

h=6.63·10⁻³⁴ J·s

CENA 2

NR 1 1974

POPULARNY MIESIĘCZNIK MATEMATYCZNO-FIZYCZNY



SPIS TREŚCI

Czy liczby rzeczywiste
są rzeczywiste
prof. dr Roman Sikorski str. 2

Zadania str. 4

Laboratorium w domu
— Widzimy
w podczerwieni
dr Jan Gaj str. 5

Delta z wizytą
w Instytucie Fizyki
Jądrowej w Krakowie str. 6

Jak zidentyfikowaliśmy
hiperjądro podwójne
*prof. dr Janusz
Zakrzewski* str. 8

Algorytmy (I)
dr Andrzej Skowron str. 11

Nasi na Olimpiadzie str. 13

Na pytanie co to jest
cząstka elementarna
odpowiada
*prof. dr Grzegorz
Białkowski* str. 14

Sześć zadań
— jedno rozwiązanie
dr Maciej Bryński str. 16

Wielościanny gwiazdziste str. 17

W następnym numerze:
Dlaczego Niemcy nie zdążyli
wynaleźć bomby atomowej
Sztuka wygrywania

„Delta”
matematyczno-fizyczny
miesięcznik popularny
Polskiego Towarzystwa Matematycznego
i Polskiego Towarzystwa Fizycznego
Wydawnictwo Polskiej Akademii Nauk
„Ossolineum”
Warszawa

Komitet Redakcyjny
prof. dr G. Białkowski
doc. dr A. Blikle
prof. dr A. Hryniewicz
doc. dr B. Iwaszkiewicz
prof. dr J. Janik
doc. dr J. Jatzak
prof. dr Leon Jeśmanowicz —
przewodniczący
prof. dr Z. Krygowska
prof. dr K. Leibler
mgr W. Łucznik
mgr A. Mąkowski
prof. dr A. Pełczyński
prof. dr Arkadiusz Piekara —
wiceprzewodniczący
prof. dr J. Rayski
prof. dr W. Rubinowicz
prof. dr A. Schinzel
prof. dr Z. Semadeni
prof. dr M. Subotowicz
dr A. Wakulicz
doc. dr W. Zawadowski

Redaguje Kolegium w składzie:
mgr J. Bednarczuk — sekr. red.
doc. dr T. Hofmokr — z-ca red. nacz.
dr M. Kordos — red. nacz.
dr Z. Płochocki
opracowanie graficzne
art. graf. K. Dobrowolski

Adres Redakcji
ul. Śniadeckich 8,
00-656 Warszawa,
PTM

Zakład Narodowy im.
Ossolińskich — Wydawnictwo.
Wrocław, Oddział w Warszawie.
Nakład 30000 egz. Objętość 2 ark.
wyd.; 2,50 ark. druk.;
papier offsetowy III kl., 80g, 61x86.
Wydrukowano w Drukarni im.
Rewolucji Październikowej,
Warszawa, ul. Mińska 65.
Nr zam. 1668/73. R-29

ŻYCIE

w coraz większym stopniu
karmi nas pojęciami,
które nie tak dawno jeszcze
były używane wyłącznie
przez matematyków i fizyków
w ich działalności zawodowej.

TO DOBRZE.

Odpowiada to bowiem miejscu,
jakie obecnie zostało zajęte
przez matematykę i fizykę
w całokształcie ludzkiej działalności.

TO ŹLE,

bo pokarm to niejednokrotnie niestrawny,
a więc mimo niewątpliwych wartości
nie przynoszący większych korzyści
konsumentowi.

KOMPUTER I LASER

to pojęcia nie należące dziś już
tylko do świata nauk ścisłych
-ich znajomość to nie mniej niezbędny
obecnie element ogólnego wykształcenia
od informacji, co to jest heksametr,
czy też, o co toczyła się Wojna Dwóch Róż.

MATEMATYKA I FIZYKA

stając się składnikami ogólnej kultury
muszą być jednak zrozumiałe.
Używając nazw zaczerpniętych
z „arsenału” tych nauk
musimy w każdym przypadku wiedzieć

CO TO JEST?

Aby to ułatwić każdemu,
Polskie Towarzystwo Matematyczne
i Polskie Towarzystwo Fizyczne
wydają matematyczno-fizyczny
miesięcznik popularny

DELTA.

Czy liczby rzeczywiste są rzeczywiste?

prof. dr
Roman SIKORSKI,
członek rzeczywisty PAN



Liczby naturalne są niewątpliwie naturalne. Liczby całkowite niewątpliwie zasługują na nazwę „całkowite”. Liczby wymierne należałoby może nazywać liczbami mierzącymi lub wymierzającymi, bowiem wszystkie pomiary wykonujemy w praktyce w liczbach wymiernych, zresztą nie tylko pomiary: wszelkie rachunki na konkretnych liczbach wykonywane są w praktyce wyłącznie w obrębie liczb wymiernych. Po co więc wprowadzać szersze, lecz znacznie trudniejsze pojęcie liczb rzeczywistych, skoro liczby wymierne wystarczają w rachunkach? Definicja liczb rzeczywistych nastrocza zawsze pewne trudności, wskutek tego w podręcznikach szkolnych jest raczej przemycana, niż precyzyjnie formułowana.

Pojęcie liczby naturalnej jest łatwe do przyswojenia. Tak łatwe, tak swojskie, że wydaje się, iż liczby naturalne są wzięte bezpośrednio z otaczającego świata materialnego. Zapominamy na ogół, że łatwo napisać, nawet na małym kawałku papieru, nazwę liczby naturalnej n , która jest nierealizowalna w świecie materialnym, tzn. dla której trudno podać przykład n -elementowego zbioru przedmiotów realnych. Liczba $n = 1000^{1000^{1000}}$ jest przykładem takiej nierealizowalnej liczby naturalnej. Początkowe liczby naturalne są łatwo wyobrażalne, nie można tego powiedzieć o liczbach bardzo dużych. Niemniej zbiór N wszystkich liczb naturalnych można łatwo zdefiniować.

Musi zawierać liczbę 1. Jeśli zawiera liczbę n , to musi zawierać także liczbę $n+1$. I nic więcej, tzn. jest najmniejszym zbiorem o powyższych dwu własnościach. Przytoczona definicja zbioru liczb naturalnych przypomina dowcip o pakowaniu do pustej walizki chusteczek do nosa. Oczywiście można tam włożyć jedną chusteczkę. Wiadomo z praktyki, że jeśli włożyliśmy do walizki n chusteczek, to $n+1$ -sza też da się załadować. Zatem do walizki można włożyć tyle chusteczek, ile jest liczb naturalnych, czyli nieskończenie wiele!

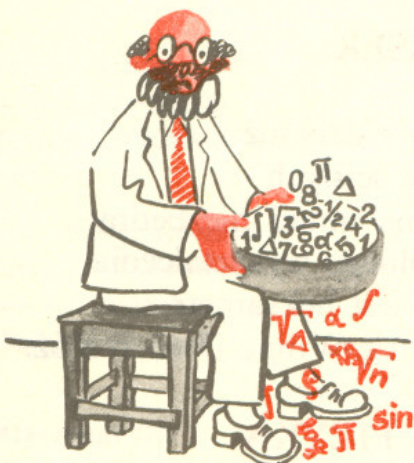
Możemy sobie wyobrazić, że matematyk ma taką abstrakcyjną walizeczkę N zawierającą wszystkie liczby naturalne. W drugiej dodatkowej walizeczce nosi ich „odbicia lustrzane w zerze”, tzn. liczby całkowite ujemne. Musi się jeszcze zdecydować, do której walizeczki włożyć liczbę zero. Zdania są podzielone, jedni lubią zaliczać zero do liczb naturalnych, inni tego nie lubią. Rzecz w istocie nie warta przysłowiowego funta kłaków. Kłócić się o zero? O matematyczne „nic”? Nie warto!

Oprócz walizeczek z liczbami całkowitymi matematyk ma także maszynkę do precyzyjnego siekania tych liczb. Mówiąc poważniej, matematyk z liczb całkowitych łatwo konstruuje liczby wymierne, tzn. „ułamki” m/n , gdzie m jest liczbą całkowitą, a n — liczbą naturalną (nie zerem!). Pewne z tych ułamków należy uznać za równe. Na ułamkach tych można w naturalny sposób zdefiniować podstawowe działania arytmetyczne: dodawanie i mnożenie, oraz wtórne działania odwrotne: odejmowanie i dzielenie (nie przez zero!). Można je też uporządkować, tzn. wprowadzić relację mniejszości $x < y$. Intuicyjnie łatwo sobie „wyobrazić” liczbę wymierną. Łatwo bowiem wyobrazić sobie n -tą część czegoś, to znaczy liczbę $1/n$; łatwo też wyobrazić sobie m takich części, tzn. liczbę m/n , z ewentualną zmianą znaku, czyli z „odbiciem lustrzanym w zerze”.

Pocziwe liczby wymierne! Tak bardzo są użyteczne! Wszystkie transakcje handlowe, bankowe, wszelkie pomiary, wszelkie rachunki techniczne są na nich oparte. Z punktu widzenia czystej praktyki jest ich za dużo, bo nieskończenie wiele. W praktyce do rachunków używa się tylko skończenie wielu liczb wymiernych (trudno byłoby oszacować ile). Tak już jednak jest z matematykami. Jeśli coś tworzą, czynią to na ogół w pełnej ogólności, wskutek tego na wyrost, na ogół więcej niż jest to potrzebne w praktyce.

Niestety zbiór liczb wymiernych ma wady. Wprawdzie jest gęsty, tzn. dla dowolnych dwu liczb wymiernych istnieje trzecia położona między nimi. Jest jednak dziurawy, przy tym jego dziury są również rozmieszczone w sposób gęsty, tzn. między dowolnymi dwiema liczbami wymiernymi znajduje się zawsze dziura. Te dziury biorą się nie ze starości zbioru, nie wskutek przetarcia zbioru w wyniku tak częstego używania liczb wymiernych w praktyce. Po prostu taka jest matematyczna natura tego zbioru. Zbiór liczb wymiernych ma postać bardzo gęstego jednowymiarowego sita.

Musimy wyjaśnić, co rozumiemy przez dziurę w zbiorze liczb wymiernych (zawodowi matematycy mówią „luka” zamiast „dziura”). Dziurą nazywamy taki



Przyjęliśmy milcząco, że p i q są wszystkimi pierwiastkami równania. Może się zdarzyć, że $p = q$ i równanie ma jeszcze inny pierwiastek, prócz liczby $p (= q)$. Należałoby postąpić następująco: warunki zadania oznaczają, że $p^2 + p \cdot p + q = 0$ i $q^2 + pq + q = 0$. Jest więc $q = -2p^2$. Z drugiego równania wynika więc, że $4p^4 - 2p^3 - 2p^2 = 0$ czyli $2p^2(2p^2 - p - 1) = 0$, skąd $p = 0$ lub $p = 1$ lub $p = -\frac{1}{2}$.

Odpowiadające wartości q to 0 , -2 i $-\frac{1}{2}$.

W porównaniu z metodą przedstawioną na str. 8 otrzymaliśmy dodatkowo

rozwiązanie $p = q = -\frac{1}{2}$.



podział zbioru W wszystkich liczb wymiernych na dwa niepuste podzbiory W_1 i W_2 , że po pierwsze każda liczba ze zbioru W_1 jest mniejsza od każdej liczby zbioru W_2 , oraz — po drugie — w zbiorze W_1 nie ma liczby największej, a w zbiorze W_2 nie ma liczby najmniejszej. Mówimy, że taka dziura leży między liczbami wymiernymi w_1 i w_2 ($w_1 < w_2$), jeśli w_1 należy do W_1 , a w_2 należy do W_2 . Na przykład, zaliczmy do W_1 wszystkie ujemne liczby wymierne i te nieujemne, których kwadrat jest mniejszy od 2, a do W_2 zaliczmy wszystkie dodatnie liczby wymierne, których kwadrat jest większy od 2. Podział zbioru W na zbiory W_1 , W_2 jest dziurą (luką) w zbiorze liczb wymiernych. Bardzo łatwo sprawdzić, że dziura ta leży między liczbami 1 i 2. Bez trudu można by wyznaczyć bardziej dokładnie położenie tej dziury. Prosty rachunek dowodzi, że leży ona między 1,41 a 1,42. Oczywiście można wyznaczyć jej położenie jeszcze dokładniej.

Istnienie dziur w zbiorze liczb wymiernych jest źródłem wielu kłopotów. Obrazowo można powiedzieć, że przez te dziury wycieka treść matematyczna wielu pięknych twierdzeń, zwłaszcza tych o bardziej subtelnej strukturze. Zbiór liczb wymiernych jest świetny do rachunków na liczbach konkretnych, ale zły dla wielu celów teoretycznych, dla bardziej skomplikowanych działań na liczbach, niż dodawanie, odejmowanie, mnożenie i dzielenie, zwłaszcza jeśli chce się wykonywać te działania w sposób dokładny, a nie przybliżony. Już pierwiastkowanie nie jest wykonalne w tym zbiorze. Spróbujcie określić tak pożyteczną funkcję jak $\log x$ ($x > 0$) tak, by zarówno x jak i $\log x$ były liczbami wymiernymi — nic z tego nie wyjdzie. W jednym i drugim przypadku czuje się po prostu brak liczb, zbiór liczb wymiernych jest za mały, by wykonywać w nim logarytmowanie lub pierwiastkowanie (dokładne, a nie przybliżone).

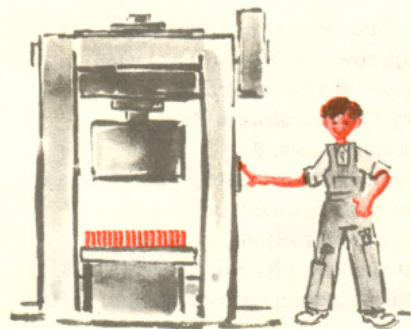
Matematyk bardzo nie lubi, gdy pewne działania, wyglądające na naturalne lub pożyteczne, są niewykonalne. W wielu przypadkach usuwa niewykonalność działań przez odpowiednie rozszerzenie zbioru przedmiotów, na których działania mają być wykonane, lub które mają być wynikiem tych działań. Można by przytoczyć wiele przykładów — nawet z najnowszej matematyki. Przytoczymy tu tylko jeden: rozszerzenie zbioru liczb wymiernych do zbioru liczb rzeczywistych.

Rozszerzenie to wykonuje się w sposób następujący. Przed każdą dziurą w zbiorze liczb wymiernych matematyk kładzie kołek do zatkania jej (liczbę wymierną wygodnie jest interpretować jako punkt na osi liczbowej; wówczas kołek do zabicia dziury też można wyobrażać sobie jako punkt na tej osi). Następnie jednym uderzeniem młotka matematyk wbija wszystkie kołki. Zwracam uwagę na fakt, że matematyk jednym aktem woli wbija od razu wszystkie kołki we wszystkie dziury! Nie należy wyobrażać sobie procesu wbijania w ten sposób, że najpierw numeruje wszystkie dziury liczbami naturalnymi, a potem chodzi kolejno od n -tej dziury do $n+1$ -szej i zabija je kołkami. Takie postępowanie byłoby niemożliwe, można bowiem udowodnić, że dziur w zbiorze liczb wymiernych jest tak dużo, że nie można ich ponumerować wszystkimi liczbami naturalnymi.

Wszystkie liczby wymierne można ponumerować kolejnymi liczbami naturalnymi, ale dziur w tym zbiorze — nie! Dziwne, nieoczekiwane, ale prawdziwe. Jakkolwiek ponumerowalibyśmy te dziury wszystkimi liczbami naturalnymi, zawsze znalazłoby się jeszcze nieskończenie wiele nieponumerowanych dziur! Zawodowi matematycy mówią, że jakiś zbiór jest przeliczalny, jeśli wszystkie jego elementy można ponumerować kolejnymi różnymi liczbami naturalnymi, a jest nieprzeliczalny, jeśli tego nie można zrobić. Zbiory nieprzeliczalne są znacznie większe, znacznie bogatsze w elementy niż zbiory przeliczalne. Zbiór liczb wymiernych jest przeliczalny, a zbiór wszystkich dziur w tym zbiorze jest nieprzeliczalny. Widać stąd, że w zbiorze liczb wymiernych jest więcej dziur niż liczb! Czyż można mieć zaufanie do takiego zbioru? Nic dziwnego, że wiele treści matematycznej wycieka przezeń.

Powróćmy do rozszerzenia liczb wymiernych do rzeczywistych. Białe kołeczki nazwiemy liczbami niewymiernymi, a całość, tzn. zarówno liczby wymierne jak i niewymierne nazwiemy liczbami rzeczywistymi zgodnie z powszechnie ustaloną terminologią. Czytelnik łatwo się domysli, że kołeczek, który zatkał dziurę, podaną jako jedyny przykład ilustrujący to pojęcie, oznaczać będziemy symbolem $\sqrt{2}$.

Jak w każdej innej konstrukcji matematycznej, tak i w tym przypadku należy wyróżnić dwie strony tego samego zadania: 1) intuicyjne wyjaśnienie celu i metody konstrukcji oraz 2) precyzyjny opis jej wykonania z zachowaniem najwyższych kryteriów ścisłości współczesnej matematyki. Opisane powyżej wbijanie kołeczków w dziury to tylko intuicyjny opis konstrukcji, wyjaśnienie jej celu. Precyzyjny



opis konstrukcji — to zupełnie inne zagadnienie. Zagadnienie — powiedziałbym — dosyć niewdzięczne. Znamy dwie metody „wbijania kołeczków”, mianowicie metodę Dedekinda i metodę Cantora. Obydwie są bardzo precyzyjne i obydwie mają tę samą wielką wadę: zaciemniają mniej istotnymi szczegółami technicznymi podstawową, jasną i prostą intencję konstrukcji. Dlatego nie przytoczymy tu żadnej z nich. Wspomnimy tylko o zasadniczej różnicy między tymi metodami. Oczywiście kołeczki muszą być z czegoś zrobione, z jakiegoś tworzywa, naturalnie z jakiegoś abstrakcyjnego tworzywa pojęć matematycznych. Otóż metody Cantora i Dedekinda różnią się głównie materiałem, z którego zrobione są kołeczki. W metodzie Dedekinda kołeczkiem zatykającym dziurę jest sama dziura! Dziurę zatyka się nią samą! Można powiedzieć, że w metodzie tej koszty zużycia materiałów zostały doprowadzone do minimum, do zera!

Konstrukcja została wykonana, kołeczki są wbite. Pozostało nam sprawdzić, czy robota została rzetelnie wykonana, czy przypadkiem w trakcie wbijania nie powstały jakieś nowe dziury. Na szczęście wszystko jest w absolutnym porządku, zbiór liczb rzeczywistych jest całkowicie szczelny. Przytoczoną definicję dziury można wprawdzie sformułować w odniesieniu do liczb rzeczywistych, nie ma jednak potrzeby wprowadzania takiego pojęcia, po prostu w ogóle nie ma takich dziur.

Okazuje się, że zbiór liczb rzeczywistych ma własności jeszcze lepsze, niż zbiór liczb wymiernych. Można ten zbiór uporządkować, można uogólnić podstawowe działania arytmetyczne, dodawanie, odejmowanie, mnożenie i dzielenie na liczby rzeczywiste, ale ponadto można w tym zbiorze wykonywać bez żadnych ograniczeń wiele innych pożytecznych operacji matematycznych, jak pierwiastkowanie, potęgowanie, logarytmowanie itp., których wykonalność w dziedzinie liczb wymiernych była mocno utrudniona, jeśli nie wręcz niemożliwa.

Okazało się ponadto, że w oparciu o pojęcie liczby rzeczywistej można zbudować całą analizę matematyczną, olbrzymią gałąź współczesnej matematyki. U podstaw wszystkich twierdzeń tej części matematyki leży „szczelność” zbioru liczb rzeczywistych (matematycy zawodowi używają przymiotnika „zupełny” zamiast „szczelny”). Twierdzenia analizy matematycznej przestają być prawdziwe, jeśli zbiór liczb rzeczywistych zastąpić przez zbiór liczb wymiernych. To właśnie mieliśmy na myśli, mówiąc żartobliwie o przeciekaniu wiedzy matematycznej przez dziury zbioru liczb wymiernych. Pojęcie liczby rzeczywistej jest niezbędne dla całej matematyki teoretycznej. Jest również niezbędne dla formowania ogólnych metod matematyki stosowanej aż do momentu, gdy w grę wchodzi przybliżone rachunki na konkretnych liczbach. Wtedy powracamy do bardziej elementarnych liczb wymiernych.

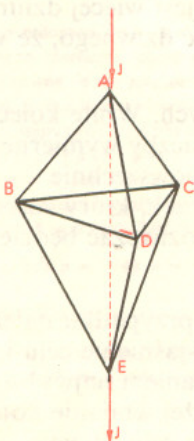
Niewątpliwie pojęcie liczby wymiernej jest prostsze niż pojęcie liczby rzeczywistej. Dla laika liczba rzeczywista, wprowadzona metodą Cantora lub Dedekinda, wydaje się być tworem dość mistycznym, wydaje się być znacznie mniej rzeczywistą — w potocznym znaczeniu tego słowa — niż liczba wymierna. Dla zawodowego matematyka liczba rzeczywista jest podstawowym narzędziem pracy, jest równie rzeczywista jak inne pojęcia matematyczne. Liczby rzeczywiste są równie rzeczywiste jak liczby wymierne, jedne i drugie są bowiem poprawnie zdefiniowanymi pojęciami istniejącymi w mózgu matematyka. Jedne i drugie mają ten sam typ rzeczywistości.



Zadania

redaguje dr Jędrzej JĘDRZEJEWSKI

redaguje mgr Andrzej MAKOWSKI

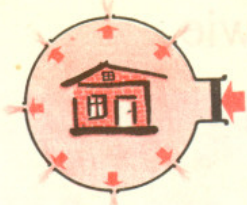


- F1.** Z prostych odcinków identycznego drutu oporowego zbudowano obwód elektryczny przedstawiony na rysunku. Prąd o natężeniu I dopływa do węzła A i wypływa z węzła E wzdłuż prostej AE , która jest prostopadła do płaszczyzny trójkąta równobocznego BCD utworzonego przez węzły B, C i D . Bryłę $ABCDE$ można uważać za dwa połączone podstawami ABC prawidłowe ostrosłupy trójkątne, każdy o innej wysokości. Wykazać, że w nieograniczonym ośrodku jednorodnym, w każdym punkcie odcinka AE wartość indukcji magnetycznej B wynosi zero. Rozwiązanie na str. 15.

- M1.** Udowodnić, że boki a, b, c trójkąta ABC spełniają równość $b^2 = a(a+c)$ wtedy i tylko wtedy, gdy $\angle ABC = 2\angle BAC$. (z pracy D. Rameshwara Rao) Rozwiązanie na str. 9.
- M2.** Znaleźć liczby rzeczywiste p i q , które są pierwiastkami równania $x^2 + px + q = 0$ (ze względu na zmienną x). (V. Gutenmacher) Rozwiązanie na str. 8.
- M3.** Dla jakich wartości rzeczywistych m układ równań

$$\begin{aligned} x^2 + y^2 &= m \\ -x^2 + y &= 2 \end{aligned}$$

ma dokładnie jedno rozwiązanie? (R. Hajłasz) Rozwiązanie na str. 10.



Laboratorium w domu

redaguje dr Jan GAJ

WIDZIMY W PODCZERWIENI

Kiepski dowcip, powie niejeden z Was. Kto nam kupi noktowizor? Mówię poważnie. Każdy, kto nie ma dwóch lewych rąk, może sam zarejestrować obraz w podczerwieni. Spróbujemy? No to do dzieła.

MATERIAŁY

Będą nam potrzebne: 1) gorący przedmiot jako źródło podczerwieni — wygodna jest lutownica, może być zagrzany w płomieniu duży klucz, 2) dwa słoiki np. po dżemie, 3) kawałek folii aluminiowej — najlepsza dostępna w handlu do celów gospodarstwa domowego, może być z opakowania czekolady, 4) kawałek folii polietylenowej — z torebki, 5) sznurek, 6) przedmiot o gładkiej kulistej powierzchni jako matryca zwierciadła wklęsłego — np. duża soczewka wypukła (może być z kondensora od powiększalnika Krokus, ewentualnie duży kulisty klosz od lampy), 7) sprzęt pomocniczy do ustawienia układu (np. dwie linijki i dwa stosy książek oraz latarka kieszonkowa).

IDEA DOŚWIADCZENIA

Pokazuje ją rysunek 1. Promieniowanie podczerwone wysyłane przez źródło (A) odbija się od zwierciadła wklęsłego (B) tworząc obraz źródła na ekranie (C) ewaporografu — przyrządu, który pozwoli nam ten obraz zobaczyć.

ZWIERCIADŁO WKŁĘSŁE

Na szyjkę słoika zakładamy folię aluminiową i zawiązujemy mocno sznurkiem. Następnie wciskamy w folię matrycę o kulistej powierzchni. Zwierciadło gotowe (rys. 2). Dla sprawdzenia próbujemy wytworzyć nim obraz świecącej żarówki na kartce papieru. Niezbyt ostry? No cóż, nie jest to przyrząd od Zeissa. Podobnej jakości możemy oczekiwać po naszym obrazie w podczerwieni — nie liczymy na zarejestrowanie drobnych szczegółów.

EWAPOROGRAF

Do drugiego słoika nalewamy trochę wody o temperaturze pokojowej i zawiązujemy sznurkiem na jego szyjce kawałek folii polietylenowej (rys. 3). Po chwili folia pokryje się od wewnątrz mgiełką wodną. To już jest ewaporograf, a jego ekranem jest oczywiście zaparowana folia. W miejscach, gdzie padnie odpowiednio silne promieniowanie podczerwone, folia ogrzeje się i zaparowanie zniknie — powstanie widoczny obraz.

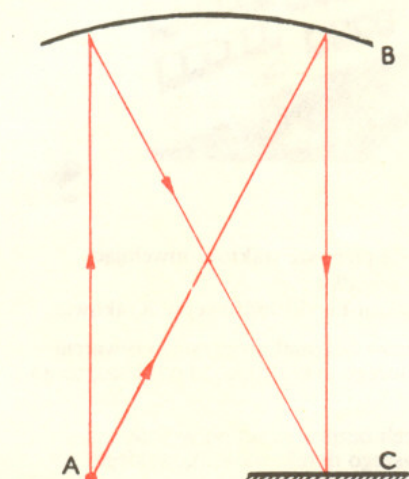
ZESTAWIAMY UKŁAD

Na stole umieszczamy blisko siebie lutownicę (lub inny gorący przedmiot) i ewaporograf, a nad nimi zwierciadło zgodnie z rysunkiem 1 w taki sposób, aby obraz lutownicy wytworzył się na ekranie ewaporografu. Wygodnie jest zastąpić na razie lutownicę płaską latarką kieszonkową i ustawić zwierciadło na odpowiedniej wysokości. Zwierciadło (słoik do góry dnem) można ustawić brzegami na dwóch linijkach opartych końcami na dwóch stosach książek. Po wstawieniu gorącego przedmiotu w miejsce latarki obraz pojawia się w ciągu kilkunastu sekund do paru minut, zależnie od warunków doświadczenia.

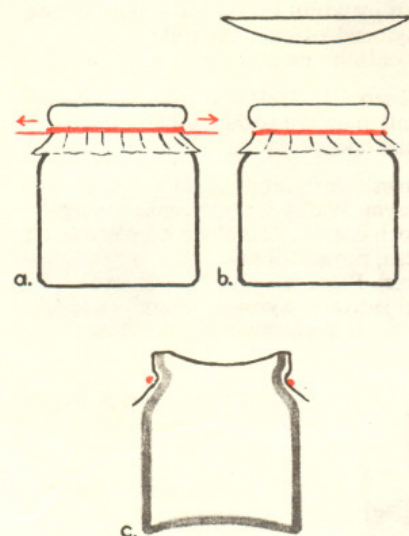
UWAGI DODATKOWE

Przedmiot gorący nie powinien być błyszczący i gładki, bo będzie słabo emitował podczerwień. Woda w ewaporografie nie powinna być za zimna (zaparowanie się nie pojawi) ani za ciepła (nie będzie ustępowało łatwo pod wpływem ciepła). Przyłożenie palca do folii na parę sekund powinno powodować zniknięcie zaparowania w tym miejscu.

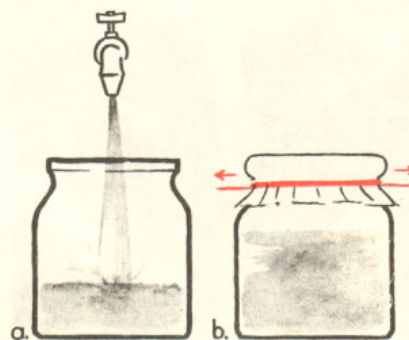
Doświadczenie można wzbogacić badając przepuszczalność różnych materiałów (szkło, folia polietylenowa) przez wstawianie ich w bieg promieni. Powodzenia! Bardzo prosimy o listy z opisem jak udało się doświadczenie lub z propozycjami udoskonaleń.



Rys. 1



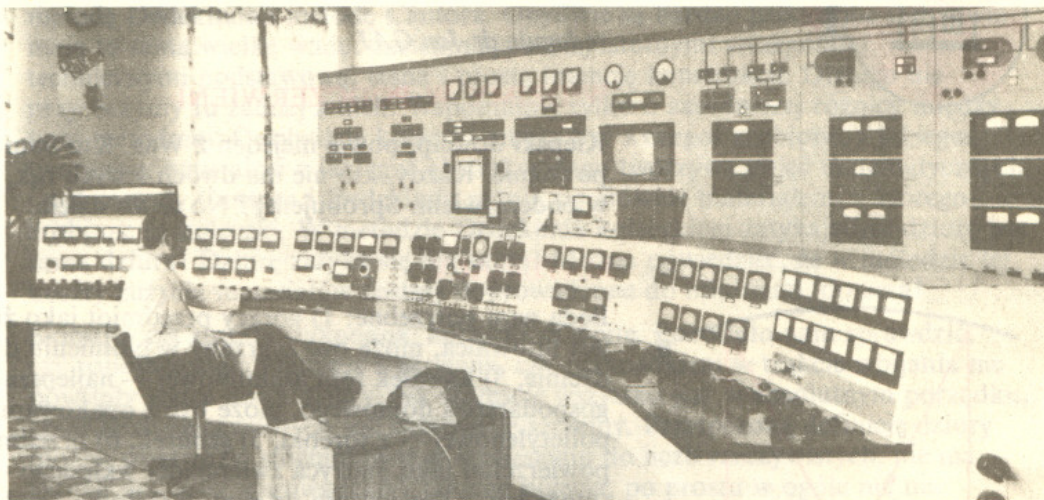
Rys. 2



Rys. 3

Delta z wizytą w Instytucie Fizyki Jądrowej w Krakowie

Pulpit sterowniczy cyklotronu



Dyrektor Instytutu Fizyki Jądrowej w Krakowie prof. dr Andrzej Hryniewicz

W 1955 r. na polach bronowickich pod Krakowem pojawiły się pierwsze traktory niwelujące teren. Było to rozpoczęcie budowy drugiego obok Warszawy ośrodka badawczego polskiej atomistyki. Nosi on obecnie nazwę Instytutu Fizyki Jądrowej w Krakowie.

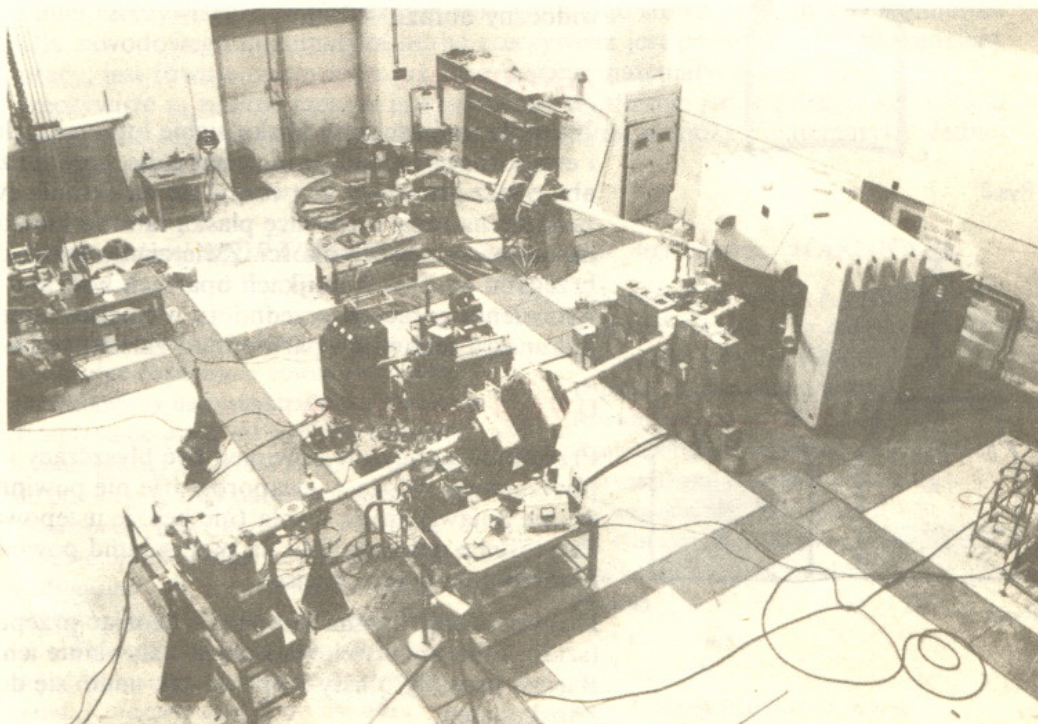
W trzy lata później, 20 listopada 1958 r., Prezes Rady Ministrów dokonał uroczystego otwarcia największego polskiego akceleratora cząstek naładowanych: dużego cyklotronu, przyspieszającego deuterony do energii 13 milionów eV.

Fizycy krakowscy niskich energii, którzy swoje prace rozpoczęli bezpośrednio po wojnie pod kierunkiem prof. dr Henryka Niewodniczańskiego, pierwszego dyrektora krakowskiego Instytutu — zyskali nowe, znakomite urządzenie badawcze.

Do prac na dużym cyklotronie byli oni dobrze przygotowani. Mieli bowiem już na swoim koncie udaną konstrukcję pierwszego polskiego cyklotronu o średnicy nabiegunków elektromagnesu 48 cm, który został uruchomiony w piwnicy dawnego budynku Instytutu Fizyki UJ w sylwestrową noc 1956 r. Dlatego też rozpoczęcie prac badawczych na nowym cyklotronie nastąpiło w rekordowo krótkim czasie. Już w kwietniu 1959 r. przeprowadzono na nim pierwsze pomiary.

Bombardując węgiel wiązką deuteronów przyspieszonych do energii 13 MeV w wyniku tzw. reakcji jądrowej strippingu [$^{12}\text{C}(\text{d}, \text{n})^{13}\text{N}$] uzyskiwano neutrony. Wynik pomiaru rozkładu kąтового i polaryzacji neutronów dostarczył cennych informacji odnośnie budowy jądra.

Od uruchomienia dużego cyklotronu minęło już 15 lat. Krakowski Instytut dokonał w tym czasie dużego skoku stając się znanym w świecie ośrodkiem badawczym. Ważnym wydarzeniem w jego historii było wejście w skład Instytutu zespołu fizyków wysokich energii. Zespół ten od pierwszych lat powojennych, pod kierunkiem prof. dr Mariana Mięśowicza, prowadził prace nad promieniowaniem kosmicznym w Akademii Górniczo-Hutniczej. Prace te stały się podstawą znanych obecnie w świecie badań tej grupy w dziedzinie fizyki jądrowej wysokich energii i cząstek



Sala wiązek cząstek naładowanych z dużego cyklotronu

elementarnych. Z nich również wywodzi się problematyka Instytutu Fizyki i Techniki Jądrowej AGH i obszerna dziedzina zastosowań metod jądrowych w geologii, górnictwie i hutnictwie.

Obecnie krakowscy fizycy wysokich energii zdobyli dostęp do największych akceleratorów cząstek naładowanych. Prowadzą eksperymenty na radzieckim synchrotronie w Sierpuchowie k/Moskwy, przyspieszającym protony do energii 76 miliardów eV .

Prowadzą prace na dużym akceleratorze w CERN (Genewa) oraz rozpoczynają doświadczenia z użyciem aktualnie największego urządzenia do przyspieszania cząstek naładowanych, które znajduje się w Batawii (USA). Akcelerator ten dostarcza wiązkę protonów o energii ok. 400 miliardów eV (GeV).

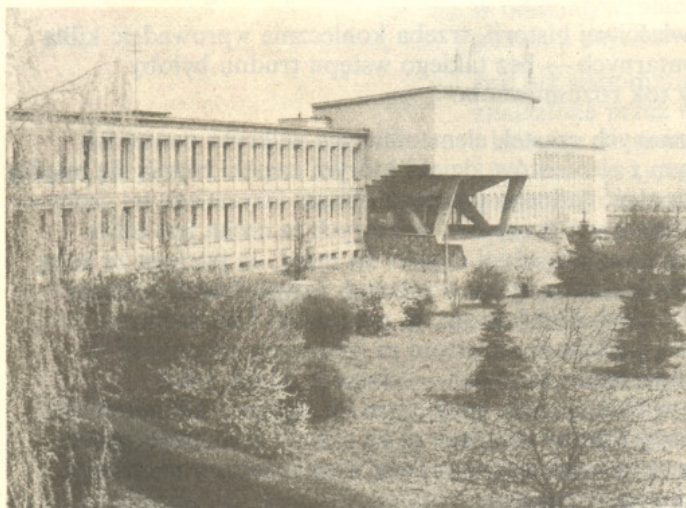
W Instytucie Fizyki Jądrowej w Krakowie pracuje obecnie 650 osób. Posiada on siedem zakładów naukowych i dwie samodzielne pracownie. Poza fizyką jądrową prowadzi też prace w dziedzinie badania fazy skondensowanej materii metodami jądrowymi oraz rozwija szeroko prace stosowane.

Jaka przyszłość?

Na to pytanie Dyrektor Instytutu, prof. dr Andrzej Hryniewicz przedstawia opracowany wspólnie z udziałem całego środowiska krakowskich fizyków program rozwoju Instytutu.

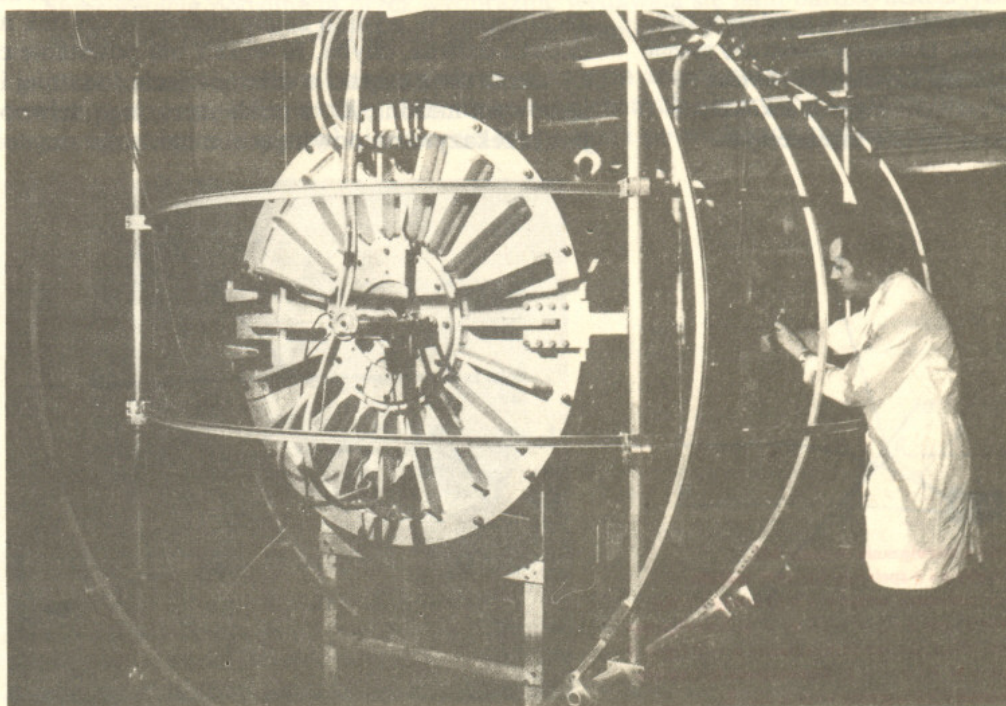
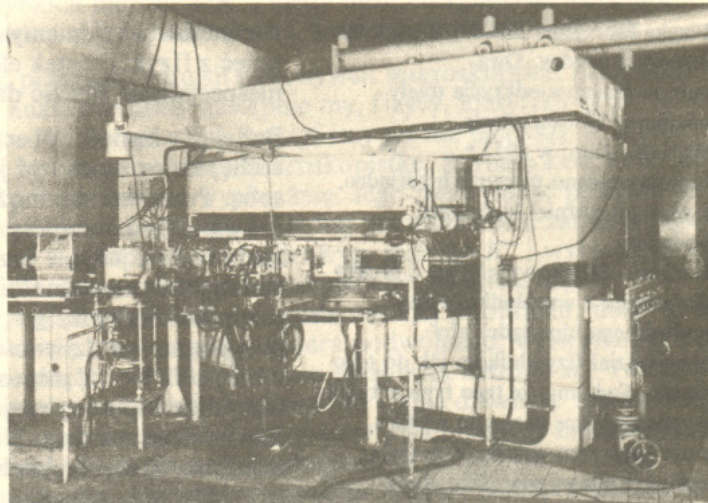
Ważne miejsce zajmuje w nim budowa w Krakowie nowego cyklotronu izochronicznego, przyspieszającego protony do energii 90 milionów eV . W oparciu o to urządzenie i istniejący potencjał krakowscy naukowcy zamierzają rozwinąć interesujące prace w mało zbadanym dotychczas obszarze energii cząstek, jak również nowe kierunki zastosowań metod fizyki jądrowej dla potrzeb energetyki jądrowej, biologii, medycyny i rolnictwa.

J. H



Ogólny widok Instytutu Fizyki Jądrowej w Krakowie

Widok dużego cyklotronu



Spektrometr promieniowania beta
zbudowany w Instytucie Fizyki
Jądrowej w Krakowie

Jak zidentyfikowaliśmy hiperjądro podwójne

Prof. dr Janusz ZAKRZEWSKI



Fizyka hiperjąder jest dziś obszerną dziedziną wiedzy. Dwa najpoważniejsze odkrycia miały miejsce w Uniwersytecie Warszawskim: w 1952 r. zaobserwowano pierwsze hiperjądro, a w 10 lat później hiperjądro podwójne. Redakcja Deltę zwróciła się z prośbą do jednego z współodkrywców podwójnego hiperjądra aby opowiedział Czytelnikom jak do tego doszło. Sądzymy, że tego typu artykuły przedstawiające drogę do wyniku, a nie tylko sam wynik, pozwalają lepiej zrozumieć specyfikę pracy badawczej. Prosimy bardzo o propozycje, jakimi osiągnięciami współczesnej fizyki zająć się w tym dziale. Dołożymy starań aby spełnić życzenia Czytelników również w zakresie osiągnięć dokonanych przez fizyków zagranicznych.

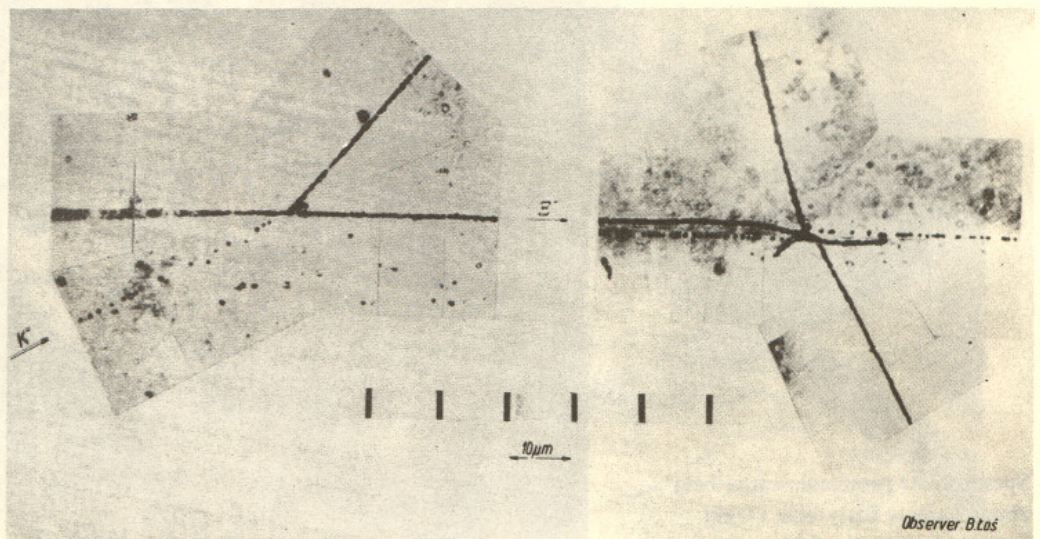
Chciałbym opisać tu historię odkrycia hiperjądra podwójnego dokonanego w roku 1962 w Warszawie. Interesowała nas wtedy sprawa tworzenia hiperjąder w oddziaływaniach szybkich mezonów K^- , przy czym zespół nasz, złożony z Mariana Danysza, Krystyny Garbowskiej, Jerzego Pniewskiego, Tadeusza Pniewskiego i autora niniejszego artykułu, uczestniczył w badaniach prowadzonych w ramach szeroko zakrojonej współpracy międzynarodowej, zwanej Europejską Współpracą K^- . Oprócz ośrodka warszawskiego, w owym czasie w skład tej grupy wchodziły zespoły z Bristolu, Brukseli, Dublinu, Genewy i Londynu — w sumie 23 fizyków i co najmniej drugie tyle techników. Badania w zakresie fizyki wielkich energii i cząstek elementarnych — to praca naprawdę zespołowa, wymagająca udziału wielu specjalistów z różnych ośrodków, oparta na wykorzystaniu akceleratorów dostarczających cząstek wielkich energii. Badania w tej dziedzinie muszą mieć charakter międzynarodowy.

W tym wypadku wykorzystywaliśmy blok specjalnej emulsji fotograficznej złożony z wielu warstw, naświetlonych wiązką mezonów K^- uzyskaną z akceleratora protonowego w Międzynarodowym Ośrodku Badań Jądrowych pod Genewą. Warstwy te, po wywołaniu i utrwaleniu, podzielone między laboratoria Europejskiej Współpracy K^- , przeglądali przy użyciu mikroskopów technicy-mikroskopisci, notując zderzenia mezonów K^- z jądrami składników emulsji fotograficznej.

Ale zanim przejdziemy do właściwej historii, trzeba koniecznie wprowadzić kilka pojęć z fizyki cząstek elementarnych — bez takiego wstępu trudno byłoby niespecjaliście śledzić dalszy tok rozumowania.

Badając własności obecnie znanych cząstek elementarnych wykryto pewną ich cechę, która może być jednym z elementów, dzięki którym cząstki różnią się między sobą. Ponieważ nie można znaleźć żadnego prostego obrazu fizycznego tej cechy — nazwano ją dziwnością. Istnieje wiele procesów, w których ta cecha jest ściśle zachowana i jedynie przekazywana od jednej cząstki do innej, co pociąga za sobą zmianę indywidualności tej cząstki. Cząstkom o różnych dziwnościach możemy przypisać liczby 0, ± 1 , ± 2 , ... itd. Nukleonom, tj. protonowi i neutronowi, przypisana jest dziwność zero, podobnie mezonom π (pionom), natomiast mezonowi K^- (K minus) i hiperonowi lambda — dziwność minus jeden. Wszystkie hiperony mają dziwność ujemną, przy czym np. hiperon Λ minus ma dziwność minus dwa. Nukleony i piony uważane są za cząstki niedziwne; jądra atomowe, których składnikami są nukleony, stanowią więc struktury niedziwne.

By cechę „dziwność” uczynić mniej zagadkową, dodajmy, że w procesach, zwanych przez fizyków szybkimi, do których należy np. zderzenie, dziwność jest zachowana. W wyniku zatem zderzenia cząstek niedziwnych nie może powstać cząstka dziwna, chyba, że wraz z nią powstanie cząstka o dziwności przeciwnej, t.j. naraz dwie cząstki o dziwnościach dodatniej i ujemnej. Przy oddziaływaniu jądrowym mezonu K^- z nukleonem, jego dziwność ujemna zatem nie ginie, lecz jest przekazywana powstającej z nukleonu cząstce dziwnej — hiperonowi.



Próba rozwiązania zadania M2.
Liczby p i q spełniają warunki zadania (na podstawie wzorów Viety) wtedy i tylko wtedy, gdy $p+q = -p$ i $pq = q$. Rozwiązując ten układ równań otrzymujemy $p = 0$ i $q = 0$ lub $p = 1$ i $q = -2$. Rozwiązanie to zawiera jednak błąd logiczny — zob. str. 2.



Rozwiązanie zadania M1.

Załóżmy, że $b^2 = a(a+c)$.

Uwzględniając tę równość mamy ze wzoru cosinusów

$$\cos A = \frac{b^2 + c^2 - a^2}{2bc} = \frac{a^2 + ac + c^2 - a^2}{2bc} = \frac{a+c}{2b}$$

$$\cos B = \frac{a^2 + c^2 - b^2}{2ac} = \frac{a^2 + c^2 - a^2 - ac}{2ac} = \frac{c(c-a)}{2ac} = \frac{c-a}{2a}$$

oraz ze wzoru na cosinus podwojonego kąta

$$\cos 2A = 2\cos^2 A - 1 = \frac{a^2 + 2ac + c^2 - 2b^2}{2b^2} = \frac{a^2 + 2ac + c^2 - 2a^2 - 2ac}{2a(a+c)} = \frac{c^2 - a^2}{2a(a+c)} = \frac{(c-a) \cdot (c+a)}{2a(a+c)} = \frac{c-a}{2a}$$

Mamy więc równość $\cos 2A = \cos B$, skąd wobec nierówności $0 < 2A < 2\pi$, $0 < B < \pi$ wynika, że $B = 2A$ lub $2\pi - 2A = B$. Jeżeli $2\pi - 2A = B$, to $2\pi - A + C = A + B + C = \pi$, skąd $A = C + \pi > \pi$, co jest niemożliwe. Jest więc $B = 2A$.

Jeżeli w trójkącie zachodzi równość $B = 2A$, to na podstawie twierdzenia sinusów i cosinusów

$$\frac{b}{a} = \frac{\sin B}{\sin A} = \frac{\sin 2A}{\sin A} = \frac{2 \sin A \cos A}{\sin A} = 2 \cos A = \frac{b^2 + c^2 - a^2}{bc}$$

Z równości tej wynika, że $b^2 c = ab^2 + ac^2 - a^3$

skąd $b^2(c-a) = a(c^2 - a^2)$ czyli

$b^2(c-a) = a(c-a) \cdot (c+a)$

skąd $c-a = 0$ lub $b^2 = a(c+a)$.

Jeżeli $c-a = 0$, to $C = A$, a ponieważ

$B = 2A$, więc $\pi = A+B+C = A+2A+A$,

skąd $C = A = \frac{\pi}{4}$, $B = \frac{\pi}{2}$.

Wówczas na mocy twierdzenia Pitagorasa

i równości $a = c$ mamy $b^2 = a^2 + c^2 = a(a+c)$.

W każdym więc przypadku otrzymamy tę równość.

Jednym z takich hiperonów jest cząstka lambda, bardzo dla nas ważna, gdyż może ona w pewnych warunkach stać się składnikiem jądra atomowego. Jądro takie ma jednostkową dziwność ujemną, jest więc samo jądrem dziwnym i nazywa się hiperjądrem. Możliwość wiązania hiperonu lambda w materii jądrowej odkryta została w roku 1952, w Uniwersytecie Warszawskim, przez Mariana Danysza i Jerzego Pniewskiego. Zapoczątkowali oni w ten sposób nowy dział fizyki — fizykę hiperjadr, która stała się odtąd specjalnością polską. Prace w tej dziedzinie prowadzone są w Warszawie do dnia dzisiejszego, nie tylko przy użyciu emulsji fotograficznej, ale również metodami licznikowymi.

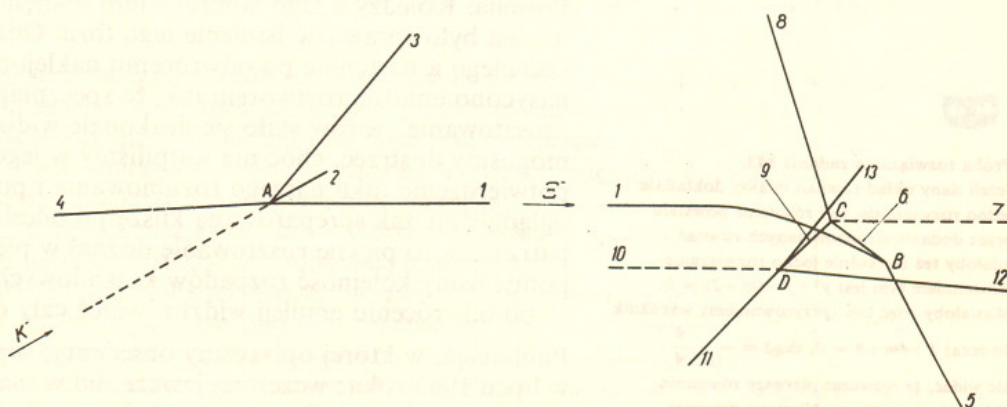
Z powyższych rozważań wynika ważny wniosek: w oddziaływaniach mezonów K^- mogą powstawać hiperjadra, skoro tworzone są hiperony lambda mogące ulec związaniu we fragmentach jądrowych. Gdy w grę wchodzi szybkie mezony K^- , na sto takich oddziaływań obserwuje się około czterech przypadków tworzenia hiperjadr. Jakie są ich dalsze losy? Otóż hiperjadra są nietrwałe, mogą ulegać rozpadowi na kilka fragmentów, mogą między innymi emitować mezony pi (jest to swoista „promieniotwórczość pionowa”). W rozpadzie takim, który jest procesem uważanym za powolny, dziwność ginie: jądro dziwne — hiperjądro — zamienia się na cząstki niedziwne, gdyż mezon pi i inne fragmenty jądrowe mają dziwność równą zero.

Czy mogą powstawać jadra o dziwności podwójnej, zawierające nie jeden, a dwa hiperony lambda? Pytanie to postawili sobie fizycy wkrótce po odkryciu hiperjadr pojedynczych, zawierających jeden hiperon lambda związany w strukturze jądrowej. Takie hiperjadra, zawierające dwa związane hiperony, nazywamy hiperjadrami podwójnymi o dziwności minus dwa. Można ich było oczekiwać właśnie w oddziaływaniach szybkich mezonów K^- , kiedy duża energia tych mezonów pozwala uzyskać dodatkową dziwność ujemną wraz z dziwnością dodatnią.

Jak już wspomniałem, przeglądu emulsji dokonywali mikroskopisci; przypadki znalezione przez nich analizowaliśmy następnie my, fizycy, klasyfikując je i kierując do pomiarów. W ten sposób zobaczyliśmy po raz pierwszy w zimie 1962 roku przypadek, zanotowany przez mikroskopistkę, Barbarę Łoś, który od pierwszego rzutu oka pobudził naszą wyobraźnię. Jego niezwykłość polegała na tym, że z bardzo małej objętości emulsji, do której prowadził ślad zatrzymującej się cząstki, wytworzonej w oddziaływaniu jądrowym mezonu K^- , wybiegały, oprócz innych cząstek, ślady dwóch mezonów pi! Interpretacja nasuwała się natychmiastowo, jak w olśnieniu, choć wymagała wielkiej wyobraźni: mógł to być dwukrotny rozpad z emisją mezonu pi hiperjadra podwójnego, wytworzonego przez hiperon ksi minus, powstający w oddziaływaniu jądrowym mezonu K^- (wraz z mezonem K^0 o dziwności dodatniej). Hiperon ten mając dziwność równą minus dwa, w oddziaływaniu z jądrem zwyczajnym tworzy dwa hiperony lambda. Jeśli hiperony te zostaną związane w tym samym fragmencie jądrowym — może powstać hiperjądro podwójne. Ale wtedy, gdy zobaczyliśmy nasz przypadek, to wszystko było czystą spekulacją — tak mogło być, ale nie musiało. Choć interpretacja podsunęta przez wyobraźnię wydawała się nam bardzo atrakcyjna, wymagała jednak potwierdzenia na drodze pomiarów i dokładnej analizy.

I tak zaczął się kilkumiesięczny okres, w którym zespół nasz starał się sprawdzić proponowaną interpretację, dokonując niezliczonych obserwacji i pomiarów. Każdy wynik przyjmowany był niezmiernie krytycznie i po wielokroć sprawdzany.

Mikrofotografia oraz rysunek schematyczny przypadku tworzenia i kaskadowego rozpadu hiperjadra podwójnego. Mezon K^- zderzył się z jądrem emulsji fotograficznej (gwiazda A) wytwarzając hiperon ksi (tor 1). Hiperon ten został pochłonięty przez jądro (gwiazda B) z emisją hiperberylu dwulambdowego (tor 6) rozpadającego się (gwiazda C) z emisją mezonu pi (tor 7) i hiperberylu jednolambdowego (tor 9). Rozpada się on również z emisją mezonu pi (tor 10) i innych cząstek naładowanych (tory 11, 12 i 13).



Hiperon ksi powstaje w wyniku przekształcenia się jednego z nukleonów jądra.



Próba rozwiązania zadania M3. Jeżeli dany układ równań miałby dokładnie jedno rozwiązanie, to i równanie powstałe przez dodanie stronami danych równań miałoby też dokładnie jedno rozwiązanie. Równaniem tym jest $y^2 + y - (m+2) = 0$. Musiałoby więc być (przyrównujemy wyróżnik do zera) $1 + 4m + 8 = 0$, skąd $m = -\frac{9}{4}$, ale widać, że wówczas pierwsze równanie układu jest sprzeczne. Niestety, pierwsze zdanie tego rozwiązania zawiera błąd logiczny. (zob. str. 16).

Do późnej nocy, dzień w dzień, wpatrywaliśmy się kolejno w mikroskop starając się dociec, z której spośród trzech gwiazd B , C i D (patrz rysunek) wybiegały tory cząstek oraz wielokrotnie mierząc zasięgi tych cząstek i kąty ich emisji. Zadanie nie było łatwe, gdyż gwiazdy te leżały w objętości emulsji nie przekraczającej trzech stumilionowych części milimetra sześciennego ($\sim 3 \cdot 10^{-8} \text{ mm}^3$). Prawdopodobieństwo przypadkowego nałożenia się obserwowanych trzech gwiazd w tak małej objętości emulsji oceniliśmy jako znacznie mniejsze od jednej dziesięciomiliardowej ($\ll 10^{-10}$). Po ustaleniu pochodzenia wszystkich torów, w szczególności stwierdzeniu, że jeden z mezonów π pochodzi z gwiazdy C (tor 7) a drugi — z gwiazdy D (tor 10), ciąg zdarzeń przedstawiał się następująco. Mezon K^- , początkujący cały proces, zderzył się z jądrem emulsji (gwiazda A) wytwarzając hiperon ksi (tor 1). Hiperon ten został pochłonięty przez jądro, prawdopodobnie węgla (gwiazda B), z emisją hiperjądra podwójnego (tor 6) rozpadającego się kaskadowo (gwiazda C) z emisją mezonu π (tor 7) i hiperjądra pojedynczego (tor 9). To ostatnie rozpadało się również z emisją mezonu π (tor 10) i innych cząstek naładowanych (tory 11, 12 i 13). Szczegółowa analiza takiego ciągu zdarzeń doprowadziła do wniosku, iż hiperjądrem podwójnym musiał być dwulambdowy beryl dziesięć (z mniejszym prawdopodobieństwem beryl jedenaście). Rozważyliśmy też szereg innych interpretacji konkurencyjnych, które mogłyby wyjaśnić obserwowany ciąg zdarzeń — gwiazdy A , B , C i D — odrzucając je w końcu jako niespójne z obserwacjami.

I tak pozostała w końcu ta pierwsza interpretacja, intuicyjna — tylko teraz potwierdzona przez pomiar i rygorystyczne rozumowanie.

Zdarzyła się rzecz przez nikogo nie przeczuwana: w Warszawie zidentyfikowano pierwsze hiperjądro podwójne, podobnie jak dziesięć lat wcześniej — również w Warszawie! — pierwsze hiperjądro pojedyncze. Jak się okazało w dalszych badaniach, prawdopodobieństwo obserwacji takiego przypadku w naszych warunkach eksperymentalnych jest bardzo małe, mniejsze niż jedno na milion oddziaływań mezonów K^- ($< 10^{-6}$). Do chwili obecnej znany jest jeszcze drugi podobny przypadek hiperjądra podwójnego, znaleziony w Stanach Zjednoczonych w kilka lat po naszym odkryciu, potwierdzający w całej pełni nasze wnioski. Nawiasem mówiąc, dopiero wtedy odetchnęliśmy z ulgą...

Kiedy już upewniliśmy się w interpretacji przypadku, wiosną 1963 roku zawiadomiliśmy naszych współpracowników z Europejskiej Współpracy K^- o odkryciu. Odbyliśmy szereg rozmów telefonicznych, wysłaliśmy listy przekazując naszym kolegom wszystkie dane pomiarowe i prosząc o sprawdzenie naszych rachunków i rozumowania. W tym czasie nie mieliśmy jeszcze do dyspozycji w Warszawie elektronicznych maszyn cyfrowych, wszystkie rachunki przeprowadzaliśmy za pomocą zwykłych maszyn do liczenia, co zajęło bardzo wiele godzin. Nasi koledzy zagranicą mieli już dostęp do maszyn cyfrowych i sprawdzili przy ich użyciu w dwóch ośrodkach nasze rachunki, uzyskując te same wyniki.

Jednym z kłopotów przy interpretacji całego zdarzenia była trudność doszukania się krótkiego toru, kryjącego się częściowo pod innym torem, właściwie niemożliwego do zaobserwowania wobec znajdującego się przypadkowo pod nim znaczka współrzędnych umieszczonego na spodzie emulsji tuż przy szkle (ciemne tło na mikrofotografii). Mimo to nie wątpiliśmy w jego istnienie, skoro jego br czyniłby całe zdarzenie zupełnie niezrozumiałym. W tym czasie ośrodkiem o największej sławie w zakresie metod emulsyjnych było laboratorium Uniwersytetu Bristolskiego, kierowane przez laureata Nagrody Nobla profesora Cecila F. Powella. Koledzy z tego laboratorium podjęli się tak spreparować emulsję, by można było sprawdzić istnienie tego toru. Odklejono emulsję od podkładu szklanego a następnie po odwróceniu naklejono powtórnie. Starto znaczek oraz nasycono emulsję roztworem tak, że spęczniała wielokrotnie i wtedy całe „rusztowanie” torów stało się doskonale widoczne. Pojawił się tor, którego nie mogliśmy dostrzec, choć nie wątpiliśmy w jego istnienie. Było to piękne potwierdzenie toku naszego rozumowania i poprawności interpretacji. Kiedy oglądaliśmy tak spreparowaną kliszę po odesłaniu do Warszawy, jeden z nas patrząc na to piękne rusztowanie doznał w pierwszej chwili szoku, że jednak pomyliliśmy kolejność rozpadów kaskadowych. Zapomniał zapewne z wrażenia, że po odwróceniu emulsji widzi również cały obraz odwrotnie...

Publikacja, w której opisaliśmy obserwację hiperjądra podwójnego, ukazała się w lipcu 1963 roku; wcześniej jeszcze, bo w marcu tegoż roku, referowaliśmy naszą pracę na dwóch konferencjach międzynarodowych, w Genewie i St. Cergue, poświęconych fizyce struktur jądrowych i hiperjader.

Na czym polegało znaczenie naszego odkrycia? Potwierdziło ono możliwość wiązania dwóch hiperonów Λ w materii jądrowej i umożliwiło pierwszą ocenę wzajemnego oddziaływania, na innej drodze nieosiągalną. Odkrycie to pokazało zarazem, że hiperjądra podwójne można wytwarzać na drodze wychwyty jądrowego hiperonów Λ . Należy się spodziewać, że fizyka hiperjader podwójnych rozwinie się dopiero po skonstruowaniu wiązek takich hiperonów (co nie jest rzeczą łatwą ze względu na krótki czas ich życia). Poszukiwanie hiperjader podwójnych wytwarzanych w oddziaływaniach szybkich mezonów K^- jest zbyt pracochłonne. A że znaleźliśmy tak rzadki przypadek? Może mieliśmy szczęście...

Krótki kurs informatyki Algorytmy cz. I

dr Andrzej SKOWRON

W rozważaniach naszych nie będziemy chwilowo dążyć do ścisłej definicji algorytmu ani do formalizacji zapisu algorytmów. Celem naszym będzie wyrobienie u Czytelnika przekonania, że algorytm jest uściśleniem przepisu postępowania prowadzącego do zamierzonego celu.

Spróbujemy wymienić kilka charakterystycznych cech (nieformalnego) pojęcia algorytmu.

- 1) Algorytm określa się przez podanie przepisu postępowania, pamięci oraz sterowania.
- 2) Pamięć określa się jako zbiór, którego elementy nazywamy stanami.
- 3) Sterowanie (którego rolę może pełnić człowiek lub urządzenie) umożliwia przeprowadzenie obliczeń, a mianowicie w oparciu o przepis postępowania wyznacza jednoznacznie czynność, którą należy wykonać na aktualnym stanie pamięci oraz określa jednoznacznie następny stan pamięci i następną czynność do wykonania (jeśli jeszcze jakąś czynność należy wykonać).
- 4) Dla każdego algorytmu określony jest sposób wprowadzania do pamięci elementów zbioru nazywanego zbiorem danych oraz sposób odczytywania wyników z pamięci.

Aby wyjaśnić bliżej sens powyższych sformułowań rozważmy dobrze znany algorytm Euklidesa, który pozwala dla dowolnych liczb naturalnych a i b wyznaczyć ich największy wspólny dzielnik oznaczany przez NWD (a, b). Jeśli chcemy wyznaczyć NWD (a, b), gdzie a, b są liczbami naturalnymi, postępujemy według następujących reguł:

- I Jeśli $a = b$, to NWD (a, b) = a .
- II Chcąc znaleźć NWD (a, b) gdy $a > b$, dzielimy a przez b i wyznaczamy resztę r z tego dzielenia. Z kolei dzielimy b przez r i wyznaczamy nową resztę r_1 . Następnie dzielimy r przez r_1 i wyznaczamy nową resztę r_2 itd., aż dojdziemy do reszty 0. Ostatnia różna od zera reszta będzie największym wspólnym dzielnikiem liczb a i b .
- III Jeśli $a < b$, to postępujemy jak w II zamieniając rolami a i b . Przyjmujemy, że liczby naturalne będą przedstawiane w zapisie dziesiętnym i jeśli nie zajdzie potrzeba, nie będziemy odróżniać liczby naturalnej od jej zapisu dziesiętnego.

Okaże się niżej, że wygodnie jest wyróżnić nazwy miejsc, w których zapisujemy liczby naturalne. Miejsca, w których będą zapisywane liczby naturalne będziemy nazywać x, y, z, u . Przez x, y, z, u będziemy oznaczać liczby naturalne zapisane odpowiednio w miejscach o nazwie x, y, z, u . Powiemy, że x jest zawartością x , y jest zawartością y , itd. Jeśli np. w miejscu o nazwie x nie zapisano żadnej liczby, to przyjmujemy, że zawartość x jest pusta i będziemy umownie przyjmować jako zawartość x symbol $-$.

Stanami pamięci naszego algorytmu będą ciągi (x, y, z, u) , a pamięć jest zbiorem wszystkich takich ciągów.

Możemy teraz bardziej szczegółowo rozpisać sposób postępowania prowadzący do wyznaczenia największego wspólnego dzielnika dwu dowolnych liczb naturalnych. Czynności, które należy wykonać (poczynając od oznaczonej numerem 1) są następujące:

1. Umieść liczbę naturalną a w miejscu o nazwie x (symbolicznie tę czynność oznaczmy $a \rightarrow x$). Jako następną wykonaj czynność oznaczoną numerem 2.
2. Umieść liczbę naturalną b w miejscu o nazwie y (symbolicznie tę czynność oznaczmy $b \rightarrow y$). Jako następną wykonaj czynność oznaczoną numerem 3.
3. Sprawdź, czy zawartość x jest większa od zawartości y (symbolicznie tę czynność oznaczmy $x > y$). Jeśli tak, to jako następną wykonaj czynność oznaczoną numerem 5; jeśli nie — to 4.
4. Umieść zawartość x w miejscu o nazwie u (symbolicznie tę czynność zapisujemy $x \rightarrow u$). Zakładamy, że po wykonaniu tej czynności zawartość x nie zmieni się. Jako następną wykonaj czynność oznaczoną numerem 7.
5. $y \rightarrow u$. Jako następną wykonaj czynność oznaczoną numerem 10.

Wykonując czynności 1, 2 wprowadzamy dane, które będą przekształcone.

Badamy, z którym z przypadków I, II, III mamy do czynienia.

Po wykonaniu tej czynności wynik jest zapisany w miejscu o nazwie u.

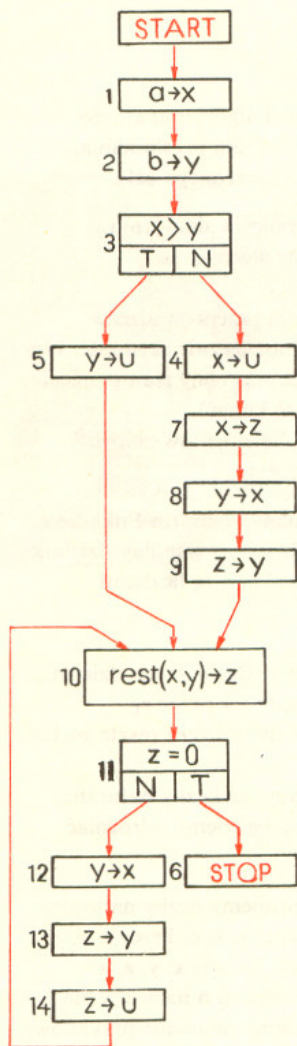
Te czynności wykonujemy, gdy $x \leq y$ (patrz 4).

Po ich wykonaniu zawartością x jest y, a zawartością y jest x.

Sprawdzenie, czy już otrzymaliśmy jako kolejną resztę 0.

Czynności przygotowujące do nowego dzielenia.

Jako dzielną należy wziąć poprzedni dzielnik, a jako dzielnik poprzednią resztę, a więc zawartość z.



Sieć działań dla algorytmu Euklidesa

6. Nie zmieniamy zawartości żadnego z miejsc. Następna czynność nie jest określona. Symbolicznie tę czynność zapisujemy STOP.
7. $x \rightarrow z$. Jako następną wykonaj czynność oznaczoną numerem 8.
8. $y \rightarrow x$. Jako następną wykonaj czynność oznaczoną numerem 9.
9. $z \rightarrow y$. Jako następną wykonaj czynność oznaczoną numerem 10.
10. Umieść resztę z dzielenia zawartości x przez zawartość y w miejscu o nazwie z (symbolicznie tę czynność zapisujemy $\text{rest}(x, y) \rightarrow z$). Jako następną wykonaj czynność oznaczoną numerem 11.
11. Sprawdź, czy zawartość z jest równa 0 (symbolicznie tę czynność zapisujemy $z = 0$). Jeśli tak, to jako następną wykonaj czynność oznaczoną numerem 6; jeśli nie — to 12.
12. $y \rightarrow x$. Jako następną wykonaj czynność oznaczoną numerem 13.
13. $z \rightarrow y$. Jako następną wykonaj czynność oznaczoną numerem 14.
14. $z \rightarrow u$. Jako następną wykonaj czynność oznaczoną numerem 10.

Często posługujemy się graficznym przedstawieniem wyżej opisanego sposobu postępowania. Otrzymany rysunek nazywamy siecią działań (flow-diagramem) rozważanego algorytmu. Na rysunku podajemy sieć działań dla naszego algorytmu. Czytelnik łatwo zauważy, jak otrzymaliśmy ten rysunek na podstawie powyższych rozważań.

W tabelce wskazujemy, jak będą się zmieniały zawartości x, y, z, u i jakie będziemy wykonywali czynności, gdy $a = 124$, $b = 36$ i rozpoczynamy od wykonania czynności oznaczonej numerem 1.

Numer wykonywanej czynności	Zawartości				Numer czynności następnej
	x	y	z	u	
	(po wykonaniu kolejnej czynności)				
1	124	—	—	—	2
2	124	36	—	—	3
3	124	36	—	—	5
5	124	36	—	36	10
10	124	36	16	36	11
11	124	36	16	36	12
12	36	36	16	36	13
13	36	16	16	36	14
14	36	16	16	16	10
10	36	16	4	16	11
11	36	16	4	16	12
12	16	16	4	16	13
13	16	4	4	16	14
14	16	4	4	4	10
10	16	4	0	4	11
11	16	4	0	4	6
6	16	4	0	4	*

* oznacza, że wykonaliśmy ostatnią czynność przewidzianą siecią działań

W oparciu o naszą sieć działań przekształciliśmy zawartości x, y, z, u w nowe zawartości x, y, z, u i wyznaczyliśmy numer następnej czynności. Rolę sterowania algorytmu pełni w naszym przypadku człowiek. Ciąg, którego kolejne wyrazy są kolejnymi wierszami w tabeli (po usunięciu z niej ostatniej kolumny) jest przykładem obliczenia naszego algorytmu. Podaj inne przykłady obliczeń tego algorytmu!

Liczba NWD (124, 36) jest umieszczona w miejscu o nazwie u po zakończeniu podanego wyżej obliczenia. Widać, że jeśli dla dowolnych liczb naturalnych a, b wykonamy czynności zgodnie z naszą siecią działań (zaczynając od czynności oznaczonej numerem 1 i kończąc wykonywanie czynności po wykonaniu 6) otrzymamy NWD (a, b) jako zawartość u. Tę własność rozważanego algorytmu nazywamy masowością.

Zauważmy jeszcze, że jakkolwiek czynność wykonujemy i jakiegokolwiek są zawartości x, y, z, u, to następna czynność jest wyznaczona jednoznacznie (jeśli jest określona) oraz następny stan pamięci też jest wyznaczony jednoznacznie. Mówimy, że rozważany algorytm jest jednoznaczny (deterministyczny).

Na zakończenie proponujemy Czytelnikowi wykonanie następujących ćwiczeń:

- Ćwiczenie 1. Podaj sieć działań dla wybranej metody rozwiązywania układu trzech równań liniowych o trzech niewiadomych przy założeniu, że współczynniki i wyrazy wolne w tych równaniach są liczbami naturalnymi.
- Ćwiczenie 2. Podaj sieć działań dla wybranej metody wyszukiwania książek w katalogu.
- Ćwiczenie 3. Podaj sieć działań na wykonanie tortu świątecznego.
- Ćwiczenie 4. Czy w każdym z ćwiczeń 1, 2, 3 można dla danej metody podać tylko jedną sieć działań? Prosimy o listy z propozycjami rozwiązań.

Nasi na Olimpiadzie



Zadania olimpijskie

XV Międzynarodowa Olimpiada Matematyczna
Moskwa 9 lipca 1973 r.

1. Punkt O leży na prostej l ; $\overrightarrow{OP_1}, \overrightarrow{OP_2}, \dots, \overrightarrow{OP_n}$ są wektorami o długości l , przy czym wszystkie punkty P leżą na płaszczyźnie zawierającej prostą l i po jednej stronie tej prostej.

Udowodnić, że jeżeli n jest liczbą nieparzystą, to $|\overrightarrow{OP_1} + \overrightarrow{OP_2} + \dots + \overrightarrow{OP_n}| > l$, gdzie $|\overrightarrow{OM}|$ — długość wektora $\overrightarrow{OM}$.

2. Rozstrzygnąć, czy istnieje skończony zbiór M złożony z punktów przestrzeni, nie zawarty w jednej płaszczyźnie i taki, że dla każdego dwóch punktów $A, B \in M$ istnieją takie dwa inne punkty $C, D \in M$, że proste AB i CD są równoległe i różne.

3. Znaleźć minimalną wartość sumy $a^2 + b^2$, gdzie a i b są liczbami rzeczywistymi, dla których równanie $x^4 + ax^3 + bx^2 + ax + 1 = 0$ ma co najmniej jeden pierwiastek rzeczywisty. Na rozwiązywanie zadań przeznaczono 4 godziny.

10 lipca 1973 r.

4. Saper ma sprawdzić, czy na polu o kształcie trójkąta równobocznego (łącznie z brzegiem) znajdują się miny. Zasięg działania jego aparatu wykrywającego jest równy połowie wysokości trójkąta. Saper wyrusza z jednego jego wierzchołka. Jaką drogę powinien wybrać, by była ona najkrótsza i by zbadał całe pole?

5. Niepusty zbiór G funkcji zmiennej rzeczywistej x postaci $f(x) = ax + b$, gdzie a i b są liczbami rzeczywistymi, $a \neq 0$ spełnia następujące warunki:

1) jeżeli $f, g \in G$, to $g \cdot f \in G$ gdzie $(g \cdot f)(x) = g(f(x))$

(zbiór jest domknięty ze względu na superpozycję czyli składanie funkcji),

2) jeżeli $f \in G$, gdzie $f(x) = ax + b$, to funkcja odwrotna $f^{-1} \in G$,

gdzie $f^{-1}(x) = \frac{x-b}{a}$,

3) dla każdej funkcji $f \in G$ istnieje takie x_f , że $f(x_f) = x_f$. Udowodnić, że istnieje takie k , że $f(k) = k$ dla każdej funkcji $f \in G$.

6. Niech $a_1, a_2, \dots, a_n$ będą danymi liczbami dodatnimi i q daną liczbą rzeczywistą, dla której zachodzi nierówność $0 < q < 1$. Znaleźć liczby rzeczywiste $b_1, b_2, \dots, b_n$ spełniające warunki:

a) $a_k < b_k$ dla każdego k od 1 do n ,

b) $q < \frac{b_{k+1}}{b_k} < \frac{1+q}{q}$ dla każdego k od 1 do $n-1$,

c) $b_1 + b_2 + \dots + b_n < \frac{1+q}{1-q} (a_1 + a_2 + \dots + a_n)$.

Na rozwiązywanie zadań przeznaczono 4 godziny.

WYNIKI

	Liczba nagród			Liczba punktów
	I	II	III	
ZSRR	3	2	3	254
Węgry	1	2	5	215
NRD	—	3	4	188
Polska	—	2	4	174
Wielka Brytania	1	—	5	164
Francja	—	3	1	153
Czechosłowacja	—	1	4	149
Austria	—	—	6	144
Rumunia	—	1	3	141
Jugosławia	—	—	5	137
Szwecja	—	1	1	99
Holandia	—	—	2	96
Bulgaria	—	—	1	96
Finlandia	—	—	2	86
Mongolia	—	—	1	64
Kuba	—	—	1	42

Reprezentacja Kuby była mniej liczna — 5 osób, wobec 8 reprezentantów innych krajów.

Redakcja zwróciła się z prośbą o opisanie swoich wrażeń z Olimpiady do jednego z jej uczestników — Piotra Bermiana. Oto jego relacja.

Miedzy 7 i 16 lipca br. braliśmy udział w XV Międzynarodowej Olimpiadzie Matematycznej w Moskwie. Oprócz drużyny polskiej było jeszcze 15 innych drużyn, w tym dwie uczestniczące po raz pierwszy — Francja i Finlandia.

Gdy w sobotę wieczorem przylecieliśmy do Moskwy, zostaliśmy zawiezieni do Hotelu Uniwersyteckiego, nowoczesnego wieżowca opodal zabudowań Uniwersytetu im. Łomonosowa (opodal tj. trochę ponad kilometr, w Moskwie oznacza to bardzo blisko). Pokoje dostaliśmy dwuosobowe, dość obszerne, z łazienkami. W ogóle locum było bardzo ładne, a sąsiedztwo przedstawicieli kilkunastu narodów w jednym budynku sprzyjało ożywionym kontaktom międzynarodowym.

Nazajutrz, czyli w niedzielę, odbyła się przejażdżka po mieście. Moskwa sprawia wrażenie swoim ogromem — o ile w Warszawie przebycie dziesięciu kilometrów w linii prostej z reguły wyprowadza na wolną przestrzeń, w Moskwie można przejechać całe dziesiątki kilometrów wśród wysokiej zabudowy.

W poniedziałek rozpoczęło się rozwiązywanie zadań. Zaraz po śniadaniu załadowaliśmy się do autobusów i pojechaliśmy z hotelu do szkoły, gdzie odbywały się zawody. Po wysłuchaniu krótkiego przemówienia, w którym członek radzieckiej Akademii Nauk Pedagogicznych prof. Markuszewicz dodał nam otuchy, twierdząc, że wszyscy wykonali już jedno zadanie — dojechali do miejsca (a trzeba przyznać, że rozwiązanie było znowu nie tak krótkie, a droga do wyniku wcale kręta), zasiedliśmy w ośmiu salach, gdzie czekały na nas koperty z tekstami zadań w ojczystym języku.

Pierwszego dnia zdecydowanie najtrudniejsze było zadanie trzecie. Żaden pomysł nie daje w nim natychmiastowego rozwiązania, trzeba się w nim było wykazać także sprawnością w liczeniu. Natomiast zadanie drugie było łatwe, ale zawierało „mały haczyk”: niektórzy dowodzili, że zbiór M nie istnieje, a wśród nich był m.in. Paweł Kröger, gwiazda poprzedniej Olimpiady.

Następnego dnia sytuacja była podobna. Też było zadanie zdecydowanie najtrudniejsze — szóste, gdzie trzeba było wpaść na jeden, ale za to skomplikowany pomysł — i zadanie z „małym haczykiem”. W tym przypadku (zadanie 4) haczyk polegał na tym, że trzeba było udowodnić m.in. pewien „oczywisty fakt” (że saper po przejściu znalezionej już optymalnej drogi rzeczywiście zbada całe pole). Zawodnicy polscy stracili na tym po 1—2 punkty.

Od momentu zakończenia pisania zadań bujnie rozkwitało międzynarodowe życie towarzyskie. Po pierwsze wypytywaliśmy się wzajemnie, ile zadań zrobiła każda drużyna (jak się okazało, niektórzy podawali zbyt dużo, za to Węgrzy byli przesadnie skromni) i jakie były rozwiązania. Nieco później zaczęliśmy rozmawiać także na inne tematy, a po paru dniach matematyka była wspomnianą sporadycznie.

Ciekawą stroną życia towarzyskiego były „imprezy sportowe”, jak mecz piłki nożnej Polska — Jugosławia czy zawody brydżowe. Smutne nieco, że co mogliśmy, tośmy przegrali. I tak wynik meczu Polska — Jugosławia wyniósł 2:4, w meczu brydża sportowego Wielka Brytania — Polska zajęliśmy drugie miejsce, a w zawodach brydżowych siedmiu par, pary polskie zajęły miejsca 5 i 7 (zwyciężyli Francuzi, ponadto udział w zawodach wzięły: jeszcze jedna para Francuzów, jedna Szwedów, jedna Holendrów i mieszana para angielsko-fińska). Mimo niewielkich sukcesów wyniosłem z tych spotkań bardzo miłe wspomnienia.

Najważniejszą pozycją programu były wycieczki. Poza wzmiankowaną już przejażdżką po Moskwie odbyła się wycieczka statkiem po rzece Moskwie, zwiedziliśmy Galerię Puszkina, byliśmy w posiadłości hrabiów Szeremetiewów w miejscowości Archangielskoje (do nich należało około połowy pięknej kolekcji obrazów impresjonistów z Galerii Puszkina), zwiedziliśmy zespoły klasztorów i cerkwi w Zagorsku (ikony Rublowa!) i Rostowie Wielkim oraz, z tłumaczką Gałą, zwiedziliśmy Kreml i Galerię Tretiakowską.

Najświetniejszym akcentem pobytu w Moskwie była uroczystość rozdania nagród. W Pałacu Pionierów po przemówieniach i oklaskach odbyło się wręczenie dyplomów, przerywane gęstymi owacjami. Wieczorem w hotelu odbyła się uroczysta kolacja. Grała muzyka, Walery Gusiew, nawspierał chybą śpiewak wódu matematyków (jeden z organizatorów) śpiewał romanse cygańskie, wszyscy tańczyli, szampa lał się strumieniami... powiedzmy strumyczkami. Nikt nie mógł sobie przypomnieć równie wesołego zakończenia Olimpiady.

Zadaliśmy też dwa pytania Przewodniczącemu Delegacji Polskiej mgr Andrzejowi Mąkowskiemu i jego zastępcy dr Maciejowi Bryńskiemu.

Co Panowie mogą powiedzieć o wynikach tegorocznej Międzynarodowej Olimpiady Matematycznej?

W Olimpiadzie uczestniczyło 125 uczniów z 16 krajów, w tym tylko jedna uczennica (z Holandii). Jury ogłosiło jak zwykle jedynie listę zdobywców nagród I, II i III stopnia. Spośród uczestniczących w Olimpiadzie naszych 8 uczniów nagrody II stopnia zdobyli Grzegorz Andrzejczak z XII LO w Łodzi i Piotr Berman z XIV LO w Warszawie, nagrody zaś III stopnia Wojciech Banaszczyk z XII LO w Łodzi, Maciej Lewenstein z XIV LO w Warszawie, Adam Smólski z VI LO w Warszawie, Jerzy Weyman z IV LO w Toruniu. Potwierdziła się w ten sposób zasada, że uczeń polski nie zdobywa I nagrody dwukrotnie. Gdyby zsumować liczby punktów zdobytych przez uczestników poszczególnych drużyn, to nasza ekipa znalazłaby się na wysokim, czwartym miejscu.

Jak Panowie oceniają przygotowanie naszych uczniów do Olimpiady?

Zadanie 5, w którym występują pojęcia matematyki współczesnej nie sprawiło naszym uczniom kłopotów, natomiast zadania bardziej tradycyjne, wymagające długich przekształceń i wyobraźni rachunkowej (takie jak zad. 3 i 6) wypadły w naszej drużynie słabo.

Na pytanie co to jest cząstka elementarna odpowiada

prof. dr Grzegorz BIAŁKOWSKI



Pytanie postawione w tytule można rozumieć dwojako: po pierwsze — co chcielibyśmy rozumieć przez cząstkę elementarną, a po drugie — co dziś obejmujemy tą nazwą. W pierwszym znaczeniu chodziłoby więc o podanie definicji, która by mogła nam dostarczyć kryterium rozpoznawania cząstek elementarnych, w drugim zaś o ustalenie stanu faktycznego, do którego doszło w wyniku wieloletnich badań teoretycznych i eksperymentalnych.

Pozornie mogłoby się wydawać, że odpowiedzi na oba pytania powinny się pokrywać. Od nas bowiem tylko zależy, jakie nazwy nadaje się obiektom fizycznym. Nie zawsze jednak tak jest. Najlepiej o tym świadczy przykład nazwy „atom”.

Jak wiadomo, pochodzi ona z języka greckiego (od słowa „atomos”, czyli „niepodzielny”) i została wprowadzona w V w. p.n.e. przez dwu filozofów greckich — Demokryta i jego nauczyciela — Leukippa. Pragnęli oni sobie odpowiedzieć na podstawowe pytanie nurtujące ludzkość od wieków a nawet tysiącleci, a mianowicie, jaka jest najgłębsza, najbardziej podstawowa struktura rzeczywistości. W owych czasach można było udzielić takiej odpowiedzi tylko na drodze spekulacji filozoficznych, gdyż nauki ścisłe były dopiero w powijkach. Leukippos i Demokryt założyli, że wszystko co jest, jest materialne (a więc i dusze!) i że materia nie jest podzielna nieograniczenie, to znaczy, że istnieją jakieś jej najdrobniejsze cząstki („atomy”), z których składa się cały wszechświat. Łączenie się i rozdzielanie atomów filozofowie ci uważali za istotę wszelkich zmian w przyrodzie. Ogromna większość późniejszych filozofów, a w tym takie autorytety jak Platon i Arystoteles, zwalczała hipotezę atomistyczną, która w ciągu wielu wieków nie znalazła należytego oddźwięku.

Do jej ponownego podjęcia przyczyniły się wyniki badań przyrodniczych, a szczególnie obserwacja, że poszczególne pierwiastki łączą się w związki chemiczne zawsze w pewnych ustalonych proporcjach. Fakt ten najłatwiej było wyjaśnić na gruncie hipotezy atomistycznej. Tak więc atom wkroczył raz jeszcze w początku XIX w. do historii myśli ludzkiej, tym razem jednak jako pojęcie podległe weryfikacji eksperymentalnej. W ciągu następnych dziesięcioleci wykryto wiele odmian atomów i uważano je za najbardziej podstawowe elementy materii, niepodzielne — zgodnie z nazwą wprowadzoną przez starożytnych myślicieli.

Jak jednak wiemy, atomy są podzielne. Świadczy o tym zjawisko jonizacji, a jeszcze dobitniej — promieniotwórczości. Nazwa „atom” przestała więc dobrze charakteryzować obiekty tą nazwą objęte. Mimo to jednak, wskutek wieloletniej tradycji nie zaniechano stosowania jej po dzień dzisiejszy. „Atom” Demokryta nie ma, jak stąd widać, wiele wspólnego z „atomem” współczesnej fizyki.

Nie zaginęła przez to jednak idea Demokryta. W poszukiwaniu tych właśnie podstawowych składników materii sformułowano pojęcie „cząstki elementarnej”. Początkowo, w latach 1910—1920 można było przypuszczać, że cały wszechświat zbudowany jest wyłącznie z trzech rodzajów cząstek: protonów, elektronów i fotonów. Protony i elektrony mogły tworzyć jądra atomowe (należałoby wziąć A protonów i $A-Z$ elektronów aby uzyskać jądro o liczbie masowej A i atomowej Z), inne elektrony krążyłyby wokół tych jąder, a fotony byłyby potrzebne jako cząstki („kwanty”) pola elektromagnetycznego zapewniającego istnienie sił między elektronami i protonami. Co więcej, zarówno protony jak i elektrony są cząstkami trwałymi („niepodzielnymi”), a więc pasowałyby do koncepcji Demokryta.

Rzeczywistość okazała się jednak znacznie bardziej skomplikowana. Po pierwsze, Dirac przewidział istnienie antycząstek i rzeczywiście antycząstki zostały wykryte, najpierw pozyton jako antycząstka elektronu, a następnie, już w latach powojennych, antyproton. Przekonano się, że elektron i pozyton spotkawszy się, mogą „anihilować”, to znaczy zmieniać się w pewną liczbę fotonów (najczęściej 2 lub 3). Z punktu widzenia koncepcji Demokryta nie jest to zrozumiałe, gdyż „atomy” nie powinny ginąć ani powstawać.

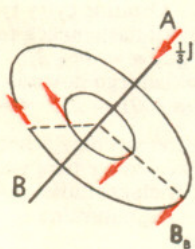
Drugim odkryciem podważającym tę koncepcję było odkrycie neutronu. Okazało się, że neutrony, bliźniaczo podobne do protonów, różnią się od nich tym, że nie są trwałe. W wyniku rozpadu przechodzą one w układ trzech cząstek — proton, elektron i neutrino. Poza tym jednym faktem neutrony są równie podstawowe i elementarne jak protony i trudno sobie wyobrazić, że jedno z nich są „atomami”

a drugie nie są. Dalsze badania, prowadzone począwszy od lat pięćdziesiątych, doprowadziły do wykrycia wielu innych cząstek, które są w ogromnej większości tworami nietrwałymi i to tak bardzo, że średnie czasy życia niektórych z nich (a nawet większości) są krótsze od 10^{-20} s! Obecnie znamy ponad 100 rodzajów tych cząstek. Ginąc, cząstki te zmieniają się w inne cząstki, czasem bardziej, czasem nawet mniej od nich trwałe. Ze względów historycznych wszystkie te obiekty nazywamy cząstkami elementarnymi. Czy jednak nazwa ta, podobnie jak to było w wypadku atomów, nie mija się z rzeczywistością? Odpowiedź na to pytanie nie jest łatwa.

Jakie bowiem kryterium „elementarności” należałoby przyjąć? Trwałość? Nie, bo są cząstki tak elementarne jak neutron, które są nietrwałe. Natomiast istnieją takie twory złożone jak na przykład wiele jąder atomowych, które są trwałe. Brak struktury wewnętrznej? Nie, bo na przykład proton nie jest na pewno cząstką bez struktury. Przedstawia się on, poglądowo mówiąc, jak chmura materii rozrzedzająca się ku brzegom.

Pozostaje więc jedyne możliwe kryterium: cząstki elementarne to te, które są niezbędne do wyjaśnienia własności wszystkich form materii, a więc i innych cząstek, same zaś nie są już przez nic wyjaśniane. Innymi słowy, kryterium to miałoby charakter raczej teoretyczny niż eksperymentalny. Gdybyśmy znali teorię, w której ze znajomości pewnej liczby rodzajów cząstek i ich własności moglibyśmy wydedukować istnienie i własności wszystkich innych cząstek, to mielibyśmy kryterium elementarności. Takiej teorii jednak nie ma, przynajmniej na razie. A priori rysują się trzy konkurencyjne możliwości. Pierwsza, że cząstki elementarne („atomy”) są to niektóre ze znanych już obiektów — może protony, neutrony, elektrony i coś jeszcze? Druga, że wszystkie znane (i może nieznane) cząstki są równie elementarne. Wedle tej koncepcji wszystkie cząstki są sobie nawzajem potrzebne i nawzajem tłumaczą się teoretycznie (jest to tzw. hipoteza demokracji cząstek). Wreszcie trzecia możliwość to ta, że żadna ze znanych cząstek nie jest elementarna, i że należy zejść o piętro niżej, aby taką cząstkę, czy cząstki wykryć. Może „atomami” są kwarki? Może jeszcze coś innego? Nauka nie daje jeszcze odpowiedzi na te pytania. Jest jednak oczywiste, że sprawa ta będzie nadal przedmiotem wytężonych badań i nie straci nic ze swej pasjonującej aktualności. Wszystko po to, aby uzasadnić przekonania dwu mędrców greckich, którzy mawiali: „z tych samych liter powstaje zarówno tragedia jak i komedia”.

Kwarki — hipotetyczne cząstki o wartościach ładunku mniejszych niż ładunek elektryczny elektronu.



Rys.1



Rozwiązanie zadania F1. Z symetrii obwodu wynika, że potencjały w węzłach B, C, D są równe, a więc odcinkami obwodu BC, CD i DB prąd nie płynie. Pole magnetyczne wytwarzają tylko dwie trójki odcinków prądu AB, AC, AD i EB, EC, ED oraz przewody doprowadzające. Doprowadzenia prądu są współosiowe z odcinkiem AE, nie wytwarzają więc pola magnetycznego w żadnym jego punkcie. Każdy prostoliniowy przewód z prądem (odcinek prądowy) wytwarza w jednorodnym ośrodku pole magnetyczne, którego linie są okręgami leżącymi w płaszczyznach prostopadłych do tego odcinka. Środki okręgów wyznaczają prostą, na której leży odcinek prądowy (rys. 1). Z symetrii obwodu wynika, że każda z gałęzi ABE, ACE i ADE płynie taki sam prąd i na mocy prawa Kirchhoffa dla węzła A, jego natężenie wynosi $I_1 = \frac{1}{3}I$. Odcinek prądowy AB w każdym

punkcie osi AE wytwarza pole o wektorze indukcji B_B prostopadłym do tej osi i do odcinka prądowego (rys. 2). Wektory indukcji wytworzone przez odcinki prądowe AC i AD są takie same co do wartości, także prostopadłe do osi AE (rys. 2a) i tworzą ze sobą kąty 120° (rys. 2b). Wypadkowy wektor indukcji wynosi więc zero w każdym punkcie osi AE.

Tę samą właściwość posiada podobne pole magnetyczne trzech pozostałych odcinków prądowych EB, EC i ED.

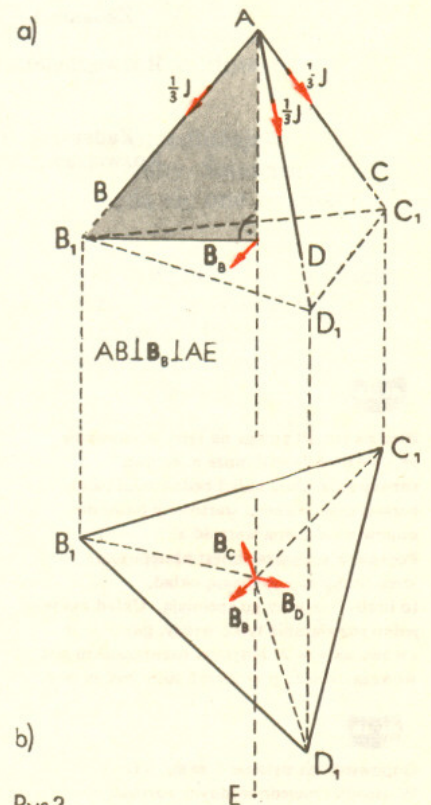
Pytania:

1) W którym miejscu rozmowienie byłoby błędne gdyby ośrodek był niejednorodny? (Odpowiedź na str. 16).

2) Czy wektor

indukcji B znika na całej prostej AE, a więc czy $B = 0$ także wewnątrz przewodów?

Odpowiedź na pytanie 2. Prąd płynący przewodem o skończonej grubości możemy wyobrazić sobie jako zbiór elementarnych prądów płynących równolegle do przewodnika. Jeżeli rozkład tych prądów ma symetrię osiową (zachodzi to w szczególności gdy gęstość prądu jest stała w każdym przekroju przewodnika), to w każdym punkcie tej osi znika pole magnetyczne. Przy takim założeniu wektor $B = 0$ na całej prostej AE.



Rys.2

Sześć zadań – jedno rozwiązanie

dr Maciej BRYŃSKI

Rozwiążemy sześć bardzo prostych zadań. Sformułowanie każdego zadania będzie inne, w każdym zadaniu będziemy zajmowali się innymi obiektami: w jednym przekształceniami płaszczyzny, w innym liczbami, w jeszcze innym dniami tygodnia. Zwróćmy jednak baczniejszą uwagę na rozwiązania tych zadań. Są one bardzo podobne. Może to podobieństwo pozwoli na wprowadzenie schematu, który będzie obejmował rozwiązania wszystkich sześciu zadań? Rozpatrzmy obrót płaszczyzny dookoła ustalonego punktu O o kąt 45° . Jakim przekształceniem płaszczyzny jest 100-krotne złożenie tego obrotu?

Zadanie 1.

Rozwiązanie:

Ponieważ $100 = 12 \cdot 8 + 4$, a ośmiokrotne złożenie danego obrotu jest przekształceniem tożsamościowym, więc w wyniku 100-krotnego złożenia obrotu o 45° otrzymamy przekształcenie identyczne z obrotem o $4 \cdot 45^\circ = 180^\circ$.

Zadanie 2.

Rozwiązanie:

Jaka najmniejsza wielokrotność liczby 6 jest podzielna przez 8? Ile różnych reszt z dzielenia przez 8 można otrzymać rozpatrując wszystkie wielokrotności liczby 6? Liczba $6n$ jest podzielna przez 8, gdy n jest podzielne przez 4. Wobec tego najmniejszą liczbą spełniającą warunki zadania jest 24. Ponieważ $6 = 0 \cdot 8 + 6$, $2 \cdot 6 = 1 \cdot 8 + 4$, $3 \cdot 6 = 2 \cdot 8 + 2$, $4 \cdot 6 = 3 \cdot 8$, więc z pewnością są co najmniej cztery różne reszty: 0, 2, 4, 6. Dla stwierdzenia, że są to wszystkie możliwe reszty, przeprowadzimy następujące rozumowanie: zauważyliśmy poprzednio, że $6n$ przy n podzielnym przez 4 dzieli się bez reszty przez 8; gdy $n = 4k + 1$, to $6n = 6(4k + 1) = 24k + 6$ — ta liczba daje przy dzieleniu przez 8 resztę 6; gdy $n = 4k + 2$, to $6n = 24k + 12$ — tu otrzymamy resztę 4; gdy zaś $n = 4k + 3$, to $6n = 24k + 18$ — tu reszta jest 2. Widzimy więc, że przy kolejnych wielokrotnościach liczby 6, reszty z dzielenia przez 8 powtarzają się cyklicznie: 0, 6, 4, 2. Reszt jest zatem dokładnie cztery.

Zadanie 3.

Rozwiązanie:

Niech T oznacza przesunięcie płaszczyzny o wektor w , którego długość wynosi $\frac{1}{8}$. Ile razy należy powtórzyć to przekształcenie, by otrzymać przesunięcie o wektor, którego długość jest liczbą całkowitą? Co można powiedzieć o n -krotnym złożeniu przesunięcia T ?

Szukamy takiej liczby naturalnej n , by $\frac{1}{8} \cdot n$ było liczbą całkowitą. Oczywiście n musi być liczbą podzielną przez 8. Najmniejszą taką liczbą jest właśnie 8. Jeśli n nie jest liczbą podzielną przez 8; $n = 8k + r$ ($r = 1, 2, \dots, 7$), to n -krotne złożenie przesunięcia T jest przesunięciem o wektor, którego długość wynosi $\frac{1}{8}(8k + r) = k + \frac{r}{8}$. Długość ta różni się od liczby całkowitej o $\frac{r}{8}$.

Zadanie 4.

Rozwiązanie:

Każdy z 30 uczestników obozu harcerskiego kolejno pełni godzinną wartę. Najmłodszy z nich pierwszego dnia pełni wartę w godzinach 12—13. W jakich godzinach przypadnie jego szósta kolejna warta?

Interesująca nas zmiana warty nastąpi po 150 godzinach ($150 = 30 \cdot 5$, a wszyscy harcerze muszą pełnić wartę pięciokrotnie). Doba ma 24 godziny; $150 = 6 \cdot 24 + 6$, a więc następna warta wypadnie po upływie sześciu dni o 6 godzin później od pierwszej warty. Będzie to więc godzina 18. Nietrudno zauważyć, że każda następna warta wypada mu o sześć godzin później niż poprzednia; godziny jego dyżurów są 12—13, 18—19, 0—1, 6—7 (oczywiście każdego dnia będzie miał najwyżej jedną wartę).

Zadanie 5.

Rozwiązanie:

1 stycznia 1974 r. wypada we wtorek. Jakim dniem tygodnia będzie 31 stycznia 1974 r? Jakim dniem tygodnia będzie setny dzień 1974 r?

Od 1 stycznia do 31 stycznia upływie 30 dni. Ponieważ $30 = 4 \cdot 7 + 2$, więc numer dnia tygodnia zmieni się o 2; interesującym nas dniem będzie czwartek. Podobnie $99 = 14 \cdot 7 + 1$, więc setnym dniem 1974 r. będzie środa (10.IV.74).

Zadanie 6.

Rozwiązanie:

Jaka jest ostatnia cyfra liczby 2^{1001} ? Rozpatrzmy kilka kolejnych potęg liczby 2: 2, 4, 8, 16, 32, 64, 128, 256, ... Ostatnie cyfry tych liczb powtarzają się cyklicznie 2, 4, 8, 6, 2, 4, 8, 6, 2, ... przy czym mamy tu następującą regułę: ostatnia cyfra liczby 2^n jest 2, gdy $n = 4k + 1$, 4 — gdy $n = 4k + 2$, 8 — gdy $n = 4k + 3$, 6 — gdy $n = 4k$. (Skrupulatny Czytelnik nie omieszcza przeprowadzić dokładnego dowodu indukcyjnego tego faktu). Wobec tego liczba 2^{1001} ma ostatnią cyfrę 2, bo $1001 = 250 \cdot 4 + 1$.

Czytelnik, który cierpliwie doczytał do tego miejsca, bez wątpienia zauważył wspólny schemat wszystkich sześciu rozwiązań. W każdym bowiem zadaniu mieliśmy do czynienia z taką sytuacją, że pewna wielokrotność wyjściowego elementu była mu równa; następne wielokrotności powtarzały się cyklicznie. Zbudujmy prosty model ilustrujący tę sytuację: Ustalmy liczbę naturalną n i w zbiorze $0, 1, \dots, n-1$ określmy działanie $\oplus$ według następującej reguły:

$a \oplus b =$ reszta z dzielenia $(a + b)$ przez n .

Zbiór $0, 1, \dots, n-1$ wraz z określonym wyżej działaniem nazywamy grupą cykliczną C_n . Działanie w grupie C_n jest określone w ten sposób, że każdy element tej grupy jest wielokrotnością elementu 1. $2 = 1 \oplus 1$, $3 = 1 \oplus 1 \oplus 1$, ..., $n-1 = \underbrace{1 \oplus 1 \oplus \dots \oplus 1}_{(n-1) \text{ razy}}$, $0 = \underbrace{1 \oplus 1 \oplus \dots \oplus 1}_n$;

następne wielokrotności powtarzają się cyklicznie.

Powróćmy jeszcze na chwilę do rozwiązanych poprzednio zadań:

1. Powtarzanie obrotu o 45° odpowiada dodawaniu elementu 1 grupy C_8 do siebie.

$\underbrace{1 \oplus 1 \oplus \dots \oplus 1}_{100 \text{ razy}} = 4$

2. W grupie C_8 $6 \oplus 6 = 4$, $6 \oplus 6 \oplus 6 = 2$, $6 \oplus 6 \oplus 6 \oplus 6 = 0$, $6 \oplus 6 \oplus 6 \oplus 6 \oplus 6 = 6$, następne wielokrotności powtarzają się cyklicznie. Równość $6 \oplus 6 \oplus 6 \oplus 6 = 0$ jest równoważna temu, że $6 + 6 + 6 + 6$ dzieli się przez 8.

3. Ponieważ interesują nas tylko części ułamkowe długości wektora przesunięcia, będącego wielokrotnością przesunięcia T , więc możemy dodawać wielokrotności wektora w tak, jak



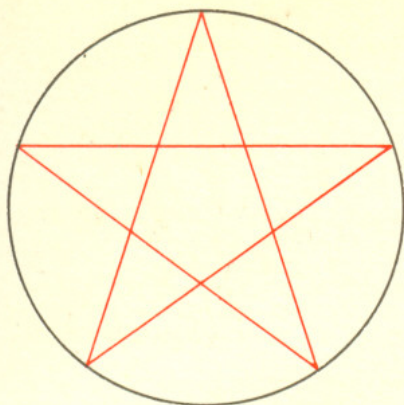
Błąd na str. 10 polega na tym, że równanie $y^2 + y - (m+2) = 0$ może mieć dwa rozwiązania, choć układ będzie miał jedno rozwiązanie; jednej z wartości y może nie odpowiadać żadna wartość x .

Poprawne rozwiązanie jest następujące: jeżeli liczby x, y spełniają układ, to liczby $-x, y$ też go spełniają. Układ ma więc jedno rozwiązanie tylko wtedy, gdy $x = 0$ i wówczas $y = 2$. Jedynym rozwiązaniem jest wówczas $x = 0, y = 2$ czyli musi być $m = 4$.

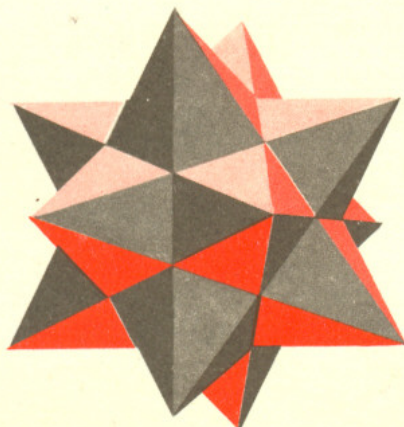


Odpowiedź na pytanie 1 ze str. 15.

W ośrodku niejednorodnym wartości bezwzględne indukcji B_A, B_B, B_C nie muszą być równe.

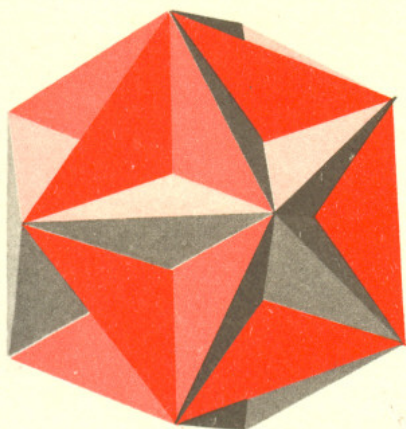


pięciokąt foremny gwiazdzisty

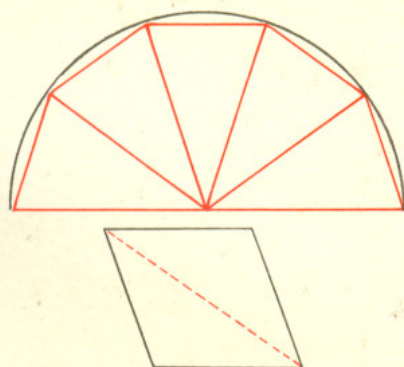


dwunastościan gwiazdzisty mały

Łamana zwyczajna to taka, która (mówiąc po prostu) nie przecina się sama ze sobą.



dwunastościan wielki



elementy C_8 : $k \cdot w + l \cdot w = (k \oplus l) \cdot w$. Stąd dla $n = 8k + r$ ($0 \leq r < 8$) część ułamkowa długości wektora $n \cdot w$ wynosi $r \cdot \frac{1}{8}$.

4.

W grupie C_{24} (dość ma 24 godziny)

$$\underbrace{1 \oplus 1 \oplus \dots \oplus 1}_{150 \text{ razy}} = 6.$$

Oznacza to, że termin następnej warty wypada o 6 godzin później niż warty poprzedniej. W tej samej grupie $6 \oplus 6 \oplus 6 \oplus 6 \oplus 6 = 6$, więc szósta warta wypadnie o 6 godzin później niż pierwsza.

5.

Numery dni tygodnia należy dodawać tak, jak elementy grupy C_7 . W tej grupie

$$\underbrace{1 \oplus 1 \oplus \dots \oplus 1}_{30 \text{ razy}} = 2 \quad \text{a} \quad \underbrace{1 \oplus 1 \oplus \dots \oplus 1}_{99 \text{ razy}} = 1$$

6.

Stwierdził, że ostatnie cyfry liczby 2^n powtarzają się cyklicznie przy zmianie n o 4. Można więc zastąpić wykładnik n przez element

$$\underbrace{1 \oplus 1 \oplus \dots \oplus 1}_n \text{ grupy } C_4, \quad \text{a} \quad \underbrace{1 \oplus 1 \oplus \dots \oplus 1}_{1001 \text{ razy}} = 1.$$

Stąd ostatnią cyfrą jest 2.

Jak więc widzimy, rozwiązanie każdego z naszych zadań polegało na wykonaniu pewnych działań w odpowiednio dobranej grupie cyklicznej C_n .

Wielościany gwiazdziste

Jeśli przy definiowaniu wielokąta zrezygnujemy z warunku, aby łamana tworząca go była zwyczajna, otrzymamy nową klasę wielokątów foremnych, tzw. gwiazdzistych.

Co otrzymamy, jeśli pójdziemy dalej i dopuścimy, aby takie wielokąty były ścianami wielościanów, rezygnując przy tym z analogicznego warunku, który ścianom wielościanów nie pozwalał przecinać się ze sobą? Otóż wtedy do grona znanych nam wielościanów foremnych (czworościan, sześciąt, ośmiościan, dwunastościan i dwudziestościan foremny) dojdą jeszcze cztery.

Zauważmy, że jeżeli w pięciokącie foremnym odpowiednio przedłużymy boki, to otrzymamy pięciokąt foremny gwiazdzisty. Wykorzystamy to przy konstrukcji pierwszego wielościanu. Weźmy więc dwunastościan foremny, którego ścianami są pięciokąty i przedłużmy jego krawędzie aż do ich przecięcia. Otrzymamy w ten sposób dwunastościan gwiazdzisty mały. Pamiętajmy czym są jego ściany! Ponieważ łamana, która tworzy pięciokąt gwiazdzisty, nie rozcina płaszczyzny na dwie rozłączne figury, więc nie możemy tu mówić o wielokącie, jako o części płaszczyzny ograniczonej łamaną. Tak więc wielościan ten jest zbudowany wyłącznie z odcinków.

Drugi z kolei wielościan, dwunastościan wielki, otrzymamy przedłużając odpowiednio, również w dwunastościanie foremnym, jego ściany. Ścianami jego będą „normalne” pięciokąty, ale jak już wspominaliśmy, będą się one ze sobą przecinały.

Ponieważ obydwa te wielościany są bardzo efektowne, warto wykonać ich modele. Model pierwszego, zgodnie z poprzednimi uwagami, należałoby wykonać właściwie z drutu, patyczków itp. Byłoby to jednak trudniejsze technicznie, dlatego skonstruujemy model taki jak na rysunku, pamiętając jednak, że naprawdę liczą się tylko krawędzie. Nie będziemy podawać całej jego siatki, gdyż łatwiej jest wykonać dwanaście ostrosłupów i następnie połączyć je np. taśmą samolepiącą. Siatkę jednego ostrosłupa tworzy połowa dziesięciokąta foremnego, czego łatwy dowód pozostawiamy Czytelnikowi.

Przy drugim modelu będziemy się musieli nieco więcej napracować. Można wprawdzie wyciąć siatkę i skleić wielościan, jednak wpłynęłoby to ujemnie na jego trwałość. Dlatego proponujemy wyciąć trzydzieści przystających rombów o kącie ostrym 72° , naciąć i zgiąć wzdłuż dłuższej przekątnej, a następnie, kierując się rysunkiem, skleić model.

Jeśli, jako punkt wyjścia weźmiemy dwudziestościan foremny, w podobny sposób otrzymamy dwa pozostałe wielościany.

Uważny Czytelnik po sklejeniu obu modeli na pewno zauważy, jaki zachodzi związek pomiędzy tymi wielościanami.

J. B.