większości na część widma światła widzialnego oraz promieniowanie podczerwone. Zastosowanie komórek **termofotowoltaicznych** (z ang. *thermophotovoltaic*, (TPV)) może teoretycznie pozwolić na osiągnięcie wyższej sprawności półprzewodnikowych ogniwa niż tzw. granicę Shockleya-Queissera, wynosząca 33,7% dla ogniwa jednowarstwowych.

Uzyskanie lepszego wyniku jest możliwe dzięki spożyciu fali o długościach do tej pory odrzucanych, a czynnikiem dodatkowym jest brak ruchomych części. Są bezgłośnie oraz nie wymagają częstej konserwacji.



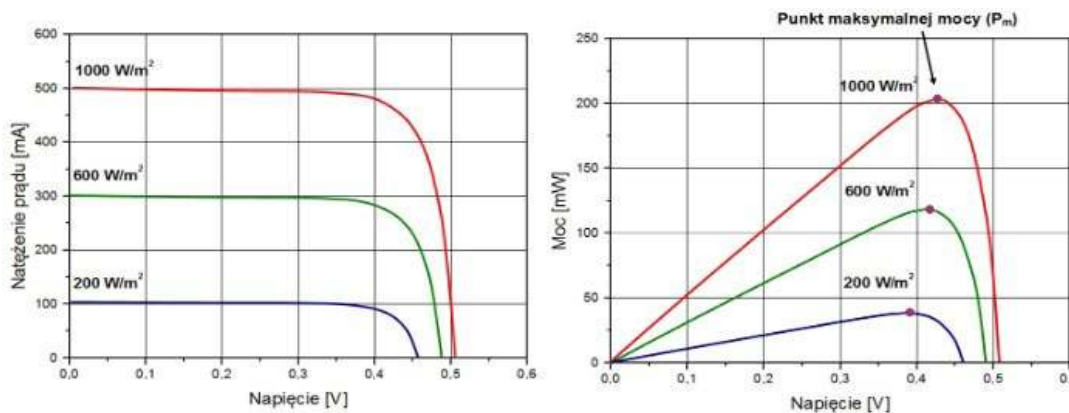
Rysunek 2. Schemat budowy ogniwa fotowoltaicznego. Źródło: <https://docplayer.pl/11383078-Badanie-ogniwa-fotowoltaicznego.html>

²⁰ https://en.wikipedia.org/wiki/Russell_Ohl

²¹ https://en.wikipedia.org/wiki/Solar_cell

²² <https://mlodytechnik.pl/technika/3635-ogniwa-sloneczne-historia>

²³ <https://www.ekoradcy.pl/blog/fotowoltaika-historia-rozwoju-technologii>



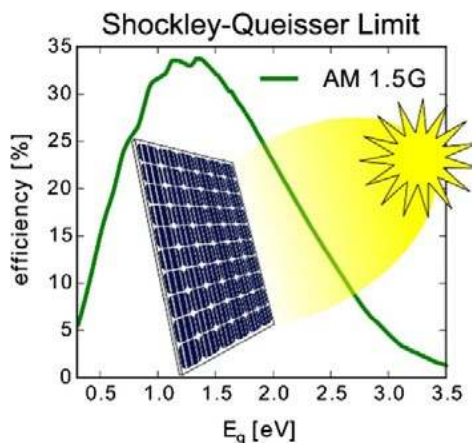
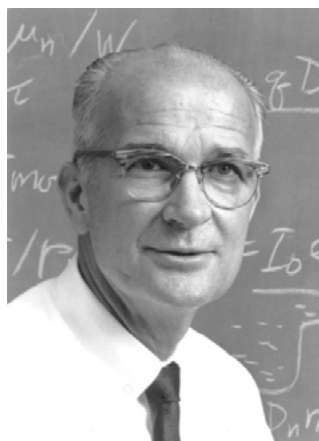
Rysunek 3. Charakterystyki prąd-napięcie oraz moc-napięcie przy różnym oświetleniu ogniwa fotowoltaicznego. Źródło: <https://docplayer.pl/11383078-Badanie-ogniwa-fotowoltaicznego.html> i Mikulski K., Ogniwa słoneczne w prostych badaniach w uczniowskim laboratorium, „Fizyka w Szkole” nr 4 lipiec/sierpień 2009, s.27-47

Co to jest granica Shockleya-Queissera?²⁴

W obszernej literaturze wskazano, że granica Shockleya-Queisser to maksymalna teoretyczna wydajność ogniwa fotowoltaicznego wykorzystującego pojedyncze złącze <p-n>²⁵. Została ona po raz pierwszy obliczona przez W. Shockleya i H. J. Queissera w Shockley Semiconductor Laboratory²⁶ w 1961 r. Podane obliczenie to jedno z najważniejszych wkładów naukowych w tej dziedzinie.

W opracowaniu Rühla pt. „Tabelaryczne wartości granicy Shockleya-Queissera dla jednozłączowych ogni słonecznych” czytamy: „Maksymalna wydajność konwersji światła na energię elektryczną η ogniwa słonecznego z jednym złączem dla danego widma oświetlenia jest znana jako szczegółowa granica równowagi lub granica Sho-

ckleya-Queissera. Na podstawie szczegółowych rozważań bilansowych ... przedstawili w 1961 r. ciało doskonale czarne o temperaturze powierzchni $T_s = 6000$ K.²⁷ Przyjęli oni, że w idealnym ogniwie fotowoltaicznym jedyną ścieżką rekombinacji, której nie można sprowadzić do zera, jest rekombinacja radiacyjna, która wyznacza górną granicę czasu życia nośników mniejszościowych. Przy generowaniu par elektron-dziura przyjęto, że fotony o energii poniżej pasma wzbronionego nie oddziałują z ogniwem słonecznym, natomiast fotony o energii powyżej pasma wzbronionego są przekształcane w pary elektron-dziura z wydajnością kwantową 100 %. ... obliczyli granicę wydajności dla ogniwa słonecznego z jednym złączem przy temperaturze ogniwa $T_c = 300$ K.,²⁸



Fotografia 13. William Bradford Shockley (1910- 1989) amerykański fizyk. Laureat Nagrody Nobla z fizyki w 1956 r. „Za badania nad półprzewodnikami i odkrycie efektu tranzystora”.

Fotografia 14. Hans-Joachim Queisser (ur.1931) fizyk ciała stałego, współautor pracy nad ogniwami słonecznymi z 1961 r. z podaniem granicy Shockleya-Queissera. Rysunek 4. Wartości teoretycznej maksymalnej sprawności konwersji światła na energię elektryczną jednozłączowych ogni słonecznych. Obliczenia dla ogniwa w temperaturze 25°C i oświetlonego widmowym natężeniem promieniowania AM 1,5G zgodnie z normą ASTM G173-03”. Źródło: https://pl.frwiki.wiki/wiki/William_Shockley <https://upload.wikimedia.org/wikipedia/commons/9/9d/Hans-JoachimQueisser1995.jpg> https://ars.els-cdn.com/content/image/1-s2.0-S0038092X16001110-fx1_lrg.jpg

²⁴ <https://www.universitywafer.com/what-is-the-shockley-queisser-limit.html> https://hmn.wiki/pl/Shockley%E2%80%9393Queisser_limit#cite_note-shockley-1 <https://www.sciencedirect.com/science/article/abs/pii/S0038092X16001110>

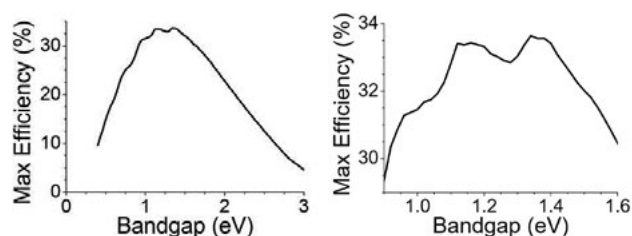
²⁵ https://pl.frwiki.wiki/wiki/Junction_P-N

²⁶ https://en.wikipedia.org/wiki/Shockley_Semiconductor_Laboratory

²⁷ http://metronu.ulb.ac.be/npauly/art_2014_2015/shockley_1961.pdf

²⁸ <https://www.sciencedirect.com/science/article/abs/pii/S0038092X16001110>

²⁹ Obecnie widmo słoneczne na Ziemi jest definiowane przez American Society for Testing and Materials (ASTM International Standard), które definiuje w dokumencie ASTM G173-03 dwa ziemskie rozkłady widmowe natężenia napromienienia

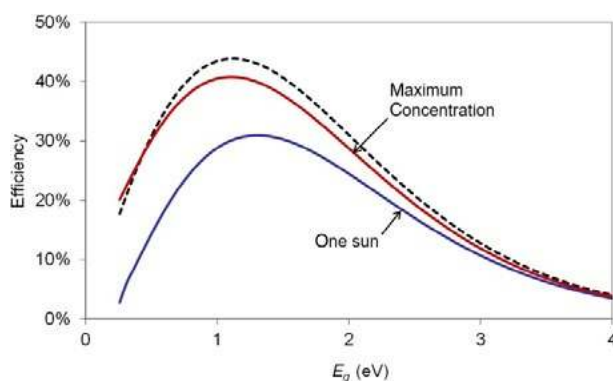


Rysunek 5. Granica Shockleya-Queissera powiększona w pobliżu obszaru szczytowej wydajności. Źródło: <https://upload.wikimedia.org/wikipedia/commons/a/a4/ShockleyQueisserZoomedIn.svg>

Rysunek 6. Granica Shockleya-Queissera dla wydajności ogniwa słonecznego bez koncentracji promieniowania słonecznego. Krzywa jest falująca z powodu pasm absorpcji w atmosferze. W pracy³⁰ widmo słoneczne było aproksymowane gładką krzywą, widmem ciała doskonale czarnego o temperaturze 6000 K. Źródło: JOURNAL OF APPLIED PHYSICS VOL. 32, KI1MBER 3 M. A. R. ell, 1961 https://en.wikipedia.org/wiki/Shockley%E2%80%93Queisser_limit

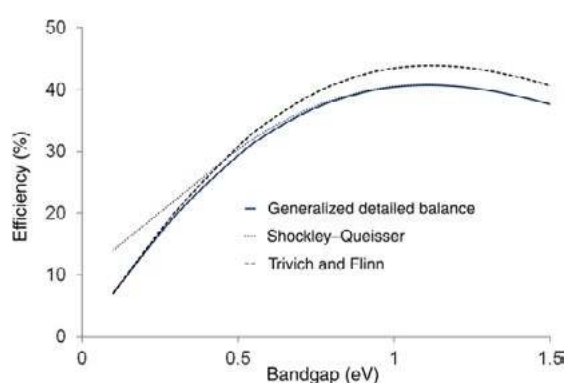


Fotografia 15. Pierre Aigrain (1924–2002) był francuskim fizykiem i sekretarzem ds. Badań Francuskiej Akademii Nauk. Źródło: <https://physicstoday.scitation.org/doi/10.1063/1.1620842>



Rysunek 7. Głównie SQ (linie pełne) i TF (linia przerywana) ograniczają wydajność ogniw słonecznych jako funkcje pasma wzbronionego. Wyniki SQ wynikają z równania dla światła słonecznego modelowanego promieniowaniem ciała doskonale czarnego o temperaturze $T_S = 6000\text{ K}$, ogniwem słonecznym i emitowanym promieniowaniem o temperaturze $T_0 = 300\text{ K}$. Źródło: <https://wires.onlinelibrary.wiley.com/doi/full/10.1002/wene.430>

Rysunek 8. Porównanie sprawności uzyskanych z uogólnionego bilansu szczegółowego ze zwykłym bilansem SQ oraz z granicą wydajności TF. Źródło: <https://wires.onlinelibrary.wiley.com/doi/full/10.1002/wene.430>



Ograniczenie sprawności 33,7%, przy założeniu, że pojedyncze złącze $\langle p-n \rangle$ z pasma zabronionego od 1,34 eV, oznacza, że z całej mocy zawartej w świetle słonecznym padającym na idealne ogniwo słoneczne (około 1000 W/m^2), tylko 33,7% można zamienić na energię elektryczną (337 W/m^2). Materiał ogniw słonecznych, krzem, ma pasmo wzbronione 1,1 eV, dające wydajność 32%. Współczesne ogniwa słoneczne zapewniają 24% sprawności konwersji. Straty wynikają głównie z problemów praktycznych, takich jak odbicie od przedniej powierzchni i blokowanie światła przez cienkie przewody na jej powierzchni. Dla komórek przy nieskończonej liczbie warstw, granica wynosi 86,8% przy zastosowaniu skoncentrowanego światła słonecznego.³¹

To Henry Herbert Kolm (1924–2010)³² skonstruował elementarny system TPV w MIT w 1956 r. Jednak Pierre Aigrain jest cytowany jako wynalazca na podstawie wykładów, które wygłosił na MIT w latach 1960–1961, co doprowadziło do badań i rozwoju.³³

W 2022 r. MIT³⁴ i NREL³⁵ ogłosiły konstrukcję urządzenia o sprawności 41%. Absorber wykorzystywał wiele warstw półprzewodnikowych III-V dostrojonych do absorbowania fotonów w różnym zakresie, ultrafioletowym, widzialnym i podczerwonym. Złoty reflektor przetwarzał niewchłonięte fotony. Urządzenie pracowało w temperaturze 2400°C , w której wolframowy emiter osiąga maksymalną jasność.³⁶

Pierwsze i najczęściej stosowane oszacowanie maksymalnego teoretycznego ogniwa słonecznego, które nie

³⁰ William Shockley i Hans J. Queisser (marzec 1961). „Szczegółowy limit bilansowy wydajności ogniw słonecznych pn Junction”. Journal of Applied Physics. 32(3): 510–519 http://metronu.ulb.ac.be/npauly/art_2014_2015/shockley_1961.pdf

³¹ Obliczenia: William Shockley i Hans J. Queisser, „Detailed Balance Limit of Efficiency of pn Junction Solar Cells”, Journal of Applied Physics, tom 32, s. 510–519 (1961), czyli „The Shockley–Limit Queissera oblicza się, badając ilość energii elektrycznej, która jest pobierana na padający foton”

³² https://en.wikipedia.org/wiki/Henry_Kolm#cite_note-1

³³ Nelson, RE (2003). „Krótka historia rozwoju termofotowoltaiki”. Nauka i technologia półprzewodników. 18(5): S141–S143. Bibcode: 2003SeScT...18S.141N. doi: 10.1088/0268-1242/18/5/301. S2CID 250921061. <https://iopscience.iop.org/article/10.1088/0268-1242/18/5/301>

³⁴ Massachusetts Institute of Technology (MIT) to prywatny uniwersytet badawczy w Cambridge w stanie Massachusetts, założony w 1861 r. Kluczowa rola w rozwoju nowoczesnych technologii i nauki. Prestiżowa i wysoko oceniana instytucja akademicka na świecie https://en.wikipedia.org/wiki/Massachusetts_Institute_of_Technology

³⁵ Narodowe Laboratorium Energii Odnawialnej (NREL) w USA specjalizuje się w badaniach i rozwoju energii odnawialnej, efektywności energetycznej, integracji systemów energetycznych i zrównoważonym transporcie. https://en.wikipedia.org/wiki/National_Renewable_Energy_Laboratory

³⁶ https://en.wikipedia.org/wiki/Thermophotovoltaic_energy_conversion#cite_note-0-13

zależy od szczegółów działania ogniw słonecznych, pochodzi od Trivicha i Flinna z 1955 r. (TF)³⁷. Model TF obejmuje tylko straty energii nośnika spowodowane szybką termalizacją po absorpcji światła i pomija energię cieplną nośników powyżej pasma wzbronionego. Podczas, gdy półprzewodnik może absorbować wszystkie padające fotony o energii powyżej jego pasma wzbronionego E_g , każda para elektron-dziura może wnieść do mocy wyjściowej tylko energię równą pasmu wzbronionemu.

O budowie nowych urządzeń

Naukowcy z MIT opracowali *termofotowoltaiczne* urządzenie składające się z absorbera nagrzewanego przez promieniowanie słoneczne oraz emitera przekształcającego zgromadzone ciepło na promieniowanie podczerwone, odbierane przez komórki ogniwa słonecznego. Konstrukcja dwuwarstwowego materiału absorber-emiter była kluczem do wykorzystania szerszego spektrum promieniowania. Absorber zbudowany został z wielościennych nanorurek węglowych (z ang. *multiwalled carbon nanotubes*, MWNT)³⁸, a selektywny emiter z jednowymiarowych kryształów fotonicznych Si/SO₂ (z ang. *1D photonic crystals*).

Nowe odkrycie dotyczy jeszcze eksperymentalnej technologii **termofotowoltaiki**. To niejako **pochodna zwykłej fotowoltaiki**, która **nie produkuje energii z samych promieni słonecznych, a za ich pomocą produkuje energię cieplną i ją przekształca w energię do użytku**. Jest to efektywniejsze niż **zwykła fotowoltaika**, która z fotonów zebranych w promieniach słonecznych generuje ok. 20% energii użytkowej. Ogniwa *termofotowoltaiczne* osiągają aż 40% według badań.³⁹

Warstwy o grubości poniżej 550 μm , zostały naniesione na fotowoltaiczną, półprzewodnikową komórkę o prze-

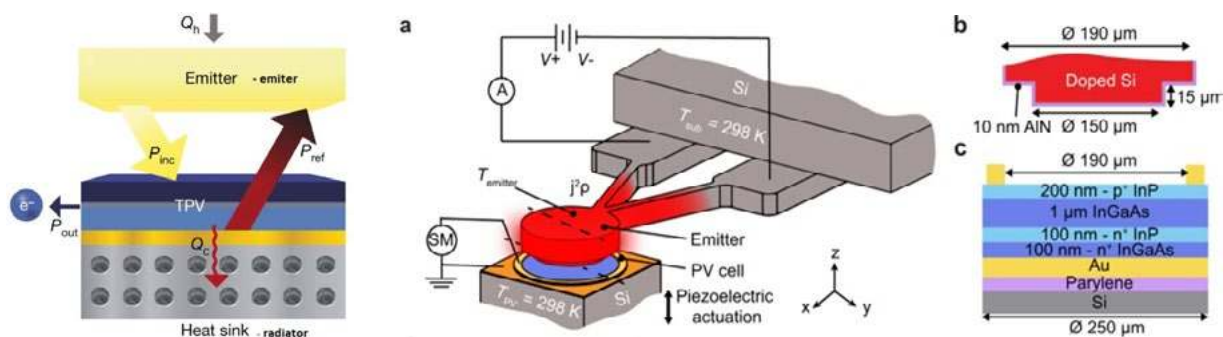
rwie energetycznej około 0,55 eV i wymiarach 1×1 cm. Podczas ekspozycji na strumień promieniowania o wartości 75 W/cm², nanorurki węglowe nagrzewały się do temperatury 962°C i przekazywały energię cieplną fotonicznym kryształom. Następnie, jednowymiarowe struktury emitowały niewidzialne fale podczerwone o określonej długości, dostosowanej do pasma wzbronionego ogniwa fotowoltaicznego. Powodowało to wzbudzenie elektronów walencyjnych w półprzewodniku.

Naukowcy z MIT zastosowali inne materiały. Wpłynęło to na uzyskanie stabilności. Wynik był kilkakrotnie niższy w porównaniu do znanych elementów fotowoltaicznych stosowanych obecnie (7-16%).

Okazało się, że w przypadku ogniw *termofotowoltaicznych*, znaczenie ma wielkość pola powierzchni w stosunku do objętości. Badana prowadzone przez naukowców na komórce o polu 1 cm² i kilku milimetrach grubości, wykazały generowanie dużych strat ciepła. Trwają prace nad urządzeniem o wymiarach 10×10 cm, które mają spowodować osiągnięcie sprawności na poziomie 20%. Udoskonalone ogniwa TPV mogłyby konwertować energię nawet z 80% skutecznością.

Ogniwa, które są urządzeniami dwuzłączowymi wykonanymi z materiałów półprzewodnikowych III-V⁴⁰ z elektronicznymi przerwami wzbronionymi między 1,0 eV a 1,4 eV. Elementy wykorzystują reflektory na tylnej powierzchni, aby skierować bezużyteczne promieniowanie z pasmem wzbronionym z powrotem do źródła. Urządzenia są zoptymalizowane pod kątem źródeł ciepła w temperaturach 1900–2400°C.

Pierwsze TPV zostały wykonane ze zintegrowanego odbłyśnika tylnej powierzchni i źródła wolframu emitującego w temperaturze 2000°C. Urządzenia te miały



Rysunek 9. Energia padająca na TPV może zostać zamieniona na energię elektryczną, odbite z powrotem do emitera lub termalizowana tracone z powodu nieefektywności ogniwa i tylnego reflektora. Źródło: <https://www.nature.com/articles/s41586-022-04473-y/figures/1>

Rysunek 10. a) Schemat układu eksperymentalnego zastosowanego do pomiarów termofotowoltaicznych bliskiego pola. Emiter (Si) ma mesę, która jest ogrzewana Joule'em (rozpraszanie ciepła mierzone ilościowo za pomocą amperomierza „A”) do 1270K poprzez przyłożenie napięcia bipolarnego (V+, V-) do dwóch wiązek. Epitaksjalnie wyhodowane ogniwo fotowoltaiczne (PV) InGaAs jest przesuwane w kierunku emitera za pomocą siłownika piezoelektrycznego, aby systematycznie kontrolować rozmiar szczeliny. Generowana energia elektryczna jest określana ilościowo za pomocą miernika źródła (SM). Substrat emitera i ogniwo fotowoltaiczne mają temperaturę ~ 298K.

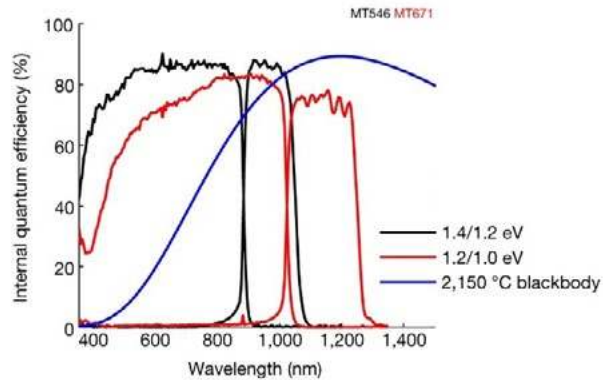
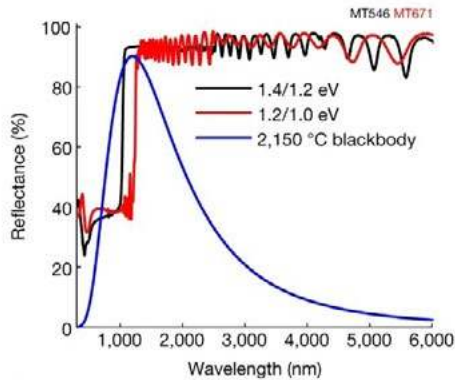
b, c) Profile przekrojowe emitera i ogniwa fotowoltaicznego na przekrojach wzdłuż czarnych przerywanych linii na ryc. 10a. Źródło: Mittapally, R., Lee, B., Zhu, L. i in. Termofotowoltaika bliskiego pola do wydajnej konwersji ciepła na energię elektryczną przy dużej gęstości mocy. Nat Commun 12, 4364 (2021) <https://doi.org/10.1038/s41467-021-24587-7>

³⁷ Trivich, D. i Flinn, P. (1955) Maksymalna efektywność przetwarzania energii słonecznej w procesach kwantowych. W F. Daniels & J. Duffie (red.), *Badania energii słonecznej* (str. 143) <https://wires.onlinelibrary.wiley.com/doi/full/10.1002/wene.430>

³⁸ Więcej na „Nanorurki – właściwości i zastosowania”, <https://materiałyinżynierskie.pl/nanorurki/>

³⁹ https://geekweek.interia.pl/technologia/news-termofotowoltaika-moze-generowac-prad-bez-przerwy-to-prawdzi,nId,6328657#utm_source=paste&utm_medium=paste&utm_campaign=chrome

⁴⁰ https://upload.wikimedia.org/wikipedia/commons/7/78/Pasmowa_teoria_przewodnictwa-schemat.svg



Rysunek 11a. Odbicie tandemów 1,4/1,2 eV i 1,2/1,0 eV. Dla odniesienia pokazano widmo ciała doskonale czarnego (blackbody) o temperaturze 2150°C, które jest średnią temperaturą emitera w aplikacji <thermal energy grid storage (TEGS)> sieciach magazynowania energii cieplnej.

Rysunek 11b. Wewnętrzna wydajność kwantowa (IQE) tandemów 1,4/1,2 eV i 1,2/1,0 eV.

Źródło: <https://www.nature.com/articles/s41586-022-04473-y/figures/2>

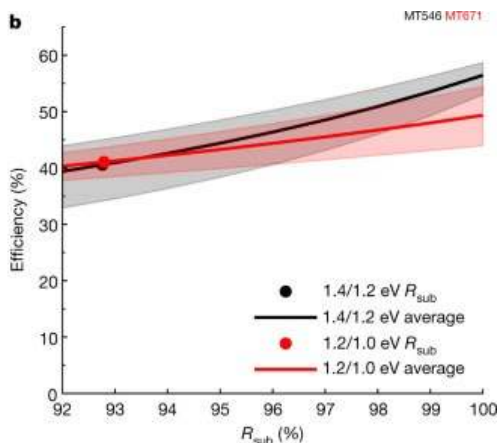
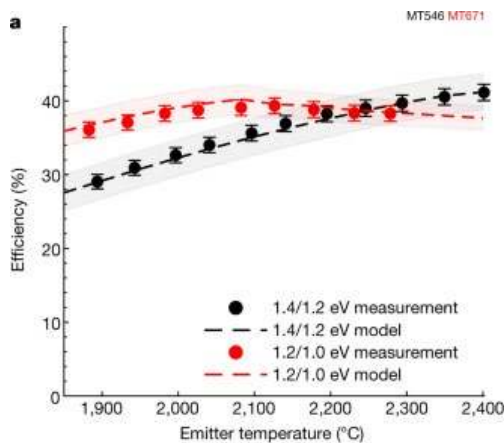
wydajność 29% i pomimo późniejszych postępów TPV nie przekraczały 32%. Działały w temperaturach poniżej 1300°C. Teoria TPV przewiduje, że ich wydajność może przekroczyć 50%.

Ogniwa TPV, opracował zespół kierowany przez A. Henry i A. LaPotin z Wydziału Inżynierii Mechanicznej MIT. Mają maksymalną wydajność 41,1% i działają z gęstością mocy 2,39 W/cm², przy użyciu źródła ciepła emitującego przy 2400°C. Urządzenia zostały wyprodukowane przez naukowców z National Renewable Energy Lab (NREL) przy użyciu techniki zwanej epitaksją metałoorganiczną z fazy gazowej.⁴¹

Co jest przyczyną takiej wydajności?

Wysoka wydajność komórek wynika z kombinacji czynników, mówi LaPotin⁴² i przekazuje informacje: „Pierwszym z nich jest zastosowanie ogniów wielozłączowych,

które pozwalają nam efektywniej przetwarzać różne pasma energetyczne widma padającego poprzez zmniejszenie tzw. strat termicznych”. Dalej wyjaśnia: „Drugim jest użycie materiałów z pasmem wzbronionym wyższym niż te zwykle stosowane w TPV, wraz z wyższymi temperaturami źródła ciepła”. Henry mówi, że „... urządzenie zespołu po raz pierwszy osiągnęło sprawność TPV na poziomie 40% i pierwszy raz, kiedy jakiegokolwiek silnik cieplny w stanie stałym wykazał sprawność wyższą niż średnia wydajność wytwarzania energii w oparciu o turbiny w USA. Porównanie z przeciętną turbiną jest kluczowe, ponieważ turbiny mają obecnie niemal monopol na produkcję energii na dużą skalę dzięki niskim kosztom i wydajności. Po raz pierwszy w historii inna technologia wykazała podobną wydajność, niższy koszt i skalowalność, dzięki czemu może konkurować z silnikami cieplnymi opartymi na turbinach” – mówi. Zdaniem Henry’ego⁴³ technologia TPV zasługuje na większą uwagę.



Rysunek 12a. Wydajność TPV mierzona przy różnych temperaturach emitera w zakresie od około 1900°C do 2400°C. Słupki błędów to niepewność pomiaru wydajności. Linie przerywane to przewidywania modelu. Zacięzione regiony to niepewność przewidywań modelu.

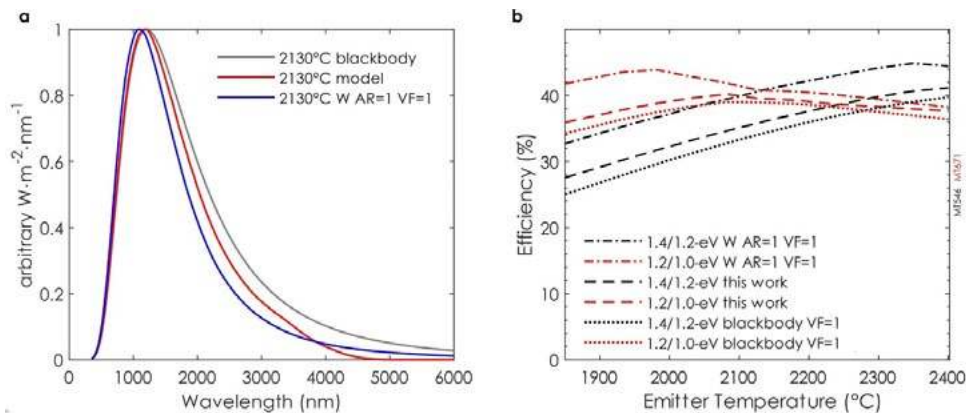
Rysunek 12b. Przewidywana wydajność tandemów 1,4/1,2 eV i 1,2/1,0 eV jako ważony współczynnik odbicia pasma wzbronionego jest ekstrapolowany przy założeniu emitera W z AR = 1 i VF = 1 i temperatura ogniwa 25°C. Linie ciągłe pokazują średnią wydajność w zakresie temperatur roboczych TEGS od 1900°C do 2400°C. Zacięzione paski to maksymalna i minimalna wydajność w zakresie temperatur.

Źródło: <https://www.nature.com/articles/s41586-022-04473-y/figures/2>

⁴¹ Pierwszy z projektów ogniów zespołu wykorzystuje górne i dolne złącza wykonane z AlGaInAs i GaInAs hodowanych na podłożu GaAs. W tym projekcie AlGaInAs ma pasmo wzbronione 1,2 eV, a GaInAs pasmo wzbronione 1,0 eV. Drugi łączy górną komórkę GaAs 1,4 eV dopasowaną do sieci z dolną komórką GaInAs 1,2 eV niedopasowaną do sieci.

⁴² Alina LaPotin <https://orcid.org/0000-0002-9002-4174> ; Wydział Inżynierii Mechanicznej, Massachusetts Institute of Technology, Cambridge, MA, USA; Źródło: <https://www.nature.com/articles/s41586-022-04473-y>

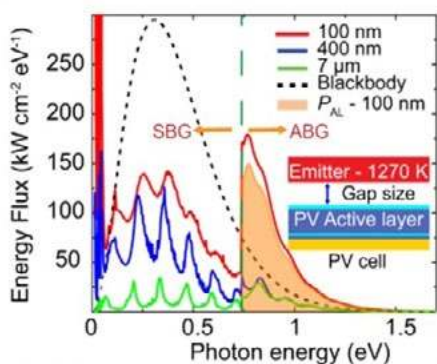
⁴³ Asegun Henry <https://orcid.org/0000-0002-9097-4882> ; Wydział Inżynierii Mechanicznej, Massachusetts Institute of Technology, Cambridge, MA, USA <https://www.nature.com/articles/s41586-022-04473-y>



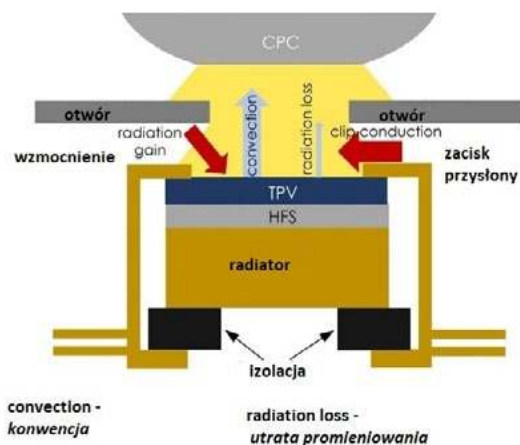
Rysunek 13a. Względne porównanie kształtu widma w pośredniej temperaturze badania (2130°C). Czerwona krzywa przedstawia modelowane widmo zgodnego z pomiarem. Szara krzywa pokazuje porównanie z kształtem widma ciała doskonale czarnego (blackbody) przy tej samej temperaturze emitera. Niebieska krzywa przedstawia porównanie z widmem opisanym literaturą emisją wolframu. Wszystkie krzywe są normalizowane według ich pików, celem porównania w kształtach widm.

Rysunek 13b. Kształt widma, w którym scharakteryzowano komórki (czerwona krzywa) jest podobny do widma ciała doskonale czarnego (szara krzywa), szczególnie powyżej pasma wzbronionego. Porównanie modelowanej wydajności TPV w widmie w tej pracy z emiterami, które można włączyć do systemu TPV. Widoczne, że wydajność mierzona w widmie żarówki w tej pracy zapewnia odpowiednią charakterystykę wydajności TPV w rzeczywistym podsystemie TPV. Temperatura ogniwa wynosi 25°C.

Źródło: <https://www.nature.com/articles/s41586-022-04473-y/figures/8>



Rysunek 14. Widmowy transfer energii z gorącego emitera termicznego o temperaturze 1270K do ogniwa fotowoltaicznego (PV) o temperaturze 300K jest wykresiany jako funkcja energii fotonu dla trzech rozmiarów szczeliny. Czarna przerywana linia to granica promieniowania ciała doskonale czarnego między dwiema półnieskończonymi płytkami przy 1270K i 300K. Poprawa transferu powyżej pasma wzbronionego, gdy przerwę zmniejsza się do 100 nm. Obszar zacieniony na pomarańczowo reprezentuje transfer energii promieniowania (PAL) z emitera do warstwy aktywnej InGaAs, która napędza generowanie nośników ładunku. Zielona linia przerywana reprezentuje pasmo wzbronione ogniwa fotowoltaicznego, podczas gdy SBG reprezentuje obszar pasma wzbronionego. Źródło: Mittapally, R., Lee, B., Zhu, L. i in. Termofotowoltaika bliskiego pola do wydajnej konwersji ciepła na energię elektryczną przy dużej gęstości mocy. Nat Commun 12, 4364 (2021). <https://doi.org/10.1038/s41467-021-24587-7>



Rysunek 15. Schemat (bez skali) przedstawiający nieużyteczne (szkodliwe) przepływy ciepła w eksperymencie. Źródło: na podstawie <https://www.nature.com/articles/s41586-022-04473-y/figures/9>

Elementy komunikatów naukowców

Szczegóły wytwarzania ogniw TPV, podano w metodach, odnoszących się do tych dwóch układów, scharakteryzowanych za pomocą ich pasm wzbronionych: 1,4/1,2 eV i 1,2/1,0 eV.⁴⁴ Pomiary współczynnika odbicia pokazano na ryc. 11a, a wewnętrzną wydajność kwantową podano na ryc. 11b.

Jak przedstawia się sprawność TPV pokazano na rysunkach 12 i 13.

Na rysunkach 14 i 15 zilustrowano fizyczny mechanizm wzmocnienia NF.

Energia słoneczna przeszła długą drogę w ciągu ostatnich 200 lat. Od obserwacji właściwości światła do znalezienia nowych sposobów przekształcania go w energię. Nowa technologia nie wykazuje żadnych oznak spowolnienia, a wręcz rozwija się w niespotykanym dotąd tempie.

Będąc na bieżąco z najnowszymi wiadomościami i postępami w zakresie pozyskania energii słonecznej, powinniśmy szybko zdecydować się, czy energia słoneczna jest dla nas odpowiednia.

Wiadomo że: „koszty energii stale rosną, coraz więcej osób szuka energii słonecznej jako bardziej wydajnej opcji zarówno dla swoich domów, jak i firm. Oprócz korzyści ekonomicznych energii słonecznej, energia słoneczna nie powoduje zanieczyszczeń i nie ma negatywnego wpływu na środowisko, co czyni ją atrakcyjną opcją dla osób zainteresowanych”⁴⁵.

dr Kazimierz Mikulski
Maksymilianowo

⁴⁴ Alina LaPotin, Kevin L. Schulte, Myles A. Steiner, Kyle Buznicki, Colina C. Kelsalla, Daniela J. Friedmana, Erica J. Tervo, Ryan M. Francja, Michelle R Young, Andrzej Rohskopf, Shomik Verma, Evelyn N. Wang & Asegun Henryk, *Sprawność termofotowoltaiczna 40%* <https://www.nature.com/articles/s41586-022-04473-y> <https://physicsworld.com/a/thermophotovoltaic-cells-top-40-percent-efficiency/>

⁴⁵ Jak obliczyć swoje szczytowe godziny słoneczne <https://www.solarpowerauthority.com/solar-news/>

O problemach z zasadami i nie tylko...

(Miniatura dydaktyczna)

Waldemar Reńda

Czy można coś nowego powiedzieć o tak dobrze znanych prawach fizycznych, jakimi są zasady dynamiki Izaaka Newtona? Pewnie niewiele, a jednak spróbuję.

Po pierwsze: nazwałem je prawami. Tak też nazwał je I. Newton. Czy słusznie? Otóż, jak podaje encyklopedia, „prawo fizyczne to twierdzenie odnoszące się do zjawisk fizycznych, dostatecznie uzasadnione doświadczalnie i oparte na rozumowaniu dedukcyjnym. Prawa fizyczne mają charakter ilościowy i sprowadzają się do podania zależności funkcyjnej między wielkościami fizycznymi.”¹

Zgodnie z tą definicją, II zasada dynamiki Newtona może być uważana za prawo fizyczne. Przy czym – niestety – najczęściej formułuje się ją jako równoczesną funkcję dwóch zmiennych: siły oraz masy ciała, na które ta siła działa. Nie jest to poprawne z punktu widzenia matematyki. Dlatego niektórzy autorzy podręczników sformułowanie tej zasady dzielą na dwie części, traktując obie te zależności jako dwie osobne proporcje. Pierwsza z nich jest zależnością wektorową przyspieszenia od działającej na ciało siły ($\mathbf{a} \sim \mathbf{F}_{\text{wyp}}$), w której odwrotność masy ciała spełnia rolę współczynnika proporcjonalności (zapis wektorowy wyczerpuje warunek zgodności obu wektorów co do ich kierunku i zwrotu). Dru-

ga zaś część opisuje zależność **wartości** przyspieszenia od masy ciała przy nie zmieniającej się wartości siły ($a \sim 1/m$).

Inni autorzy przyjmują za słuszną tylko pierwszą wersję. I jest to ujęcie najbliższe postaci, którą sformułował Newton. Pisał on:² „Zmiana ruchu jest proporcjonalna do działającej siły i zachodzi wzdłuż linii działania siły.”³ W tym sformułowaniu pojawia się określenie: „zmiana ruchu”. Słowa te sprawiły, że wśród fizyków trwa dyskusja nad tym, co ono oznacza. W jednej z definicji zawartej w „Principiach” Newton pisze: „Ilość ruchu jest jego miarą, daną przez prędkość i ilość materii.” A ponieważ masę ciała definiuje on jako „ilość materii” zawartej w ciele, należałoby zatem przypuszczać, że owa „zmiana ruchu” to raczej wielkość, którą dziś nazywamy pędem ciała. Co sugerowałoby, że wg Newtona, zasada ta dotyczy zależności pomiędzy zmianą pędu ciała i działającą na nie siłą, czyli odpowiada obecnej postaci pędowej tej zasady: $\Delta \mathbf{p} = \mathbf{F} \cdot \Delta t$.⁴ (Zapis wektorowy ustala zgodność co do kierunku oraz zwrotów wektorów $\mathbf{F}$ i $\Delta \mathbf{p}$.) Jak widać i ta postać jest funkcją podwójną, bo owa zmiana pędu jest zarówno proporcjonalna do działającej na ciało siły (wypadkowej)⁵, jak i do czasu jej działania. O czym jednak Newton w swoich „Principiach” nie wspomina.

Warto tu jeszcze zauważyć i to, że nazwa: „zmiana pędu” nie wywołuje wątpliwości znaczeniowych

¹ Por.: *Fizyka. Ilustrowana encyklopedia dla wszystkich*, pod redakcją A. Januszajtisa i J. Langer, Wydawnictwa Naukowo-Techniczne, Warszawa 1987, s. 213.

² Por.: A. K. Wróblewski. *Historia fizyki*, PWN, Warszawa 2006, s. 129. Zob. też: E. M. Rogers, *Fizyka dla dociekliwych, Część II*, PWN, Warszawa 1986, s. 260.

³ Obecnie zasady dynamiki najczęściej formułowane są w postaci implikacji: „Jeżeli..., to...” I tę formę uważam za dydaktycznie korzystną. Natomiast, omawiając te zasady, warto także przytoczyć sformułowania newtonowskie. Zob. też: W. Reńda, *Dynamika w gimnazjum i liceum*, Fizyka w Szkole, 2/2013.

⁴ I. Newton nie używał wzorów.

⁵ W pismach I. Newtona nie występuje to pojęcie.

w przypadku ruchu **opóźnionego** ciał. Uczniowie bowiem z niechęcią przyjmują w tym przypadku nazwę „przyspieszenie”.⁶ Dlatego stosuję tu nazwę: „opóźnienie”, co jest zgodne z potocznym rozumieniem tego słowa, a równocześnie nie ma negatywnego wpływu na strukturę pojęciową tego działu fizyki. Problemy związane z tego typu ruchem są jednak na tyle duże, że w szkole podstawowej w zasadzie nie analizuje się tego rodzaju ruchu, a jedynie się o nim wspomina.⁷

I jeszcze jedno: najczęściej zasadę tę ilustrujemy wzorem: $\mathbf{F} = m \cdot \mathbf{a}$, co sugeruje, że siła jest tu proporcjonalna do przyspieszenia, a nie odwrotnie. Czy słusznie? Nie minimy się z prawdą twierdząc, że jeżeli ciało o masie m chcemy nadać przyspieszenie $\mathbf{a}$, to na owo ciało musimy działać siłą proporcjonalną do oczekiwanego przyspieszenia, a także działać na nie zgodnie z wektorem tego przyspieszenia. Poza tym postać tę wykorzystujemy do zdefiniowania jednostki siły: $[F] = [m] \cdot [a] = 1 \text{ kg} \cdot \text{m/s}^2 = 1 \text{ N}$. Tak więc: 1N (niuton) to siła, która ciału o masie 1kg nadaje przyspieszenie 1 m/s^2 .

I tu mała dygresja dotycząca pojęcia masy ciała. I. Newton definiował masę jako ilość materii zawartej w ciele. Dziś mówi się, że masa jest miarą bezwładności ciała, co wymaga zdefiniowania owej bezwładności. I tu musimy sięgnąć do **I zasady dynamiki**. W obecnej formie zasada ta formułowana jest jako szczególny przypadek II zasady. Można ją zatem uważać za prawo fizyczne postaci: „jeżeli $F = 0$, to $a = 0$, czyli: $v = 0$ lub $\mathbf{v} = \text{constans}$ ”, co oznacza, że ciało jest w spoczynku lub porusza się ruchem jednostajnym prostoliniowym.

Zasadę tę nazywa się również **zasadą bezwładności**. Niestety, w takim sformułowaniu trudno doszukać się tej cechy ciał. Tu lepszą postacią byłaby ta, którą podał sam Newton, a mianowicie: „Każde ciało trwa w stanie spoczynku lub ruchu jednostajnego po linii prostej, dopóki siły przyłożone nie zmuszą ciała do zmiany tego stanu.”⁸ Ciało wykazuje zatem cechę zwaną bezwładnością, której miarą jest owa masa. W tej postaci mamy raczej do czynienia z ogólną prawidłowością, a więc z zasadą.⁹ Ba, jest ona nawet pokrewna z zasadami zachowania, w tym ściśle związana jest z zasadą zachowania pędu.

O kolejności wprowadzania tych zasad

Jeżeli lekcje dynamiki zaczniemy od stwierdzenia, że dynamika jest to dział fizyki o związkach ruchu z działającymi na ciało siłami, to w sposób naturalny pojawia się najpierw I zasada dynamiki, a następnie II zasada. Jeżeli jednak zaczniemy ten dział od definicji siły jako miary wzajemnego oddziaływania ciał, to pojawia się konieczność omówienia najpierw III zasady dynamiki. Kolejność

pozostałych zasad będzie zależała od tego, jak potraktujemy I zasadę; czy jako zasadę bezwładności, czy też szczególny przypadek II zasady. O ile wiem, to ten drugi wariant jest częściej stosowany przez nauczycieli.

A teraz o trzeciej zasadzie dynamiki Newtona

Izaak Newton sformułował ją tak: „Każdemu działaniu odpowiada równe i przeciwnie skierowane przeciwdziałanie.”¹⁰ Dziś najczęściej zasadę tę formułuje się w postaci implikacji: „Jeżeli jedno ciało działa na drugie siłą, to i drugie oddziałuje na pierwsze siłą równą co do wartości, lecz przeciwnie skierowaną”, co w zasadzie odpowiada wersji Newtona. Z tym, że dodał on jeszcze: „...wzajemne działania dwóch ciał są zawsze równe i przeciwnie skierowane.”¹¹ Niektórzy autorzy przyjmują, że to drugie zdanie jest wystarczającą formą tej zasady i można się z tym zgodzić.

Ale uwaga. W wielu podręcznikach bywa takie sformułowanie: „Ciała działają **na siebie** siłami...” Otóż, żadne ciało nie może działać **na siebie**! Należy zatem mówić lub pisać tak: „Ciała oddziałują **wzajemnie** siłami...”

I jeszcze jeden problem

W jednym z podręczników analizę sił działających na leżącą na stole książkę zaczęto od siły, którą stół działa na nią. Czy słusznie? Wątpię. Wydaje mi się, że w tym przypadku analizę oddziaływań należy rozpocząć od działającej na ciało siły ciężkości,¹² która z kolei wywołuje nacisk książki na powierzchnię stołu, co w konsekwencji powoduje wystąpienie **reakcji** stołu na dolną okładkę książki. W opisanym tu przykładzie występuje łańcuch przyczynowo-skutkowy, którego nie należy dowolnie odwracać.

Użyłem tu nazwy „reakcja” na siłę, z jaką stół oddziałuje na książkę. Otóż jest to zgodne z może najkrótszą i dość popularną formą III zasady dynamiki, a mianowicie, że „reakcja równa jest akcji”.¹³ To prawda, bo zwykle jedna z sił jest akcją, a druga reakcją na akcję. Na przykład w czasie tenisowego serwisu akcją jest siła działająca na piłkę ze strony rakiety. Natomiast siła z jaką ta piłka działa na powierzchnię kortu jest też akcją. Z tym że nie musi to być wzajemne oddziaływanie ciał sprężystych, bo np. siła wyporu jest reakcją na siłę, z jaką pływające po niej ciało działa na wodę.

Oczywiście, że nie zawsze można ustalić, która z sił jest akcją, a która reakcją. Tak jest np. w przypadku oddziaływania ładunków elektrycznych czy biegunów magnetycznych, a także w oddziaływaniu grawitacyjnym ciał. Jednakże tam, gdzie jest to możliwe, należy używać tych nazw, gdyż odpowiada to intuicyjnemu rozumieniu przez uczniów tych pojęć. Ponadto w mechanice mówi się o si-

⁶ W tym przypadku przyspieszenia tego nie wolno nazwać przyspieszeniem ujemnym.

⁷ W szkole podstawowej należy wyłącznie stosować skalarny opis ruchu.

⁸ A. K. Wróblewski. *Historia fizyki...* op. cit., s. 130. Zasady dynamiki w formie newtonowskiej znajdziemy też w popularnym podręczniku: R. Resnick, D. Halliday, *Fizyka Tom 1*, PWN, Warszawa 1999, s. 96 i następne.

⁹ Dotyczy to także III zasady dynamiki Newtona.

¹⁰ Zasada ta powszechnie zwana jest zasadą akcji i reakcji.

¹¹ A. K. Wróblewski. *Historia fizyki...* op. cit., s. 129.

¹² Niektórzy autorzy nazywają ją siłą grawitacji. Czy słusznie? Piszę o tym w dalszej części tej pracy.

¹³ A nie odwrotnie!

łach nacisku i siłach reakcji (np. reakcja podpory czy reakcja więzów). A przecież podręczniki fizyki są takie same zarówno dla uczniów liceum jak i technikum.

W niektórych podręcznikach zamiast nazwy „siła reakcji” pojawia się nazwa „siła sprężystości”. Otóż siły sprężystości są siłami **wewnętrznymi**, a tu są siły **zewnętrzne** wzajemnego nacisku.

Słów parę o wypadkowej

Po pierwsze należy zwrócić uwagę na to, że **wypadkowa** nie jest kolejną siłą działającą na ciało, lecz postulowaną siłą, która może zastąpić siły składowe bez zmiany skutków ich działania. Dlatego też jej wektor warto rysować innym kolorem lub kreską przerywaną. Nie należy też mówić/pisać: „siła wypadkowa” ale „wypadkowa sił”.

Bywa, że analizując siły działające na ciało, ich wektory rysuje się tak, że przykładają się je do jednego punktu, a mianowicie do środka ciężkości tego ciała. I tu popełnia się nierzadko poważne błędy, bo nie można np. przyczepić tam wektora siły tarcia. O ile bowiem wektor można przesunąć wzdłuż prostej, wzdłuż której ta siła działa, to przesunięcie równoległe siły zmienia skutek jej działania.¹⁴ Poza tym uczniowie niechętnie widzą takie zabiegi, gdy wektor danej siły przykładamy nie tam, gdzie w rzeczywistości występuje dane oddziaływanie.¹⁵

Dodam, że nie zawsze istnieje wypadkowa dwóch sił. Tak jest np. w przypadku tzw. pary sił.¹⁶ W tym przypadku działanie takiej pary można zrównoważyć jedynie drugą parą sił o tym samym co do wartości momencie, lecz o przeciwnym zwrocie.

Warto tu jeszcze zauważyć, że wypadkowa sił nie jest ich prostą sumą wektorową. W przypadku dodawania przez doczepianie znajdujemy wprowadzić wartość tej siły i zwrot, ale najczęściej nie jest prawidłowo usytuowana jej prosta działania. Zatem tak znaleziony wektor nie spełnia wszystkich warunków wypadkowej. Wszystkie te cechy można uzyskać poprzez kolejne dodawanie wektorów metodą równoległoboku. Dotyczy to jednak jedynie sił, które można przesunąć do wspólnego punktu zaczepienia. W przypadku sił równoległych problem jest bardziej złożony. W tym przypadku należy skorzystać z zasady, że ciało jest w spoczynku, gdy suma wektorowa sił składowych równa jest zero i równocześnie zeruje się suma wektorowa ich momentów względem dowolnie wybranego punktu.



Foto – Dreamstime

Kilka uwag o sile ciężkości, sile grawitacji, ciężarze oraz o środku ciężkości ciała

We wspomnianym podręczniku R. Resnicka i D. Halliday’a czytamy: „Ciężarem ciała nazywamy siłę grawitacyjną z jaką Ziemia przyciąga to ciało.”¹⁷ Czy słusznie? Jeżeli *ciężarem ciała* nazwiemy tę siłę, którą naciskamy na łazienkową wagę sprężynową,¹⁸ to o równości tej siły i siły grawitacji można w zasadzie mówić jedynie na biegunach Ziemi.¹⁹ W każdym innym miejscu na Ziemi są to inne siły. I różnią się nie tylko wartością, ale też kierunkiem działania.²⁰ Wynika to stąd, że Ziemia jest układem nieinercyjnym. A w tym układzie występują siły bezwładności.²¹ Ponadto każde ciało znajduje się na dnie atmosfery ziemskiej i działa na nie siła wyporu aerostatycznego.²² Tak więc siła ciężkości jest wypadkową działającej na ciało siły grawitacji, siły odśrodkowej bezwładności oraz owego wyporu.²³ Czy te dodatkowe siły są duże? Zwykle są niewielkie, bo np. na równiku stanowią one ok. 0,5% działającej na nas siły grawitacji. Mimo to proponuję jednak odróżniać pojęcie siły ciężkości od siły grawitacji.

A **ciężar** ciała? Tu jest pewien problem, bo sama nazwa wskazuje, że jest to jakaś cecha ciała związana z jego masą. Ba, nawet potocznie oba te pojęcia się utożsamia, bo waga łazienkowa podaje wynik w kilogramach masy. A przecież mierzy ona ową wspomnianą wyżej siłę ciężkości. Proponuję więc unikać tej nazwy. Jeżeli jednak użyjemy tego pojęcia, to należy go dokładnie i **poprawnie** zdefiniować. Podobnie jest i z pojęciem siły ciężkości.

¹⁴ Może bowiem pojawić się wówczas moment siły.

¹⁵ Przypomnę, że siła to wielkość fizyczna będąca miarą wzajemnego oddziaływania ciał.

¹⁶ Parę sił tworzą dwie równoległe i równe co do wartości siły, ale mające przeciwne zwroty. Przy czym ich proste działania oddalone są o $d > 0$. Nie zrównoważona para sił wywołuje obrót ciała, a zjawisko to opisują zasady dynamiki ruchu obrotowego.

¹⁷ R. Resnick, D. Halliday, *Fizyka*, op cit... s. 102.

¹⁸ Obecnie wykorzystuje się tu m.in. zjawisko piezoelektryczne.

¹⁹ Tu siła odśrodkowa bezwładności związana z ruchem obrotowym Ziemi jest równa zero.

²⁰ Zob.: W. Reñda, *Szerokość geograficzna czy geocentryczna*, *Fizyka w Szkole*, 6/2018.

²¹ $F_b = -m \cdot a_{układu}$, przy czym znak minus nie oznacza tu „ujemności” siły, ale fakt, iż siła bezwładności ma zwrot przeciwny do zwrotu wektora przyspieszenia. Niestety, w jednym z artykułów w *Fizyce w Szkole* spotkałem stwierdzenie, że siła bezwładności ma wartość: $-m \cdot a$, co jest karygodnym błędem, bo wartość wektora nie może być ujemna!

²² Np. na osobę o masie 70 kg działa siła wyporu aerostatycznego o wartości ok. 1N, ale w przypadku balonu jej wartość może być większa od ciężaru jego gondoli i powłoki.

²³ Występują tu jeszcze inne siły, jak siły związane z ruchem mimosładowym ziemi, czy oddziaływaniem grawitacyjnym Księżyca i Słońca, ale są one znacznie mniejsze od tych, które wyżej wymieniałem. Zob.: W. Reñda, *Pływy na Ziemi i w Kosmosie*, *Fizyka w Szkole*, 1/2015.

Siła ciężkości to siła, którą mierzy spoczywający względem Ziemi siłomierz (dynamometr), na którym swobodnie zawieszone jest dane ciało po uwzględnieniu siły wyporu aerodynamicznego. Wartość owej siły możemy również obliczyć ze wzoru: $F_c = m \cdot g$ pod warunkiem, że znamy masę ciała i wartość przyspieszenia ziemskiego g w punkcie pomiaru. Niestety, wartość ta zależy zarówno od szerokości geograficznej, jak i wysokości nad poziom morza. Wiemy, że na poziomie morza, na 45° szer. geogr. $g = 9,80665 \text{ m/s}^2$ (wartość standardowa). Natomiast na równiku i na poziomie morza $g = 9,7805 \text{ m/s}^2$. Dodam, że wartość tę mierzy się grawimetrem, którego działanie opiera się na precyzyjnym pomiarze okresu drgań wahadła fizycznego.

Pozostaje jeszcze problem, gdzie przyłożyć ów wektor siły ciężkości.²⁴ Odpowiedź, że w środku **ciężkości** ciała (a jednak „ciężkości”!). I tu pojawia się kolejny problem, jak ów punkt zdefiniować? Jest z tym kłopot od kiedy usunięto z programu statykę ciał, bo programy o tym nie mówią, a nie możemy obejść się bez tego pojęcia dlatego, że wektor siły ciężkości trzeba gdzieś zaczepić. Podobnie jest z wektorem nacisku ciała na podłoże i reakcji podłoża na ciało. A wiemy, że w przypadku wspomnianej książki, oba wektory muszą być przyłożone pod owym środkiem ciężkości. Kiedyś położenie środka ciężkości wyznaczano doświadczalnie na lekcji fizyki w klasie VI i badano warunki równowagi ciał. Teraz pozostaje jedynie stwierdzenie, że dla niektórych ciał, mających np. postać kuli lub prostopadłościanu, będzie to środek geometryczny owych brył, ale pod warunkiem, że jest to ciało jednorodne. W innych przypadkach położenie tego punktu należałoby ustalić doświadczalnie.

A jaka jest definicja owego punktu?

Zwykle mówi się, że jest to punkt zaczepienia wypadkowej sił działających na poszczególne cząsteczki danego ciała. Matematyczne ustalenie położenia tego punktu jest możliwe jedynie w kilku prostych przypadkach. Zamiana jego na „środek masy” też nie rozwiązuje problemu, bo definicja środka masy jest jeszcze bardziej kłopotliwa niż definicja środka ciężkości i równie trudna do ustalenia.²⁵ Tak więc musimy poświęcić nieco czasu lekcyjnego na przybliżenie uczniom tego problemu i to bez względu na to, czy owo pojęcie znajduje się w *Podstawie*, czy też nie.²⁶ Dodam tylko, że siła nacisku ciała na podłoże jak i reakcja na nią jest też wypadkową sił działających na poszczególne elementy oddziałujących powierzchni, przy czym zawsze obie te siły działają prostopadle do owej po-

wierzchni. Podobnie jest też z siłą wyporu. Tu również można określić punkt, który nazwiemy środkiem wyporu. Pokrywa się on ze środkiem ciężkości cieczy wypartej przez dane ciało.²⁷

A siła grawitacji?

Wartość tej siły należałoby obliczać ze wzoru: $F_{gr.} = m \cdot \gamma$. Gdzie γ to wartość **przyspieszenia grawitacyjnego** w punkcie pomiaru. Na równiku:

$$\gamma = 9,7805 \text{ m/s}^2 + 0,0339 \text{ m/s}^2 = 9,8144 \text{ m/s}^2.$$

W innych miejscach jest to wartość wektora będącego sumą geometryczną wektorów przyspieszenia dośrodkowego a_d w punkcie pomiaru oraz przyspieszenia ziemskiego g .²⁸

Wartość owego przyspieszenia możemy oczywiście obliczyć też ze wzoru: $\gamma = GM_Z/r^2$. Gdzie G to stała powszechnej grawitacji, M_Z to masa Ziemi, zaś r to odległość punktu pomiaru od środka Ziemi. Trzeba jednak znać tę odległość.

Dodam, że masy ciał w powyższym, jak i w poprzednim wzorze, to tzw. **masy grawitacyjne**. Natomiast masa występująca we wzorze: $F = m \cdot a$, to **masa bezwładna**. Na szczęście obie te masy są sobie równe, co udowodniono doświadczalnie.

Warto dodać, że wektor siły bezwładności zaczepiamy w środku masy ciała, oraz że ciała traktujemy jako ciała rozciągłe, a nie jako punkty materialne.²⁹

A teraz o masie

Co można powiedzieć o tym pojęciu? I. Newton masę ciała definiował jako *ilość materii* w ciele. Ta definicja bliska jest intuicyjnemu rozumieniu masy jako budulca materii, jako jedna z form tej materii. W tym przypadku nie można podać jej definicji.³⁰

W dynamice pojęcie masy jest jedną z **wielkości fizycznych**. A wielkość fizyczna to mierzalna cecha ciała, zjawiska lub procesu fizycznego. Jej wartość uzyskuje się w wyniku pomiaru, którego wynikiem jest liczba mianowana. Natomiast pomiar wielkości fizycznej polega na porównaniu wartości danej wielkości fizycznej z jej jednostką. Jednostka wielkości fizycznej to określona wartość danej wielkości. Niektórym jednostkom odpowiadają pewne wzorce – np. wzorec masy przechowywany w Międzynarodowym Biurze Miar i Wąg w Paryżu. Podobnie jest z kilkoma innymi wzorcami jednostek wielkości przyjętych za podstawowe.³¹

²⁴ W Podstawie programowej występuje nazwa: siła ciężkości a nie grawitacji.

²⁵ Zob.: *Fizyka. Ilustrowana encyklopedia dla...* op. cit. s. 274. Dodam, że oba te punkty pokrywają się tylko w jednorodnym polu grawitacyjnym.

²⁶ Nie wszystkie pojęcia, które są stosowane w nauczaniu fizyki muszą się znaleźć w *Podstawie*. O tym powinna decydować potrzeba dydaktyczna ich stosowania.

²⁷ Zob.: *Fizyka. Ilustrowana encyklopedia dla...* op. cit. s. 10.

²⁸ Zob.: W. Reńda, *Szerokość geograficzna czy geocentryczna*, Fizyka w Szkole, 6/2018 oraz W. Reńda, *Ciężar a siła ciężkości*, Fizyka w Szkole, 2/2014.

²⁹ Punktem materialnym nazywamy punkt geometryczny mający nieskończenie małą masę lub całe ciało, jeżeli jego rozmiary nie mają wpływu na przebieg i opis badanego zjawiska. Dodam, że punkt geometryczny ma nieskończenie małe rozmiary, ale nie zerowe.

³⁰ Uważam, że masa jako atrybut materii, podobnie jak energia czy ładunek, to niedefiniowalne pojęcia podstawowe. Zob. też twierdzenie Gödel’a o nierozstrzygalności.

³¹ Patrz układ SI. Obecne wzorce oparte są o prawa i stałe fizyczne. Niestety są zbyt skomplikowane, by omawiać je na lekcjach fizyki w szkole średniej.



Wróćmy jednak do masy i jej pomiaru. Do tego celu wykorzystuje się oddziaływanie grawitacyjne. Dawniej masę ciał mierzono np. przy użyciu wagi laboratoryjnej przez porównanie wartości sił ciężkości działających na odważniki i na dane ciało. Jeżeli można było pominąć siły wyporu, to pomiar można było uważać za wiarygodny. Stosując jednak wagę sprężynową lub dynamometr, mierzymy nie masę ciała, ale działającą na ciało siłę ciężkości. Jeżeli chcemy uzyskać wartość masy ciała, to musimy skorzystać ze wzoru: $m = F_c/g$, gdzie g to wartość przyspieszenia ziemskiego w punkcie pomiaru, której na ogół nie znamy. Zwykle tego typu wagi wyskalowane są już w kilogramach masy, ale dokonuje się tego w miejscu skalowania, a nie w miejscu późniejszego pomiaru. Oczywiście na ogół występujące tu niepewności pomiaru nie są duże i zwykle można je pominąć.

Do pomiaru masy ciał może być wykorzystane zjawisko bezwładności. Wówczas pomiar masy polega na porównaniu sił działających na dane ciało w układzie nieinercyjnym. Jeżeli więc $a_{d2} = a_{d1}$, to ze wzoru: $F_b = m \cdot a_{układu}$ otrzymamy proporcję: $m_2/m_1 = F_2/F_1$. Jednakże tej metody nie stosuje się w praktyce życia codziennego. Natomiast stosowana jest w wielu badaniach naukowych.

I ostatni problem

Ruch jest względny i wszystkie te problemy należałoby omawiać uwzględniając ten fakt. W przypadku, gdy układem odniesienia jest układ inercyjny,³² to nie ma większych problemów. Gdy zaś tym układem jest układ nieinercyjny, to posługujemy się pojęciem siły bezwładności³³ oraz zasadą d'Alemberta. Dokładniejsze omówienie tego problemu wykracza jednak poza ramy tej pracy. Dodam jedynie, że opis pewnych zjawisk – np. chwilowych zmian ciężaru ciała w windzie – najłatwiej jest wyjaśnić uczniom z pozycji układu **nieinercyjnego**, gdyż sami doświadczają owych efektów. Dodam też, że w opisie tych efektów – w porównaniu z opisem dokonany z układu inercyjnego (zewnętrznego) – występuje jedynie zamiana sił akcji na siły reakcji lub odwrotnie.³⁴

Kończąc, chciałbym zaapelować do nauczycieli, by starali się możliwie jak najpełniej rozumieć treści nauczania jako elementy struktury pojęciowej danego działu i całej fizyki. Tego typu analiz nie znajdziemy w podręcznikach. Mogłyby pojawić się w poradnikach metodycznych, ale – niestety – brak na rynku tego typu wydawnictw.

Waldemar Reñda
Olkuś, wrzesień 2023 r.

³² Układ inercyjny definiujemy zwykle jako taki, w którym spełniona jest I zasada dynamiki Newtona. Nie można jednak stwierdzić, że istnieje idealny układ inercyjny. My tylko postulujemy jego istnienie. Jeżeli jednak efekty bezwładnościowe są pomijalnie małe, to możemy przyjąć, że układ, w którym badamy dane zjawisko, jest układem inercyjnym.

³³ Mówi się, że siły bezwładności są siłami pozornymi. Siła jest tylko pojęciem służącym do opisu oddziaływań mechanicznych ciał. Jeżeli zatem widzimy zjawisko, które możemy uznać za wynik jakiegoś oddziaływania, to możemy również użyć pojęcia siły do jego opisu. Zatem to nie siła jest pozorna, ale pozorne może być owo oddziaływanie. Mówi się też o nierozróżnialności efektów grawitacyjnych i bezwładnościowych obserwowanych w układzie nieinercyjnym. Otóż ów brak możliwości odróżnienia tych efektów wystąpi jedynie w przypadku, gdy nasz układ znajduje się w jednorodnym polu grawitacyjnym.

³⁴ Zob.: W. Reñda, *Siły bezwładności*, Fizyka w Szkole, 3/2020.

„Doświadczenia skutków rzeczy pod zmysły podpadających ...” – włoskowatość, aerodynamika i rozpraszanie światła

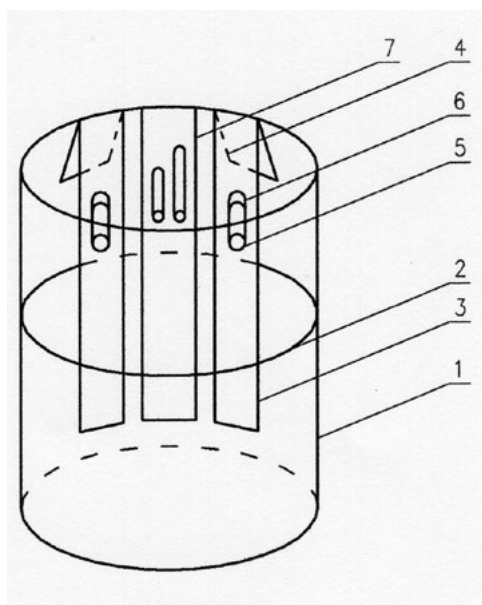
Stanisław Bednarek

W artykule jest opisanych kilka doświadczeń, dotyczących zjawisk zaliczanych do różnych działów fizyki. Doświadczenia te mają wspólną cechę, polegającą na tym, że mogą być łatwo przeprowadzone przy użyciu przedmiotów codziennego użytku. Ważną cechą zjawisk, których dotyczą opisane doświadczenia są również ich widowiskowe i interesujące wyniki. Dlatego te doświadczenia warto zaproponować do wykonania zarówno przez nauczyciela podczas lekcji, jak również samodzielnie przez uczniów.

Włoskowatość na kolorowo

Zjawisko włoskowatości kojarzy się zwykle z cienkimi rurkami szklanymi. Tym razem stanie się inaczej, a poza tym wartością dodaną doświadczeń będzie jeszcze jeden efekt, polegający na rozdzielaniu związków chemicznych o różnych barwach. Ten efekt ma ważne zastosowanie w chemii analitycznej i nazywa się chromatografią bibulową. Do doświadczenia zostanie wykorzystany niewielki słoik, kilka pasków bibuły filtracyjnej, kolorowe pisaki i ciecz (Rys. 1). Paski powinny mieć szerokość ok. 1,5 cm i długość ok. 10 cm. Zamiast bibuły filtracyjnej może być biały papier, np. do drukarek. Papier nie powinien być zbyt twardy i gładki.

Fragmenty pasków papieru należy zagiąć, tak żeby dało się je zawiesić na górnym brzegu słoika. Dolny koniec paska, znajdujący się wewnątrz słoika, powinien dotykać jego dna, a drugi koniec zwisać na zewnątrz. Na kawałku paska, który będzie umieszczony wewnątrz słoika, należy



Rys. 1. Budowa układu do chromatografii bibulowej; 1 – słoik, 2 – ciecz, 3 – pasek bibuły filtracyjnej lub papieru, 4 – zagięta część paska, 5 – plamki barwnika, 6 – smugi barwnika, 7 – pasek kontrolny.



Fot. 1. Barwne smugi otrzymane w jednym z doświadczeń z chromatografią bibulową.

wykonać w tym samym miejscu kilka kropek, używając do tego celu pisaków. Kropki powinny być wykonane w tym samym miejscu, tzn. jedna na drugiej, w wyniku przyłożenia i przytrzymania końców pisaków o różnych barwach przez kilkadziesiąt sekund.

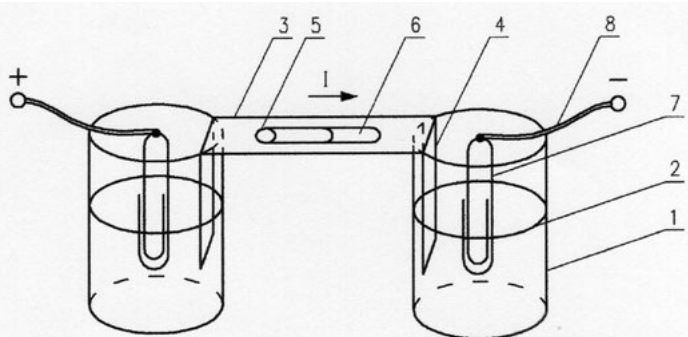
Można też przygotować pasek kontrolny, na którym kropki o różnych barwach będą obok siebie, ale w tej samej odległości od końca paska. Po tym paski należy zawiesić na górnym brzegu słoika i ostrożnie nalać do niego cieczy. W przypadku pisaków wodoodpornych będzie to alkohol etylowy, a przypadku zwykłych pisaków wystarczy woda. Górny poziom cieczy powinien być ok. 1 cm poniżej kropek.

Układ doświadczalny, przygotowany zgodnie z podanym opisem należy pozostawić w spokojnym miejscu na kilkadziesiąt minut. Po tym czasie na paskach będą dobrze widoczne barwne smugi skierowane od kropek pionowo w górę (fot. 1). Smugi o różnych barwach mają różne długości. Na pasku kontrolnym smugi znajdują się obok siebie. Na paskach, na których kropki były wykonane w tym samym miejscu, smugi o jednych barwach wyglądają tak, jakby wychodziły ze smug o innych barwach. Różna długość barwnych smug wynika z różnej szybkości rozchodzenia się substancji, tworzących poszczególne barwy w wykonanych kropkach.

Decydujące znacznie ma tutaj efekt włoskowatości, zachodzący w mikroszczelinach między włóknami zawartymi w papierze [1]. Dzięki temu możliwe jest rozdzielanie związków chemicznych, wchodzących w skład mieszanin i ewentualne poddanie ich dalszym, dokładniejszym badaniom, np. analizie spektralnej [2]. Jak wspomniano wcześniej, ta metoda nazywa się chromatografią bibulową. Do opracowania tej metody w istotny sposób przyczynił się rosyjski chemik i biolog Michaił Cwiet, który również podał prawidłową interpretację jej wyników i nadał nazwę. Cwiet prowadził badania w Warszawie na Politechnice i w Uniwersytecie w pierwszych latach XX w. Nomen omen, warto dodać, że słowo Cwiet w języku rosyjskim oznacza kolor.

Do rozdzielania mieszanin różnych substancji chemicznych można też wykorzystać przepływ prądu elektrycznego przez układ zbudowany podobnie, jak do chromatografii. Zjawisko, które będzie wówczas zachodziło nazywa się elektroforezą. Układ przeznaczony do przeprowadzenia elektroforezy jest pokazany na Rys. 2. W dwóch małych naczyniach szklanych – mogą to być zlewki lub kieliszki, są umieszczone zagięte końce paska bibuły lub papieru z kropką zaznaczonym na nim przy użyciu kilku pisaków o różnych barwach. Pasek jest podobny, jak poprzednio, kropka zaznaczona na środku jego poziomej części. Po wykonaniu kropki pasek należy zwilżyć wodnym roztworem soli kuchennej.

Ten sam roztwór wlewa się też do obu naczyń do takiej wysokości, żeby oba zagięte końce paska były w nim zanurzone. Na ścianki naczyń trzeba jeszcze nałożyć duże spinacze biurowe, których dolne końce też powinny być za-



Rys. 2. Budowa układu do elektroforezy; 1 – zlewka lub kieliszek, 2 – elektrolit, 3 – pasek bibuły filtracyjnej lub papieru, 4 – zagięta część paska, 5 – plamki barwnika, 6 – smugi barwnika, 7 – duży spinacz biurowy, 8 – przewód, 1 – natężenie prądu.

nurzone w roztworze. Do górnych części tych spinaczy są przymocowane końce przewodów w izolacji, przyłączone do źródła prądu stałego o napięciu kilku V. Tym źródłem mogą być 2-3 okrągłe baterijki, tzw. „paluszki”, umieszczone w odpowiednim pojemniku, który kosztuje kilka zł. i jest dostępny w sklepach z artykułami elektronicznymi.

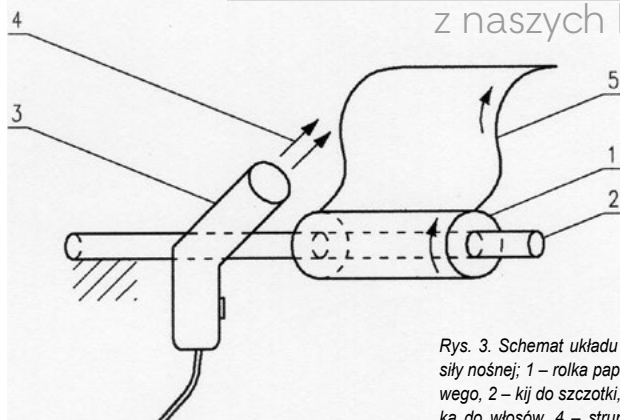
Wodny roztwór soli kuchennej jest elektrolitem i umożliwia prąd elektryczny przez zwilżony nim pasek. Podczas tego przepływu są również przenoszone cząsteczki substancji, które zostały użyte do zaznaczenia kropki. Szybkość tego transportu jest różna dla różnych cząsteczek i zależy m.in. od ich właściwości fizykochemicznych. Podczas przepływu prądu elektrycznego w pasku istnieje pole elektryczne, które oddziałuje na cząsteczki i powoduje ich ruch. Skutkiem tego po pewnym czasie można zaobserwować smugi o różnych barwach, co wskazuje na rozdzielanie różnych substancji, podobnie jak w przypadku chromatografii bibułowej.

Należy wspomnieć, że czasem pole elektryczne, wpływające na ruch cząsteczek w celu ich rozdzielania mieszanin, zastępuje się polem magnetycznym. Zachodzi wtedy zjawisko nazywane magnetoforesą, ale jest trudniejsze do zaobserwowania, ponieważ wymaga dość silnych pól magnetycznych, utrzymywanych przez czas kilkudziesięciu minut.

Aerodynamika papieru toaletowego

Siła nośna ma podstawowe znaczenie w aerodynamice, ponieważ zabezpiecza przed spadaniem heterodyny, czyli obiekty latające mające średnią gęstość większą od gęstości powietrza. Wytworzenie siły nośnej można w uproszczeniu wyjaśnić stosując prawo Bernoulliego. Gdy odpowiednio wyprofilowany płat, np. skrzydło, jest opływane przez strumień powietrza, to ciśnienie powietrza nad płatem jest mniejsze, niż pod nim. Różnica tych ciśnień scalaowana po powierzchni płata daje siłę nośną. Żeby taka różnica się pojawiła, prędkość powietrza na płacie musi być większa, niż pod nim. W tym celu płatom nadaje się odpowiedni profil i są one bardziej wypukłe po stronie górnej, niż od dołu. Strumień powietrza nad płatem ma wtedy do pokonania dłuższą drogę, niż pod nim i musi poruszać się wolniej. Dzięki temu pojawia się siła nośna.

Do wytworzenia siły nośnej będą potrzebne suszarka do włosów i rolka papieru toaletowego. Rolka może być umieszczona na wieszaku bez obudowy, ale lepiej nałożyć ją luźno na poziomo zamocowany pręt, np. kij do szczotki przywiązany do oparcia krzesła. Suszarkę można zastąpić inną, niewielką dmuchawą, np. opalarką, oczywiście z wy-



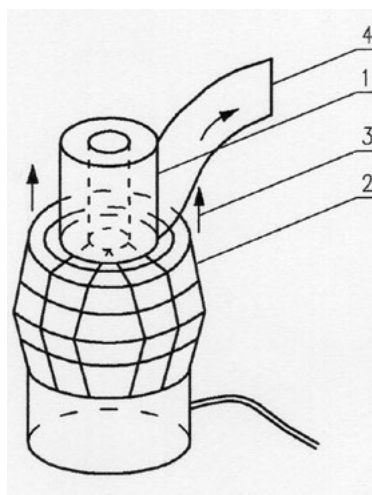
Rys. 3. Schemat układu do badania siły nośnej; 1 – rolka papieru toaletowego, 2 – kij do szczotki, 3 – suszarka do włosów, 4 – strumień powietrza, 5 – unoszony fragment papieru toaletowego.



Fot. 2. Rozwijający i unoszący się papier toaletowy.

łączonym nagrzewaniem. Z rolki trzeba usunąć zewnętrzną warstwę papieru, zabezpieczającą przed rozwijaniem i oderwać przyklejony koniec papieru. Strumień powietrza z dmuchawy należy skierować na rolkę ukośnie ku górze. Prędkość powietrza powinna być zgodna z kierunkiem nawinięcia papieru. Obserwuje się wówczas obrót rolki i odwijanie papieru, który jest unoszony ku górze wraz ze strumieniem powietrza (Fot. 2). Rolka zachowuje się podobnie, jak wirnik turbiny i chociaż nie ma łopatek to jest wprawiana w ruch obrotowy przez oddziaływanie strumienia powietrza z powierzchnią papieru za pośrednictwem sił lepkości.

Inny wariant tego doświadczenia przedstawia Rys. 4. W tym wariantcie rolka papieru toaletowego została ustawiona pionowo na środku osłony śmigła wentylatora



Rys. 4. Inny wariant układu do badania siły nośnej; 1 – rolka papieru toaletowego, 2 – osłona śmigła wentylatora, 3 – strumień powietrza, 4 – unoszony fragment papieru toaletowego.



Fot. 3. Inna wersja doświadczenia pokazującego odwijanie i unoszenie papieru toaletowego na małym wentylatorze.

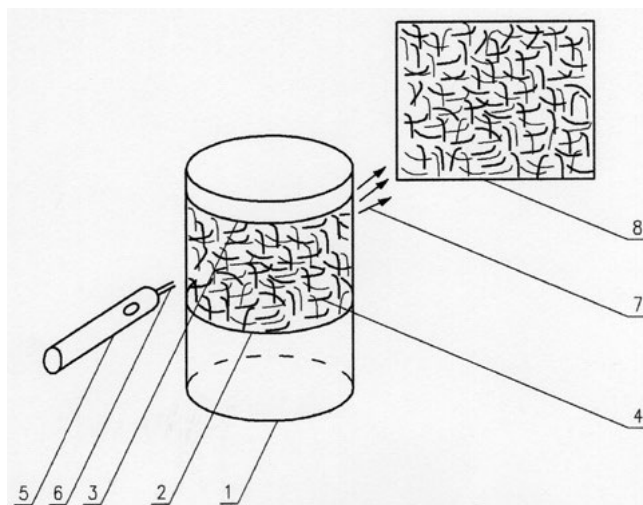
wykonanej z drutów. Średnica osłony powinna wynosić ok. 30–40 cm, a osłona ustawiona poziomo, tak żeby strumień powietrza opływał rolkę z boków. Można też wykorzystać wentylator o mniejszej średnicy z podobną osłoną śmigła, ale wtedy średnica rolki papieru też powinna być mniejsza. Najprostszym sposobem jest wykorzystanie w tym celu już używanej rolki, na której pozostało niewiele papieru. Rolka nie może być jednak zbyt mała, żeby w całości nie została uniesiona ku górze. W tym wariancie doświadczenia rolka nie obraca się, natomiast obserwuje się odwijanie papieru w kierunku radialnym i unoszenie odwiniętej części ku górze (Fot. 3).

Rozpraszanie światła na pianie

Do doświadczenia zostanie wykorzystany niewielki słoik, który należy częściowo napełnić płynem do mycia naczyń lub szamponem, szczelnie zamknąć zakrętką i kilka razy mocno nim potrząsnąć. (Rys. 5). W ten sposób nad powierzchnią płynu w słoiku zostanie wytworzona piana. Słoik z pianą trzeba ustawić w odległości 1–2 m przed zawieszonym pionowo ekranem lub białą ścianą w zaciemnionym pomieszczeniu i skierować na pianę wiązkę światła ze wskaźnika laserowego o dowolnej barwie. Na ekranie pojawi się wówczas interesujący obraz.

Ten obraz jest wynikiem oddziaływania światła z pianą, która ma strukturę komórkową. Ścianki komórek piany mają kształt wielokątów, najczęściej pięcio- lub sześciokątów i są utworzone z cienkich warstewek cieczy, ograniczonych błonami powierzchniowymi. Światło padając na te ścianki ulega załamaniu, odbiciu, a następnie dyfrakcji i interferencji oraz jest też w niewielkim stopniu pochłaniane.

Skutkiem tych złożonych oddziaływań na ekranie pojawia się obraz o interesujących walorach poznawczych i estetycznych (Fot. 4). W wyniku działania siły ciężkości ciecz wewnątrz ścianek, rozdzielających komórki spływa w dół i ścianki pękają. W ten sposób komórki łączą się ze sobą i piana powoli zanika. Dlatego obraz na ekranie też ulega zmianom (Fot. 5). Jeszcze ciekawsze efekty można uzyskać



Rys. 5. Schemat układu do badania rozpraszania światła na pianie; 1 – słoik, 2 – płyn do mycia naczyń lub szampon, 3 – zakrętka, 4 – piana, 5 – wskaźnik laserowy, 6 – wiązka światła padającego, 7 – światło rozproszone, 8 – ekran lub ściana do obserwacji efektów rozpraszania.

po skierowaniu na komórki piany wiązki światła laserowego o innej barwie, np. zielonej lub czerwonej (Fot. 6).

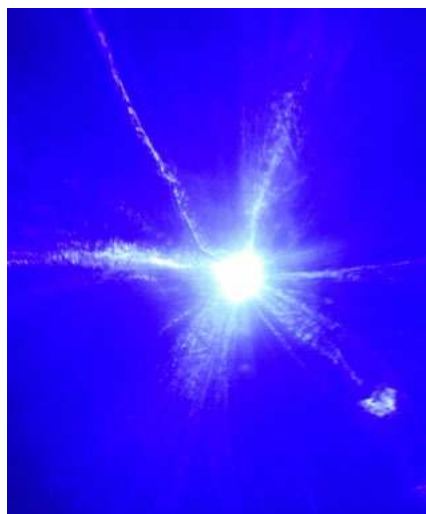
Można też skierować w ten sam obszar piany światło laserowe jednocześnie z trzech wskaźników o barwie czerwonej, zielonej i niebieskiej. To doświadczenie dobrze nadaje się na temat konkursu dla uczniów, polegającego na otrzymaniu najładniejszego obrazu piany przy użyciu wiązek światła laserowego.

Stanisław Bednarek

Wydział Fizyki i Informatyki Stosowanej
Uniwersytetu Łódzkiego

LITERATURA

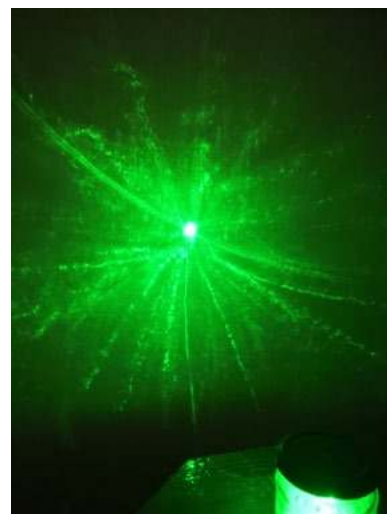
- [1] S. Szczeniowski, Fizyka Doświadczalna, cz. II, Ciepło i fizyka drobinowa, Państwowe Wydawnictwo Naukowe, Warszawa 1971, s. 288.
- [2] G. W. Ewing, Metody instrumentalne w analizie chemicznej, Państwowe Wydawnictwo Naukowe, Warszawa 1967, s. 385.
- [3] S. Szczeniowski, Fizyka Doświadczalna, cz. III, Elektryczność i magnetyzm, Państwowe Wydawnictwo Naukowe, Warszawa 1972, s. 184.
- [4] S. Szczeniowski, Fizyka Doświadczalna, cz. I, Mechanika i akustyka, Państwowe Wydawnictwo Naukowe, Warszawa 1972, s. 509.
- [5] S. Szczeniowski, Fizyka Doświadczalna, cz. IV, Optyka, Państwowe Wydawnictwo Naukowe, Warszawa 1971, s. 216.
- [6] Sojecki, Optyka, Wydawnictwa Szkolne i Pedagogiczne, Warszawa 1985, s. 173.



Fot. 4. Przykład obrazu otrzymanego w wyniku oświetlenia piany wiązką niebieskiego światła laserowego.



Fot. 5. Inny przykład obrazu otrzymanego w wyniku oświetlenia wiązki niebieskiego światła laserowego piany użytej do wykonania fot. 4 po upływie 20 min.



Fot. 6. Przykład obrazu otrzymanego w wyniku oświetlenia piany wiązką zielonego światła laserowego.

W dolnej części fot. 5 i 6 jest widoczna górna część słoika z pianą wytworzoną z płynu do mycia naczyń, fotografie autora.

Wybrane zagadnienia z mechaniki nieba

– część IV

Zadania cz. 2

Marcin Wesołowski

Zadanie 7. Scałkować równanie ruchu prostoliniowego w zagadnieniu 2 ciał. Założyć dodatkowo, że energia całkowita układu jest równa zero. W chwili początkowej odległość wzajemna ciał jest równa r_0 .

Rozwiązanie:

Na podstawie treści zadania oraz wcześniejszych rozważań możemy zapisać, że:

$$m\ddot{r} = -\frac{GMm}{r^2}. \quad (7.1)$$

Zauważmy, że możemy nieco uprościć równanie (7.1) dzieląc go obustronnie przez m , wówczas mamy:

$$\ddot{r} = -\frac{GM}{r^2}. \quad (7.2)$$

Postępując analogicznie jak w pierwszej części rozważań związanych z mechaniką nieba pomożemy obustronnie równanie (7.2) przez czynnik r , a następnie scałkujemy to równanie względem czasu dt :

$$\ddot{r} = -\frac{GM}{r^2} \quad / \cdot r, \quad (7.3)$$

$$\int \ddot{r} r dt = -GM \int \frac{r}{r^2} dt, \quad (7.4)$$

Wówczas równanie (7.4) możemy zapisać jako:

$$\frac{1}{2}(\dot{r})^2 = \frac{GM}{r} + C, \quad (7.5)$$

Jednocześnie pamiętając, że: $\dot{r} = \frac{dr}{dt} = v$, wówczas równanie (7.5) możemy wyrazić jako:

$$\frac{1}{2}v^2 - \frac{GM}{r} = C. \quad (7.6)$$

Zgodnie z założeniem $C > 0$, wówczas na podstawie równania (7.6) mamy:

$$v = \pm \sqrt{\frac{2GM}{r}}, \quad (7.7)$$

$$v = \frac{dr}{dt} = \pm \sqrt{\frac{2GM}{r}}. \quad (7.8)$$

Zatem na mocy równania (7.8) słuszna jest zależność:

$$r^{\frac{1}{2}} dr = \pm \sqrt{2GM} dt. \quad (7.9)$$

Całkujemy równanie (7.9) w granicach od $r_0(t_0)$ do $r(t)$ otrzymujemy:

$$\frac{\frac{3}{2}r^{\frac{3}{2}} - \frac{3}{2}r_0^{\frac{3}{2}}}{\frac{3}{2}} = \pm \sqrt{2GM}(t - t_0), \quad (7.10)$$

lub

$$\pm(t - t_0) = \frac{2}{3\sqrt{2GM}} \left(r^{\frac{3}{2}} - r_0^{\frac{3}{2}} \right). \quad (7.11)$$

Uwaga: Zauważmy, że do wyznaczenia lewej strony zależności (7.10) wykorzystaliśmy następujący wzór wynikający z rachunku całkowego:

$$\int r^{\frac{1}{2}} dr = \int \sqrt{r} dr = \frac{2}{3} r^{\frac{3}{2}}.$$

Przeprowadź teraz następującą dyskusję:

1. Jeżeli $v_0 = \dot{r}_0 > 0$, $t = t_0 = 0$, $v_0 = +\sqrt{\frac{2GM}{r_0}}$, wtedy

$$(t - t_0) = \frac{2}{3\sqrt{2GM}} \left(r^{\frac{3}{2}} - r_0^{\frac{3}{2}} \right).$$

Wtedy gdy t rośnie, to r również rośnie.

2. Jeżeli $v_0 = \dot{r}_0 < 0$, $t = t_0 = 0$, $v_0 = -\sqrt{\frac{2GM}{r_0}}$, wtedy

$$-(t - t_0) = \frac{2}{3\sqrt{2GM}} \left(r^{\frac{3}{2}} - r_0^{\frac{3}{2}} \right),$$

lub

$$(t - t_0) = \frac{2}{3\sqrt{2GM}} \left(r_0^{\frac{3}{2}} - r^{\frac{3}{2}} \right),$$

Wtedy gdy t rośnie, to r maleje.

Odpowiedź: Wnioski wynikające z całkowania równania ruchu w zagadnieniu dwóch punktów materialnych zostały omówione w dyskusji.

Zadanie 8. Wyznaczyć prędkość, z jaką podczas całkowitego zaćmienia Słońca, cień Księżyca porusza się po powierzchni Ziemi, nie uwzględniając poprawki wynikającej z orbitalnego ruchu Ziemi. Dla uproszczenia przyjąć, że zaćmienie obserwowane jest w południe na równiku, oraz, że oś Ziemi jest prostopadła do płaszczyzny orbity Księżyca. Kierunki obrotu Ziemi wokół własnej osi i ruchu orbitalnego Księżyca są zgodne. Odległość Księżyca od Ziemi wynosi $r = 3.8 \cdot 10^5$ km. Promień Ziemi $R_Z = 6.4 \cdot 10^3$ km. Przyjąć, że miesiąc księżycowy (gwiazdowy) jest równy 27^d . W rachunkach uwzględnić, że odległość Ziemi od Słońca jest znacznie większa niż odległość Księżyca od Ziemi.

Rozwiązanie:

Przesuwanie się cienia Księżyca związane jest z obrotem Ziemi wokół własnej osi (Δx_1) oraz z ruchem Księżyca po orbicie (Δx_2). Po czasie Δt przesunięcia te są równe:

$$\Delta x_1 = \frac{2\pi R_Z}{T_Z} \Delta t, \quad T_Z = 1^d, \quad (8.1)$$

$$\Delta x_2 = \frac{2\pi r}{T_K} \Delta t, \quad T_K = 28 T_Z. \quad (8.2)$$

Kierunek ruchu orbitalnego Księżyca i kierunek ruchu spinowego Ziemi są zgodne, więc wypadkowe przesunięcie cienia jest równe. Wówczas na podstawie równań (8.1) oraz (8.2) możemy zapisać, że:

$$\Delta x = \Delta x_2 - \Delta x_1 = 2\pi \left(\frac{r}{T_K} - \frac{R_Z}{T_Z} \right) \Delta t. \quad (8.3)$$

Wykorzystując równanie (8.3) możemy teraz wyznaczyć prędkość przesuwania się cienia Księżyca, która wynosi:

$$v = \frac{\Delta x}{\Delta t} = 2\pi \left(\frac{r}{T_K} - \frac{R_Z}{T_Z} \right) = 0.52 \text{ km/s}. \quad (8.4)$$

Odpowiedź: Poszukiwana wartość prędkości przesuwania się cienia Księżyca wynosi 0.52 km/s.

Zadanie 9. Planeta porusza się po orbicie eliptycznej o dużej półosi a i mimośrodku e wokół Słońca. Wykorzystując prawa zachowania energii mechanicznej i momentu pędu, obliczyć wartości tych wielkości fizycznych. Masę Słońca M i stałą grawitacji G przyjmij jako znane.

Rozwiązanie:

Z treści zadania wynika, że planeta porusza się po orbicie eliptycznej więc na samym początku, aby określić dwa podstawowe parametry orbitalne (perycentrum i apocentrum). W tym miejscu wyjaśnijmy te dwa ważne pojęcia:

- **perycentrum** to punkt położony na orbicie eliptycznej, który jest najbliższy środkowi masy rozważanego układu, zwykle ciała centralnego o masie znacznie większej od mniejszego ciała obiegającego. Inaczej mówiąc jest to punkt największego zbliżenia, który jest odpowiednikiem perygeum;
- **apocentrum** to punkt położony na orbicie eliptycznej, który jest najbardziej odległy od środka masy rozważanego układu. Inaczej mówiąc jest to punkt największego oddalenia, który jest odpowiednikiem apogeum.

Aby wyznaczyć wspomniane dwa punkty wykorzystajmy równanie biegunowe orbity eliptycznej, które możemy zapisać jako:

$$r = \frac{p}{1 + e \cdot \cos \varphi}, \quad (9.1)$$

gdzie p jest parametrem orbity eliptycznej ($p = a \cdot (1 - e^2)$), a e jest jej mimośrodem. Wówczas równanie (9.1) jest równe:

$$r = \frac{a \cdot (1 - e^2)}{1 + e \cdot \cos \varphi}. \quad (9.2)$$

W trakcie ruchu planety po orbicie eliptycznej kąt φ będzie zmieniał się w zakresie od 0° do 180° . Wówczas na podstawie równania (9.2) mamy:

- dla **perycentrum**

$$r_1 = \frac{a \cdot (1 - e^2)}{1 + e \cdot \cos 0} = \frac{a \cdot (1 - e)(1 + e)}{1 + e} = a \cdot (1 - e) = q, \quad (9.3)$$

- dla **apocentrum**

$$r_2 = \frac{a \cdot (1 - e^2)}{1 + e \cdot \cos 180} = \frac{a \cdot (1 - e)(1 + e)}{1 - e} = a \cdot (1 + e) = Q. \quad (9.4)$$

Korzystając z zasady zachowania energii mechanicznej i momentu pędu możemy zapisać:

$$\frac{mv_1^2}{2} - \frac{GMm}{r_1} = E, \quad (9.4)$$

$$2mEr^2 + 2GMm^2r - k^2 = 0. \quad (9.5)$$

Zauważmy, że równanie (9.5) jest równaniem kwadratowym, które możemy uprościć do następującej postaci:

$$\alpha r^2 + \beta r + \gamma = 0, \quad (9.6)$$

gdzie słuszne są następujące zależności:

$$r_1 + r_2 = -\frac{\beta}{\alpha}, \quad r_1 r_2 = \frac{\gamma}{\alpha}, \quad r_1 + r_2 = 2a. \quad (9.7)$$

Niech $E = -\frac{GMm}{2a}$, następnie w oparciu o równania (9.3) oraz (9.4) mamy:

$$a(1 - e) \cdot a(1 + e) = -\frac{k^2}{2mE}. \quad (9.8)$$

Wykorzystując powyższe zależności oraz wykonując proste przekształcenia otrzymujemy:

$$a^2(1 - e^2) = -\frac{k^2}{2m \left(-\frac{GMm}{2a} \right)} = \frac{k^2 a}{GMm^2}. \quad (9.9)$$

Przekształcając równanie (9.9) łatwo można pokazać, że:

$$k = \sqrt{GMm^2 a (1 - e^2)}. \quad (9.10)$$

Odpowiedź: Poszukiwane równania zostały wykazane.

Zadanie 10. W jaki sposób anomalia rzeczywista zależy od czasu w ruchu parabolicznym w zagadnieniu 2 ciał (równanie Barkera)?

Rozwiązanie:

Wykorzystując rozważania przedstawione w I i II części artykułu związanego z mechaniką nieba możemy zapisać że:

$$r = \frac{p}{1 + e \cdot \cos \varphi}. \quad (10.1)$$

Zauważmy, że równanie (10.1) jest określane jako równanie opisujące krzywą stożkową, gdzie p jest parametrem elipsy ($p = a(1 - e^2)$), gdzie a jest dużą półosią orbity, a e jest jej mimośrodem:

$$r = \frac{p}{1 + e \cdot \cos \varphi} = \frac{a(1 - e^2)}{1 + e \cdot \cos \varphi}. \quad (10.2)$$

Żałujemy dla uproszczenia, że ciało znajduje się w peryhelium orbity wówczas $\varphi = 0$, więc równanie (10.2) jest równe:

$$r = \frac{a(1-e^2)}{1+e \cdot \cos \varphi} = \frac{a(1-e)(1+e)}{1+e} = a(1-e) = q. \quad (10.3)$$

Ponieważ rozważamy ruch ciała, dla którego torem jest parabola, więc na mocy wcześniejszych rozważań $e = 1$, wówczas z równania (10.2) otrzymamy:

$$p = r \cdot (1 + e \cdot \cos \varphi). \quad (10.4)$$

Wówczas na podstawie powyższych założeń mamy:

$$\begin{aligned} p &= r \cdot (1 + e \cdot \cos \varphi) = r + r \cdot e \cdot \cos \varphi = r + r \\ &= q + q = 2q. \end{aligned} \quad (10.5)$$

Na podstawie równania (10.5) możemy zapisać, że:

$$r = \frac{2q}{1 + \cos \varphi}. \quad (10.6)$$

Zauważmy, że z wzorów trygonometrycznych wynika, że:

$1 + \cos \varphi = 2 \cdot \cos^2 \frac{\varphi}{2}$, wówczas równanie (10.6) dane jest jako:

$$r = \frac{q}{\cos^2 \frac{\varphi}{2}} = q \cdot \sec^2 \frac{\varphi}{2}. \quad (10.7)$$

Upraszczając równanie (10.7) wykorzystaliśmy funkcję trygonometryczną secans, która jest odwrotnością cosinusa. Aby określić ruch po paraboli należy wykorzystać równanie określające całość pól w postaci:

$$r^2 \dot{\varphi} = c, \quad (10.8)$$

Następnie wykorzystajmy zależność na moment pędu: $c = K = \sqrt{G \cdot M \cdot a \cdot (1 - e^2)}$, wówczas możemy zapisać, że $c^2 = p \cdot \mu$, gdzie p jest parametrem elipsy a $\mu = GM$. Takie podejście sprawia, że równanie (10.8) możemy wyrazić jako:

$$r^2 \dot{\varphi} = \sqrt{2q\mu}. \quad (10.9)$$

W oparciu o powyższe zależności możemy zapisać, że:

$$\left(\frac{q}{\cos^2 \frac{\varphi}{2}} \right)^2 \frac{d\varphi}{dt} = \sqrt{2q\mu}, \quad (10.10)$$

lub

$$\frac{q^2}{\cos^4 \frac{\varphi}{2}} d\varphi = \sqrt{2q\mu} dt. \quad (10.11)$$

Następnie dla wygody dalszych obliczeń rozpiszmy lewą stronę równania (10.11) do następującej postaci:

$$\frac{q^2}{\cos^2 \frac{\varphi}{2}} \cdot \frac{\sin^2 \frac{\varphi}{2} + \cos^2 \frac{\varphi}{2}}{\cos^2 \frac{\varphi}{2}} d\varphi = \sqrt{2q\mu} dt, \quad (10.12)$$

lub

$$\frac{1}{2} \sqrt{2q\mu} dt = q^2 \left(1 + \tan^2 \frac{\varphi}{2} \right) \frac{d\frac{\varphi}{2}}{\cos^2 \frac{\varphi}{2}}. \quad (10.13)$$

W wyniku całkowania równania (10.13) otrzymujemy, że:

$$\sqrt{\frac{\mu}{2}} \frac{(t-t_0)^{\frac{3}{2}}}{q^{\frac{3}{2}}} = \tan \frac{\varphi}{2} + \frac{1}{3} \tan^3 \frac{\varphi}{2}. \quad (10.14)$$

Uwaga = w powyższych obliczeniach wykorzystaliśmy następujące zależności:

$$\frac{\varphi}{2} = y, \quad d\frac{\varphi}{2} = dy, \quad \frac{d\frac{\varphi}{2}}{\cos^2 \frac{\varphi}{2}}, \quad (10.14a)$$

$$\int \frac{d\frac{\varphi}{2}}{\cos^2 \frac{\varphi}{2}} = \int \frac{dy}{\cos^2 y} = \tan y, \quad (10.14b)$$

$$\begin{aligned} \int \frac{\tan^2 \frac{\varphi}{2}}{\cos^2 \frac{\varphi}{2}} d\frac{\varphi}{2} &= \int \frac{\tan^2 y}{\cos^2 y} dy = \int \tan^2 y d(\tan y) \\ &= \frac{\tan^3 y}{3} = \frac{1}{3} \tan^3 \frac{\varphi}{2}. \end{aligned} \quad (10.14c)$$

Odpowiedź: Równanie (10.14) nosi nazwę równania Barkera i wyraża ono związek między anomalią prawdziwą i czasem jaki upłynął od chwili przejścia ciała przez peryhelium swojej orbity.

Zadanie 11. Kometa Halley'a, dla której parametry orbitalne wynoszą odpowiednio $e = 0.967281$, $a = 17.047$ au zbliża się do Słońca. W jakim okresie czasu, w ciągu jednego obiegu, kometa znajdzie się bliżej Słońca niż Ziemia. Obliczyć szukany odstęp czasu przybliżając rzeczywistą orbitę komety orbitą paraboliczną o tym samym peryhelium. Przyjąć, że Ziemia porusza się wokół Słońca po orbicie kołowej o promieniu $a_0 = 1$ au, a okres obiegu Ziemi wokół Słońca wynosi $T_0 = 1$ rok.

Rozwiązanie:

Aby rozwiązać to zagadnienie wykorzystajmy w tym miejscu równanie opisujące krzywą stożkową:

$$r = \frac{p}{1 + e \cdot \cos \varphi}. \quad (11.1)$$

Zgodnie z treścią zadania mamy przybliżyć rzeczywistą orbitę komety na orbitę paraboliczną (dla której na podstawie wcześniejszych rozważań przyjmujemy, że $e = 1$), wówczas mamy:

$$r = \frac{p}{1 + \cos \varphi}, \quad (11.2)$$

gdzie tak jak poprzednio p jest parametrem elipsy. Jeżeli założymy, że rozpatrywana kometa znajduje się w peryhelium orbity to w oparciu o równanie (11.2) otrzymujemy:

$$r = \frac{a \cdot (1 - e^2)}{1 + 1 \cdot \cos 0} = \frac{p}{2}. \quad (11.3)$$

Następnie przekształćmy równanie (11.1) w taki sposób aby można było wyznaczyć na jego podstawie parametr p . Dodatkowo wykorzystajmy założenia wynikające z treści zadania, wówczas mamy:

$$p = r \cdot (1 + e \cdot \cos \varphi) = r + r \cdot e \cdot \cos \varphi \\ = r + r = q + q = 2q. \quad (11.4)$$

Po podstawieniu odpowiednich wartości mamy, że:

$$p = 2q = 1.1155 \text{ au}. \quad (11.5)$$

Wykorzystajmy teraz III prawo Keplera, niech:

$$T^2 = \frac{4\pi^2}{GM} a^3, \quad (11.6)$$

lub

$$GM = \frac{4\pi^2 a^3}{T^2} = \frac{4\pi^2 a_0^3}{T_0^2}. \quad (11.7)$$

Następnie niech słuszną będzie zależność:

$$GM = \mu = 4\pi^2 \left[\frac{1 \text{ au}^3}{1 \text{ rok}^2} \right].$$

Aby wyznaczyć poszukiwany czas wykorzystajmy do tego celu równanie Barkera, które możemy wyrazić jako:

$$t - t_0 = \frac{1}{3}(r+p) \sqrt{\frac{2r-p}{\mu}}. \quad (11.8)$$

Następnie w oparciu o zależność (11.8) możemy zapisać:

$$\tau = 2(t - t_0) = \frac{2}{3}(r+p) \sqrt{\frac{2r-p}{\mu}} \\ = \frac{2}{3}(a_0+p) \sqrt{\frac{2a_0-p}{4\pi^2}}. \quad (11.9)$$

Upraszając prawą stronę tożsamości (11.9) możemy zapisać, że:

$$\tau = \frac{1}{3 \cdot \pi} (a_0 + p) \sqrt{2a_0 - p} \\ = 0.211101 \text{ lat} \approx 0.21 \text{ lat}. \quad (11.10)$$

Odpowiedź: Poszukiwany okres odstępu czasu wynosi w przybliżeniu 0.21 lat.

Zadanie 12. Sztuczny satelita Ziemi wystrzelony został z równika w kierunku Ziemi. Znaleźć taki stosunek promienia orbity satelity do promienia Ziemi, przy którym satelita periodycznie przelatuje nad miejscem, skąd został wystrzelony, dokładnie co dwie doby. Promień Ziemi $R_Z = 6400 \text{ km}$.

Rozwiązanie:

Dane: $t = 2T_0 = 48^h$, $T_0 = 24^h$, $R_Z = 6400 \text{ km}$, $r = ?$

Rozważmy dwa następujące przypadki, niech

a) $\omega_s > \omega_z$ wówczas mamy:

$$\omega_s t = \omega_z t + 2\pi. \quad (12.1)$$

Podzielmy teraz obustronnie równanie (12.1) przez t wówczas mamy:

$$\omega_s = \frac{\omega_z t + 2\pi}{t} = \omega_z + \frac{2\pi}{t}. \quad (12.2)$$

Podstawmy teraz warunki wynikające z treści zadania wówczas mamy:

$$\omega_s = \frac{2\pi}{T_0} + \frac{2\pi}{2T_0} = \frac{2\pi}{T_0} + \frac{\pi}{T_0} = \frac{3\pi}{T_0}. \quad (12.3)$$

Następnie porównajmy ze sobą dwie siły dośrodkową i grawitacji co możemy zapisać jako:

$$m_s \omega_s^2 r = \frac{GM_Z m_s}{r^2}. \quad (12.4)$$

Uprośćmy teraz równanie (12.4) dzieląc je obustronnie przez m_s :

$$\omega_s^2 r^3 = GM_Z. \quad (12.5)$$

Wykorzystując równanie (12.3) lewą stronę równania (12.5) możemy wyrazić jako:

$$\frac{9\pi^2 r^3}{T_0^2} = GM_Z, \quad (12.6)$$

jednocześnie pamiętając, że: $g = \frac{GM_Z}{R_Z^2}$, wówczas otrzymamy:

$$\frac{9\pi^2 r^3}{T_0^2} = g R_Z^2. \quad (12.7)$$

Po prostych przekształceniach równania (12.7) możemy zapisać, że:

$$\frac{r}{R_Z} = \sqrt[3]{\frac{T_0^2 g}{9\pi^2 R_Z}} \approx 5. \quad (12.8)$$

W drugim przypadku mamy, że:

b) $\omega_s < \omega_z$ wówczas możemy zapisać, że:

$$\omega_s t = \omega_z t - 2\pi. \quad (12.9)$$

Postępując jak w podpunkcie a) łatwo można pokazać, że:

$$\omega_s = \frac{\pi}{T_0}. \quad (12.10)$$

Dokonajmy teraz również analogicznego porównania siły dośrodkowej i grawitacji jak w podpunkcie a), wówczas mamy:

$$m_s \omega_s^2 r = \frac{GM_Z m_s}{r^2}, \quad (12.11)$$

Przeprowadzając analogiczne rozumowanie jak w podpunkcie a) otrzymujemy, że:

$$\frac{r}{R_Z} = \sqrt[3]{\frac{T_0^2 g}{\pi^2 R_Z}} \approx 10.57. \quad (12.12)$$

Odpowiedź: Poszukiwane stosunki promieni dane są odpowiednio zależnościami (12.8) oraz (12.12).

dr hab. Marcin Wesołowski, prof. UR

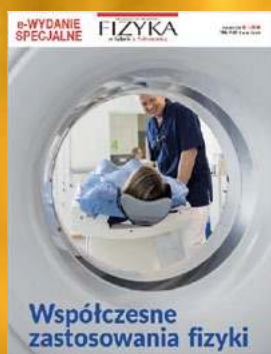
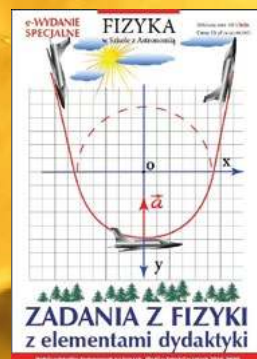
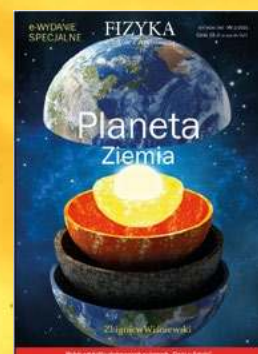
Uniwersytet Rzeszowski, Kolegium Nauk Przyrodniczych, Instytut Nauk Fizycznych,
Centrum Innowacji i Transferu Wiedzy Techniczno-Przyrodniczej
Uniwersytetu Rzeszowskiego.

Cyfrowe wydania specjalne

Fizyki w Szkole

Tylko w wersji PDF!

Już od 10 zł!
Wysyłamy na adres
mailowy!



Szczegóły i formularz zamówienia na www.aspress.com.pl/wydania-specjalne/

eprasa.pl c0287c7302

Prenumerata 2024

Otwiera dostęp do:

- ▶ artykułów o najważniejszych odkryciach i wydarzeniach w fizyce
- ▶ publikacji o zastosowaniu fizyki w różnych dziedzinach życia
- ▶ zadań, doświadczeń i eksperymentów
- ▶ nowości z astronomii i astrofizyki



**PRZEDŁUŻ
LUB
ZAMÓW**

Do wyboru:

- ☐ Wersja drukowana
lub cyfrowa – pliki PDF
- ☐ Prenumerata roczna
i półroczna

Szczegóły i formularz zamówienia na www.aspress.com.pl/prenumerata/

eprasa.pl c0287c7302