

$h=6.63 \cdot 10^{-34}$ J·s $\pi=3.1415$

CENA 5 ZŁ

NR 9 1976

deia
POPULARNY MIESIĘCZNIK MATEMATYCZNO-FIZYCZNY



SPIS TREŚCI

Czy mechanika kwantowa jest teorią kompletną?	
<i>Prof. dr Grzegorz Białkowski</i>	str. 1
Półwiecze	
<i>Doc. dr Jerzy Kijowski</i>	str. 4
Zadania	str. 7
Iloczyn skalarny	
<i>Mgr Krzysztof Nowiński</i>	str. 8
Mała Delta	str. 10
Prawdopodobieństwo?	
<i>Dr Tadeusz B. Iwiński</i>	str. 13
Cudowny wynik pewnego doświadczenia	str. 15
Laboratorium w domu	
<i>Dr Jan A. Gaj</i>	str. 16

Zachęcamy naszych Czytelników do uczestniczenia w popularno-naukowych odczytach organizowanych przez Polskie Towarzystwo Matematyczne i Centrum Obliczeniowe PAN. Zbliżająca się jesienna sesja odczytów będzie poświęcona statystyce.

Zabiorą głos:

- doc. Ryszard Zieliński z Instytutu Matematycznego PAN, który przedstawi metody statystyczne aproksymacji stochastycznej (15 X br);
- dr Roman Zmyślony z Instytutu Matematycznego PAN, który nakreśli historię i aktualny stan badań w zakresie modeli liniowych (29 X br);
- dr Andrzej Kozek z Instytutu Matematycznego PAN, który będzie mówił o pewnych zaskakujących faktach w teorii testowania hipotez, związanych z porównywalnością testów (12 XI br);
- dr Adam Kowalski z Centrum Obliczeniowego PAN, który dokona krytycznego przeglądu metod taksonomicznych (26 XI br).

Wszystkie odczyty będą się odbywać zgodnie z tradycją w piątki o godzinie 13.15 w sali 1044 Pałacu Kultury i Nauki w Warszawie.

Zainteresowanym radzimy nawiązać kontakt z sekretariatem odczytów (Centrum Obliczeniowe PAN, 00-901 Warszawa — PKiN) w celu uzyskania bliższych informacji oraz zawiadomień o następnych spotkaniach i w celu przekazania swoich uwag i sugestii.

„Delta”
matematyczno-fizyczny miesięcznik
popularny
Polskiego Towarzystwa
Matematycznego i Polskiego
Towarzystwa Fizycznego
wydawany przy poparciu
Ministerstwa Oświaty i Wychowania

Komiteta Redakcyjnego
doc. dr J. Bartke
prof. dr Grzegorz Białkowski —
przewodniczący
doc. dr A. Bączyński
doc. dr B. Gleichgewicht
doc. dr K. Goebel
doc. dr B. Iwaszkiewicz
doc. dr T. Iwiński
prof. dr A. Januszajtis
prof. dr Leon Jeśmanowicz —
wiceprzewodniczący
mgr H. Kaczorek
prof. dr B. Karczewski
prof. dr M. Kuczma
mgr A. Mąkowski
prof. dr Z. Pawlak
prof. dr A. Piekara

prof. dr Z. Semadeni
prof. dr J. Stankowski
prof. dr M. Subotowicz
doc. dr S. Turnau
doc. dr J. Wdowczyk

Redaguje Kolegium w składzie:
doc. dr T. Hofmokr — z-ca red. nac.
dr T. B. Iwiński
dr M. Kordos — red. nac.
mgr K. Prażmowski — red. techn. graf.
doc. dr M. Święcki
D. Tys — sekr. red.
Adres Redakcji
ul. Hoża 69 pok. 151,
00-681 Warszawa,

Zakład Narodowy im.
Ossolińskich — Wydawnictwo
Wrocław, Oddział w Warszawie
Nakład 20 000 egz. Objętość 2 ark.
wyd.; 2,50 ark. druk.;
papier offsetowy III kl. 80 g, 61×86
Wydrukowano w Drukarni im.
Rewolucji Październikowej
Warszawa, ul. Mińska 65.
Nr zam. 814/76 J-113

Wydano z pomocą finansową Polskiej Akademii Nauk

WARUNKI PRENUMERATY Cena prenumeraty rocznej zł 60, — cena prenumeraty półrocznej zł 30, —

Prenumeratę na kraj przyjmują Oddziały RSW „Prasa—Książka—Ruch” oraz urzędy pocztowe i doręczyciele — w terminach:
— do 25 listopada na styczeń, I kwartał, I półrocze roku następnego i na cały rok następnego.
— do dnia 10 miesiąca, poprzedzającego okres prenumeraty na pozostałe okresy roku bieżącego.
Jednostki gospodarki społecznej, instytucje i organizacje społeczno-polityczne składają zamówienie w miejscowych Oddziałach RSW „Prasa—Książka—Ruch”.
Zakłady pracy i instytucje w miejscowościach, w których nie ma Oddziałów RSW oraz prenumeratorzy indywidualni, zamawiają prenumeratę w urzędach pocztowych lub u doręczycieli.
Prenumeratę ze zleceniem wysyłki za granicę, która jest o 50% droższa od prenumeraty krajowej, przyjmuje RSW „Prasa—Książka—Ruch”, Centrala Kolportażu Prasy i Wydawnictw, ul. Towarowa 28 00-958 Warszawa, konto PKO nr 1531-71 w terminach podanych dla prenumeraty krajowej.

Sprzedaż numerów bieżących i uprzednich

Instytucje państwowe i społeczne, zakłady pracy, szkoły i czytelnicy indywidualni mogą nabywać „DELTA”:

w Księgarni Ośrodka Rozpowszechniania Wydawnictw Naukowych PAN.
Sprzedaż gotówkowa i wysyłkowa, numerów bieżących i archiwalnych; płatność gotówką, przelewem lub za zaliczeniem pocztowym.

Adres: ORPAN 00-901 Warszawa, Pałac Kultury i Nauki, konto PKO nr 1-6-100312

w Księgarni Ossolineum, Rynek 8, 50-106 Wrocław

w Głównej Księgarni Naukowej, Krakowskie Przedmieście 7, 00-068 Warszawa

w Księgarni Naukowej, ul. Podwale 6, 31-118 Kraków

Orders for this periodical from abroad can be placed with „Ars Polona” Krakowskie Przedmieście 7 00-068 Warszawa, Poland or with

— Kubon & Sagner, Inhaber Otto Sagner, D8 Munchen 34, Postfach 68, Bundesrepublik Deutschland.

— Earls Court Publications Ltd., 130 Shephard Bush Centre, London W 12, Great Britain.

— Licos Commissionaria Sansoni, Via Lamarmora 45, 50 121 Firenze, Italia.

Cena 1 egzemplarza zł 5, — nr indeksu 35723/35550

Czy mechanika kwantowa jest teorią kompletną?

Prof. dr Grzegorz BIAŁKOWSKI

Niewątpliwie tą teorią fizyczną, która najbardziej współdecyduje o obliczu fizyki współczesnej, jest mechanika kwantowa. Szereg jej twierdzeń i postulatów trudno zrozumieć w świetle naszego doświadczenia codziennego. Toteż właściwie od momentu jej powstania — a upływa właśnie pięćdziesiąt lat, odkąd przyjęła postać niemal definitywną — budziła ona wiele wątpliwości i prowokowała do sprzeciwu. Do dziś szereg kwestii interpretacyjnych nie znalazło jeszcze ostatecznego wyjaśnienia i nie brak fizyków, którzy się spodziewają, że stanie się to w końcu powodem do sformułowania nowej, lepszej teorii. Jakkolwiek nikt w zasadzie nie wątpi, że mechanika kwantowa, podobnie jak wszystkie pozostałe teorie fizyczne, zostanie z czasem zastąpiona przez inną koncepcję, której ona sama stanie się tylko jakąś wersją graniczną czy przybliżoną, to przecież trzeba powiedzieć, że na razie brak jest przesłanek fizycznych do sformułowania takiej koncepcji. Po prostu mechanika kwantowa, jak dotąd, przechodzi zwycięsko rozmaite testy eksperymentalne, wobec czego zastrzeżenia, wysuwane przeciwko tej teorii mają swoje źródło nie w trudnościach czysto naukowych, ale, można powiedzieć, w jej ogólnym klimacie pojęciowym, a może nawet filozoficznym.



W artykule tym chcę poruszyć jedną z podstawowych kwestii podnoszonych przez oponentów mechaniki kwantowej, a mianowicie — czy teoria ta jest kompletna?

Pytanie to w wersji bardziej rozwiniętej można by wypowiedzieć następująco: czy w mechanice kwantowej znajdują swoje odzwierciedlenie wszystkie elementy rzeczywistości danej w doświadczeniu?

Łatwo dostrzec, która cecha mechaniki kwantowej prowokuje do postawienia takiego pytania. Jak wiadomo — a mówi się o tym w setkach i tysiącach mniej lub bardziej udanych artykułów i książek popularnych — mechanika kwantowa nie jest teorią deterministyczną w sensie fizyki klasycznej. W mechanice, stworzonej przez Galileusza i Newtona, znając położenie i prędkość (lub położenie i pęd) ciała materialnego, które dla uproszczenia potraktujemy jako punkt materialny, można przewidzieć cały przyszły ruch tego ciała (a więc podać położenie i pęd w każdej chwili późniejszej), a także ustalić, jaki był ruch tego ciała w przeszłości. Wszystko to, oczywiście, przy założeniu, że znane są całkowicie wszystkie siły działające na ciało w przeszłości i w przyszłości. Inaczej z mechaniką kwantową; zgodnie bowiem z prawami kwantowymi nie jest spełniony już ów wstępny warunek, a mianowicie nie jest możliwy jednoczesny dowolnie dokładny pomiar położenia i prędkości żadnego obiektu kwantowego. Nie jest więc też możliwe precyzyjne przewidywanie ruchu takiego obiektu. Całkowita informacja, na którą nam pozwala mechanika kwantowa sprowadza się do prawdopodobieństwa, że przyszłe zachowanie obiektu kwantowego będzie takie a takie. Tak więc na to, aby sprawdzać prawa kwantowe musimy mieć do dyspozycji bardzo wiele identycznych obiektów. Zachowanie całego takiego zespołu potrafimy określić jednoznacznie, wiemy bowiem, jaka część tego zespołu znajdzie się po chwili t w określonym dowolnie małym obszarze przestrzeni. Nie wiemy tylko, które indywidua z tego zespołu to uczynią. Znanym przykładem prawa kwantowego jest prawo rozpadu promieniotwórczego: wiemy, że po czasie t określona część atomów danego pierwiastka rozpadnie się, ale nie umiemy przewidzieć, które to będą atomy.

W tym sensie mechanika kwantowa jest teorią statystyczną. Jest to jednak teoria statystyczna szczególnego typu. Przecież w fizyce klasycznej także są znane teorie statystyczne! Weźmy choćby gaz zamknięty w jakimś naczyniu. Wiemy, że w warunkach równowagi termodynamicznej dwie równe co do objętości części tego naczynia będą zawierać jednakową liczbę cząsteczek gazu. Nie wiemy jednak, które cząsteczki znajdują się w której z dwu połówek naczynia. Sytuacja pozornie przypomina prawo rozpadu: można podać taki czas, w którym rozpadnie się połowa atomów i w ten sposób podzielić wszystkie atomy na dwie równe części — te, które się w tym czasie rozpadną, i te, które się nie rozpadną, podobnie jak podzieliłiśmy cząsteczki gazu według kryterium, w której połowie naczynia się znajdują.



Jednakże w rzeczywistości w obu teoriach sytuacja jest zupełnie inna: Otóż w klasycznej fizyce statystycznej znamy prawa rządzące zachowaniem pojedynczych cząsteczek (są nimi z założenia prawa mechaniki newtonowskiej), a nasza niewiedza co do tego zachowania jest spowodowana po pierwsze niemożliwością śledzenia ruchu wielu trylionów obiektów, a po drugie brakiem potrzeby, aby to czynić: wystarczy nam znać właśnie tylko pewne wielkości średnie, które ujawniają się fenomenologicznie na przykład jako temperatura gazu, czy też jego ciśnienie.

Tak więc rzeczywisty kompletny opis stanu gazu musiałby zawierać informacje dotyczącą N wektorów położenia i N wektorów pędu (N — liczba cząsteczek gazu), co jest liczbą ogromną, podczas gdy opis statystyczny ogranicza się do kilku potrzebnych liczb. Rodzi się więc niemal automatycznie pytanie, czy mechanika kwantowa nie jest także teorią statystyczną w tym właśnie sensie, czy nie operuje ona tylko jakimiś wielkościami średnimi, za którymi stoi wiele, może nawet bardzo wiele pomijanych zmiennych, nie ujawniających się w równaniach mechaniki kwantowej. Takie zmienne, oczywiście hipotetyczne, nazywa się parametrami ukrytymi. Można by, trochę złośliwie, powiedzieć, że parametry te są „ukryte” nie tylko dlatego, że nie pojawiają się w równaniach kwantowych, ale także dlatego, że nikt jak dotąd nie podał rozsądnej sugestii, jakiej natury miałyby być owe parametry.

Pozostaje jednak sensowne pytanie, czy przyroda nie domaga się wprowadzenia jakichś parametrów ukrytych. Na ten temat wypowiedzieć się może jedynie doświadczenie. W ramach mechaniki kwantowej wykazano jednak (jest to tzw. twierdzenie von Neumanna), że wprowadzenie parametrów ukrytych nie da się pogodzić z prawami kwantowymi, wymagałoby więc ono radykalnej przebudowy całej teorii. Szczególnie jedna z zasad kwantowych stoi tu na przeszkodzie, a mianowicie tzw. zasada superpozycji stanów.

Aby dobrze zrozumieć, o co tutaj chodzi, przypomnijmy sobie, że stan dowolnego obiektu kwantowego opisujemy pewną funkcją zespoloną, zwaną funkcją falową. Zasada superpozycji mówi nam, że jeśli mamy dwie funkcje falowe opisujące dopuszczalne stany określonego obiektu kwantowego, to dowolna kombinacja liniowa tych dwu funkcji także reprezentuje dopuszczalny stan tego obiektu. Natomiast prawdopodobieństwo znalezienia obiektu w danym stanie jest równe kwadratowi modułu odpowiedniej funkcji falowej.

Chcąc sobie zdać sprawę z „nieklasyczności” zasady superpozycji rozpatrzmy dwa różne obiekty: elektron „klasyczny” i elektron „kwantowy”. Przypuśćmy, że mamy pewien określony stan początkowy a , z którego możliwe jest przejście do stanu końcowego c . Niech prawdopodobieństwo takiego przejścia wynosi $P(a, c)$. Zapominamy przy tym na chwilę, że w mechanice klasycznej prawdopodobieństwo to może mieć tylko dwie wartości, a mianowicie 0 lub 1, różnice z mechaniką kwantową są bowiem znacznie głębsze. Niech teraz przejście $a \rightarrow c$ dokonuje się przez jeden z wielu (N) stanów pośrednich b_i ($i = 1, \dots, N$). Oznaczmy prawdopodobieństwo przejścia $a \rightarrow b_i \rightarrow c$ przez $P(a, b_i, c)$. Jest chyba oczywiste, że

$$(1) \quad P(a, c) = \sum_{i=1}^N P(a, b_i, c) = \sum_{i=1}^N P(a, b_i) P(b_i, c) ?$$

I tak jest naprawdę w fizyce klasycznej, natomiast może tak nie być w fizyce kwantowej. Podając powyższy wzór założyliśmy bowiem milcząco, że jest rzeczą obojętną dla przebiegu procesu, czy dokonaliśmy pomiaru, zmierzającego do ustalenia, czy rzeczywiście nasz elektron przeszedł przez określony stan pośredni b_i . Okazuje się, że nie jest to rzeczą obojętną. Powyższy wzór opisuje rzeczywistość fizyczną tylko wtedy, gdy pomiar taki naprawdę został wykonany. W przeciwnym razie obowiązuje wzór inny, a mianowicie

$$(2) \quad \psi(a, c) = \sum_{i=1}^N \psi(a, b_i, c),$$

gdzie symbolem ψ oznaczyliśmy odpowiednią funkcję falową.

Weźmy konkretny przykład. Niech a odpowiada elektronowi w pewnym stanie emitowanemu z jakiegoś źródła, a c — elektronowi padającemu w określonym punkcie na ekran, który może być kliszą fotograficzną. Niech stan b_i odpowiada przejściu przez i -tą szczelinę w pewnej przesłonie, umieszczonej między źródłem i ekranem, i dla uproszczenia niech $N = 2$. $P(a, b_1, c)$ mierzymy, zasłaniając

Rozwiązanie zadania M 97. Niech q będzie ilorazem rozpatrywanego ciągu geometrycznego, r zaś różnicą ciągu arytmetycznego. Mamy

$$s = xq, z = xq^2, b = a + r, c = a + 2r$$

skąd

$$x^b y^c z^a = x^{a+r} (xq)^{a+2r} (xq^2)^a =$$

$$= x^{a+r} x^{a+2r} q^{a+2r} x^a q^{2a} = x^{3a+3r} q^{3a+2r},$$

$$x^c y^a z^b = x^{a+2r} (xq)^a (xq^2)^{a+r} =$$

$$= x^{a+2r} x^a q^a x^{a+r} q^{2a+2r} = x^{3a+3r} q^{3a+2r},$$

$$\text{a więc } x^b y^c z^a = x^c y^a z^b.$$

Nieco prościej można rozwiązać to zadanie wykorzystując związek $y^2 = xz$:

$$\frac{x^b y^c z^a}{x^c y^a z^b} = x^{b-c} y^{c-a} z^{a-b} = x^{-r} y^{2r} z^{-r} =$$

$$\left(\frac{y^2}{xz}\right)^r = 1.$$

drugą szczelinę i wyznaczając zaczerpnienie kliszy w wybranym punkcie. Podobnie mierzymy $P(a, b_2, c)$. Jeżeli doświadczenie będziemy wykonywać, przepuszczając systematycznie elektrony tylko przez jedną z dwóch szczelin, to otrzymamy na kliszy obraz zgodny z wzorem (1). Jeżeli jednak odsłoniemy obie szczeliny, a więc zrezygnujemy z pomiaru b_i , to obraz na kliszy będzie zupełnie inny — zgodny z wzorem (2). Na kliszy ujawni się wówczas obraz interferencyjny pochodzący od nakładania się dwu „fal”, dwu składników we wzorze (2). Nakładanie to nie jest wynikiem interferencji fal odpowiadających dwu różnym elektronom. Robiono bowiem doświadczenia z elektronami wypuszczanymi pojedynczo w pewnych odstępach czasu, a jednak obraz interferencyjny był taki sam. Tak więc zasada superpozycji jest jednym z najistotniejszych elementów, różniących fizykę klasyczną od kwantowej.

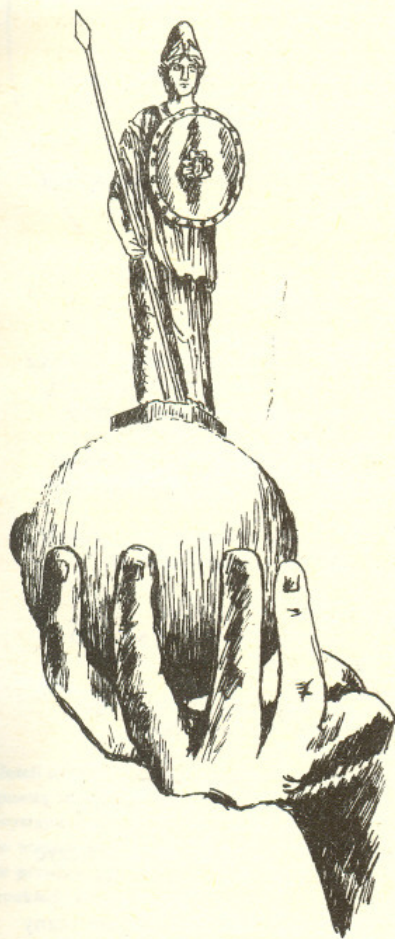
Jak wiadomo, do przeciwników mechaniki kwantowej należał również Einstein. (Z faktu tego do dziś czerpią otuchę rozmaici maniacy, którzy chcą poprawić mechanikę kwantową nie rozumiejąc jej i sądzą, że znajdują się w tym samym obozie co i Einstein). W r. 1934 Einstein opublikował wraz ze swymi współpracownikami Podolskim i Rosenem artykuł wymierzony właśnie przeciw mechanice kwantowej. Rozpatrują oni układ fizyczny, składający się z dwu podukładów, które oddziałują ze sobą tylko w pewnym czasie t zawartym między $t = 0$ i $t = T$. Przed powstaniem oddziaływania ($t < 0$) oba układy są niezależne. Funkcja falowa całego układu jest wtedy po prostu iloczynem funkcji falowych podukładów, gdyż prawdopodobieństwa znalezienia podukładu a w stanie a_1 i podukładu b w stanie b_1 są niezależne. Sytuacja zmienia się po włączeniu oddziaływania: funkcja falowa całego układu przestanie być iloczynem funkcji falowych podukładów, a będzie zależeć w pewien bardziej ogólny sposób od zmiennych charakteryzujących stan obu części. W chwili $t = T$ oddziaływanie wygasa; mimo to jednak funkcja falowa nie staje się automatycznie iloczynem funkcji falowych obu podukładów, lecz trwa w tej właśnie ogólniejszej postaci. Tę ogólną funkcję falową, zgodnie z zasadą superpozycji, można przedstawić w postaci kombinacji liniowej iloczynu funkcji falowych obu podukładów z osobna,

$$(3) \quad \psi(a, b) = \sum_i \psi(a_i) \psi(b_i).$$

Wyobraźmy sobie teraz, że w chwili $t > T$ wykonujemy na podukładzie b pomiar, który dowodzi nam, że podukład ten znalazł się w stanie b_i . Tym samym z całej kombinacji liniowej (3) pozostaje jeden tylko, i -ty wyraz, i zarazem zostaje ustalone, że podukład a znajduje się w stanie a_i , mimo, że na podukładzie tym nie wykonaliśmy żadnego pomiaru. Rzeczywiście, nie oddziałaliśmy wcale na układ a , gdyż z założenia układ b nie oddziałuje już z układem a i nie może mu przekazać żadnej informacji typu „robią na mnie pomiar”.

Co więcej, jak argumentują Einstein, Podolsky i Rosen, można na układzie b wykonać dwa pomiary wielkości niewspółmierzalnych, jak np. położenia i pędu. Jeden pomiar zakłóca stan ustalony przez drugi, ale oczywiście tylko dla tego podukładu, na którym pomiar ten był wykonany, czyli podukładu b . Natomiast układu a to nie dotyczy. Tak więc o tym samym układzie a , na którego stan, po ustaniu oddziaływania $a-b$ nie możemy wpłynąć przez pomiar na układzie b , otrzymujemy dwie zupełnie różne, i nawet — zgodnie z prawami kwantowymi nie współzależne informacje. Einstein i współautorzy wyciągają stąd wniosek, że opis rzeczywistości w ramach mechaniki kwantowej nie jest kompletny. Bohr, Heisenberg i inni przedstawiciele tzw. szkoły kopenhaskiej na ten zarzut (tzw. paradoks Einsteina, Podolsky'ego i Rosena) odpowiedzieli, że funkcja falowa nie jest opisem samej rzeczywistości, a tylko naszej informacji o tej rzeczywistości. Nic dziwnego, mówią oni, że pomiar podukładu b , zmieniając naszą informację, wpływa na funkcję falową podukładu a . Nic też dziwnego, że dowiadujemy się o tym natychmiast, jakby z nieskończoną prędkością.

Nietrudno zauważyć, że taka interpretacja funkcji falowej może być pobudką do pewnych rozważań filozoficznych, w których pojęcie przedmiotu fizycznego zacierza się, a na jego miejscu pojawia się pojęcie przedmiotu „dla nas”, być może jakiegoś tylko wytworu naszej świadomości. Z taką konsekwencją zapewne nie każdy (w tym także i Einstein) mógłby się pogodzić. W tej sytuacji nie należy wykluczać konieczności jakiejś przebudowy mechaniki kwantowej, m.in. na przykład przez ograniczenie roli zasady superpozycji. W tej chwili jednak, jak już mówiłem, brak przeciw tej zasadzie argumentów eksperymentalnych. Mechanika kwantowa, ze wszystkimi wątpliwościami interpretacyjnymi, które budzi, nadal świetnie opisuje wyniki doświadczeń. A to może w końcu jest najważniejsze.



Za początek nowej, kwantowej ery w fizyce przyjmuje się zazwyczaj rok 1900. Wtedy to właśnie fizyk niemiecki Max Planck odkrył tak zwany kwant działania, którego wielkość, równa

$$h = 6,6255 \cdot 10^{-27} \text{ erg} \cdot \text{sek},$$

jest uniwersalną stałą fizyczną charakteryzującą procesy zachodzące w świecie obiektów mikroskopowych. Przedstawiona dalej hipoteza Plancka wyjaśniała początkowo jedno tylko zjawisko fizyczne — rozkład energii promieniowania wysyłanego przez rozgrzane ciało fizyczne między różne częstotliwości tego promieniowania. Rozkład ten badano doświadczalnie, otrzymując w rezultacie krzywe w rodzaju przedstawionej linią ciągłą na rys. 1. Wielkość $E(\nu, \Delta\nu)$ równa iloczynowi $u(\nu) \cdot \Delta\nu$ (pole zakreskowane na rysunku) jest ilością energii wypromieniowanej przez badane ciało w ciągu jednostki czasu, w zakresie częstotliwości od ν do $\nu + \Delta\nu$.

Zagadnienie to w teoretycznie wyidealizowanej formie, mianowicie przy założeniu, że promieniowanie elektromagnetyczne pozostaje w równowadze termodynamicznej z wysyłającym je ciałem, otrzymało nazwę zagadnienia promieniowania ciała doskonale czarnego. Według fizyki klasycznej — termodynamiki i elektrodynamiki — rozkład promieniowania ciała doskonale czarnego powinna opisywać krzywa dana wzorem Rayleigha-Jeansa

$$u(\nu) = \frac{8\pi^2\nu^2 kT}{c^3}$$

(linia przerywana na rys. 1), co stało w rażącej sprzeczności z wynikami doświadczalnymi. Zresztą wzór ten był sam w sobie paradoksalny, bowiem całkowita energia wypromieniowana

w jednostce czasu, równa $\int_0^\infty u(\nu) d\nu$ byłaby nieskończona. Paradoks ten nosi nazwę katastrofy w ultrafiolecie.

Jak pokazał Planck, promieniowanie ciała doskonale czarnego można opisać poprawnie, jeśli założyć, że atomy ciała drgające z częstotliwością ν mogą przyjmować taki jedynie stan ruchu, by w każdej chwili ich energia była równa całkowitej wielokrotności wielkości $h \cdot \nu$. Promieniując lub pochłaniając promieniowanie atom taki może zmieniać swoją energię skokowo — znów o wielokrotność „kwantu energii” $h \cdot \nu$. Z założenia tego Planck wyprowadził swój słynny wzór

$$u(\nu) = \frac{8\pi^2\nu^2}{c^3} \frac{h\nu}{\exp(h\nu/kT) - 1},$$

który doskonale zgadzał się z wynikami doświadczalnymi.

Zakaz przyjmowania energii różnej od $n \cdot h \cdot \nu$ (n — liczba naturalna) był niezrozumiałym, wręcz paradoksalnym dodatkiem do doskonale niemal pięknego gmachu fizyki dziewiętnastowiecznej. Tymczasem coraz to nowe zjawiska świata obiektów mikroskopowych dawały się opisać przy założeniu takich niezrozumiałych, administracyjnych ograniczeń (zjawisko fotoelektryczne — Einstein, 1905; ciepło właściwe ciał w niskich temperaturach — Einstein, 1907; wreszcie model atomu — Bohr, 1913). Prawdziwa eksplozja zakazów i nakazów nastąpiła w optyce, gdzie przy pomocy takich pojęć jak „stany dozwolone” czy „przejścia wzbronione” wyjaśniano z powodzeniem strukturę widm atomowych. W systemie tym istniało wiele procesów fizycznych dających się opisać przez aparat matematyczny i pojęciowy teorii, jednak z niewiadomych powodów zabronionych przez jej część „administracyjną”. Ten stan rzeczy, zwany obecnie starszą teorią kwantów, nie mógł zadowolić fizyków. Poszukiwano teorii, w której obserwowane zakazy wynikałyby w konieczny sposób z samego sposobu opisu zjawisk.

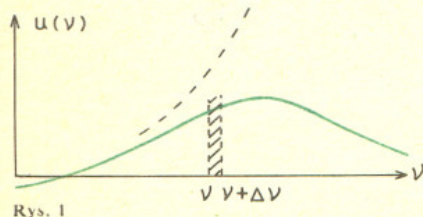
W roku 1925 ukazała się praca młodego podówczas fizyka z Getyngi, Wernera Heisenberga „Über quantentheoretische Umdeutung kinematischer und mechanischer Beziehungen”, a następnie dwie prace Heisenberga oraz Maxa Born’a i Pasquala Jordana zatytułowane „Zur Quantenmechanik”. Autorzy zaproponowali opis ruchu cząstek mikroskopowych w zupełnie nowym języku, różnym od mechaniki klasycznej. Z formalnego punktu widzenia nową teorię otrzymywali autorzy zastępując wielkości fizyczne takie jak położenie, pęd czy energia cząstki przez nieskończone macierze. Taka „mechanika macierzowa” pozwalała obliczyć poziomy energetyczne atomu wodoru czy oscylatora kwantowego (te ostatnie — nota bene —

wynosiły $\left(n + \frac{1}{2}\right) \cdot h \cdot \nu$ a nie — jak w hipotezie Plancka — $n \cdot h \cdot \nu$), jednak sens matematyczny i pojęciowy wprowadzonych macierzy był jeszcze dość długo niejasny.

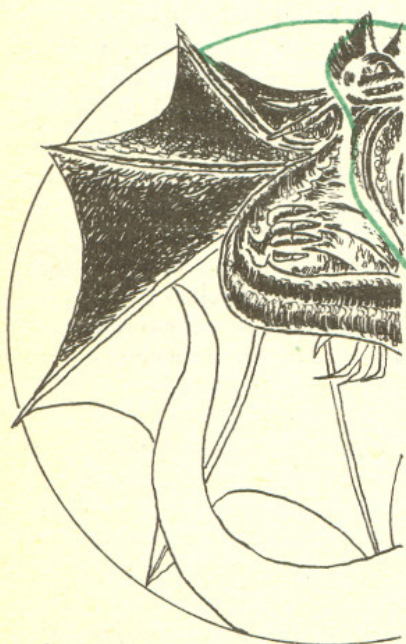
Tymczasem w roku 1926 pojawiła się seria prac Erwina Schrödingera, profesora na Politechnice w Zurychu, który proponował zastąpienie mechaniki klasycznej jeszcze inną teorią, nazwaną wkrótce mechaniką falową. Punktem wyjścia rozważań Schrödingera była zauważona już dawno analogia między mechaniką a optyką geometryczną.

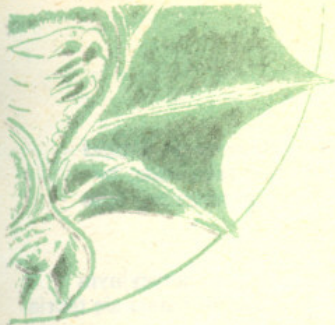
Optyka geometryczna to bardzo nietypowa teoria fizyczna. Właściwie — to bardziej matematyka niż fizyka. Opisuje ona bieg promieni świetlnych w sytuacjach, gdy własności optyczne ośrodka bardzo mało zmieniają się na odległościach rzędu wielu długości fali świetlnej. Kształtem promieni rządzi zasada Fermata. Głosi ona, że spośród wszystkich linii łączących dwa wybrane punkty ośrodka promień świetlny pobiegnie po tej, dla której tak zwana droga optyczna jest najmniejsza. Droga optyczna jest całką ze współczynnika załamania:

$$I = \int_{x_0}^{x_1} n(x) dx.$$



Rys. 1





Na przykład w ośrodku jednorodnym, gdy $n(x) = n = \text{const}$, droga optyczna I jest równa iloczynowi stałej n przez długość badanej linii. Oznacza to, że w tej sytuacji światło biegnie po liniach najkrótszych — prostych.

Stosując zasadę Fermata możemy projektować nawet bardzo skomplikowane urządzenia optyczne (obiektywy!) nie interesując się zupełnie fizyczną naturą światła.

W mechanice klasycznej istnieje ścisły odpowiednik zasady Fermata. Jest to tak zwana zasada Maupertuis-Lagrange'a. Głosi ona, że cząstka o masie m i energii całkowitej E , poruszająca się z punktu x_0 do x_1 pod wpływem pola sił o potencjale V wybiera taki tor, na którym całka

$$I = \int_{x_0}^{x_1} \sqrt{\frac{2}{m} (E - V(x))} ds$$

przybiera wartość minimalną. Wielkość $\sqrt{\frac{2}{m} (E - V(x))}$ odgrywa tu jak gdyby rolę

współczynnika załamania w punkcie x dla cząstek o energii E .

Optyka geometryczna nie nadaje się do opisu zjawisk interferencji i dyfrakcji, w których ujawnia się mikroskopowa struktura światła. Być może — rozumował Schrödinger — mechanika klasyczna jest właśnie „optyką geometryczną” prawdziwej, mikroskopowej mechaniki. Taki punkt widzenia sugerowały zresztą wcześniejsze prace młodego fizyka francuskiego, Louis de Broglie'a, który postulował rozważanie tak zwanych „fal materii”. Schrödinger postawił sobie za cel znalezienie takiej „optyki fizycznej fal materii”, dla której „optyką geometryczną” byłaby mechanika klasyczna.

Zakładając możliwie najprostszą postać praw takiej optyki otrzymał równanie różniczkowe, zwane obecnie równaniem Schrödingera, którego rozwiązania opisywały na przykład poziomy energetyczne elektronu w atomie wodoru.

W kilka miesięcy później Schrödinger zauważył, że jego teoria — z pozoru tak różna od podejścia Heisenberga — jest w istocie rzeczy równoważna mechanice macierzowej.

Równoważność tę zbadali dokładniej P. A. M. Dirac i P. Jordan. Okazało się, że wszystkim obiektom matematycznym występującym w jednej z tych teorii można było przyporządkować odpowiednie obiekty drugiej i to w taki sposób, by równania opisujące prawa ruchu w obu teoriach przeszły wzajemnie na siebie.

Sytuacja była jednak bardzo dramatyczna. Rozwiązując równania Schrödingera lub Heisenberga otrzymywano na przykład częstotliwości linii widmowych prostych pierwiastków, jednak ani macierze Heisenberga, ani tajemnicza funkcja falowa ψ Schrödingera nie miały jasnej interpretacji fizycznej. Równania opisywały jakieś drgania — co gorsza zespolone, bowiem funkcja falowa przybiera wartości zespolone — ale nie wiadano co tam mianowicie drga. Wkrótce Schrödinger zauważył — ciągle jeszcze w tym samym, przełomowym roku 1926 — że z jego równania wynika prawo zachowania całki (rozciągniętej po całej przestrzeni) z kwadratu modułu funkcji falowej:

$$(1) \quad \frac{d}{dt} \int |\psi(x, t)|^2 dx = 0.$$

Zaproponował zatem, by $|\psi(x, t)|^2$ uważać za gęstość ładunku elektrycznego, dzięki czemu równanie (1) wyrażałoby prawo zachowania całkowitego ładunku. Elektron byłby więc grudką naładowanej materii o gęstości $|\psi|^2$.

Taka interpretacja nie dała się jednak utrzymać. Przeczyło jej zbyt wiele faktów. Omówimy tu najprostszy z nich, tzw. rozpyływanie się paczki falowej.

Otóż w nieobecności sił zewnętrznych rozwiązania równania Schrödingera rozpyływają się po całej przestrzeni niemal jak gaz doskonały, któremu dać dostatecznie dużo miejsca na rozprężenie. Istnieje na przykład rozwiązanie, dla którego rzekoma gęstość ładunku ma postać

$$|\psi(x, t)|^2 = \frac{1}{\sigma \sqrt{2\pi}} \left(1 + \frac{\hbar^2}{4m^2\sigma^4} t^2 \right)^{-\frac{1}{2}} \exp \left\{ \frac{-x^2}{2\sigma^2} \frac{1}{1 + \frac{\hbar^2}{4m^2\sigma^4} t^2} \right\},$$

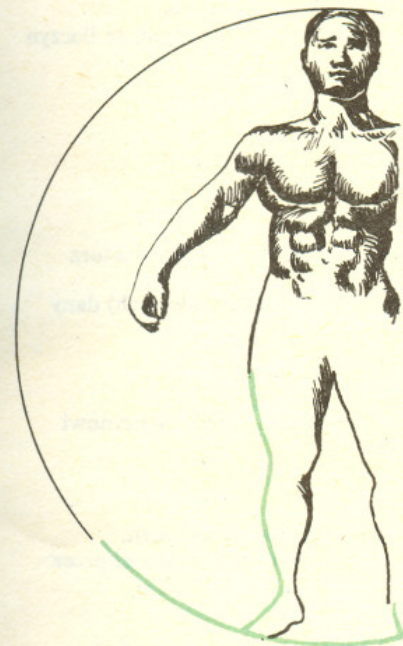
gdzie m jest masą elektronu, $\hbar = \frac{h}{2\pi}$, zaś σ jest dowolną stałą o wymiarze długości. Poznajemy tu rozkład normalny Gaussa o dyspersji równej

$$\sigma \cdot \sqrt{1 + \frac{\hbar^2}{4m^2\sigma^4} t^2}.$$

Gdyby nawet σ (dyspersja w chwili $t = 0$) była mikroskopowo mała, to i tak po dostatecznie długim czasie możemy uzyskać dowolnie, makroskopowo wielką dyspersję. Dokonując pomiarów na takim elektronie, rozmytym na przykład na obszarze stu kilometrów, powinniśmy wykrywać zawsze tylko tę część ładunku, która leży w zasięgu naszej aparatury. Tymczasem nigdy nie zarejestrowano części ładunku elementarnego. We wszystkich znanych nam pomiarach ładunek elektronu pojawia się zawsze jako całość. Widać to dobrze w doświadczeniach z dyfrakcją elektronów na kryształach, w których najlepiej manifestują się „falowe” własności materii (patrz „Delta” 10/75).

Linia falista na rys. 2 przedstawia wartości funkcji $|\psi|^2$ w różnych punktach ekranu. Wartości te w uderzający sposób zgadzają się z gęstością zaczernienia emulsji fotograficznej, którą pokryto ekran, proporcjonalną do ilości elektronów padających w danym punkcie. Zaczernienie to składa się jednak z dużej ilości pojedynczych plamek, z których każda świadczy o zarejestrowaniu jednego, całego, mikroskopowo małego elektronu.

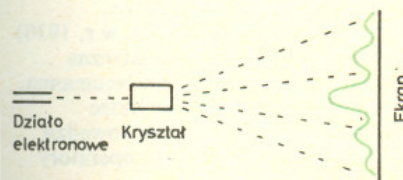
Zasadniczego przewrotu, który ostatecznie ukonstytuował mechanikę kwantową, zmieniając całkowicie sposób fizycznego opisu mikroświata dokonała hipoteza Maxa Borna o probabilistycznej interpretacji funkcji falowej. Zaproponował on — ciągle jeszcze w roku 1926 — by gęstość $|\psi(x, t)|^2$ interpretować jako gęstość prawdopodobieństwa zarejestrowania elektronu (ale całego elektronu) w chwili t w punkcie x .



Jeżeli układ jest scharakteryzowany przez pewną zmienną x , to prawdopodobieństwo wystąpienia w tym układzie fluktuacji, w wyniku których zmienna może zmieniać się w przedziale $(x, x + dx)$ jest dane przez rozkład Gaussa

$$d\omega = \frac{1}{\sqrt{2\pi\sigma^2}} \exp \left\{ -\frac{(x-x_0)^2}{2\sigma^2} \right\} dx,$$

gdzie kwadrat dyspersji (wariancji) $\sigma^2 = (x-x_0)^2$ równa się średniemu odchyleniu kwadratowemu od wartości średniej x_0 .



Rys. 2

(2)

$$P(x(t) \in \mathcal{D}) = \int_{\mathcal{D}} |\psi(x, t)|^2 dx$$

równa jest prawdopodobieństwu zarejestrowania elektronu w chwili t w obszarze $\mathcal{D}$. Dzięki równości (1) funkcja falowa daje się unormować, to znaczy pomnożyć przez taką stałą, by prawdopodobieństwo zarejestrowania elektronu gdziekolwiek (tzn. w całej przestrzeni) było stale równe jedności.

W tym momencie Rubikon został przekroczony. Fizyka mikroświata zerwała z tradycyjnym językiem, ukształtowanym dzięki potocznemu, makroskopowemu doświadczeniu i stała się teorią probabilistyczną. Dziwna to jednak probabilistyka — jakże różna od klasycznej teorii prawdopodobieństwa, której uczymy się w szkole. „Teatrem działań” jest tu (podajemy współczesną wersję teorii) tak zwana przestrzeń Hilberta — abstrakcyjna przestrzeń wektorowa nad ciałem liczb zespolonych, na ogół nieskończenie wymiarowa. Jako konkretne przedstawienie takiej przestrzeni można wziąć na przykład zbiór schrödingierowskich funkcji falowych. W przestrzeni Hilberta określony jest iloczyn skalarny, który parze wektorów ψ i φ przyporządkowuje liczbę zespoloną oznaczaną zazwyczaj (ψ, φ) . Zakłada się przy tym, że iloczyn skalarny jest:

1° liniowy w pierwszym czynniku:

$$(\alpha\psi_1 + \beta\psi_2, \varphi) = \alpha(\psi_1, \varphi) + \beta(\psi_2, \varphi)$$

2° zmienia swą wartość na liczbę sprzężoną przy zamianie kolejności czynników:

$$(\varphi, \psi) = \overline{(\psi, \varphi)}$$

3° dodatnio określony, to znaczy $(\psi, \psi) \geq 0$, przy czym równość zachodzi tylko dla wektora zerowego.

W przedstawieniu schrödingierowskim iloczyn skalarny dwu wektorów (funkcji falowych) dany jest całką z iloczynu:

$$(\psi, \varphi) = \int \psi(x) \bar{\varphi}(x) dx$$

(całka rozciągnięta na całą przestrzeń).

Przestrzeń Hilberta ma wiele cech znanych nam przestrzeni euklidesowych. Dzięki iloczynowi skalarnemu można na przykład określić pojęcie długości wektora:

$$\|\psi\| = \sqrt{(\psi, \psi)}.$$

Można też mówić o prostokątności wektorów, gdy ich iloczyn skalarny równy jest zeru.

Ogólniej — można określić kąt $\angle(\psi_1, \psi_2)$ między dwoma kierunkami reprezentowanymi przez wektory ψ_1 i ψ_2 . Przyjmujemy mianowicie

$$(3) \quad \cos \angle(\psi_1, \psi_2) = \frac{|(\psi_1, \psi_2)|}{\|\psi_1\| \cdot \|\psi_2\|}.$$

Przez „kierunek” rozumiemy tutaj całą jednowymiarową podprzestrzeń złożoną ze wszystkich wektorów proporcjonalnych do danego. Kierunek można reprezentować na przykład wektorem o długości jednostkowej: $\|\psi\| = 1$.

Wybór takiego reprezentanta oznacza, w języku funkcji falowych, spełnienie warunku unormowania gęstości prawdopodobieństwa, bowiem

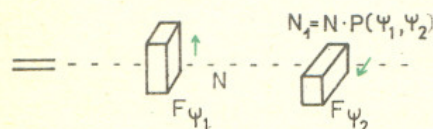
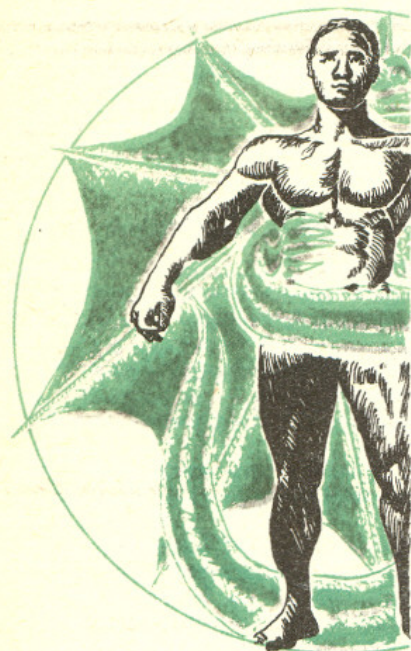
$$\|\psi\| = \int |\psi(x)|^2 dx.$$

W ten sposób mechanika kwantowa stała się jak gdyby geometrią przestrzeni Hilberta. Stany opisywanego układu fizycznego (a przynajmniej tak zwane stany czyste) odpowiadają właśnie kierunkom w przestrzeni Hilberta. Każdym dwu stanom można przypisać liczbę dodatnią zawartą między zerem a jednością, zwaną prawdopodobieństwem przejścia. Jest ona równa kwadratowi kosinusa kąta zawartego między nimi. Reprezentując oba stany (kierunki) wektorami unormowanymi ($\|\psi_1\| = \|\psi_2\| = 1$) otrzymujemy zgodnie z (3) następującą wartość prawdopodobieństwa przejścia:

$$P(\psi_1, \psi_2) = |(\psi_1, \psi_2)|^2.$$

Aby zrozumieć sens fizyczny liczby $P(\psi_1, \psi_2)$ wyobraźmy sobie filtr F_{ψ_1} (rys. 3), który z każdej wiązki elektronów przepuszcza tylko te, które są w stanie ψ_1 . Niech — dla ustalenia uwagi — ψ_1 oznacza stan, w którym spin elektronu jest skierowany w górę, w kierunku osi z . Filtrując dowolny strumień elektronów przez F_{ψ_1} , otrzymujemy wiązkę w stanie czystym, w której wszystkie spiny są skierowane w górę. Niech teraz ψ_2 oznacza stan, odpowiadający innej, ale ustalonej konfiguracji spinu zaś F_{ψ_2} — odpowiedni filtr. Filtr ten jest oczywiście tym samym urządzeniem fizycznym co F_{ψ_1} , tylko nieco obróconym przestrzennie. Jak się okazuje, pewna część elektronów wiązki przefiltrowanej przez F_{ψ_1} (a więc będących w stanie ψ_1) przejdzie również przez F_{ψ_2} , co oznacza, że są one w stanie ψ_2 . Jak się obecnie wydaje, w przypadku pojedynczego elektronu nie potrafimy nigdy przewidzieć czy przejdzie on przez F_{ψ_2} czy nie. Natomiast prawdopodobieństwo takiego przejścia jest równe właśnie $P(\psi_1, \psi_2)$.

Taka kwantowa teoria prawdopodobieństwa wykazuje wiele zadziwiających własności, nie znanych w klasycznym rachunku prawdopodobieństwa. Trudno omówić je w krótkim artykule. Warto jednak wspomnieć, że ostateczny kształt pojęciowy zawdzięcza ona w dużej mierze P. A. M. Diracowi, którego piękna książka „Principles of Quantum Mechanics” (I wyd. w r. 1930) była źródłem natchnienia dla wielu pokoleń fizyków i matematyków. Struktura matematyczna nowej teorii, przez długi czas nie w pełni zrozumiała, została zbadana przez Johna von Neumanna, który właśnie pod koniec lat 20-tych był asystentem wielkiego Dawida Hilberta w Getyndze i tam zetknął się osobiście z twórcami mechaniki kwantowej. Właśnie von Neumann wprowadził pojęcie abstrakcyjnej przestrzeni Hilberta, zinterpretował „macierze” Heisenberga jako operatory w tej przestrzeni i zbadał ich własności, co dało początek nowemu działowi analizy funkcjonalnej.



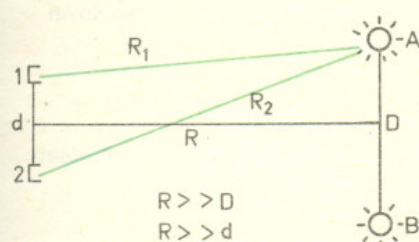
Rys. 3

Ciekawe, że sam Hilbert w swej teorii równań całkowych, opublikowanej dwadzieścia lat przed powstaniem mechaniki kwantowej, wprowadził pojęcie tak zwanego widma formy kwadratowej. Jak się okazało, energia atomu wyraża się właśnie przez formę kwadratową, której widmo (w sensie Hilberta) odpowiada obserwowanemu w spektroskopie widmu promieniowania wysyłanego przez świecące atomy. Genialna intuicja Hilberta objawiła się jeszcze raz w roku 1925, gdy Born i Jordan próbowali zasięgnąć jego rady w sprawie pewnych trudności matematycznych związanych z macierzami nieskończonymi. Hilbert odpowiedział wówczas, że jedynym miejscem, gdzie takie macierze pojawiają się w sposób naturalny, jest teoria równań. Nakłaniał więc fizyków do szukania odpowiedniego równania różniczkowego — właśnie późniejszego równania Schrödingera — jednak Born i Jordan nie potraktowali wówczas tej rady poważnie. Mechanika kwantowa, a nawet jej wersja relatywistyczna podana w roku 1928 przez Diraca, nie rozwiązała bynajmniej wszystkich trudności opisu mikroświata. Wiele kłopotów filozoficzno-pojęciowych nastroczał fakt, że funkcję falową potrafimy przypisać jedynie zespołom statystycznym — na przykład zbiorowi cząstek produkowanych przez dany akcelerator, przy ustalonych parametrach makroskopowych, takich jak napięcia przyspieszające, szerokości szczelin kolimatorów itp. Natomiast próby zdefiniowania funkcji falowej pojedynczego mikroobiekta — na przykład przypadkowo zbląkanej cząstki promieniowania kosmicznego — prowadziły do paradoksalnych wniosków — słynne są paradoksy Einsteina-Rosena-Podolsky'ego (patrz artykuł G. Białkowskiego) oraz Schrödingera. Jak dotychczas jednak próby zastąpienia mechaniki kwantowej czymś innym nie powiodły się. Wyjaśniając ogromną ilość zjawisk z tak różnych dziedzin jak optyka, teoria ciała stałego, fizyka jądrowa jest ona obecnie „wzorcową” teorią mikroskopową. Na jej obraz i podobieństwo budujemy teorie opisujące systemy nawet dużo bardziej skomplikowane niż mechaniczne układy kilku cząstek (pole elektromagnetyczne). Kryje się w tym jednak ogromna ilość zagadek i niejasności dotychczas nie rozwiązanych.



Zadania

Redaguje dr Andrzej ZIEMIŃSKI



F 33. Dwa niezależne monochromatyczne źródła światła A i B (patrz rysunek) emitują fotony, które są następnie rejestrowane przez detektory 1 i 2. Odległości między źródłami oraz między detektorami są znacznie mniejsze niż odległości źródeł od detektorów. Należy znaleźć prawdopodobieństwo jednoczesnego (w czasie Δt) zarejestrowania koincydencji fotonów przez oba detektory. Pokazać, że badając zależność liczby koincydencji od odległości między detektorami można zmierzyć odległość między dwoma odległymi źródłami światła.

(A. Para)

Wskazówka: Amplitudę rejestracji przez detektor 1 (w czasie Δt) fotonu wysłanego ze źródła A można zapisać w postaci $f_{A1} = ce^{ikr}$, gdzie r — odległość detektora od źródła, i — jednostka urojona, zaś c i k — stałe, przy czym stała c jest na ogół zespolona. Oznacza to, że prawdopodobieństwo zarejestrowania w detektorze 1 fotonu pochodzącego ze źródła A wynosi $P_{A1} = |f_{A1}|^2$ ($|a|^2 = a \cdot \bar{a}$, $\bar{a}$ — liczba zespolona sprzężona z a).

Przed rozwiązywaniem zadania Czytelnik powinien przeczytać artykuły G. Białkowskiego i J. Kijowskiego.

Rozwiązanie na str. 17

Redaguje mgr Andrzej MAKOWSKI

M 97. Udowodnić, że jeżeli liczby dodatnie x, y, z są kolejnymi wyrazami ciągu geometrycznego, liczby zaś a, b, c — ciągu arytmetycznego, to

$$x^b y^c z^a = x^c y^a z^b.$$

Rozwiązanie na str. 2

M 98. W pewnym państwie, w którym jest $n \geq 2$ miast, wybudowano sieć dróg łączących parami te miasta, przy czym

1° drogi łączące różne pary miast nie krzyżują się (ale mogą być skrzyżowania dwupoziomowe, jednak bez możliwości zjechania z jednej drogi na inną),

2° z każdego miasta można dojechać do każdego innego jedną tylko trasą.

Udowodnić, że omawiana sieć składa się z $n-1$ odcinków łączących parami miasta.

Rozwiązanie na str. 14

M 99. Udowodnić, że każda potęga o wykładniku naturalnym dowolnego wyrazu ciągu arytmetycznego o pierwszym wyrazie a^2 i różnicy $2a$, gdzie a jest liczbą całkowitą, jest wyrazem tego ciągu.

Rozwiązanie na str. 14

Jednym z podstawowych wzorów trygonometrycznych jest *twierdzenie kosinusów* podające zależność między bokami trójkąta a jednym z jego kątów:

$$c^2 = a^2 + b^2 - 2ab \cos C.$$

Na formułę tę można patrzeć jako na uogólnienie twierdzenia Pitagorasa (do którego sprowadza się, gdy kąt C jest prosty, czyli $\cos C = 0$).

Dla bliższego zbadania geometrycznego i algebraicznego sensu twierdzenia kosinusów użyteczny jest zapis wektorowy, w którym boki trójkąta będą reprezentowane przez wektory $\mathbf{a}, \mathbf{b}$ i $\mathbf{a}-\mathbf{b}$.

Oznaczając jak zwykle przez $|\mathbf{a}|$ długość wektora $\mathbf{a}$ otrzymamy wzór:

$$|\mathbf{a}-\mathbf{b}|^2 = |\mathbf{a}|^2 + |\mathbf{b}|^2 - 2|\mathbf{a}||\mathbf{b}|\cos \angle(\mathbf{a}, \mathbf{b}).$$

Nazwijmy wyrażenie $|\mathbf{a}||\mathbf{b}|\cos \angle(\mathbf{a}, \mathbf{b})$ *iloczynem skalarnym wektorów $\mathbf{a}$ i $\mathbf{b}$* , oznaczmy je w skrócie symbolem $(\mathbf{a}, \mathbf{b})$ i zbadajmy pewne szczególne własności tej funkcji.

Łatwo sprawdzić, kładąc we wzorze kosinusów $\mathbf{a} = \mathbf{b}$, że

$$(0) \quad (\mathbf{a}, \mathbf{a}) = |\mathbf{a}|^2,$$

co pozwala przepisać ten wzór w postaci

$$(\mathbf{a}-\mathbf{b}, \mathbf{a}-\mathbf{b}) = (\mathbf{a}, \mathbf{a}) + (\mathbf{b}, \mathbf{b}) - 2(\mathbf{a}, \mathbf{b}),$$

a to zaczyna przypominać wzór na kwadrat różnicy. Nieco bliższe sprawdzenie przekona nas, że wyrażenie $(\mathbf{a}, \mathbf{b})$ ma wiele formalnych cech iloczynu (mówimy, że jest dwuliniowe), mianowicie:

$$(1) \quad (\mathbf{a}, \mathbf{b}+\mathbf{c}) = (\mathbf{a}, \mathbf{b}) + (\mathbf{a}, \mathbf{c}),$$

$$(2) \quad (\mathbf{a}, \mathbf{b}) = (\mathbf{b}, \mathbf{a})$$

oraz gdy λ jest liczbą rzeczywistą, to

$$(3) \quad (\lambda \mathbf{a}, \mathbf{b}) = \lambda(\mathbf{a}, \mathbf{b}).$$

Ponadto z równości (0) wynika, że

$$(4) \quad (\mathbf{a}, \mathbf{a}) \geq 0 \text{ i } (\mathbf{a}, \mathbf{a}) = 0 \Leftrightarrow \mathbf{a} = \mathbf{0},$$

a wyrażenie $\sqrt{(\mathbf{a}, \mathbf{a})}$ jest długością wektora $\mathbf{a}$.

Wynika stąd dalej, że odległość euklidesowa końców wektorów $\mathbf{a}$ i $\mathbf{b}$ zaczepionych we wspólnym początku (nazywamy ją odległością tych wektorów) wyraża się wzorem

$$\varrho(\mathbf{a}, \mathbf{b}) = \sqrt{(\mathbf{a}-\mathbf{b}, \mathbf{a}-\mathbf{b})}.$$

Znając własności (0)–(3) i wiedząc, że wersory (wektory kierunkowe) osi układu współrzędnych w n -wymiarowej przestrzeni euklidesowej są prostopadłe i mają długość 1, możemy również wyprowadzić wzór analityczny na iloczyn skalarny:

$$(\mathbf{x}, \mathbf{y}) = x_1 y_1 + \dots + x_n y_n = \sum_{i=1}^n x_i y_i,$$

gdzie $\mathbf{x} = (x_1, \dots, x_n)$, $\mathbf{y} = (y_1, \dots, y_n)$ w danym układzie współrzędnych.

Zastanówmy się, co by było, gdybyśmy zaczęli rozważać „uogólniony iloczyn skalarny”, tzn. dowolną funkcję $(\mathbf{x}, \mathbf{y})$ przyporządkowującą parom wektorów liczby i spełniającą warunki (1)–(4). Okaże się, że taki iloczyn analitycznie będzie się wyrażał wzorem

$$(\mathbf{x}, \mathbf{y}) = \sum_{i,j=1}^n a_{ij} x_i y_j$$

(Układ liczb a_{ij} będzie spełniał pewne warunki, z których jeden: $a_{ij} = a_{ji}$ wynika bezpośrednio z wzoru (2), a inne są nieco bardziej skomplikowane. Geometrycznie sprowadzają się one do tego, że wszystkie kule muszą być podobnymi elipsoidami). Powstaje pytanie, co właściwie uzyskujemy wprowadzając taką funkcję. Odpowiedź jest nieco zaskakująca: uzyskujemy mianowicie „całość” geometrii euklidesowej. Możemy bowiem przy pomocy takiej funkcji określić:

— nową długość wektora wzorem

$$||\mathbf{x}|| \stackrel{\text{df}}{=} \sqrt{(\mathbf{x}, \mathbf{x})},$$

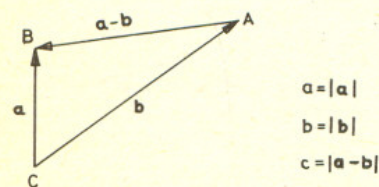
(a więc i pewną nową metrykę, dlaczego?);

— nowy kosinus kąta między wektorami — wzorem

$$\cos \angle(\mathbf{x}, \mathbf{y}) \stackrel{\text{df}}{=} \frac{(\mathbf{x}, \mathbf{y})}{||\mathbf{x}|| \cdot ||\mathbf{y}||}$$

(a więc kąt wypukły między wektorami $\mathbf{x}, \mathbf{y}$);

— nową prostopadłość wektorów, (tym samym również prostych, płaszczyzn itp) — mówiąc, że $\mathbf{x} \perp \mathbf{y}$ wtedy i tylko wtedy, gdy $(\mathbf{x}, \mathbf{y}) = 0$.



Można rozważać sytuacje jeszcze ogólniejsze, w których wektory mnoży się przez liczby zespolone (por. artykuł J. Kijowskiego). Warunek (2) zastępuje się wtedy warunkiem

$$(2') \quad (\mathbf{x}, \mathbf{y}) = \overline{(\mathbf{y}, \mathbf{x})}$$

(kreska jest symbolem sprzężenia liczby zespolonej).

Np. na płaszczyźnie $\mathbf{e}_1 = (1, 0)$ jest wersorem pierwszej osi, a $\mathbf{e}_2 = (0, 1)$ — wersorem drugiej osi. Zatem jeśli $\mathbf{x} = (x_1, x_2)$ i $\mathbf{y} = (y_1, y_2)$ to $\mathbf{x} = x_1 \cdot \mathbf{e}_1 + x_2 \cdot \mathbf{e}_2$, $\mathbf{y} = y_1 \cdot \mathbf{e}_1 + y_2 \cdot \mathbf{e}_2$. Z własności iloczynu skalarnego wynika, że

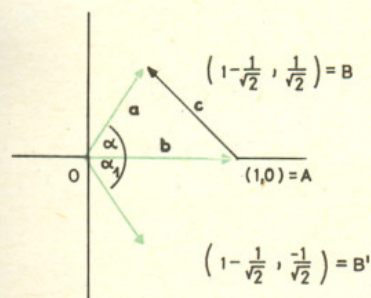
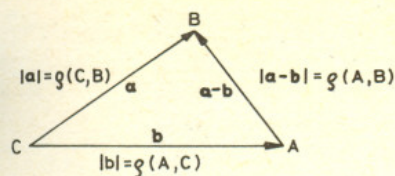
$$(\mathbf{x}, \mathbf{y}) = x_1 y_1 (\mathbf{e}_1, \mathbf{e}_1) + x_1 y_2 (\mathbf{e}_1, \mathbf{e}_2) + x_2 y_1 (\mathbf{e}_2, \mathbf{e}_1) + x_2 y_2 (\mathbf{e}_2, \mathbf{e}_2) = x_1 y_1 + x_2 y_2$$

bo $(\mathbf{e}_1, \mathbf{e}_1) = 1 = (\mathbf{e}_2, \mathbf{e}_2)$ oraz $(\mathbf{e}_1, \mathbf{e}_2) = 0$ (bo są prostopadłe).

Przykład: niech $(\mathbf{x}, \mathbf{y}) = \frac{1}{4} (x_1 - x_2)(y_1 - y_2) + \frac{1}{9} (x_1 + x_2)(y_1 + y_2)$.

Wtedy $||\mathbf{x}||^2 = \frac{1}{4} (x_1 - x_2)^2 + \frac{1}{9} (x_1 + x_2)^2$, skąd wynika, że np. „kula jednostkowa”, czyli $\{\mathbf{x}: ||\mathbf{x}|| < 1\}$ będzie elipsą

o długościach półosi 1 i $\frac{3}{2}$, której dłuższa oś leży na prostej $x_1 = x_2$.



Mamy: $\rho(A, O) = \rho(O, B) = 1$, $\rho(A, B) =$
 $= \left| 1 - \frac{1}{\sqrt{2}} - 1 \right| + \left| \frac{1}{\sqrt{2}} - 0 \right| = \sqrt{2}$.

Zatem wzór kosinusów jest następujący:

$$(\sqrt{2})^2 = 1^2 + 1^2 - 2 \cdot 1 \cdot 1 \cdot \cos \alpha.$$

Stąd $2\cos \alpha = 0$, $\alpha = -\pi/2$. Analogicznie dla α' .

Dla β mamy $\rho(B, B') = \sqrt{2}$, stąd również $\cos \beta = 0$!

Tu można by wysunąć jedno zastrzeżenie: przecież jeżeli mamy np. na płaszczyźnie dowolną metrykę (a wiele przykładów różnych metryk na płaszczyźnie podaje np. M. Moszyńska w artykule o przestrzeniach metrycznych, Delta 1/1975), to wypisany na początku wzór kosinusów, w którym liczby a , b i c traktujemy jako znane odległości trzech danych punktów, pozwala bezpośrednio wyznaczyć kosinus kąta między bokami o długościach a i b .

Takie podejście nie prowadzi na ogół niestety do pozytywnych rezultatów. Przy pewnych metrykach przestaje bowiem być prawdziwe twierdzenie, że dwa wektory prostopadłe do trzeciego są równoległe. Na przykład na płaszczyźnie z metryką określoną wzorem

$$\varrho(x, y) = |x_1 - y_1| + |x_2 - y_2|$$

wektory $\left(1 - \frac{1}{\sqrt{2}}, \frac{1}{\sqrt{2}}\right)$ oraz $\left(1 - \frac{1}{\sqrt{2}}, -\frac{1}{\sqrt{2}}\right)$ byłyby przy takiej definicji kąta

prostopadłe, a jednocześnie każdy z nich byłby prostopadły do wektora $(1, 0)$! Szkic dowodu — obok.

Istnieje jednak pewne bardzo proste kryterium geometryczne pozwalające odróżniać metryki dające się określić za pomocą pewnego iloczynu skalarnego od innych metryk. Zachodzi mianowicie

Twierdzenie. Jeżeli w danej metryce każde dwie kule są (w metryce euklidesowej) jednokładnymi elipsoidami, a kule o równych promieniach są przystające, to istnieje iloczyn skalarny (x, y) wyznaczający tę metrykę.

Mamy więc sposób pozwalający na odróżnianie metryk „dobrych”, czyli takich, które można traktować jako fragment pełnej geometrii euklidesowej, od metryk „przypadkowych”.

Zauważmy teraz, że w warunkach, które powinien spełniać nasz iloczyn skalarny, występowało jedynie dodawanie wektorów i ich mnożenie przez stałe. Łatwo zauważyć, że płaszczyzna i przestrzeń trójwymiarowa nie są jedynymi zbiorami, w których rozważa się operacje tego typu. Działanie dodawania i mnożenia przez liczbę wprowadza się między innymi w zbiorach funkcji określonych na pewnym ustalonym zbiorze.

Dla przykładu weźmy pod uwagę zbiór wszystkich funkcji określonych i ciągłych na przedziale $\langle 0, 2\pi \rangle$ i przyjmujących tę samą wartość na końcach przedziału ($f(0) = f(2\pi)$).

Powstaje naturalne pytanie: czy można tak rozszerzyć pojęcie iloczynu skalarnego, by obejmowało ono również przypadek takich funkcji? Okazuje się, że wyrażenie

$$(*) \quad (f, g) = \int_0^{2\pi} f(x)g(x)dx$$

przyporządkowuje parom funkcji liczby w taki sposób, że spełnione są wszystkie warunki (1)–(4), a więc (f, g) można nazwać iloczynem skalarnym funkcji f i g .

Sprawdzenie, że rzeczywiście warunki te są spełnione, jest dość żmudne rachunkowo i pominiemy je tutaj. Ważna dla nas jest konsekwencja tego faktu: w przestrzeni funkcji ciągłych na przedziale $\langle 0, 2\pi \rangle$ można uprawiać „geometrię euklidesową”. Teraz łatwo już sobie wyobrazić, że istnieje wiele przestrzeni, w których można wprowadzać iloczyn skalarny. Nazywa się je przestrzeniami unitarnymi lub — jeśli spełniają pewne dodatkowe warunki — przestrzeniami Hilberta.

Ten dodatkowy warunek — to *zupełność*.

Mówi się, że przestrzeń metryczna jest zupełna, jeśli każdy ciąg (x_n) spełniający tzw. warunek Cauchy'ego

$$\lim_{n, m \rightarrow \infty} \rho(x_n, x_m) = 0$$

jest zbieżny w tej przestrzeni. Przykładem przestrzeni metrycznej niezupełnej jest odcinek otwarty $(0, 1)$.

Ciąg $\left(\frac{1}{n}\right)$ spełnia warunek Cauchy'ego

$$\lim_{n, m \rightarrow \infty} \left| \frac{1}{n} - \frac{1}{m} \right| = 0$$

a jednocześnie ciąg ten nie jest zbieżny w odcinku $(0, 1)$, bo jego granica — liczba zero — do tego odcinka nie należy.

Spośród wielu zastosowań tej „geometrii euklidesowej” wymienimy jedno o poważnych konsekwencjach dla analizy matematycznej:

Twierdzenie. W zdefiniowanej powyżej przestrzeni funkcji ciągłych z iloczynem skalarnym określonym wzorem (*) funkcje $\cos nx$ i $\sin nx$ oraz funkcja stała ($n \in \mathbb{N}$) są parami wzajemnie prostopadłe.

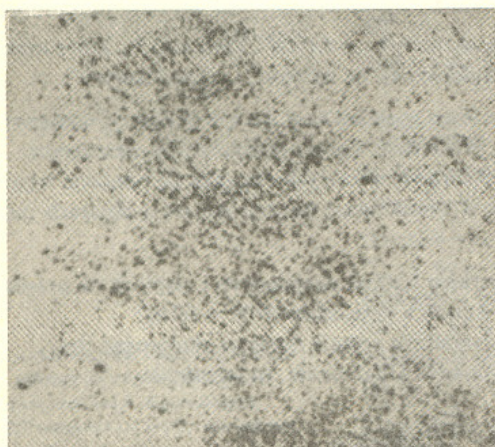
Dowód tego faktu oraz jego konsekwencje, z których m.in. wynika, że można ten układ funkcji traktować jako swego rodzaju „układ współrzędnych” w rozważanej przestrzeni, znajdzie Czytelnik w artykule W. Szlenka o szeregach Fouriera (Delta 4/1976).

Zadanie. Obliczyć długość boków i kąty trójkąta o wierzchołkach 1 , x , x^2 w przestrzeni funkcji ciągłych na przedziale $\langle 0, 1 \rangle$ z iloczynem skalarnym

$$(f, g) = \int_0^1 f(x) \cdot g(x) dx.$$

O innych zastosowaniach pisze J. Kijowski w tym numerze.

mała delta



$1,2 \times 10^4$ fotonów



$9,3 \times 10^4$ fotonów



$7,6 \times 10^4$ fotonów

Opowiadanie palnika gazowego

Gdzie te dawne dobre czasy, kiedy zmęczeni i przemarznięci myśliwi wracali z ubitym niedźwiedziem z polowania i grzali się przed ogniskiem, wielbiąc boską, życiodajną moc ognia? Co z tego zostało w Waszych miejskich mieszkaniach? Grzeją Was głupie żelazne rury — kaloryfery, świecą Wam nudne żarówki. Nawet zapalek nie chcecie już używać, bo macie jakieś nowe elektryczne pstrykawki. Zostały Wam tylko w kuchni cztery gazowe palniki. Ich słabo świecące płomyczki gotują Wam na każde zamówienie mleko nie od krowy, tylko ze sklepu, i te nędzne ochłapki mięsa, którymi się żywicie.

A musisz wiedzieć, że chociaż świecenie płomienia ludzie obserwowali od setek tysięcy lat, to kiedy osiemdziesiąt lat temu zaczęli się nad nim dokładniej zastanawiać, wywołało to jeden z największych przewrotów w rozumieniu świata.

Na czym w ogóle polega świecenie? Na tym, że z ciała świecącego wybiegają bardzo małe cząsteczki, które fizycy nazywają fotonami. Te cząsteczki mają dziwne własności. Nie mogą w ogóle być nieruchome, muszą biec z ogromną prędkością. Prędkość ta wynosi 300 000 kilometrów na sekundę. To znaczy, że w sekundę dobiegają z Ziemi prawie do Księżyca. W sekundę mogłyby też obieć całą Ziemię wzdłuż równika siedem razy dookoła. Kiedy fotony dobiegną do Twojego oka „giną” w nim, a Ty widzisz światło. Kiedy foton dobiegnie do kliszy fotograficznej, w miejscu, w które trafił, powstaje czarny punkt. Kiedy upadnie na kliszę wiele fotonów, zacernienie może przedstawić jakiś obraz i w ten sposób powstaje zdjęcie. Popatrz na fotografie obok: przedstawiają one fragmenty kliszy fotograficznej, zacernione przez różne liczby fotonów. Jeżeli jest ich mało, rysunek jest niewyraźny, bardziej przypomina rozsypany piasek, niż coś konkretnego. Dla dużej liczby widać, że jest to buzia dziewczynki.

Fotony mogą być różne i wywołują w oku różne wrażenia. Jedne powodują wrażenie barwy czerwonej, inne pomarańczowej, żółtej, zielonej, niebieskiej czy fioletowej. Są zresztą i takie, których w ogóle nie widzimy. Skąd biorą się w palniku fotony? „Rodzą się” w nim. Mogą one powstawać w każdym rozgrzanym ciele, na przykład we włóknie żarówki elektrycznej. W płomieniu gazowym sytuacja jest jednak bardzo specjalna. Jak wiesz, temperatura jest tam bardzo wysoka, a poza tym płomień jest gazem. W tak rozgrzanym gazie cząsteczki poruszają się bardzo szybko i często zderzają ze sobą.



Pod wpływem zderzeń rozpadają się na poszczególne atomy. A więc w płomieniu świecą oddzielne atomy, które się tam właśnie znalazły. Fotony „rodzą się” w atomach. A każdy atom wysyła fotony specjalne, charakterystyczne dla siebie. Posyp palący się palnik zwykłą solą, a zobaczysz, że zabarwia się na żółto. To zaczynają świecić atomy sodu, z których między innymi zbudowana jest sól kuchenna. Wsadź do palnika drucik miedziany (trzymaj go obciążkami, żebyś nie poparzył palców!). Możesz go przedtem zanurzyć na chwilę w occie, żeby efekt był wyraźniejszy. Atomy miedzi w palniku będą świecić zielono. A może pamiętasz fioletowe płomyczki dogasającego ogniska? To świecą jeszcze inne atomy — atomy potasu. Tak więc obserwując płomień palnika gazowego można dowiedzieć się wielu rzeczy o własnościach atomów.

Punkt, odcinek, trójkąt...

Chociaż są to pojęcia abstrakcyjne (bo przecież nikt nie widział ani punktu, ani odcinka), przemawiają dobrze do wyobraźni i zgadzają się ze zdrowym rozsądkiem. I aż dziw bierze, jak wiele wokół nas zjawisk, które zdają się ostrzegać: uwaga, to co wydaje się takie oczywiste, wcale nie musi być prawdziwe.

Sprawa pozornie jasna — dwa punkty, nawet bardzo bliskie, można połączyć odcinkiem punktów pośrednich. Tymczasem ... na reproduktowanej w gazecie fotografii można odszukać takie punkty, między którymi nie ma już niczego. Obejrzyjcie taką fotografię pod lupą. Składa się ona ze skończonej ilości czarnych plamek, co widać wyraźnie w powiększeniu. Mimo to odbieramy ją normalnie i na ogół rozpoznajemy, co lub kogo przedstawia.

Płaszczyzna fotografii reproduktowanej w gazecie składa się z izolowanych punktów. Matematycy powiedzieliby o niej, że jest dyskretna. Dyskretny (w matematycznym sensie) jest również obraz na ekranie telewizora, dyskretnym jest także koncert fortepianowy (pod względem wysokości dźwięków, a także ... czasu odmierzanego nutami, półnutami, ćwiartkami, ósemkami itd.)

Na dobrą sprawę żyjemy równolegle w dwóch światach: jeden jest gęsty, spójny, z płynnymi zmianami, a drugi dyskretny. Umiemy jednak jakoś je pogodzić. Zawdzięczamy to właściwości naszych oczu, uszu i mózgu, że izolowane wrażenia odbieramy jako połączone w sposób ciągły. I całe szczęście, gdyż w przeciwnym wypadku musielibyśmy się obyć bez wielu przyjemności. Na przykład w kinie nie widzielibyśmy płynnego ruchu, tylko ciąg pojawiających się i znikających nieruchomych obrazów.



Oto przykład ciekawych zależności. Figura na rysunku to koło o środku w początku układu współrzędnych i o promieniu 10.

Współrzędne punktów jego dyskretnego obrazu, czyli współrzędne grubych kropek, spełniają zależność:

$$x^2 + y^2 \leq 10^2$$

(trzeba przypomnieć sobie twierdzenie Pitagorasa).

Grubych kropek będzie tyle, ile jest rozwiązań całkowitych nierówności:

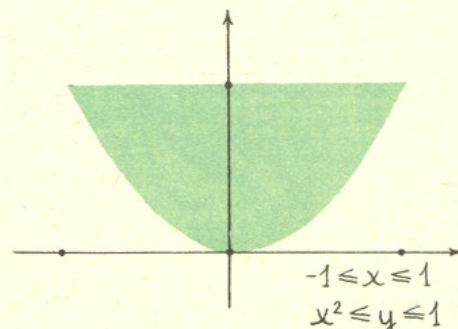
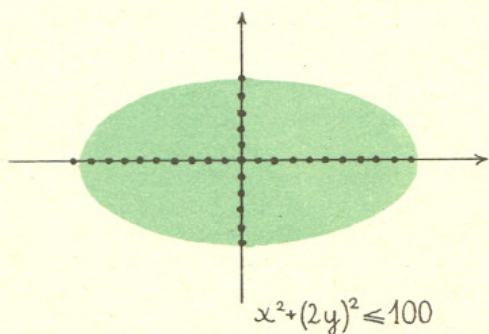
$$x^2 + y^2 \leq 100.$$

(Pamiętajcie, że rozwiązaniami mogą być nie tylko pary liczb dodatnich, ale i takie pary, w których jedna lub obie liczby są ujemne, na przykład $x = -5$ i $y = -3$ albo $x = -10$ i $y = 0$.) Kto ma cierpliwość może policzyć, że jest ich 317. Z drugiej strony pole rzeczywistego koła filmowanego przez kamerę jest równe $\pi \cdot 10^2 \approx 314$. Pole dyskretnego obrazu koła dosyć dobrze przybliża pole rzeczywistego koła... lub na odwrót. Zgodność będzie tym lepsza, im większe będą rozmiary koła. Możemy to wykorzystać w dwie strony. Ci, którzy znają przybliżoną wartość liczby π , mogą obliczyć ile mniej więcej rozwiązań całkowitych ma nierówność:

$$x^2 + y^2 \leq n \quad (n \text{ duże}).$$

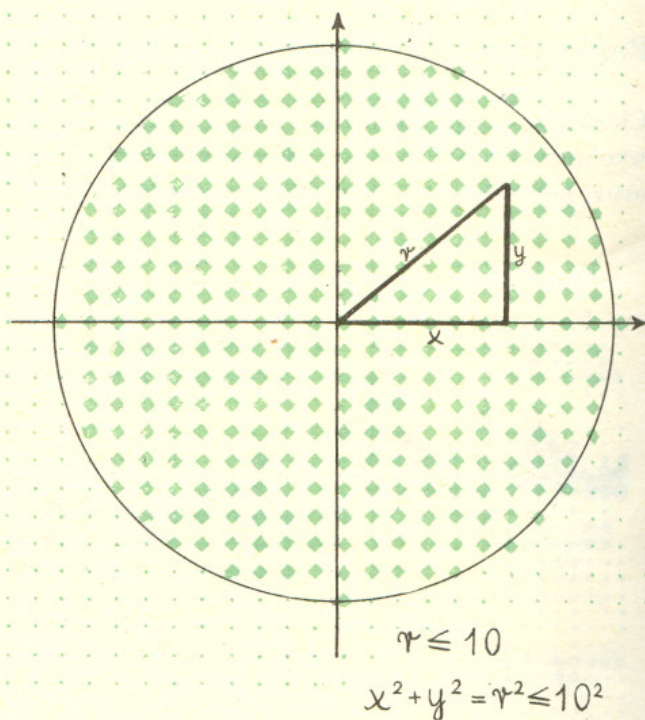
Ci, którzy umieją dobrze liczyć i mają cierpliwość, mogą spróbować wyznaczyć przybliżoną wartość liczby π (w jaki sposób?).

A jak wykorzystać metodę kropek do obliczenia przybliżonych wartości pól figur pokazanych na rysunkach poniżej?



Związki zachodzące między światem „płynnym” a dyskretnym są niekiedy bardzo interesujące. Spójrzmy na rysunek wyobrażając sobie, że jest to oglądany pod lupą ekran telewizora. Trochę upraszczając możemy przyjąć, że linią ciągłą zaznaczone są kontury przedmiotu ujmowanego przez kamerę, natomiast pogrubione kropki to obraz tego przedmiotu oglądany już na ekranie. A więc telewizyjny obraz jest figurą dyskretną. Pola takich dyskretnych figur będziemy mierzyli licząc, z ilu kropek się one składają (jest to chyba sensowna propozycja?).

W tym sensie, pole figury przedstawionej na rysunku jest równe 37. Ale uwaga, oto inny obraz. Kamera ujmuje ten sam przedmiot; tyle tylko, że obrócił się on w stosunku do pierwszego położenia. Jego telewizyjny obraz to jednak coś innego niż poprzednio, gdyż składa się z 31 kropek. Ten pierwszy eksperyment stawia trochę sceptycznie do możliwości odkrycia sensownych zależności między światem „płynnym” a dyskretnym. Podobne paradoksy są niestety regułą w mikroświecie, to znaczy wtedy, gdy rozmiary rozpatrywanych figur są względnie małe. Co innego jednak, jeśli rozpatrywać będziemy figury duże w stosunku do odstępów między kropkami.





Mechanika kwantowa wprawiała i ciągle jeszcze wprawia matematyków w zakłopotanie, dostarczając tym samym wielu interesujących problemów. Zaczęło się od Diraca, który różniczkował funkcje nieróżniczkowalne i otrzymywał sensowne wyniki. Potrzeba było lat, by rzecz uporządkować: stworzono (zob. Delta 10/1975) teorię dystrybucji, na gruncie której poczynania Diraca nabrały głębokiego matematycznego sensu.

Heisenberg swą niewinnie wyglądającą zasadą nieoznaczoności dostarczył zajęcia logikom: koniunkcja dwóch zdań, z których każde jest prawdziwe lub fałszywe może tu nie być (zob. Delta 4/1975) ani prawdziwa, ani fałszywa. Do dziś trwają poszukiwania „logiki kwantowej”. I choć sformułowano tu wiele interesujących propozycji — żadna z nich nie zdobyła sobie jeszcze pełnych praw obywatelskich.

Nie koniec na tym. Born wprowadza do mechaniki kwantowej probabilistyczną interpretację występujących w niej obiektów (np. funkcji falowej, zob. artykuł J. Kijowskiego w tym numerze). Interpretacja ta rozwiązuje fizykom pewne kłopoty pojęciowe, ale rodzi nowe kłopoty matematyczne.

Pierwszy oczywisty problem wynika z zasady nieoznaczoności. Jak wiemy — w zwykłym rachunku prawdopodobieństwa koniunkcja dwu zdarzeń jest zdarzeniem. Koniunkcja dwu „zdarzeń kwantowych” wcale „zdarzeniem kwantowym” być nie musi.

Rozpatrzmy rzecz nieco dokładniej. Wyobraźmy sobie dla uproszczenia pojedynczą cząstkę poruszającą się po prostej, oznaczmy jej położenie przez q , a pęd przez p . Dla dowolnego $\varepsilon > 0$ zdanie A : „w chwili t cząstka znajduje się w przedziale $\langle \alpha, \alpha + \varepsilon \rangle$ ” określa pewne „zdarzenie kwantowe”, bowiem teoretycznie możliwy jest pomiar położenia z dowolną dokładnością. Również dla dowolnego $\eta > 0$ zdanie B : „w chwili t pęd cząstki zawarty jest w przedziale $\langle \beta, \beta + \eta \rangle$ ” też opisuje pewne zdarzenie kwantowe, bo i pęd można mierzyć dowolnie dokładnie. Mimo to, jeśli tylko ε i η są dostatecznie małe to formalnie napisana koniunkcja $A \cap B$ zdarzeń A i B „zdarzeniem kwantowym” nie jest. Koniunkcja ta ma bowiem sens wtedy i tylko wtedy gdy

$$\varepsilon \eta \geq \hbar/2.$$

Okazuje się jednak, że kłopot matematyczny jest tu mniejszy, niż się na pierwszy rzut oka wydaje: można tak zmodyfikować elementarny rachunek prawdopodobieństwa, że zasada nieoznaczoności nie jest sprzeczna z nowym „kwantowym rachunkiem prawdopodobieństwa”. Przy tym ten nowy rachunek jest wystarczająco podobny do starego na to, aby nie razić ekstrawagancją.

W zwykłym rachunku prawdopodobieństwa, przypomnijmy, rodzina zdarzeń scharakteryzowana jest następująco:

Dany jest zbiór zdarzeń elementarnych Ω oraz pewna rodzina Z podzbiorów zbioru Ω zwanych zdarzeniami. Przy tym zakłada się, że

- 1) $\emptyset$ i Ω są zdarzeniami;
- 2) jeśli A jest zdarzeniem, to $A' = \Omega - A$ też jest zdarzeniem;
- 3) jeśli A, B są zdarzeniami, to również $A \cup B$ jest zdarzeniem.

Prawdopodobieństwo natomiast jest taką funkcją $P: Z \rightarrow \langle 0, 1 \rangle$, że

- 4) $P(\Omega) = 1$;
- 5) jeśli A i B są zdarzeniami i $A \cap B = \emptyset$, to $P(A \cup B) = P(A) + P(B)$.

Kwantowy rachunek prawdopodobieństwa wygląda podobnie. Dany jest zbiór zdarzeń elementarnych Ω , oraz pewna rodzina Z_k podzbiorów Ω zwanych zdarzeniami kwantowymi (krótko: k-zdarzeniami). Przy tym żąda się, aby spełnione były następujące warunki:

- 1) $\emptyset$ i Ω są k-zdarzeniami;
- 2) jeżeli A jest k-zdarzeniem, to A' też jest k-zdarzeniem;
- 3') jeśli A, B są k-zdarzeniami i $A \cap B = \emptyset$, to $A \cup B$ też jest k-zdarzeniem.

Prawdopodobieństwo natomiast jest taką funkcją $P: Z_k \rightarrow \langle 0, 1 \rangle$, że

- 4) $P(\Omega) = 1$;
- 5) jeśli A i B są k-zdarzeniami i $A \cap B = \emptyset$, to $P(A \cup B) = P(A) + P(B)$.

Różnica jest więc bardzo niewielka. Sprowadza się ona jedynie do zastąpienia warunku 3) warunkiem 3'), który na gruncie mechaniki kwantowej jest całkiem do przyjęcia. O ile bowiem z warunków 1), 2), 3) wynika, że jeśli A i B są zdarzeniami, to również $A \cap B$ jest zdarzeniem, to z warunków 1), 2), 3'), wynika jedynie, że jeśli A, B są k-zdarzeniami, to $A \cap B$ jest k-zdarzeniem *wtedy i tylko wtedy, gdy $A \cup B$ jest k-zdarzeniem*. Warunek 3') zakazuje więc co prawda rozpatrywać alternatywy takich k-zdarzeń, których koniunkcja nie jest k-zdarzeniem, ale, jak się okazuje, w tych problemach mechaniki kwantowej, w których obowiązuje zasada nieoznaczoności nie powstaje nigdy potrzeba rozważania takich alternatyw. Prawdopodobieństwo kwantowe ma natomiast dokładnie te same własności rachunkowe, co prawdopodobieństwo zwykłe i operowanie nim, jeśli się dobrze określi rodzinę zdarzeń kwantowych, nie nastręcza żadnych trudności (zob. zadanie na końcu artykułu).

Gdzie więc — i czy rzeczywiście — powstają zapowiedziane istotne kłopoty z probabilistyczną interpretacją mechaniki kwantowej?



Rozwiązanie zadania M 99. Należy udowodnić, że różnica

$$R = (a^2 + 2na)^k - a^2,$$

gdzie n jest liczbą całkowitą nieujemną, k — liczbą naturalną, jest wielokrotnością liczby $2a$.

$$\text{Mamy } R = a^{2k} + \binom{k}{1} a^{2(k-1)} 2na + \dots + \binom{k}{k-1} a^2 (2na)^{k-1} + (2na)^k - a^2.$$

Wszystkie składniki tej sumy prócz, być może, pierwszego i ostatniego są podzielne przez $2a$. Ponadto mamy

$$a^{2k} - a^2 = a^2(a^{2(k-1)} - 1) = a \cdot a \cdot (a-1) \cdot A_k = 2a \frac{a(a-1)}{2} A_k,$$

gdzie A_k jest liczbą całkowitą:

$$A_k = \frac{a^{2(k-1)} - 1}{a-1}$$

oraz $\frac{a(a-1)}{2}$ jest też liczbą całkowitą, gdyż

liczba $a(a-1)$ jako iloczyn dwóch kolejnych liczb całkowitych jest parzysta. R jest więc sumą liczb podzielnych przez $2a$, a więc jest liczbą podzielną przez $2a$.

Tu dygresja. „Funkcja δ ” Diraca — póki nie skonstruowano dystrybucji — w świadomości jej użytkowników była funkcją, „ale nie zawsze”. Można było np. powiedzieć, że dla $x \neq 0$ jest ona funkcją równą tożsamościowo zeru; nie miało sensu natomiast stwierdzenie, że dla $x = 0$ przyjmuje ona jakąkolwiek wartość. W gruncie rzeczy to, czy była, czy nie była funkcją zależało od fizycznej interpretacji kontekstu, w którym występowała.

Wartość matematyki jako narzędzia używanego przez przyrodników polega zaś właśnie na tym, że pojęcia matematyczne mają jednoznaczny sens, który jest niezależny od przyrodniczej interpretacji tych pojęć. Obiekt badany przez matematykę albo jest funkcją — albo nią nie jest. Trzeciej możliwości nie ma.

Matematycy musieli więc uznać, że δ funkcją nie jest i że mechanika kwantowa posługuje się pewną klasą obiektów ogólniejszą niż klasa funkcji. I dlatego stworzyli teorię dystrybucji; teorię ogólniejszą niż teoria funkcji.

Z prawdopodobieństwem jest podobnie. Jest to pojęcie matematyczne o jednoznacznie określonych własnościach. Jeśli więc coś w mechanice kwantowej jest podobne do prawdopodobieństwa, „ale nie zawsze” — to nie jest to już prawdopodobieństwo. Dlatego właśnie wykonaliśmy krok polegający na wprowadzeniu prawdopodobieństwa kwantowego. Okazuje się jednak, że jest to krok zbyt mały.

Poważne kłopoty pojawiają się bowiem dopiero wtedy, gdy wychodzi się poza elementarny rachunek prawdopodobieństwa i zaczyna rozważać zmienne losowe.

Przypomnijmy: Zmienną losową nazywa się taka funkcja $X: \Omega \rightarrow \mathbb{R}$, że dla każdego przedziału $\langle a, b \rangle \subset \mathbb{R}$ zbiór $\{\omega \in \Omega: X(\omega) \in \langle a, b \rangle\}$ jest zdarzeniem.

Zmienną losową nazywa się ciągłą, jeśli istnieje taka funkcja $f: \mathbb{R} \rightarrow \mathbb{R}$, która jest nieujemna

$$\text{ i dla każdego przedziału } \langle a, b \rangle \quad P(\{\omega \in \Omega: X(\omega) \in \langle a, b \rangle\}) = \int_a^b f(x) dx.$$

Funkcję f nazywa się gęstością rozkładu zmiennej X .

Jeśli więc (wracamy do przykładu pojedynczej cząstki) $\psi(q)$ jest funkcją falową, a — zgodnie z przyjętą interpretacją — $f(q) = |\psi(q)|^2$ jest gęstością rozkładu położenia cząstki na prostej,

to $\int_a^b |\psi(q)|^2 dq$ jest prawdopodobieństwem tego, że cząstka znajduje się w przedziale $\langle a, b \rangle$.

Interpretacja ta oznacza, że położenie q traktowane jest jako zmienna losowa. Analogicznie — pęd p jest teraz drugą zmienną losową, o gęstości $g(p) = |\Phi(p)|^2$. Przy tym funkcja Φ jest jednoznacznie wyznaczona przez funkcję ψ . Tak więc cząstka opisywana jest parą ciągłych zmiennych losowych (q, p) o znanych rozkładach $|\psi(q)|^2$ i $|\Phi(p)|^2$.

W rachunku prawdopodobieństwa parę zmiennych losowych (X, Y) nazywa się zmienną dwuwymiarową. Gęstością rozkładu dwuwymiarowej zmiennej losowej (X, Y) nazywa się taką nieujemną funkcję $h(x, y)$, że dla dowolnych liczb $a < b$ i $c < d$ zachodzi równość

$$P(\{\omega \in \Omega: (X(\omega) \in \langle a, b \rangle) \text{ i } (Y(\omega) \in \langle c, d \rangle)\}) = \int_a^b \int_c^d h(x, y) dy dx.$$

$$\text{Dowodzi się przy tym, że } f(x) = \int_{-\infty}^{+\infty} h(x, y) dy \text{ jest gęstością rozkładu zmiennej } X, \text{ a } g(y) = \int_{-\infty}^{+\infty} h(x, y) dx \text{ jest gęstością rozkładu zmiennej } Y.$$

Okazuje się, że gdyby zamiast prawdopodobieństwa wykorzystać do definicji zmiennej losowej prawdopodobieństwo kwantowe, to definicja ta pozostałaby poprawna. Nasuwa się więc przypuszczenie, że jeśli zacznie się używać kwantowych zmiennych losowych, to wszystkie ich własności formalne będą zgodne z twierdzeniami mechaniki kwantowej.

(Byłoby tak, gdyby wszystkie kłopoty z interpretacją prawdopodobieństwa wynikały z zasady nieoznaczoności. Niestety — tak nie jest).

Wróćmy do przykładu. Mamy do czynienia z parą zmiennych losowych (q, p) o znanych gęstościach rozkładu. Powstaje naturalne pytanie: Czy istnieje gęstość rozkładu zmiennej dwuwymiarowej (q, p) , tzn. funkcja nieujemna $h(q, p)$ taka, że

$$(1) \int_{-\infty}^{+\infty} h(q, p) dp = |\psi(q)|^2 \quad \text{ i } \quad \int_{-\infty}^{+\infty} h(q, p) dq = |\Phi(p)|^2;$$

(2) Spełniony jest podany wyżej związek między Φ i ψ oraz zasada nieoznaczoności;

(3) Wartości średnie obserwowalnych kwantowych zmiennych losowych obliczane zgodnie z definicją wartości średniej w rachunku prawdopodobieństwa są równe wartościom średnim tych zmiennych obliczanym w formalizmie operatorowym mechaniki kwantowej.

Odpowiedzi na to pytanie udzielił L. Cohen. Jest ona następująca:

Istnieją co prawda funkcje spełniające (1) i (2), ale nie istnieje funkcja spełniająca (1), (2) i (3).

A więc nie można posuwać się zbyt daleko w probabilistycznych interpretacjach mechaniki kwantowej. Co prawda można traktować zarówno p jak i q jako zmienne losowe, ale rozpatrując p i q łącznie — wychodzimy już z rachunku prawdopodobieństwa, para (q, p) zmienną losową nie jest.



Rozwiązanie zadania M 98. Najpierw udowodnimy, że istnieje miasto, z którego wychodzi jedna tylko droga. Wyjeżdżamy z dowolnego miasta dowolną drogą. Dojechawszy do jakiegoś miasta jedziemy dalej (jeśli to możliwe) drogą, którą jeszcze nie jechaliśmy. Jeżeli drogi takiej nie ma, to oznacza to, że z miasta tego wychodzi tylko ta droga, którą przyjechaliśmy. Zauważmy, że podróżując w opisany sposób odwiedzimy każde miasto najwyżej raz, w przeciwnym bowiem wypadku istniałaby droga zamknięta („pętla”), co przeczy warunkowi 2'. Zastosujmy teraz rozumowanie indukcyjne. Dla $n = 2$ twierdzenie jest prawdziwe, mamy bowiem wówczas jeden tylko odcinek drogi. Załóżmy, że twierdzenie jest prawdziwe dla liczby n . Wówczas, gdy mamy $n+1$ miast, na podstawie udowodnionego wyżej faktu istnieje miasto, z którego wychodzi tylko jedna droga. Odrzućmy to miasto i wychodzącą zeń drogę. Pozostaje nam sieć dróg łącząca n miast i spełniająca warunki zadania, a więc złożona z $n-1$ odcinków. Sieć dróg łącząca $n+1$ miast składa się więc z n odcinków, a więc na podstawie zasady indukcji matematycznej twierdzenie jest prawdziwe.

I dlatego nie obejdzie się tu bez nowej teorii: teorii ogólniejszej niż rachunek prawdopodobieństwa. Takiej teorii jeszcze nie ma.

Zadania. 1. Udowodnić, że jeśli Z jest zwykłą rodziną zdarzeń, P — zwykłym prawdopodobieństwem oraz $A \in Z$ ustalonym zdarzeniem takim, że $P(A) > 0$, to rodzina

$$\{B \in Z: P(A \cap B) = P(A) \cdot P(B)\}$$

(tzn. rodzina zdarzeń niezależnych od A) jest kwantową rodziną zdarzeń.

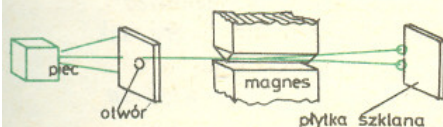
2. Pokazać, że jeśli A, B są k-zdarzeniami oraz $A \subset B$, to $B - A$ jest k-zdarzeniem.

Literatura

Tom 33 (z roku 1966) czasopisma „*Philosophy of Science*”, artykuły P. Suppesa (str. 14) i L. Cohena (str. 317).

Cudowny wynik pewnego doświadczenia

Wszystko to brzmi nieprawdopodobnie — powiecie po przeczytaniu dwóch poprzednich artykułów o mechanice kwantowej. Jak to, najbardziej fundamentalna teoria mikroświata pozwalająca przewidzieć wyniki doświadczeń z fantastyczną dokładnością nie może sobie poradzić z opisaniem losów zwykłego, swobodnego elektronu! Czy naprawdę musimy budować teorię, w której z elektronem, a raczej z informacją o nim, jeżeli nań nie patrzymy, wiążemy pewną falę (falę prawdopodobieństwa), podczas gdy w trakcie każdej obserwacji ukazuje się on nam w postaci mikroskopowej cząstki. Na tego typu wątpliwości najlepiej odpowiada zawsze doświadczenie. Opiszemy tu wyniki jednego z przełomowych dla fizyki doświadczeń, wykonanego w 1922 roku przez Sterna i Gerlacha. Idea eksperymentu jest prosta. Wiązka atomów srebra powstała przez odparowanie srebra w specjalnym piecu jest skierowana między bieguny magnesu



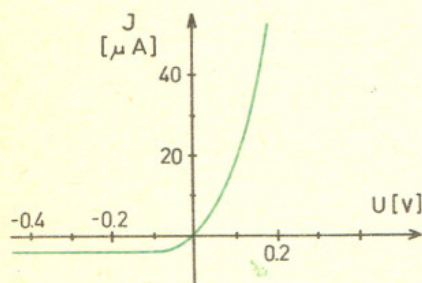
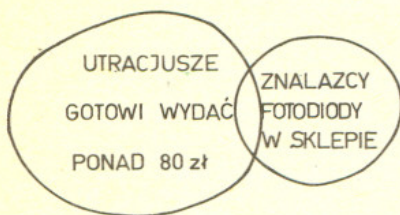
Na magnes atomowy oprócz siły odchylającej działa też siła dążąca do ustawienia go zgodnie z kierunkiem linii pola zewnętrznego. Siła ta nie zmienia jednak kąta między magnesem i kierunkiem linii sił, a tylko wywołuje precesję dookoła tego kierunku. Jest tak dlatego, że własne pole magnetyczne atomu jest związane z występowaniem własnego momentu pędu i atom stanowi mały giroskop.

Samo urządzenie Sterna-Gerlacha po wywierceniu otworu na ekranie w miejscu, gdzie pada jedna z wiązek, stanowi przyrząd do wytwarzania atomów o określonym zwrocie spinu. Można przeszedźć cały aparat pojęciowy mechaniki kwantowej ustawiając szereg takich urządzeń (filtrów), różnie obróconych, jedno za drugim i analizując wyniki wymyślonych tak doświadczeń.

Wyniki doświadczenia Sterna-Gerlacha zostały przewidziane przez tzw. starą teorię kwantów przed jego wykonaniem.

wytwarzającego silnie niejednorodne pole (patrz rysunek). Następnie atomy osadzają się na płytce szklanej. Atomy srebra podobnie, jak wszystkie inne atomy, mają własny moment pędu (spin). Powstaje on ze złożenia momentów pędu jądra atomowego i elektronów. Ponieważ wszystkie składniki atomu są naładowane, więc atom taki zachowuje się jak swego rodzaju pętla z prądem i wytwarza własne pole magnetyczne, o kierunku zgodnym z ustawieniem spinu. W niejednorodnym zewnętrznym polu magnetycznym na taki atomowy magnes działa siła odchylająca w górę lub w dół, zależna od kąta między osią magnesu, a zwrotem zmiany (gradientu) pola zewnętrznego. I tak na magnes skierowany wzdłuż gradientu działa maksymalna siła do góry, na skierowany przeciwnie maksymalna siła w dół, na magnes ustawiony prostopadłe do kierunku zmiany pola nie działa żadna siła itd. Ponieważ atomy srebra wyprodukowane w piecu stanowią zbiór chaotyczny, więc związane z nimi magnesy są losowo poustawiane i na ekranie powinniśmy otrzymać ciągłą linię wzdłuż kierunku gradientu pola. Tego wymaga fizyka klasyczna i ukształtowany w codziennym doświadczeniu rozsądek. Tymczasem Stern i Gerlach znaleźli na płytce jedynie dwa izolowane punkty. Atomy srebra utworzyły tylko dwie odrębne wiązki. Powtórzmy jeszcze raz: atomy srebra, których spiny były ustawione zupełnie chaotycznie utworzyły dwie oddzielne wiązki. Zupełnie jakby spiny atomów wiedziały (tylko skąd?), że wolno im się ustawić względem pola magnetycznego w ściśle określonych kierunkach (dwóch dla atomów srebra o spinie $1/2$). Wynik doświadczenia nie zależy od tego, jak obrócimy układ magnesów wytwarzających pole. Zawsze dostajemy dwie plamki na linii równoległej do kierunku zmiany pola. I to plamki równie intensywne. Choć wynik ten jest wprost fantastyczny, to jednak tak dzieje się w doświadczeniu. Zaczniemy teraz puszczać atomy po kolei w pewnych odstępach czasu. Pojedynczy atom nawet o znanym z poprzedniego doświadczenia ustawieniu spinu wybierze raz jedno, raz drugie ustawienie swego spinu względem kierunku zmiany pola. Tu tkwi element losowy i to w sytuacji, gdy o własnym polu magnetycznym atomu wiemy chyba wszystko. Mechanika kwantowa też oferuje nam tu tylko prawdopodobieństwo określonego zwrotu. Natomiast skwantowanie rzutu momentu pędu na każdą wyróżnioną oś jest jednoznaczny przewidywaniem w tej teorii. Pozwala ona też obliczyć z ogromną dokładnością wielkość własnego pola magnetycznego związanego ze spinem, a więc i wielkość odchylenia każdej z dwóch wiązek. Dlatego mówimy, że nie ma żadnych przesłanek doświadczalnych na to, żeby poprawiać mechanikę kwantową.

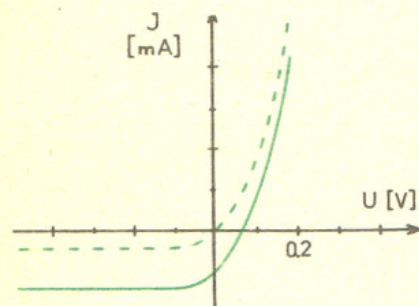
FOTODIODA



Charakterystyka prądowo-napięciowa diody półprzewodnikowej

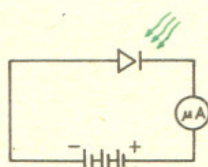
Rys. 1

Diody o własnościach prostowania prądu (także światłoczuła) może być również kontakt metal — półprzewodnik.

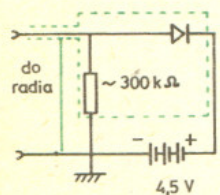


Charakterystyka prądowo-napięciowa fotodiody przy oświetleniu (linia przerywana-bez oświetlenia)

Rys. 2

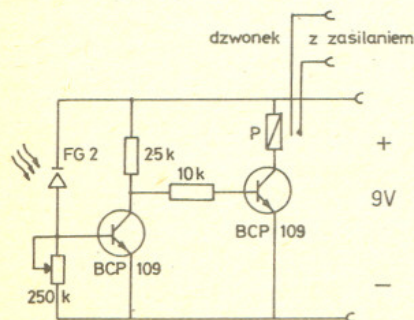


Rys. 3



Rys. 4

Czytelnik zauważy z pewnością uproszczenia w tym rozumowaniu.



Rys. 5

Dzisiejsze nasze eksperymenty będą niestety dostępne tylko dla stosunkowo niewielkiej części zbioru Czytelników „Delfy”: będzie to przekrój zbioru osób, które zdecydują się na wydatek ponad 80 zł i zbioru (znacznie, jak przypuszczam, mniej licznego) szczęśliwców, którzy natrafią na fotodiody w sprzedaży. Produkowane w Polsce fotodiody germanowe FG-2, podobnie jak i dowolne inne (także fototranzystory) świetnie nadają się do naszych doświadczeń, cała sztuka w tym, żeby je kupić.

Jedno przemawia niewątpliwie za omówieniem tego tematu w bieżącym numerze: zjawisko fotoelektryczne wewnętrzne, umożliwiające działanie fotodiody jest efektem czysto kwantowym — fizyka klasyczna nie wyjaśniła go zadowalająco. A więc, skoro wytłumaczyłem się jakoś z wyboru niedemokratycznego tematu, zadajmy sobie pytanie

JAK DZIAŁA FOTODIODA?

Trzeba zacząć od zwykłej diody (mówiliśmy już o niej w „Laboratorium w domu” z nr 8/75), znanej też pod nazwą złącza p-n. Jeżeli przyłożymy do diody napięcie w kierunku zaporowym, na pograniczu obszarów o przewodnictwie akceptorowym (p) i donorowym (n) powstaje warstwa opróżniona z nośników ładunku stanowiąca izolator i powodująca, że prąd praktycznie nie płynie. Jeżeli w warstwie tej pojawią się z jakiegokolwiek przyczyny nośniki — popłynie prąd. W zwykłej diodzie zawsze powstaje pewna bardzo mała ilość nośników, gdyż elektrony są przenoszone z pasma walencyjnego do pasma przewodnictwa kosztem energii drgań cieplnych — stąd niezerowa wartość natężenia prądu zaporowego (patrz wykres 1). Ten sam proces można wywołać oświetlając złącze p-n, potrzebne jest jednak światło — i tu bez mechaniki kwantowej obejść się nie sposób — o dostatecznie małej długości fali. Fakt ten jest jednym z dowodów eksperymentalnych istnienia w półprzewodnikach przerwy energetycznej — obszaru energii między wierzchołkiem pasma walencyjnego a dnem pasma przewodnictwa, który jest dla elektronów zabroniony. Elektron w pasmie walencyjnym może wchłonąć energię kwantu świetlnego (równą $h\nu$, gdzie $h = 6,63 \cdot 10^{-27}$ erg. sek jest stałą Plancka, a ν — częstotliwością fali świetlnej) wtedy, kiedy jest ona wystarczająco duża, by przenieść go do pasma przewodnictwa, a więc większa od szerokości przerwy energetycznej.

Bardzo to piękne — powiedzą zwolennicy konkretów, ale

JAK SIĘ TO OBJAWIA PRAKTYCZNIE?

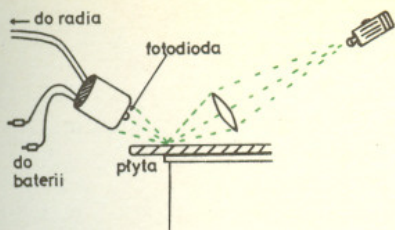
Przez wzrost natężenia prądu zaporowego przy oświetleniu (patrz wykresy 1 i 2).

Fotodiody różni się od zwykłej diody tylko tym, że jej konstrukcja umożliwia doprowadzenie światła do obszaru złącza. Aby skończyć z gołosłownością zabierzmy się do doświadczeń. Najłatwiej będzie to sprawdzić dysponując mikroamperomierzem. Łączymy prosty obwód (rys. 3) — uwaga na bieguny! i oświetlamy fotodiody np. żarówką — natężenie prądu w kierunku zaporowym rośnie. W tym miejscu większość Czytelników nie posiadających mikroamperomierza zaprotestuje: Z takim wydatkiem się nie liczyliśmy! Nic strasznego. Wiele ciekawych eksperymentów przed nami, jeżeli tylko dysponujemy radiem z wejściem adapterowym. Domyślcie się już, że ich tematem będzie

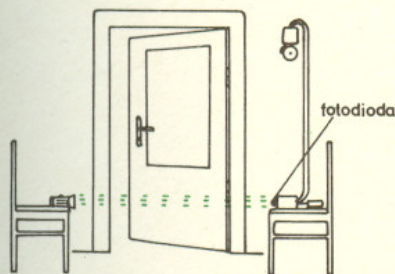
ZAMIANA ŚWIATŁA NA DŹWIĘK

Aby uzyskać silne wahania napięcia na fotodiodzie pod wpływem światła zestawiamy obwód składający się z baterijki 4,5 V, fotodiody i opornika rzędu 300 kΩ (rys. 4). Natężenie prądu w obwodzie (jest to prąd zaporowy fotodiody) wahając się przy zmianach oświetlenia spowoduje zmiany napięcia na oporniku R zgodnie z prawem Ohma: $U = IR$. Wartość R musi być więc możliwie duża (napięcie jest do niej proporcjonalne), nie na tyle jednak, aby decydować o natężeniu prądu w obwodzie. W praktyce dobieramy ją doświadczalnie tak, aby otrzymać najsilniejszy efekt (na przykład próbując 100 kΩ, 300 kΩ i 500 kΩ). Jeszcze jedna ważna uwaga: obwód powinien być ekranowany (przynajmniej fotodiody, opornik i przewód łączący je z radiem) co można zapewnić umieszczając go w metalowym pudełku połączonym z masą układu. Najprostszym doświadczeniem będzie teraz wystawienie fotodiody na światło żarówki zasilanej z sieci — w głośniku usłyszymy buczenie o częstotliwości 100 Hz, co odpowiada częstotliwości zmian temperatury włókna żarówki. Oczywiście jeśli użyjemy latarki kieszonkowej, buczenia nie usłyszymy — jeśli jednak postukać w latarkę nawet palcem, słyszy się zaskakująco melodyjny dźwięk pochodzący od drgań włókna żaróweczki.

Rozmiary „punktu” świetlnego powinny być rzędu odległości między rowkami ($\sim 0,1$ mm).



Rys. 5



Rys. 6

Zabawny eksperyment można też wykonać z adapterem i płytą. Po wprawieniu adaptera w ruch ogniskujemy światło żarówki w mały punkt na płycie i „patrzmy” na niego fotodiodą (rys. 5). W głośniku słychać dźwięki, w których w zależności od precyzji wykonania doświadczenia oraz zasobu dobrej woli można łatwiej lub trudniej rozpoznać muzykę nagrałą na płycie. Na zakończenie dla nieco bardziej zaawansowanych majsterkowiczów urządzenie specjalne, w którym

FOTODIODA ŁAPIE ZŁODZIEJA

No, może nie dosłownie, ale w każdym razie pomaga. W przejściu, którego nasza dioda ma pilnować, ustawiamy ją naprzeciwko latarki skierowanej na nią (rys. 6). Zestawiamy obwód z fotodiody, dwóch tranzystorów i czułego przełącznika P (rys. 7), który w chwili zasłonięcia wiązki światła włącza dzwonek alarmowy. Jeżeli nasza instalacja alarmowa ma działać nieco dłużej, lepiej zasilać żarówkę latarki z transformatora dzwonkowego zamiast z baterijki. Ten sam obwód z fotodiodą — już bez udziału latarki może też służyć na przykład do automatycznego zapalania światła w pokoju z nastaniem zmierzchu. Liczę na Waszą inwencję w znajdowaniu innych zastosowań fotodiody. Powodzenia w pracy!



Rozwiązanie zadania F 33

Równoczesna rejestracja fotonu w obu detektorach może być wynikiem jednego z następujących procesów:

1. Foton ze źródła A wpada do detektora 1, inny foton z tego źródła wpada do detektora 2,
2. Fotony wysłane przez źródło B rejestrowane są w obu detektorach 1 i 2,
3. Foton wysłany z A wpada do 1, wysłany z B wpada do 2,
4. Odwrotnie niż 3 (tzn. z A do 2, z B do 1).

Procesy 1 i 2 można w zasadzie odróżnić od siebie oraz od 3 i 4 (np. wyłączając lub zasłaniając jedno ze źródeł) natomiast przypadki 3 i 4 są nieodróżnialne. Amplitudę f zdarzenia „składającego się” z dwu zdarzeń o amplitudach odpowiednio f_1 i f_2 zapisujemy w postaci

$$f = f_1 \cdot f_2$$

Zatem poszczególnym procesom prowadzącym do jednoczesnej rejestracji fotonów w obu detektorach odpowiadają następujące amplitudy:

1. $f^1 = f_{A_1} f_{A_2}$,
2. $f^2 = f_{B_1} f_{B_2}$,
3. $f^3 = f_{A_1} f_{B_2}$,
4. $f^4 = f_{A_2} f_{B_1}$.

Ponieważ procesy 1 i 2 są odróżnialne, 3 i 4 — nie, więc prawdopodobieństwo koincydencji ma postać

$$P_{\text{koinc}} = |f^1|^2 + |f^2|^2 + |f^3 + f^4|^2$$

(tzn. dodajemy prawdopodobieństwa procesów odróżnialnych, amplitudy — dla procesów nieodróżnialnych). Postać poszczególnych amplitud łatwo napisać patrząc na rysunek przy zadaniu

$$f^1 = f_{A_1} f_{A_2} = c^2 e^{ik(R_1 + R_2)},$$

$$f^2 = f_{B_1} f_{B_2} = c^2 e^{ik(R_1 + R_2)},$$

$$f^3 = f_{A_1} f_{B_2} = c^2 e^{2ikR_1},$$

$$f^4 = f_{A_2} f_{B_1} = c^2 e^{2ikR_2},$$

$$P_{\text{koinc}} = f^1 \cdot \bar{f}^1 + f^2 \cdot \bar{f}^2 + (f^3 + f^4) \cdot \overline{(f^3 + f^4)} = 2|c|^4 [2 + e^{2ik(R_2 - R_1)} + e^{-2ik(R_2 - R_1)}].$$

Korzystając z faktu, że $e^{ix} + e^{-ix} = 2\cos x$ możemy to zapisać w postaci $P_{\text{koinc}} = 2|c|^4 [2 + \cos 2k(R_2 - R_1)]$.

Jeżeli $R \gg D$ oraz $R \gg d$, to różnicę obu promieni możemy wyrazić następująco

$$R_1 = \sqrt{R^2 + \frac{1}{4}(D-d)^2} = R \sqrt{1 + \left(\frac{D-d}{2R}\right)^2} \approx R \left[1 + \frac{1}{2} \left(\frac{D-d}{2R}\right)^2\right],$$

$$R_2 = \sqrt{R^2 + \frac{1}{4}(D+d)^2} = R \sqrt{1 + \left(\frac{D+d}{2R}\right)^2} \approx R \left[1 + \frac{1}{2} \left(\frac{D+d}{2R}\right)^2\right],$$

$$R_2 - R_1 \approx \frac{d \cdot D}{2R}.$$

Ostatecznie wzór na prawdopodobieństwo koincydencji przybierze postać

$$P_{\text{koinc}} = 2|c|^4 \left[2 + \cos \frac{k \cdot d \cdot D}{R}\right].$$

Wystarczy teraz rozsuwając oba detektory znaleźć takie dwa ich położenia, przy których P_{koinc} (a zatem i liczba równoczesnych zliczeń na jednostkę czasu) jest takie samo. Wówczas:

$$\frac{k d_2 D}{R} - \frac{k d_1 D}{R} = 2\pi, \quad \text{skąd} \quad D = \frac{2\pi R}{k(d_2 - d_1)}.$$

Jeżeli powyższe rozumowanie zastosować do pojedynczej gwiazdy, tzn. rozważyć pary punktów na powierzchni emitującej promieniowanie (światłone lub radiowe), to okazuje się, że po odpowiednim wysumowaniu (wycalkowaniu) po wszystkich parach punktów, uwzględniając jednocześnie fakt, że emitowane promieniowanie nie jest monochromatyczne — otrzymamy związek między prawdopodobieństwem koincydencji przy różnych odległościach wzajemnych detektorów, a rozmiarami gwiazdy.

Ten sposób pomiaru rozmiarów gwiazd został zaproponowany i wykorzystany przez Hanbury, Browna i Twissa w połowie lat 50-tych naszego stulecia („Nature” 178 (1956), 1046).



Sprostowanie

Mgr inż. Ryszard Piasecki ze Szczecina zwrócił nam uwagę na błąd w rozwiązaniu zadania F 25 w „Delfie” 1/76. Mianowicie

$\int_H^0 \frac{dh}{\sqrt{h}}$ równa się oczywiście $-2\sqrt{H}$, a nie

$-\frac{1}{2}\sqrt{H}$. Błąd jest drobny, niestety ma

poważniejsze konsekwencje. Wzór (3) powinien teraz mieć postać:

$t = \frac{t_0^2}{2t_n} \left(\ln \frac{t_0}{t_0 - 2t_n} - \frac{2t_n}{t_0} \right)$ i w warunkach podanych w zadaniu $t \rightarrow \infty$. Jedynie dla

$\frac{2t_n}{t_0} < 1$ zbiornik napelni się w skończonym

czasie, przy czym zawsze $t \geq t_n$. Przepraszamy.