

nr 60

zestyl 3

rok 2009

nr indeksu 369721

cenę 12 zł (0% VAT)

POSTĘPY FIZYKI

Dwumiesięcznik Polskiego Towarzystwa Fizycznego

Motory molekularne

Modele kosmologiczne

Technika w LHC



ISSN 0032-5430



9 770032 543004 >



Bukiet kwiatów (płyta Lippmanna 9×12 cm). Kwiaty doskonale nadają się jako przedmiot zdjęć metodą Lippmanna. Na tym zdjęciu widzimy różne kolory kwiatów i piękne odbicie światła od szklanego wazonu (dalsze zdjęcia barwne na III stronie okładki, artykuł na s. 121 – red.).

RADA REDAKCYJNA

Andrzej Kajetan Wróblewski (przewodniczący), Mieczysław Budzyński, Andrzej Dobek, Witold Dobrowolski, Zofia Gołąb-Meyer, Adam Kiejna, Józef Szudy

REDAKTOR HONOROWY

Adam Sobiczewski

KOMITET REDAKCYJNY

Jerzy Warczewski (redaktor naczelny), Ewa Lipka (sekretarz redakcji), Mirosław Łukaszewski (redaktor techniczny), Magdalena Staszal

Adres Redakcji:

ul. Hoża 69, 00-681 Warszawa, e-mail: postepy@fuw.edu.pl, Internet: postepy.fuw.edu.pl

KORESPONDENCI ODDZIAŁÓW PTF

Maciej Piętka (Białystok), Aleksandra Wronkowska (Bydgoszcz), Wojciech Gruhn (Częstochowa), Tomasz Jarosław Wąsowicz (Gdańsk), Roman Bukowski (Gliwice), Beata Kozłowska (Katowice), Aldona Kubala-Kukuś (Kielce), Małgorzata Nowina Konopka (Kraków), Elżbieta Jartych (Lublin), Michał Szanecki (Łódź), Halina Pięta (Opole), Maria Potomska (Poznań), Małgorzata Pociask (Rzeszów), Małgorzata Kuzio (Ślupsk), Janusz Typek (Szczecin), Winicjusz Drozdowski (Toruń), Aleksandra Miłosz (Warszawa), Bernard Jancewicz (Wrocław), Joanna Borgensztajn (Zielona Góra)

POLSKIE TOWARZYSTWO FIZYCZNE

ZARZĄD GŁÓWNY

Reinhard Kulessa (prezes), Jacek M. Baranowski (sekretarz generalny), Roman Puźniak (skarbnik), Przemysław Dereń, Mirosław Trociuk i Jerzy Warczewski (członkowie wykonawczy), Bolesław Augustyniak, Maria Dobkowska, Stanisław Dubiel, Henryk Figiel, Jacek Przemysław Goc, Zofia Gołąb-Meyer, Bernard Jancewicz i Ewa Kurek (członkowie)

Adres Zarządu:

ul. Hoża 69, 00-681 Warszawa, tel./fax: 022-6212668, e-mail: ptf@fuw.edu.pl, Internet: ptf.fuw.edu.pl

PRZEWODNICZĄCY ODDZIAŁÓW PTF

Eugeniusz Żukowski (Białystok), Stefan Kruszewski (Bydgoszcz), Józef Zbrozczyk (Częstochowa), Bolesław Augustyniak (Gdańsk), Bogusława Adamowicz (Gliwice), Wiktor Zipper (Katowice), Małgorzata Wysocka-Kunisz (Kielce), Stanisław Wróbel (Kraków), Jerzy Żuk (Lublin), Tadeusz Wibig (Łódź), Stanisław Waga (Opole), Roman Świetlik (Poznań), Małgorzata Klisowska (Rzeszów), Włodimir Tomlin (Ślupsk), Adam Bechler (Szczecin), Grzegorz Karwasz (Toruń), Mirosław Karpierz (Warszawa), Bernard Jancewicz (Wrocław), Marian Olszowy (Zielona Góra)

REDAKTORZY NACZELNI INNYCH CZASOPISM

WYDAWANYCH POD EGIDĄ PTF

Witold D. Dobrowolski – *Acta Physica Polonica A*, Kacper Zalewski – *Acta Physica Polonica B*, Andrzej Jamiołkowski – *Reports on Mathematical Physics*, Marek Kordos – *Delta*, Zofia Gołąb-Meyer – *Foton*, Zbigniew Wiśniewski (redaktor prowadzący) – *Fizyka w Szkole*

Czasopismo ukazuje się od 1949 r.

Wydawca: Polskie Towarzystwo Fizyczne

Dofinansowanie: Ministerstwo Nauki i Szkolnictwa Wyższego

Patronat: Wydział Fizyki Uniwersytetu Warszawskiego

Skład komputerowy w redakcji

Opracowanie okładki: Studio Graficzne etNova Piotr Zenda i Wspólnicy sp.j., tel.: 022-8735520, e-mail: etnova@etnova.pl

Druk i oprawa: „UNI-DRUK”, Warszawa, ul. Buńczuk 7b

ISSN 0032-5430

SPIS TREŚCI

M. Kostur, J. Łuczka – Motory molekularne	
Część I: Motory biologiczne	90
A. Krasieński – O modelach kosmologicznych i niektórych związanych z nimi nieporozumieniach	98
J. Kulka – Techniczne aspekty zderzacza LHC	109
M. Heller – Wielki Program Wszechświata	119
P. Ranson, R. Ouillon, J.-P. Pinan-Lucarré – W stulecie Nagrody Nobla z fizyki za rok 1908 przyznanej Gabrielowi Lippmannowi za fotografię barwną	121
NOWI PROFESOROWIE	125
ZE ZJAZDÓW I KONFERENCJI	126
WSPOMNIENIA: Jan Turkiewicz (1934–2009)	127
KRONIKA	128

Drodzy Czytelnicy,

Pięć artykułów stanowi zasadniczy trzon zeszytu 3 Postępów. Cztery z nich zapoczątkowują cztery nowe cykle. Artykuł Marcina Kostura i Jerzego Łuczki – przedstawiający mikroskopowy świat biologicznych motorów, które są odpowiedzialne za transport wewnątrzkomórkowy – jest pierwszym artykułem z cyklu, który by można nazwać „Zastosowanie fizyki statystycznej do opisu procesów biologicznych”. Artykuł Andrzeja Krasieńskiego rozpoczyna cykl, który po prostu nosi tytuł „Kosmologia”. Artykuł ks. Michała Hellera rozpoczyna cykl esejów Centrum Kopernika Badań Interdyscyplinarnych w Krakowie. Cykl ten nosi tytuł „Filozofia Przyrody”. Wszystkie teorie fizyczne mają strukturę matematyczną i zwykle jest to struktura dynamiczna. Struktury matematyczne nie tylko opisują świat fizyczny, lecz raczej należy myśleć o nich jako o Wielkim Programie, który nasz Wszechświat wykonuje. Artykuł Jana Kulki dotyczy aspektów technicznych budowy Wielkiego Zderzacza Hadronów (LHC) i opowiada o różnych sekretach i tajemnicach jego konstrukcji. Ten artykuł rozpoczyna cykl zatytułowany „Wielki Zderzaczy Hadronów (LHC) – jego fizyka i technika”. Artykuł P. Ransona, R. Ouillona i J.-P. Pinan-Lucarré’a opowiada o Nagrodzie Nobla sprzed 101 lat, która była przyznana Gabrielowi Lippmannowi za fotografię barwną.

Prócz tego polecam tradycyjne rubryki, takie jak Kronika, Wspomnienia, Ze Zjazdów i Konferencji, Nowi Profesorowie.

Jerzy Warczewski

Na okładce:

Goście Dni Otwartych w CERN-ie (6 kwietnia 2008 r.) przy detektorze eksperymentu CMS (fot. – Archiwum CERN-u). Skrót CMS oznacza Compact Muon Solenoid, co można przetłumaczyć jako zintegrowany z magnesem detektor mionów. Eksperyment CMS to jeden z 4 wielkich eksperymentów fizyki wysokich energii (ATLAS, ALICE, CMS, LHCb) zainstalowanych przy zderzaczu LHC (patrz artykuł Jana Kulki, s. 109).

Motory molekularne

Część I: Motory biologiczne

Marcin Kostur, Jerzy Łuczka

Instytut Fizyki, Uniwersytet Śląski, Katowice

Molecular motors. Part I: Biological motors

Abstract: We describe the microworld of biological motors which are responsible for intercellular transport. The physicist's point of view, i.e. the concept of the Feynman ratchet and an abstract approach to modelling of noisy transport, is presented.

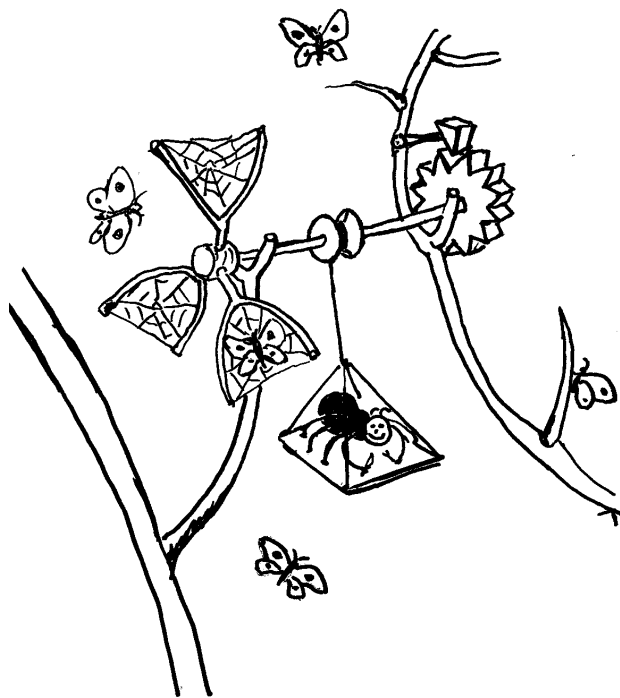
Opisujemy okiem fizyka mikroskopowy świat biologicznych motorów, które są odpowiedzialne za transport wewnątrzkomórkowy. Zaprezentowane zostaną: koncepcja zębatego Feynmanowskiej i abstrakcyjne podejście do modelowania nierównowagowego transportu stochastycznego.

Wstęp

Rok 1989 stał się rokiem przełomowym dla nanotechnologii i inżynierii molekularnej: 29 września Don Eigler, pracownik IBM w San Jose (USA), przedstawił logo swojej firmy składające się z 35 atomów ksenonu ułożonych na powierzchni niklu. Miało ono rozmiar 5 nanometrów wysokości i 17 nanometrów szerokości. Przeciętny zjadacz chleba nie zdaje sobie sprawy z wielkości tego logo i nie rozumie tego rewolucyjnego przełomu. Eigler otrzymał nanologo przy użyciu skaningowego mikroskopu tunelowego (SMT), który notabene sam zbudował. Nie byłoby tego sukcesu, gdyby nie osiągnięcie innych pracowników IBM (tym razem z Zurichu), a mianowicie Gerda Binniga i Heinricha Rohrera, którzy pierwsi zbudowali SMT w 1982 roku.

Obaj naukowcy pod koniec 1978 roku rozpoczęli badania dotyczące cienkich warstw. Aby móc kontynuować swoją pracę, potrzebowali urządzenia dającego możliwość obserwacji powierzchni w skali nanometrów. Ponieważ do tej pory nie było przyrządów, które by to umożliwiały, sami postanowili taki zbudować. Znowu potrzeba okazała się matką wynalazku. Powstał pierwszy skaningowy mikroskop tunelowy.

Osiągnięcie to zostało bardzo szybko zauważone i docenione przez Komitet Noblowski. Zaledwie 4 lata później, w 1986 roku, twórcy otrzymali nagrodę Nobla z fizyki. W tym samym roku Gerd Binnig (patrz wyżej), Calvin F. Quate i Christoph Gerber skonstruowali mikroskop sił atomowych. Obydwa mikroskopy stały się potężnymi narzędziami do budowania dowolnych struktur złożonych z określonej liczby atomów i cząsteczek, do przesuwania, zginania i obracania nanodrutów, nanorurek oraz innych nanoobjektów. To była nie tylko rewolucja technologiczna, ale również rewolucja działania. W ogromnej mierze dotychczasowy postęp cywilizacyjny był związany z minia-



Rys. Krzysztof Mazur

turyzacją, tj. wytwarzaniem urządzeń, których rozmiary wraz z upływem czasu i towarzyszącym mu wzrostem poziomu technologii stawały się coraz mniejsze. Znamienym przykładem mogą być komputery i ich elementy. Nanotechnologia wymaga odwrotnego postępowania: z małych elementów (atomów i molekuł) buduje coraz większe i coraz bardziej złożone obiekty. W takim przypadku prognozowanie i planowanie jest mniej przewidywalne. Warto uzmysłowić sobie, że np. własności układu zbudowanego ze 120 molekuł mogą być radykalnie inne niż własności podobnego układu zbudowanego ze 121 molekuł. Nowe procesy wytwarzania wymagać będą nowych koncepcji

pracy naukowej. Fizycy teoretycy zdają sobie sprawę jak trudno jest zbadać kwantowomechaniczny układ zbudowany ze 120 molekuł, który na dodatek musi „pracować” w realnych temperaturach. Dokładne uwzględnienie zjawisk kwantowych i dyssypacji dla takiego układu nie jest obecnie możliwe. A to zaledwie 120 molekuł. Prawdopodobnie w przyszłości badania naukowe w większej mierze będą przypominać rzemieślniczą pracę z zastosowaniem komputerów potężnej mocy (być może zbudowanych z qubitów i qutritów), niż pionierską pracę twórczą. Te wizjonerskie dywagacje nie są nowe. Richard P. Feynman, pierwszy wizjoner nauk ścisłych, 29 grudnia 1959 roku na posiedzeniu Amerykańskiego Towarzystwa Fizycznego, które odbywało się wówczas w Kalifornijskim Instytucie Technologicznym (Caltech), wygłosił nietypowy wykład. Jego tytuł brzmiał: „There’s Plenty of Room at the Bottom”, a podtytuł: „An Invitation to Enter a New Field of Physics”. Nie jest to odpowiednie miejsce na omawianie tego wykładu. Odsyłamy czytelnika do oryginału. Warto go przeczytać! Przytoczmy tu jednak jedno niezwykle istotne zdanie: „What I want to talk about is the problem of manipulating and controlling things on a small scale”. Tą małą skalą jest oczywiście skala atomowa. Data tego wykładu (à propos, Koledzy Fizycy, czy u nas są możliwe posiedzenia 29 grudnia?) może być uznana za dzień narodzin teoretycznych nanonauk, a jego mówca może być uważany za ojca nanofizyki.

Dzisiaj, po upływie dwudziestu lat od osiągnięcia Eiglera, wyzwaniem nie jest budowa statycznych struktur w skali nano. Wyzwaniem jest budowa dynamicznych nanoobjektów, które mogą coś robić w sposób kontrolowany. Istnieją liczne naukowe i nienaukowe doniesienia o zbudowaniu różnego rodzaju maszyn molekularnych. „Wyścig szczurów” trwa od dwóch dekad: entuzjastyczne doniesienia mówią o zbudowaniu transporterów, rotorów, podnośników, przełączników, blokad, trybów. Zaangażowani są w to naukowcy z prawie wszystkich przyrodniczych i stosowanych nauk: fizycy, chemicy, biolodzy i technolodzy wszelkiej maści. Należy mieć do tego odpowiedni dystans, aby móc odróżnić naukę od pseudonauki czy paranoi (czy wszyscy do końca potrafią odróżnić astrofizykę od astrologii?). Wydaje się, że postęp w konstrukcji najprostszych nanomaszyn, zbudowanych z molekularnych jednostek, jest zauważalny. Przykładem może być nanowinda, dla której paliwem jest światło, i która może przemieszczać się o 1,4 nm lub platforma molekularna przesuwająca się o 0,7 nm i generująca siłę rzędu 200 pN.

Natomiast planowa konstrukcja dynamicznych obiektów złożonych z molekularnych jednostek, które miałyby wykonywać określone kontrolowane i skoordynowane procesy, jest obecnie niemożliwa. Wyjątkiem są przełączniki (switches) intensywnie konstruowane i badane przez chemików. Musimy zdać sobie sprawę z tego, że nasze myślenie o maszynach molekularnych jest obciążone myśleniem „makroskopowym”, opartym na doświadczeniach z życia codziennego. Maszyny w skali nano i mezo pracują w diametralnie innym świecie. Wymieńmy kilka różnic:

A. Maszyny molekularne to obiekty zbudowane nie ze sztywnych metalowych czy plastikowych elementów, ale z materii miękkiej, z elastycznych makromolekuł, które mogą zmieniać swój kształt i być w wielu różnych stanach. Często są to łańcuchy połączone określonymi jednostkami i przerwanie łańcucha zmienia nie tylko kształt, ale i jego własności. W makroświecie, np. przecięcie metalowego miedzianego drutu nie zmienia jego własności.

B. Maszyny molekularne pracują w świecie ruchów Browna, w świecie fluktuacji, turbulencji i zaburzeń losowych; są one non stop bombardowane cząstkami otoczenia i fotonami pochodzącymi ze spontanicznej emisji i absorpcji. Niestety nie możemy tych fluktuacji wyeliminować. Ruchy Browna w mikroświecie to nie tylko problem środowiska i temperatury, w których pracuje maszyna. To głównie problem skali rozmiarów, w jakiej ona pracuje. W makroświecie boimy się fluktuacji i procesów losowych. Walczymy z nimi, starając się je maksymalnie eliminować. Często się nam to udaje (np. fluktuacje napięcia w domowej sieci elektrycznej muszą być zredukowane do takiego poziomu, aby nie zniszczył telewizora, radia czy komputera). W mikroświecie nie możemy zredukować fluktuacji. Jedyne co można zrobić, to nauczyć się współżyć z fluktuacjami i maksymalnie je wykorzystywać. Ta swoista symbioza z fluktuacjami bywa nie tylko pożyteczna, ale i często konieczna do tego, by maszyna mogła działać.

C. Częstą konsekwencją tego, o czym piszemy w punkcie B, jest zaniedbywalny wpływ bezwładności na ruch nanoobjektu. To tak, jak gdyby formalnie założyć, że jego masa jest zerowa. Jest to konsekwencja względnie dużego tarcia, które zawsze istnieje w świecie fluktuacji. W makroświecie obowiązuje II zasada dynamiki Newtona: siła działająca na ciało jest równa szybkości zmian pędu tego ciała. Ale czy tak jest zawsze? Na przykład, gdy na cząstkę w miodzie działamy niewielką siłą, to czy cząstka porusza się ruchem przyspieszonym? W skali mikro jest jak w miodzie: siła na ogół jest proporcjonalna do prędkości ciała, a nie do przyspieszenia, które jest zaniedbywalnie małe w „normalnej” skali czasowej. Parametrem determinującym relacje między siłą bezwładności i siłą tarcia jest liczba Reynoldsa. W makroświecie jej wartość waha się od setek do miliardów, co świadczy o dominującej roli sił bezwładności i masy. W układach biologicznych liczba Reynoldsa jest w granicach od setnych do milionowych części ułamka. Tak więc o sile bezwładności można zapomnieć, ale trzeba bezwzględnie pamiętać o sile tarcia.

D. Maszyny molekularne, aby wykonać swoje zadania, wymagają szczególnego rodzaju paliw napędowych, szczególnego otoczenia i innych komplementarnych obiektów. Często otoczenie musi być regulowane i dopasowywane do określonych procesów przez inne nanomaszyny.

Ponieważ dynamika w makro- i mikroświatach jest całkowicie odmienna, nie może więc dziwić, że mechanizm kontroli ruchu motorów molekularnych musi być zdecydowanie inny.

Czy wobec powyższych różnic jesteśmy w stanie zbudować molekularne maszyny? Podpowiedzią, jak to uczy-

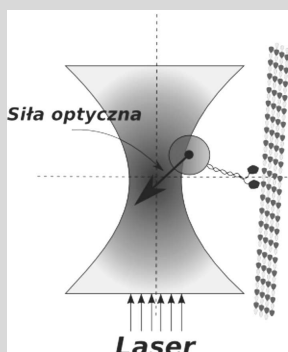
WSTAWKA 1

SZCZYPCE OPTYCZNE

Zbadanie mechanizmów funkcjonowania biologicznych motorów molekularnych wymaga specjalnych narzędzi. Jedną z najbardziej fascynujących technik jest manipulacja pojedynczymi molekułami za pomocą tak zwanych szczypiec optycznych. Umożliwiają one uwięzienie nano- i mikroobiektów dielektrycznych w środowisku wodnym. Do takich obiektów należy m.in. polistyrenowa mikrokulka, do której można chemicznie dowiązać badaną molekułę motoryczną.

Zasada działania szczypiec optycznych jest bardzo prosta. Jak wiadomo, płytka dielektryka wciągana jest do wnętrza kondensatora. Zjawisko to można wykorzystać do pułapkowania mikrocząstek. Problemem jest jednak wytworzenie pola elektrycznego w ściśle określonym miejscu i do tego w roztworze wodnym. Rozwiązaniem okazało się użycie skoncentrowanego światła, na przykład promienia lasera skoncentrowanego przez obiektyw o dużej numerycznej aperturze.

Metoda ta doskonale nadaje się do zastosowania w konwencjonalnym mikroskopie, przy czym stosując odpowiednie filtry dichroiczne wykorzystuje się tę samą drogę optyczną do pułapkowania i obserwacji obiektów.



Pułapka optyczna: kulka dielektryczna przyciągana jest do punktu o maksymalnym natężeniu światła laserowego. Do kulki można chemicznie dołączyć proteinę motoryczną i za jej pomocą manipulować nią i mierzyć siły.

Szywność pułapki optycznej rośnie wraz z mocą stosowanego lasera i może być płynnie sterowana. By „złapać” mikrometrową kulkę polistyrenową, wystarczy od kilku do kilkunastu miliwatów mocy. Pułapką można sterować mechanicznie, przemieszczając lustro na drodze optycznej promienia laserowego. Większe możliwości daje wykorzystanie spójnego charakteru światła i należy zmodyfikować front falowy tak, by pułapka miała żądany kształt. Zastosowanie ciekłokrystalicznych przestrzennych modulatorów frontu falowego pozwala na sterowanie w czasie rzeczywistym.

Pułapka optyczna jest też precyzyjnym i bezkontaktowym siłomierzem, operującym w zakresie pikoniutonów (pN). Kalibrację takiego siłomierza przeprowadza się, mierząc średnie kwadratowe odchylenie pułapkowanej cząsteczki o znanym rozmiarze pod wpływem fluktuacji termicznych.

nić, może być Natura i możemy być my sami. To Natura w trakcie swej ewolucji wykształciła perfekcyjne maszyny i motory molekularne. Występują one w niezliczonych ilościach w naszym ciele, w komórkach naszych organizmów. Krok po kroku, dzięki osiągnięciom i aparaturze współczesnych nauk przyrodniczych, w szczególności dzięki aparaturze wytworzonej na potrzeby fizyki, poznajemy cząstkowe mechanizmy działania motorów biologicznych. Są one niedoścignionym wzorcem. Wytworzone przez naukowców maszyny molekularne są na razie wielce ułomne i dalekie od maszynierii biologicznej. Pamiętajmy, że ludzie próbują konstruować maszyny molekularne za ledwie od kilkunastu lat. Natura potrzebowała milionów lat, aby dojść do dzisiejszego ideału. Należy podkreślić, że to właśnie próba zrozumienia funkcjonowania motorów biologicznych stała się inspiracją dla wielu fizyków do podjęcia badań naukowych na temat motorów molekularnych w szerokim znaczeniu tego słowa. Przyjrzyjmy się więc im bliżej.

Motory biologiczne

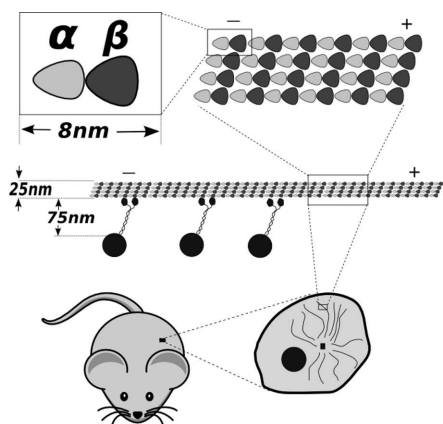
W dużym uproszczeniu komórkę biologiczną w naszym ciele można przyrównać do wielkiej aglomeracji miejskiej (coś à la Tokio). Aby mogła ona funkcjonować, musi posiadać infrastrukturę: sieć dróg, komunikację, sieć wodociagową zaopatrującą w wodę, kanalizację, energię, służby oczyszczania, straż pożarną, policję itp. Taka infrastruktura istnieje w naszych komórkach.

Dla przykładu, analogiem sieci dróg i autostrad są mikrotubule, po których pędzą biomotory.

Istnieją 3 wielkie rodziny biomotorów: miozyna (znana od 1864 roku), dyneina (odkryta w 1963 roku) i kinezyna (odkryta dopiero w 1985 roku). Motory te transportują wewnątrz komórek towary i ładunki z jednego miejsca do drugiego, biorą udział w segregacji chromosomów, podziale komórek, przemieszczaniu organelli, pęcherzyków, biorą udział w procesie transkrypcji genów. Fakt, że biomotory biorą udział w podziale komórek ma kluczowe znaczenie w medycynie chorób nowotworowych. Wiadomo przecież, że komórki nowotworowe rozmnażają się (tzn. następuje ich podział) w zawrotnym tempie. Bez pracy biomotorów proces rozmnażania takich komórek mógłby zostać zahamowany. Widać stąd, dlaczego zrozumienie mechanizmów ruchu motorów ma fundamentalne znaczenie. Praca wykonywana przez biomotory jest co prawda na poziomie mikroskopowym, ale ich działanie jest widoczne także w skali makroskopowej. Efekty pracy miozyny, pierwszego z odkrytych motorów, są widoczne w każdej chwili naszego życia. To miozyna jest odpowiedzialna za ruch dowolnej części naszego ciała: ręki, palca, powieki. Miliony, a czasem miliardy cząstek miozyny muszą w skoordynowany sposób działać jednocześnie, abyśmy mogli wykonywać typowe prace w domu, ogrodzie czy przy sportowym – czasami ekstremalnym – wysiłku fizycznym.

Autostrady wewnątrzkomórkowe, jakimi są mikrotubule, mają strukturę cylindra o długości od ułamka

mikrona do setek mikronów i średnicy 25 nm. Cylinder składa się z 13 włókien, zwanych protofilamentami (dlaczego pechowa trzynastka?). Każde włókno ma przestrzenną strukturę periodyczną zbudowaną z naprzemiennych 2 białek (z grubsza kulek): α -tubuliny i β -tubuliny. Sekwencja przestrzenna α - β - α - β ... ma okres 8 nm i cechę zwaną polarnością (biegunowością): jeden koniec jest oznaczany jako biegun plus, drugi – jako biegun minus. Do końca plus łatwiej przyczepiają się tubule. Innymi słowy, mikrotubula rośnie szybciej od strony plus. Ta cecha biegunowości ma ogromne znaczenie dla motorów biologicznych. Są one w stanie rozpoznać, w którym kierunku mają się poruszać. Czy od plusa do minusa czy odwrotnie. Mikrotubule nie są strukturami stałymi. Są to struktury dynamiczne, często zmieniające swą długość. Nie są losowo rozłożone. Na ogół radialnie rozchodzą się od centrosomu (który znajduje się w pobliżu jądra). Końce (–) są w pobliżu centrosomu, natomiast końce (+) rozchodzą się na peryferie komórki.



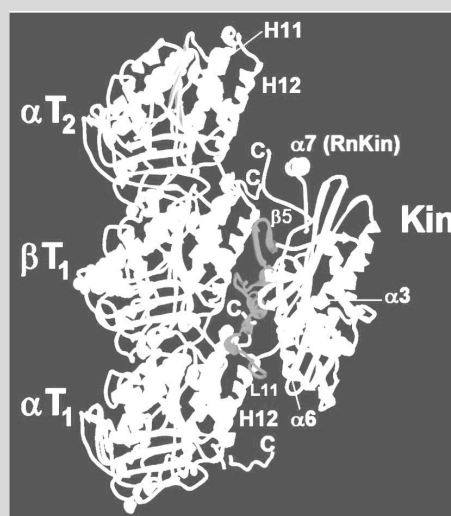
Rys. 1. We wnętrzu typowej komórki, tuż obok jądra (ciemne koło), znajduje się centrosom (mały prostokąt). Z jego otoczenia rozchodzi się sieć komórkowych autostrad-mikrotubuli w kierunku peryferii komórki. Po nich poruszają się biomotory: kinezy i dyneiny. Kinezyzna wozi towar na peryferie, dyneina mknie w przeciwnym kierunku – ku centrum. Mikrotubula jest białkiem będącym polimerem o kształcie rurki (o średnicy 25 nm) składającym się z 13 nici-protofilamentów, a każdy protofilament będący strukturą przestrzennie periodyczną o okresie 8 nm zbudowany jest z dwóch białek: α - i β -tubuliny. Mikrotubula posiada wyróżniony kierunek (bieguny minus i plus). Kierunek poruszania się kinezy i dyneiny jest ściśle związany z biegunowością (polarnością). Kinezyzna porusza się od minusa do plusa. Chodzi ona na dwóch głowach (dlaczego nie na nogach? ach ci biolodzy!). Jej długość wynosi 75 nm. Do jej końca przyczepiają się towary (np. pęcherzyki), które są transportowane tam, gdzie są potrzebne.

Badania pokazują, że dyneina zawsze transportuje towar w kierunku bieguna minus mikrotubuli, natomiast konwencjonalna kinezyzna – w kierunku końca plus. W późniejszych latach odkryto wiele innych (kilkaset) odmian biomotorów oraz dokonano ich klasyfikacji, łącząc różne motory w superrodziny. Okazało się, że istnieją takie

odmiany kinezyzny, które pędzą w kierunku przeciwnym do ruchu konwencjonalnej kinezyzny. Oczywiście od razu nasuwa się naturalne pytanie, skąd taka (bezmózgowa) kinezyzna wie, w którą stronę ma transportować towary? Oczywiście pytań jest o wiele więcej:

- Dlaczego kinezyzna się porusza?
- Jak kinezyzna się porusza wzdłuż mikrotubuli (ruchem spiralnym po cylindrze, a może tylko wzdłuż jednego protofilamentu)?
- Czy jest to ruch deterministyczny czy losowy?
- Z jaką prędkością pędzi kinezyzna?
- Gdzie i jak transportowany towar przyczepia się i odczepia się od kinezyzny?

WSTAWKA 2



Struktura atomowa otrzymana metodami rentgenowskimi fragmentu jednego z włókien mikrotubuli, tzn. protofilamentu składającego się z α - i β -tubuliny oraz kinezyzny (Kin). Odległość między najbliższymi cząstkami dwóch α -tubulin wynosi 8 nm. Zauważmy, jak skomplikowane są to struktury zawierające białka o masie kilkuset tysięcy mas atomowych. Na przykład, kinezyzna to obiekt o rozmiarze kilku nanometrów i masie 360 kDa (kilodaltonów) wykonująca 225 kroków w ciągu sekundy. Jeden krok to 8 nm. Fizycy zwykli modelować takie obiekty jako punktowe cząstki o masie M . Czasami traktują te obiekty jak kulki o promieniu R . Ilustracja pochodzi ze strony <http://www.mpasmb-hamburg.mpg.de> (za zgodą i dzięki uprzejmości prof. dr Eckharda Mandelkowa).

Te i inne pytania dotyczą też pozostałych motorów biologicznych. Odpowiedzi na nie poszukuje wiele grup badawczych na całym świecie. Niektórym wydaje się, że je znaleźli. Ale czy do końca rozumiemy te procesy? Wiadomo od dawna, że paliwem dla biomotorów jest ATP (adenozynotrójfosforan). Podczas reakcji hydrolizy ($\text{ATP} \rightarrow \text{ADP}$) uwalnia się energia rzędu 20 kT (k – stała Boltzmann, T – temperatura wewnątrz komórki), która jest źródłem energii dla motoru. Kinezyzna w jednym cyklu reakcji $\text{ATP} \rightarrow \text{ADP}$ wykonuje ruch o długości 8 nm, czyli tyle

ile wynosi okres przestrzenny mikrotubuli. Z badań wynika, że w procesie transportu kinezyzna „przyczepia się” tylko do β -tubuliny i, jak mawiają biolodzy, transport jest procesywny, tzn. kinezyzna wykonuje dziesiątki–setki kroków bez odrywania się od mikrotubuli. Kinezyzna może być źródłem siły o maksymalnej wartości 5–7 pN. Kojarząc to z odległością 8 nm otrzymamy wartość pracy 40 pN·nm na jeden cykl $\text{ATP} \rightarrow \text{ADP}$. To odpowiada sprawności 50% (maksymalna energia z hydrolizy ATP wynosi 80 pN·nm). Konwencjonalna kinezyzna porusza się z prędkością 1800 nm/s, czyli może wykonać 225 kroków na sekundę. Jak marnie wygląda przy tym człowiek, który może wykonać zaledwie kilka kroków na sekundę. Istnieją także motory znacznie szybsze od konwencjonalnej kinezyzny.

Skale w komórce biologicznej

Po części opisowej, przejdźmy do bardziej konkretnych zagadnień. Musimy zdawać sobie sprawę w jakim świecie rozpatrujemy ruch motorów biologicznych. Innymi słowy, jakie skale wielkości, odległości, czasu, energii i sił istnieją w komórkach biologicznych.

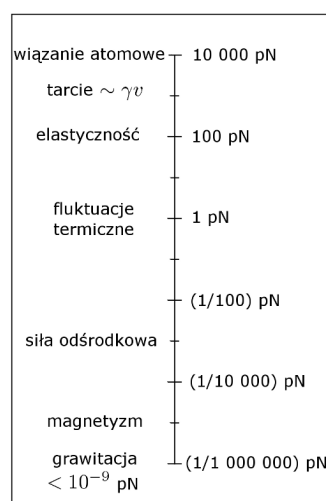
Rozmiar typowej komórki biologicznej wynosi 10–100 mikronów. Są też komórki o długości 100 cm (włókna nerwowe). Komórka przypomina zupełną kuleczkę z wody (70–90%), z zawartością molekuł (cukry proste, aminokwasy, witaminy, białka, kwasy nukleinowe, jony sodu, potasu, wapnia i wielu innych cząsteczek). Najmniejszymi składnikami są jony, następnie małe molekuly (jak glukoza), makromolekuly (jak hemoglobina) i organelle (rzędu 100 nm). W typowej komórce jest 10 000 różnych białek i średnio milion białek każdego typu, co daje w sumie 10 miliardów wszystkich białek. Typowa liczba makrocząsteczek jest rzędu 10^7 . Taka komórka-zupa, aby się nie rozlała, jest otoczona błoną komórkową o grubości 6–10 nm. Całkowita masa komórki jest rzędu 10^{-9} g.

Najważniejszym źródłem energii w komórce są reakcje hydrolizy ATP. Podczas jednej reakcji jest uwalniana energia rzędu $100 \cdot 10^{-21}$ J. Dla porównania, energia fotonu o długości 500 nm (światło zielone) wynosi $397 \cdot 10^{-21}$ J. Tego rzędu energia jest dostarczana podczas fotosyntezy. Pamiętajmy także o energii fluktuacji termicznych, $E = kT$. Przy temperaturze $T = 37$ C, energia fluktuacji $E = 4,28 \cdot 10^{-19}$ J.

Energia elektryczna może być oszacowana przy wzięciu pod uwagę potencjału V membrany, którego wartość przyjmijmy jako równą 70 mV. Wówczas $E = eV = 11 \cdot 10^{-21}$ J. Widać, że wszystkie formy energii wahają się w granicach $(1\text{--}400) \cdot 10^{-21}$ J. Skalę sił przedstawiamy na rys. 2.

Korzystając ze standardowych wyników fizyki statystycznej, możemy oszacować typowe wielkości dla cząstek w komórce. Jeżeli cząstki przemieszczają się (w wodzie) tylko ruchem dyfuzyjnym (ruchem Browna), to na pokonanie dystansu 100 mikronów potrzeba: 2,5 s dla jonu potasu, 50 s dla białka i około 3 godzin dla organeli. Jeżeli ten

dystans skrócimy do 1 mikrona, to czasy te wynoszą odpowiednio: 0,25 s, 5 ms i 1 s. Pamiętajmy, że jest to średnie kwadratowe przemieszczenie $x = \sqrt{2Dt}$, gdzie współczynnik dyfuzji $D = kT/\gamma$, współczynnik tarcia γ jest dany przez wzór Stokesa, $\gamma = 6\pi\eta R$, η jest współczynnikiem lepkości wody, a R jest liniowym rozmiarem dyfundującej cząstki. Podobnie, korzystając z zasady ekwipartycji energii, możemy oszacować średnią kwadratową prędkość cząstek, $v = \sqrt{3kT/M}$, gdzie M jest masą cząstki. Dla typowego białka otrzymamy prędkość rzędu 10 m/s. Czas korelacji prędkości takiego białka wynosi $\tau = m/\gamma = 3$ ps. To oznacza, że jakiekolwiek zmiany położenia białka mogą zostać zaobserwowane po czasie znacznie dłuższym od tego czasu korelacji, ponieważ w czasie 3 ps białko może się przesunąć zaledwie na odległość 1/300 swojego rozmiaru.



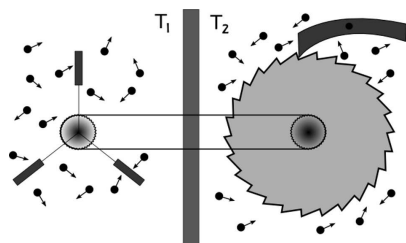
Rys. 2. Skala sił w komórce biologicznej

Mechanizm zębatkowy

Biolodzy od dawna sugerowali, że transport biomotorów można wyjaśnić za pomocą mechanizmu zębatkowego. Od końca XIX wieku zębatka służyła jako przykład obalający II zasadę termodynamiki. Marian Smoluchowski był pierwszą osobą, która pokazała, że mechanizm zębatkowy nie obala II zasady termodynamiki. Cytowany już Pierwszy Wizjoner, R.P. Feynman, wykorzystuje ten przykład w swoich słynnych *Wykładach z fizyki*, aby pokazać, że zębatka wraz z zapadką pracuje w zgodzie z II zasadą termodynamiki. Warto wspomnieć, że ignorancja w szerokim świecie jest tak wielka, że aż 99% badaczy cytuje „Feynman ratchet and pawl machine” zamiast „Smoluchowski ratchet and pawl machine”.

Na rys. 3 przedstawiony jest schemat takiego urządzenia. Lewa część urządzenia składa się z łopatek i jest umieszczona w zbiorniku z gazem o temperaturze T_1 , natomiast prawa część urządzenia składa się z zębatki wraz ze sprężynującą zapadką (sprężynka nie jest uwidoczniiona na rysunku) i jest w zbiorniku z gazem o temperaturze T_2 . Obydwie części są mechanicznie połączone w taki sposób,

że obrót łopatek przenosi się na obrót zębatki i odwrotnie. Na rysunku jest to pas transmisyjny, ale inne rozwiązania są równie dobre. Historycznie, w książkach i wielu pracach, urządzenie jest połączone osią. Aby uniknąć wizualizacji przestrzennej, zdecydowaliśmy się przedstawić własne oryginalne połączenie mechaniczne. Ponieważ jest to eksperyment myślowy, nie ma to większego znaczenia.



Rys. 3. „Smoluchowski–Feynman Brownian Ratchet and Pawl” – mechanizm zapadkowy zębatki Browna. W dwóch odizolowanych zbiornikach znajduje się gaz. W lewym zbiorniku, o temperaturze T_1 , fluktuacje pojedynczych molekuł mogą losowo poruszać łopatkami wiatraka. Wiatrak jest mechanicznie połączony w mechanizmem zębatkowym w prawym zbiorniku, którego temperatura wynosi T_2 . Zapadka umożliwia ruch zębatki tylko w jedną stronę. Czy taki układ jest perpetuum mobile – urządzeniem czerpiącym energię z termicznych fluktuacji równowagowych? Okazuje się, że nie. Błędem w rozumowaniu jest przeniesienie makroskopowej intuicji do mikroskopowego świata, w którym niemożliwe jest zbudowanie mechanizmu umożliwiającego ruch w jedną stronę. Trzeba wziąć pod uwagę, że mechanizm zapadkowy również jest pod wpływem fluktuacji termicznych. W takim przypadku, jeśli oba układy będą w tej samej temperaturze $T_1 = T_2$, cały układ nie będzie w stanie wykonywać pracy mechanicznej. Natomiast w przypadku, gdy $T_1 > T_2$, układ będzie działał jak silnik cieplny i praca będzie wykonana kosztem transportu ciepła ze zbiornika lewego do prawego.

Cząsteczki gazu w lewym zbiorniku uderzają w łopatki ze wszystkich stron i łopatka wykonuje chaotyczne ruchy Browna, obracając się losowo w jedną lub drugą stronę. Z symetrii wynika, że średnia siła wypadkowa działająca na łopatki jest zerowa i łopatka nie wykonuje obrotów średnio w jedną stronę. Sytuacja w prawym zbiorniku jest inna. Tam nie ma tej symetrii, ponieważ zęby są niesymetryczne i zapadka blokuje ruch w jedną stronę. Mogłoby się wydawać, że urządzenie będzie obracać się tylko w jedną stronę, nawet gdy temperatury w obu zbiornikach są takie same. Wówczas, po niewielkiej modyfikacji urządzenia, moglibyśmy podnosić i opuszczać małe obiekty. Oznaczałoby to, że urządzenie wykonuje pracę i to kosztem termicznych fluktuacji cząstek gazu. Oczywiście łamałoby to II zasadę termodynamiki w przypadku, gdy $T_1 = T_2$. Paradoks ten wyjaśnia Feynman w swoich *Wykładach*. Jeżeli temperatury w obu zbiornikach są różne, maszyna obraca się w jedną stronę, ale to nie przeczy niczemu.

W tym miejscu warto przytoczyć argument, że R.P. Feynman chyba gdzieś czytał lub słyszał o wywodach

M. Smoluchowskiego, które są zawarte w pracy: M. Smoluchowski, „O fluktuacjach termodynamicznych i ruchach Browna”, *Prace Matematyczno-Fizyczne*, tom 25, s. 187–263, Warszawa, 1914.

W paragrafie 46, poświęconym perpetuum mobile, M. Smoluchowski pisze: „Jeszcze przejrzystsza będzie konstrukcja, jeżeli użyjemy przyrządu, opisanego w końcowym ustępie §39, ale zamiast zwierciadła umieścimy koło zębate i połączymy z nim zatrzask dopuszczający obrót koła w jedną stronę. Automatyczne fluktuacje powinny wtedy obracać koło jednostronnie; ciągle okracanie włókna przedstawiałoby trwałe źródło pracy, które można by wyzyskać za pomocą odpowiedniego urządzenia, połączonego z górnym końcem włókna. Nie chodzi wcale o kwestię trudności technicznej przy wykonaniu przyrządu; nie ma tu granic nieprzekraczalnych.

Sądzymy jednak, że zbudowanie takiego perpetuum mobile jest niemożliwe ze względów zasadniczych; w dwóch wymienionych właśnie przypadkach właściwą przeszkodą jest zasadnicza niemożliwość sporządzenia odpowiednich zatrzasków czy wentyli. Funkcjonowanie takich przyrządów daje się do tego sprowadzić, że w zwykłych warunkach pozostają one w pozycji (względnie) minimalnej energii potencjalnej. Jeżeli w ogóle mają się poddawać pod wpływem słabych impulsów ruchu Browna, sprężyny w nich czynne muszą być tak podatne, że tym bardziej już ich własny ruch Browna musi się uwidocznić; wskutek tego zatrzask czy wentyl w ogóle nie będzie funkcjonował jednostronnie. Nastąpi rozkład wszystkich możliwych pozycji w myśl prawa (61), które musi obejmować także wszystkie tego rodzaju przyrządy mechaniczne, bez względu na szczegóły konstrukcyjne, jako szczególne przypadki”.

Wypisz, wymaluj, wywody Feynmana są takie same.

Co zębatka opisana powyżej ma wspólnego z motorami biologicznymi? Wydaje się na pierwszy rzut oka, że chyba niewiele. Ale to nie jest prawdą. Pobieźna analiza gadżetu Smoluchowskiego–Feynmana pozwala na następujące obserwacje:

1. Zębatka jest strukturą periodyczną.
2. Ta struktura periodyczna jest asymetryczna ponieważ zęby nie są symetryczne.
3. Na układ nie działa żadna siła powodująca systematyczny ruch w jedną stronę i średnia siła pochodząca od fluktuacji termicznych jest zerowa.
4. Aby pojawił się transport (tzn. obrót średnio w jedną stronę), temperatury w obu zbiornikach muszą być różne. Innymi słowy, w języku fizyki statystycznej, musi być naruszona zasada równowagi szczegółowej. Albo jeszcze inaczej, układ nie może być w stanie równowagi termodynamicznej.

Teraz możemy zauważyć pewne podobieństwa do problemu transportu biomotorów:

1. Mikrotubula jest przestrzennie strukturą periodyczną.
2. Mikrotubula ma polarność, tzn. ma wyróżniony koniec plus i koniec minus. Wynika to z tego, że tubulina α wiąże cząsteczkę GTP w sposób nieodwracalny.

Tubulina β wiąże także cząsteczkę GTP, ale w sposób odwracalny: GTP hydrolizuje do GDP. Zwróćmy uwagę na to, że gdyby protofilament był zbudowany tylko z jednej tubuliny, nie można by zrobić tego rozróżnienia i układ byłby w pełni symetryczny. Można wnioskować, że 2 różne tubuliny, tzn. α -tubulina i β -tubulina, to minimalna liczba potrzebna dla asymetryczności. Natura zminimalizowała tę liczbę (czy ma to cokolwiek wspólnego z zasadą najmniejszego działania?).

3. Wewnątrz komórki nie ma sił systematycznych i gradientów działających na motor. Motor pracuje w środowisku wodnym i fluktuacje termiczne nie powodują systematycznych ruchów.

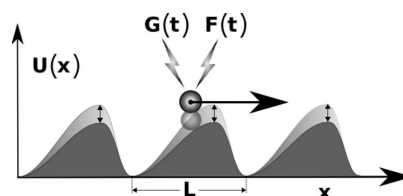
4. Transport pojawia się, ponieważ istnieje źródło energii (z reakcji ATP) i układ nie jest w stanie równowagi termodynamicznej. Zasada równowagi szczegółowej jest naruszona. Zauważmy, że źródło energii jakim jest tu reakcja chemiczna nie wyróżnia kierunku przestrzennego.

Modelowanie: abstrahowanie

Fizyk, spostrzegając podobieństwa 1–4 zarówno w gadźecie zębatkowym, jak i w biomotorach, może dokonać abstrakcji w stylu matematyka: odrzucić zbędne detale i skonstruować minimalny model motoru molekularnego. Nie musi on opisywać ani zębatki Smoluchowskiego–Feynmana ani kinezyzny poruszającej się po mikrotubuli. Oczywiście później można komplikować model w dowolny sposób, starając się przybliżyć do opisu rzeczywistego układu. Jednakże minimalny model ma ogromną zaletę: jest on prosty, zrozumiały, być może ujmuje istotę procesu i wyjaśnia najważniejsze własności transportowe. Dla przykładu, czy badając model minimalny możemy zrozumieć dlaczego jedna cząstka (np. kinezyzna) porusza się od minusa do plusa, a inna (np. dyneina) – w kierunku przeciwnym? Jaka prosta cecha cząstki może o tym decydować?

Zaczynamy więc budować najprostszy model. Po pierwsze, motor mimo swej skomplikowanej struktury i kształtu (porównaj wstawka 2) jest rozważany jako cząstka punktowa o masie M . Rozważamy jednowymiarowy ruch klasycznej (tzn. nie kwantowej) cząstki, scharakteryzowany przez jej położenie x . Biorąc pod uwagę punkt 1, zakładamy, że cząstka porusza się w przestrzeni periodycznym potencjale $U(x) = U(x+L)$, gdzie L jest okresem struktury periodycznej. Uwzględniając punkt 2, należy przyjąć, że potencjał $U(x)$ jest asymetryczny, tzn. ma złamaną symetrię odbicia $U(x) = U(-x)$. Dalej, ponieważ istnieją fluktuacje termiczne, na cząstkę działa losowa siła $F(t)$ pochodząca od zderzeń z cząstkami otoczenia (np. wody w komórce). Średnia wartość tej siły wynosi zero (patrz punkt 3). Ponieważ otoczenie nie jest próżnią, na cząstkę działa też siła tarcia proporcjonalna do prędkości cząstki v i współczynnika tarcia γ . Należy jeszcze uwzględnić źródło energii (patrz punkt 4) do napędzania cząstki. Może to być dowolna siła $G(t)$ (ale nie może to być siła pochodząca z termicznych fluktuacji). Jeżeli ta siła nie jest średnio równa zero, to w sposób trywialny otrzymamy transport cząstki, ponieważ niezerowa średnia

wartość siły wyróżnia kierunek przestrzenny. Natomiast nietrywialnym przypadkiem jest deterministyczna lub losowa siła $G(t)$ o zerowej wartości średniej. Na przykład deterministyczna siła $G(t) = A \cos(\omega t + \varphi_0)$ uśredniona po jednym okresie $T = 1/\omega$ ma wartość zero. Przykładem siły losowej może być losowy ciąg impulsów o statystyce Poissona (tzw. szum śrutowy czy poissonowski biały szum) lub szum dychotomiczny (skorelowany dwustanowy proces Markowa). Jeżeli uwzględnimy wszystko to, o czym piszemy powyżej, to dynamika cząstki Browna jest określona przez równanie Newtona w następującej postaci (patrz rys. 4): $Ma + \gamma v = -U'(x) + G(t) + F(t)$. W równaniu tym wszystkie trzy siły są średnio zerowe: siła potencjalna $U'(x)$ uśredniona po okresie przestrzennym L , siła $G(t)$ uśredniona po okresie czasowym T lub realizacjach losowych, gdy jest to siła losowa oraz siła zderzeń losowych $F(t)$ uśredniona po wszystkich realizacjach szumu termicznego. Tak więc, średnio na cząstkę działa zerowa siła, a mimo to średnia prędkość cząstki może być niezerowa i cząstka może być transportowana na dowolne odległości.



Rys. 4. Abstrakcyjny model motoru molekularnego. Działają na niego dwie siły: $F(t)$ – stochastyczna, o średniej wartości zero, pochodząca od równowagowych fluktuacji termicznych oraz $G(t)$ – nierównowagowa siła zależna od czasu o zerowej wartości średniej. Cząsteczka znajduje się w przestrzeni periodycznym potencjale $U(x) = U(x+L)$ o okresie L , który również może być modulowany w czasie przez czynnik nierównowagowy. Fizyk bada warunki, jakie muszą być spełnione, aby taki układ miał własności transportowe, tzn. aby średnia prędkość cząsteczki w stanie stacjonarnym była różna od zera.

Jeżeli zastosujemy to równanie do cząstek w komórce biologicznej (abstrahując od tego, czy jest to uprawnione, czy nie), to wówczas w odpowiednio przeskalowanym bezwymiarowym równaniu Newtona otrzymamy bezwymiarowy współczynnik tarcia $\gamma = 1$. Natomiast przeskalowana bezwymiarowa masa M będzie miliony razy mniejsza od 1. Oznacza to, że możemy pominąć cechę cząstki, jaką jest jej bezwładność scharakteryzowana przez masę. Uzasadnia to w inny sposób stwierdzenie zawarte w punkcie C Wstępu o przetłumionym charakterze ruchu. Stwierdzenie to było tam oparte na wartości liczby Reynoldsa.

W ciągu ostatnich kilkunastu lat opublikowano kilka tysięcy prac na temat transportu w układach opisanych powyższym równaniem Newtona i jego modyfikacjami.

Wydaje się, że współczesne badania naukowe, w których podstawowym równaniem jest równanie Newtona, nie mogą prowadzić do ciekawych wyników, nie mówiąc już

o odkryciach (przecież równanie Newtona jest znane od kilkuset lat!). Taka jest opinia wielu naszych kolegów. Nie chcemy być zgryźliwi, dlatego nie będziemy polemizować z takimi opiniami.

Co wynika z tych kilku tysięcy prac?

– Po pierwsze, badano takie warunki, aby w stanie stacjonarnym (dla długich czasów) można było transportować cząstki. Kluczowym warunkiem jest złamanie symetrii, czy to przestrzennej symetrii odbicia ukrytej w potencjale $U(x)$, czy to symetrii nierównowagowych fluktuacji $G(t)$, czy też czasowej symetrii deterministycznej siły $G(t)$.

– Po drugie, pokazano że może istnieć optymalna temperatura, w której średnia prędkość cząstki jest maksymalna. Innymi słowy, fluktuacje termiczne odgrywają konstruktywną rolę i pomagają w transporcie. Być może to, że optymalna temperatura naszego ciała wynosi 36,6 °C, ma jakiś związek z tym zjawiskiem (rezonans stochastyczny?).

– Po trzecie, stwierdzono że cząstki mogą przemieszczać się w przeciwnych kierunkach w zależności od parametrów modelu, w szczególności w zależności od wielkości (rozmiaru liniowego) cząstki, w zależności od masy cząstek (jeżeli nie jest to przypadek ruchu przetłumionego) lub od niewielkiej deformacji potencjału.

Można zauważyć, że najprostszy model, oparty na równaniu Newtona, przewiduje podstawowe własności transportowe motorów biologicznych. Bardzo łatwo można uzasadnić, dlaczego model ten nie jest adekwatny do opisu biomotorów, ponieważ nie uwzględnia tego i tamtego, i jeszcze czegoś innego. I to jest święta prawda. Być może Natura, mimo skomplikowanej budowy swoich elementów, poddaje się czasami prostemu, żeby nie powiedzieć trywialnemu, opisowi swego funkcjonowania i działania.

Uwagi końcowe

Przedstawiliśmy podstawowe informacje o naturalnych motorach molekularnych, które wytworzyła sama Natura. Mimo że przeprowadzono wiele badań na temat motorów biologicznych, na podstawowe pytania brak jest zadowalających odpowiedzi. Biolodzy dają do dyspozycji fizyków i chemików mnóstwo danych doświadczalnych. Wołają do nich o pomoc. I fizycy, i chemicy odpowiadają na wezwanie biologów, ale ogromna część badań poszła w innych kierunkach. Fizycy stworzyli wiele modeli tego, co dzisiaj nazywamy motorami czy zębatkami brownowskimi. Nie są to jednak modele transportu kinezyzny czy dyneiny. Chemicy poszli także w innym kierunku. Zaczęli budować innego rodzaju maszyny molekularne. Mają ogromne osiągnięcia w budowie molekularnych przełącz-

ników: cząsteczek o strukturze rotaksanu i katenanu. Paliwem dla takich przełączników może być promieniowanie elektromagnetyczne, zmiana pH środowiska, temperatury czy ładunku elektrycznego. Jednak, jak sami chemicy zaznaczają, nie należy tych maszyn utożsamiać z motorami molekularnymi. Widać też ciągle, jak mocno środowisko przyrodników jest rozproszone i jak często trudno jest znaleźć wspólny język, zrozumiały jednocześnie dla fizyków, chemików, biologów i materiałoznawców. Czasy i wyzwania wymagają interdyscyplinarnych badań. Sami fizycy czy biolodzy nie poradzą sobie.

Literatura

Aktualna wiedza o komórce biologicznej zawarta jest we wspaniałej książce [1]. Napisana jest przez wybitnych biologów molekularnych, bogato ilustrowana kolorowymi rysunkami i zdjęciami. Pełna wersja zaopatrzona jest w filmy i animacje. Z ubolewaniem przyznajemy, że nie znamy tak nowoczesnie wydanej książki z fizyki.

Książka [2] może być z łatwością czytana przez fizyków. Zawiera ona te elementy mechaniki klasycznej, fizyki statystycznej i kinetyki reakcji chemicznych, które są wykorzystywane do wyjaśniania i zrozumienia procesów w komórkach biologicznych. Zawiera mnóstwo oszacowań liczbowych wielkości fizycznych dla obiektów, które są składnikami komórek.

Praca [3] jest jednym z najnowszych artykułów przeglądowych na temat motorów molekularnych. Jest napisana przez chemików i dotyczy w głównej mierze maszyn molekularnych będących nanoprzełącznikami. Poprzedzona jest znakomitą wstępem i bogatym zbiorem literatury.

Pozycja [4] jest specjalnym wydaniem, zawierającym zaproszone, oryginalne artykuły naukowe napisane głównie przez fizyków.

Książka [5] w rozdziale 46 zawiera opis mechanizmu zapadkowego.

Szczegóły na temat struktury biomotorów i innych białek są zawarte na stronach internetowych RCSB Protein Data Bank. Polecamy także wyżej cytowaną stronę internetową: <http://www.mpasmb-hamburg.mpg.de>, która zawiera między innymi animacje ruchu kinezyzny.

- [1] Harvery Lodish i in., *Molecular Cell Biology* (W.H. Freeman and Company, New York 2004).
- [2] Jonathon Howard, *Mechanics of Motor Proteins and the Cytoskeleton* (Sinauer Associates, 2001).
- [3] E.R. Kay, D.A. Leigh, F. Zerbetto, „Synthetic Molecular Motors and Mechanical Machines”, *Angew. Chem. Int. Ed.* **46**, 72 (2007).
- [4] Special Issue containing articles on molecular motors, *Journal of Physics: Condensed Matter* **17**, zesz. 47 (2005).
- [5] R.P. Feynman, R.B. Leighton, M. Sands, *Feynmana wykłady z fizyki*, t. I, cz. 2 (PWN, Warszawa 1969).

MARCIN KOSTUR obecnie pracuje na stanowisku adiunkta w Instytucie Fizyki Uniwersytetu Śląskiego. W ostatnich latach przebywał na stażach naukowych w Niemczech i Stanach Zjednoczonych. Jego zainteresowania naukowe obejmują transport w układach z szumem, dynamikę stochastyczną, fizykę mikrocieczy oraz metody numeryczne.

JERZY ŁUCZKA jest profesorem zwyczajnym w Instytucie Fizyki Uniwersytetu Śląskiego. Jest kierownikiem Zakładu Fizyki Teoretycznej (www.zft.us.edu.pl). Jego zainteresowania naukowe dotyczą motorów molekularnych, procesów stochastycznych, kwantowych układów otwartych i fizyki układów mezoskopowych.

O modelach kosmologicznych i niektórych związanych z nimi nieporozumieniach

Andrzej Kasiński

Centrum Astronomiczne im. Mikołaja Kopernika, Polska Akademia Nauk, Warszawa

On cosmological models and some misunderstandings about them

Abstract: The article explains the need to use exact inhomogeneous models of the Universe and discusses examples of successful application of such models. The following topics are discussed: 1. It is reminded that the cosmological principle that is the basis of the classical Robertson–Walker models is merely a philosophical postulate which should be subject to observational verification (never carried out). 2. It is shown that the linearised Einstein equations used in cosmology to describe structure formation and evolution have shaky mathematical foundation. 3. The basic geometrical properties of the Robertson–Walker, Lemaître–Tolman (L–T) and Szekeres models are presented and compared with each other. 4. The following astrophysical applications of the L–T model are discussed: (a) Explanation of the dimming of type Ia supernovae without accelerated expansion of the Universe and without „dark energy”; (b) Description of formation of galaxy clusters and voids by exact models; (c) A solution of the „horizon problem” without inflationary models and their scalar fields; (d) The connection between increasing and decreasing modes of perturbation and geometrical characteristics of the models. A brief listing of the most important events in the history of inhomogeneous models is given.

W artykule uzasadniono potrzebę używania niejednorodnych modeli do opisu Wszechświata i podano przykłady skutecznych zastosowań takich modeli. Omówiono następujące tematy: 1. Przypomniano, że zasada kosmologiczna, która stanowi podstawę klasycznych modeli Robertsona–Walkera jest tylko postulatem filozoficznym, który należy poddać obserwacyjnej weryfikacji (czego dotychczas nie zrobiono). 2. Pokazano, że zlinearyzowane równania Einsteina używane w kosmologii do opisu powstawania i ewolucji struktur nie mają solidnych podstaw matematycznych. 3. Przedstawiono podstawowe własności geometryczne modeli Robertsona–Walkera, Lemaître’a–Tolmana (L–T) i Szekeresa i porównano je ze sobą. 4. Przedyskutowano następujące astrofizyczne wnioski z modelu L–T: (a) Wyjaśnienie pociemnienia supernowych typu Ia bez wprowadzania ekspansji z przyspieszeniem i bez „ciemnej energii”; (b) Opis powstawania gromad galaktyk i pustek za pomocą ścisłych modeli; (c) Rozwiązanie „problemu horyzontu” bez użycia modeli inflacyjnych i ich pól skalarnych; (d) Związek między rosnącymi i malejącymi modami zaburzenia a geometrycznymi charakterystykami modelu. W zakończeniu podano listę najważniejszych wydarzeń w historii modeli niejednorodnych.

1. Co to jest „zasada kosmologiczna”?

„Zasada kosmologiczna” wywodzi się pośrednio z odkrycia Kopernika, które można streścić tak: układ współrzędnych o początku w środku Słońca daje prostszy opis ruchu planet niż układ o początku w środku Ziemi. Potomni dorobili do tego (w wyniku dalszych odkryć) filozoficzny wniosek: Ziemia nie zajmuje uprzywilejowanej pozycji we Wszechświecie.

Później wniosek ten przybrał bardziej fundamentalistyczną formę: wszystkie położenia we Wszechświecie są równouprawnione, każdy obserwator, niezależnie od swo-

jego położenia, zaobserwowałby taki sam wielkoskalowy obraz Wszechświata.

Ta „zasada kosmologiczna” nie jest podsumowaniem wiedzy doświadczalnej, tylko aktem wiary, który pozuje na prawo fizyki. Akt ten został sformułowany na podstawie pierwszych teoretycznych modeli Wszechświata (modeli Robertsona–Walkera, R–W), jako ich uzasadnienie *ex post*.

Tak samo, jak elementy innych teorii wyprzedzających doświadczenie, zasada kosmologiczna powinna podlegać obserwacyjnej weryfikacji, co niestety nie było możliwe przez długie lata wskutek niedostatecznego zaawansowania techniki. Testowanie jej rozpoczęło się, dość przy-

padkowo i niespodziewanie dla społeczności astronomów, w roku 1978 od obserwacyjnego odkrycia pustek przez Gregory'ego i Thompsona [1].

Postęp po roku 1978 nie doprowadził do potwierdzenia zasady kosmologicznej. Przeciwnie, im dalej sięgały obserwacje, tym więcej widziano niejednorodnych struktur. Mimo to wiara w zasadę kosmologiczną przetrwała do dziś. Mówi się, że Wszechświat jest jednorodny „w odpowiednio dużej skali”, ale definicja tej „odpowiedniej skali” jest daleka od precyzji („kilka” setek megaparseków). Rozmiar „podstawowej komórki” Wszechświata, która miałaby być powtarzalna, jest tak duży, że nic pewnego nie umiemy powiedzieć o rozkładzie gęstości materii w pobliżu jej brzegu i poza jej granicami.

2. Dlaczego trzeba badać modele Wszechświata ogólniejsze niż klasyczne – przestrzennie jednorodne i izotropowe

2.1. Argumenty astronomiczne

Argumenty astronomiczne zostały krótko podane w poprzednim paragrafie. Skoro widzimy struktury o różnych rozmiarach i kształtach, trzeba użyć modeli, które mogą te struktury opisać.

Tradycyjna astronomiczna metoda opisu struktur to zaburzenia jednorodności znajdowane przez rozwiązywanie zlinearyzowanych równań Einsteina. Są dwa problemy z tą metodą:

1. Aby rozwiązanie perturbacyjne było realistyczne, jego parametry muszą być „dostatecznie małe”. Ale nie wiadomo, co to znaczy „dostatecznie”. Z punktu widzenia matematyki, kolejne rzędy rachunku perturbacyjnego są coraz lepszymi przybliżeniami dokładnej wartości funkcji, gdy wartość parametru mieści się w promieniu zbieżności szeregu przedstawiającego tę funkcję. Aby określić promień zbieżności, musimy znać ogólny wzór na n -ty wyraz szeregu. To się n i g d y nie zdarza w kosmologii – mamy tam do dyspozycji tylko pierwszy i czasami drugi wyraz szeregu.

2. W praktyce oczekuje się, że rachunek perturbacyjny da realistyczny wynik, gdy jego parametr jest mniejszy od 1. W kosmologii mamy dwa takie parametry: kontrast gęstości $\Delta\rho/\rho$ i kontrast krzywizny $\Delta R/R$.¹ O b y d w a muszą być małe. Wbrew rozpowszechnionemu mniemaniu, sam kontrast krzywizny nie jest miarą dokładności przybliżenia. Na przykład w modelu Lemaître'a–Tolmana z dodatnią energią (o którym będzie mowa dalej), maleje on z czasem niezależnie od warunków początkowych ([2], par. 18.10), podczas gdy $\Delta\rho/\rho$ rośnie.

Jak pokazuje tabela poniżej, z wyjątkiem pustek i Wielkiego Atraktora, dla innych badanych w kosmologii obiektów obecny kontrast gęstości jest wielokrotnie większy od 1. Jednostką gęstości masy używaną w tabeli jest średnia gęstość kosmiczna ρ_b .

Obiekt	Średnia gęstość [ρ_b]
gwiazda	$1,5 \cdot 10^{29}$
gromada kulista	$2 \cdot 10^5$
galaktyka	$6 \cdot 10^4$
gromada Virgo	190
Wielki Atraktor	1,6
pustka	0,1

Wskutek tych dwu problemów, wyniki rachunków perturbacyjnych nie mogą być bezpiecznie ekstrapolowane w przyszłość.

Problem zbieżności szeregu przybliżeń nakłada się na drugi, znany w teorii względności pod nazwą problemu stabilności linearyzacyjnej równań Einsteina. Wśród rozwiązań zlinearyzowanych równań Einsteina istnieją takie, które nie są przybliżeniem żadnego ścisłego rozwiązania [3–5].

Rozważanie ogólniejszych modeli nie oznacza negowania sukcesów modeli R–W. Modele niejednorodne mają uzupełnić szczegóły, których geometria R–W „nie chwyta”, jak np. powstawanie struktur albo oddziaływanie światła z zaburzeniami gęstości. Odgrywają podobną rolę, jak rozwiązania perturbacyjne, ale, dopóki wierzymy, że teoria Einsteina jest poprawna, można ich wyniki ekstrapolować dowolnie daleko w przyszłość i w przeszłość, ponieważ są dokładnymi rozwiązaniami równań pola grawitacyjnego. Modele R–W nadal pozostają „ślusne” jako pierwsze przybliżenie. Jeśli modele R–W są dla potrzeb kosmologii wystarczająco dobre, to modele ogólniejsze mogą być tylko lepsze – bo zawierają w sobie R–W jako przypadek graniczny, a oprócz tego zawierają dodatkowe swobodne parametry, które umożliwiają opis większego zbioru zjawisk.

2.2. Argumenty matematyczne i fizyczne

Poszukiwanie uogólnień znanych rozwiązań jest naturalną potrzebą ludzkiego umysłu. Jeśli rozwiązanie, które chcemy uogólnić, ma interpretację fizyczną, jest rzeczą oczywistą zastosowanie nowego rozwiązania do tej samej sytuacji.

3. Geometria modeli kosmologicznych

3.1. Modele Robertsona–Walkera (R–W)

Forma metryczna tych modeli jest dana wzorem, wynikającym ze zrobionych a priori założeń o symetrii czasoprzestrzeni:

$$ds^2 = dt^2 - S^2(t) \left[\frac{dr^2}{1 - kr^2} + r^2 (d\theta^2 + \sin^2 \theta d\varphi^2) \right], \quad (3.1)$$

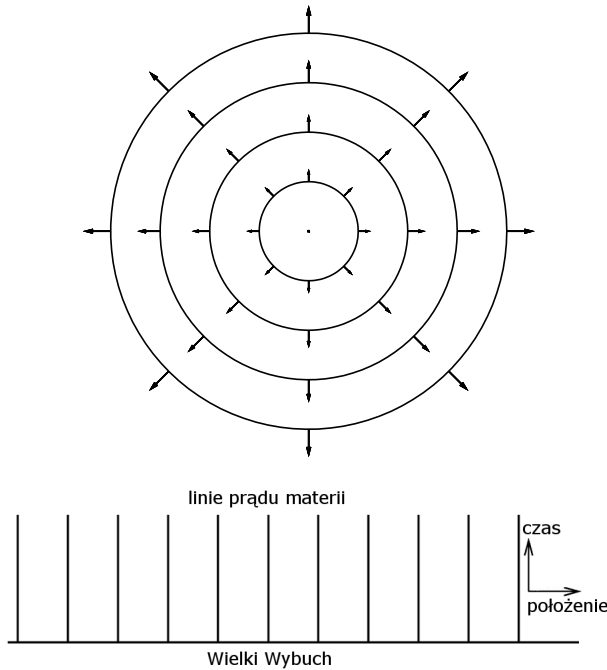
gdzie funkcja $S(t)$ jest wyznaczona przez równania Einsteina. Jeśli źródłem jest pył (materia o ciśnieniu tożsa-

¹ R jest tu średnią krzywizną trójwymiarowej przestrzeni stałego czasu. ΔR oznacza różnicę między wartością krzywizny w badanym punkcie a R . ρ jest średnią gęstością masy we Wszechświecie, a $\Delta\rho$, podobnie jak ΔR , jest różnicą między wartością lokalnie mierzoną a wartością średnią.

ściowo równym zero), to $S(t)$ spełnia następujące równanie, będące całką pierwszą równań Einsteina:

$$S_{,t}^2 = \frac{2GM}{c^2 S} - k + \frac{1}{3} \Lambda S^2. \quad (3.2)$$

Wielkości k i M są stałymi dowolnymi, Λ jest stałą kosmologiczną. Geometrię tych modeli ilustruje rys. 1.



Rys. 1. Ekspansja w modelach Robertsona-Walkera. U góry: prędkość ekspansji jest sztywno sprzężona z położeniem: jest proporcjonalna do odległości od obserwatora w centrum w ustalonej chwili czasu, ale zmienia się z czasem. U dołu: wybuch początkowy, w naturalnej kosmologicznej synchronizacji, następuje równocześnie $\Rightarrow$ wszystkie cząstki materii w każdej późniejszej chwili czasu mają ten sam wiek.

Odległość jasnościowa obserwatora umieszczonego w $(t, r) = (t_0, 0)$ od źródła światła umieszczonego w (t_e, r_e) jest zdefiniowana jako

$$D_L = r_e S(t_e) (1+z)^2, \quad (3.3)$$

gdzie z jest przesunięciem częstości promieniowania ku czerwieni zdefiniowanym przez

$$z = \frac{\text{emitowana częstość}}{\text{obserwowana częstość}} - 1 \equiv \frac{\nu_e}{\nu_o} - 1. \quad (3.4)$$

Dla modeli R-W wynosi ono:

$$z = S(t_0)/S(t_e) - 1. \quad (3.5)$$

Aby obliczyć D_L , trzeba scałkować równanie promienia biegnącego do obserwatora, $dt/S(t) = -dr/\sqrt{1-kr^2}$. W astronomii przedstawia się ten wynik za pomocą wielkości obserwowalnych, tzn. z , „stałej” Hubble’a w chwili t_0 (obecnej):

$$H_0 = S_{,t}/S|_{t=t_0} \quad (3.6)$$

oraz parametrów gęstości, krzywizny i stałej kosmologicznej

$$(\Omega_m, \Omega_k, \Omega_\Lambda) \stackrel{\text{def}}{=} \frac{1}{3H_0^2} (8\pi G\rho_0, -3k/S_0^2, \Lambda), \quad (3.7)$$

które spełniają tożsamość $\Omega_m + \Omega_k + \Omega_\Lambda \equiv 1$ (S_0 i ρ_0 są obecnymi wartościami S i średniej gęstości materii we Wszechświecie). W tych zmiennych odległość jasnościowa źródła światła umieszczonego w punkcie odpowiadającym z od obserwatora umieszczonego w $z = 0$ wynosi:

$$D_L(z) = \frac{1+z}{H_0 \sqrt{\Omega_k}} \times \sin \left\{ \int_0^z \frac{\sqrt{\Omega_k} dz'}{\sqrt{\Omega_m(1+z')^3 + \Omega_k(1+z')^2 + \Omega_\Lambda}} \right\}. \quad (3.8)$$

Wzór ten obowiązuje dla obydwu znaków Ω_k ($\sin(ix) \equiv i \sinh x$) i także w granicy $\Omega_k \rightarrow 0$.

Obecnie ulubionym modelem większości kosmologów jest przypadek szczególny (3.8) odpowiadający $k = \Omega_k = 0$.

3.2. Model Lemaître’a–Tolmana

Forma metryczna modelu Lemaître’a–Tolmana (L–T) [6,7] jest dana wzorem

$$ds^2 = dt^2 - \frac{R_{,r}^2}{1+2E(r)} dr^2 - R^2(t, r) (d\vartheta^2 + \sin^2 \vartheta d\varphi^2), \quad (3.9)$$

gdzie $R(t, r)$ spełnia równanie (wynikające z równań Einsteina):

$$R_{,t}^2 = 2E(r) + \frac{2M(r)}{R} - \frac{1}{3} \Lambda R^2. \quad (3.10)$$

Przy $\Lambda = 0$, to równanie jest identyczne z newtonowskim równaniem ruchu radialnego cieczy w sferycznie symetrycznym polu grawitacyjnym.

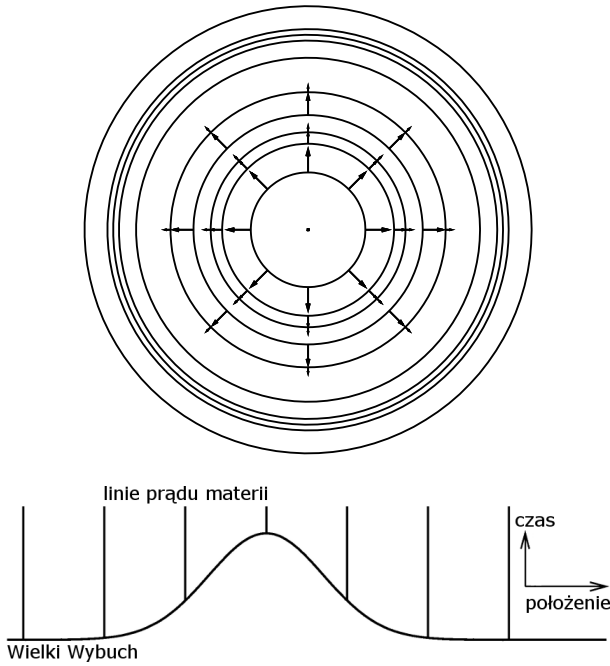
Funkcje $M(r)$ (rozkład masy) i $E(r)$ (rozkład energii) są dowolne. Całka równania (3.10) zawiera trzecią funkcję dowolną, $t_B(r)$, która opisuje „rozkład jazdy” wybuchu początkowego. Na przykład, przy $E = 0 = \Lambda$ rozwiązaniem (3.10) jest

$$R = \left(\frac{9M}{2} \right)^{1/3} (t - t_B(r))^{2/3}, \quad (3.11)$$

skąd widać, że osobliwość początkowa występuje w chwili $t = t_B(r)$ – zależnej od położenia. Geometrię modelu L–T ilustruje rys. 2.

Tylko dwie spośród funkcji (M, E, t_B) opisują fizyczne stopnie swobody, jedną z nich można ustalić wybierając współrzędną r , ponieważ równania opisujące ten model nie zmieniają się przy dowolnych transformacjach postaci $r \rightarrow r(r')$. Gęstość masy jest dana przez:

$$\kappa\rho = \frac{2M_{,r}}{R^2 R_{,r}}; \quad \kappa \stackrel{\text{def}}{=} \frac{8\pi G}{c^4}. \quad (3.12)$$



Rys. 2. Ekspansja w modelu Lemaître'a-Tolmana. U góry: prędkość ekspansji jest nieskorelowana z położeniem powłoki materii. Przestrzenny rozkład prędkości jest dowolną funkcją zmiennej radialnej. U dołu: wybuch początkowy, w naturalnej kosmologicznej synchronizacji, następuje nierównocześnie $\Rightarrow$ cząstki materii w ustalonej chwili czasu kosmologicznego mają wiek zależny od położenia. „Rozkład jazdy” wybuchu początkowego jest drugą funkcją dowolną zmiennej radialnej. Na rysunku wybrano $t_B(r) = 2e^{-r^2}$, maksimum krzywej leży w środku symetrii.

Model Friedmanna jest zawarty w modelu L-T jako następująca granica:

$$M = \text{const} \cdot r^3, \quad E = -kr^2/2, \quad t_B = \text{const}, \quad R = rS(t). \quad (3.13)$$

Porównanie (3.10) z (3.2) pokazuje, że ewolucją modeli L-T i Friedmanna rządzi to samo równanie, ale w modelu L-T każda powłoka $r = \text{const}$ ma swoje własne warunki początkowe, niezależne od innych powłok.

Ponieważ prędkość ekspansji poszczególnych powłok materii jest nieskorelowana z ich położeniem, mogą występować efekty, które nie ujawniają się w modelu Friedmanna, na przykład maksima i minima gęstości wędrujące niezależnie od cząstek materii („fale gęstości”) oraz zderzanie się różnych powłok w trakcie ewolucji. Miejsce zderzenia jest drugą osobliwością, niezależną od wybuchu początkowego. (W granicy Friedmanna wybuch początkowy i przecięcie powłok pokrywają się). Można jej uniknąć przez nałożenie pewnych nierówności na funkcje dowolne i ich pochodne; ten problem został dawno rozwiązany i nie będziemy go dyskutować [8].

3.3. Modele Szekeresa

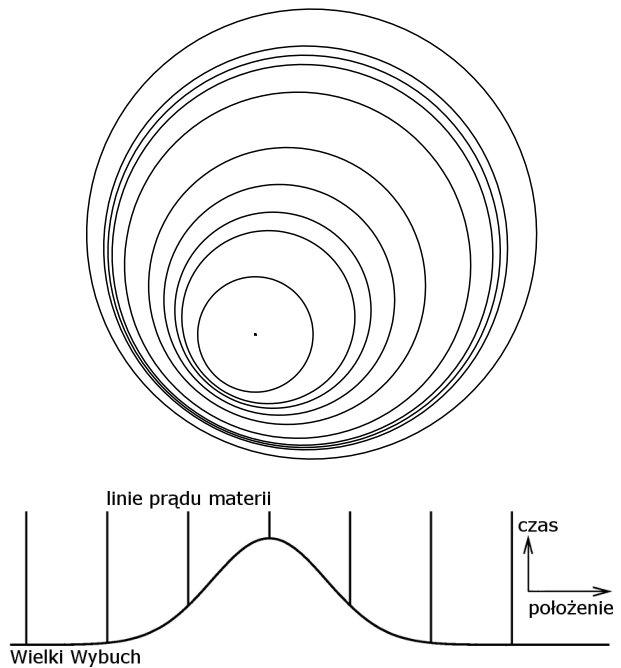
Forma metryczna modeli Szekeresa [9] jest dana wzorami:

$$ds^2 = dt^2 - \frac{(R_{,r} - RE_{,r}/\varepsilon)^2}{\varepsilon + 2E} dr^2 - \frac{R^2}{\varepsilon^2} (dx^2 + dy^2), \quad (3.14)$$

gdzie

$$\varepsilon \stackrel{\text{def}}{=} \frac{S}{2} \left[\left(\frac{x-P}{S} \right)^2 + \left(\frac{y-Q}{S} \right)^2 + \varepsilon \right]; \quad (3.15)$$

funkcja $R(t, r)$ spełnia to samo równanie ewolucji (3.10) co w modelu L-T. Funkcje (P, Q, S) są dowolnymi funkcjami zmiennej r i opisują anizotropię początkowego rozkładu prędkości; $\varepsilon = \pm 1, 0$. Geometrię modeli Szekeresa z $\varepsilon = +1$ ilustruje rys. 3.



Rys. 3. Ekspansja w modelach Szekeresa. U góry: prędkość ekspansji jest nieskorelowana z położeniem powłoki materii i jest anizotropowa: kuliste powłoki materii nie są współśrodkowe. Trzy dodatkowe funkcje dowolne modelu opisują trajektorię środków sfer w przestrzeni. U dołu: wykres „rozkładu jazdy” wybuchu początkowego w modelu Szekeresa wygląda tak samo, jak w modelu Lemaître'a-Tolmana.

Model L-T jest przypadkiem granicznym modeli Szekeresa odpowiadającym $\varepsilon = +1$ i stałym funkcjom (P, Q, S) . (Aby otrzymać (3.9) z (3.14) trzeba przetransformować współrzędne (x, y) na (ϑ, φ) ; w modelach Szekeresa używa się innych współrzędnych na sferze).

Gęstość masy w modelach Szekeresa dana jest wzorem

$$\kappa\rho = \frac{2(M_{,r} - 3ME_{,r}/\varepsilon)}{R^2(R_{,r} - RE_{,r}/\varepsilon)}. \quad (3.16)$$

Jest to sferycznie symetryczny rozkład gęstości, na który nałożono „dipol grawitacyjny”. (Wzory na część monopową i dipolową są dość skomplikowane (patrz [2,10]), więc ich nie przytaczamy. W szczególnym przypadku można otrzymać „czysty dipol”, ale wtedy gęstość masy

staje się ujemna w części przestrzeni, czego nikt nie potrafi zinterpretować fizycznie).

Rozwiązania odpowiadające $\varepsilon = +1$, $\varepsilon = 0$ i $\varepsilon = -1$ nazywają się odpowiednio kwazisferycznym, kwaziplat-skim i kwazihiperbolicznym modelem Szekeresa. Modele z $\varepsilon \leq 0$ są słabo zbadane i nawet ich interpretacja geometryczna nie jest do końca jasna (patrz [11,12]). Natomiast pierwszy okazał się bardzo ciekawy i przydatny z punktu widzenia astrofizyki (patrz [13] oraz seria prac K. Bolejko [14–17]).

4. Fizyka i astrofizyka w modelu L–T

Ponieważ model L–T jest sferycznie symetryczny, wszystkie struktury, które można w nim opisać, są też sferycznie symetryczne. Nie jest on więc używany jako model całego Wszechświata, tylko jako model pojedynczej lokalnej struktury zanurzonej we Wszechświecie Friedmanna. Można też dokonać „zszycia” większej liczby obszarów L–T z jednorodnym tłem opisywanym modelem Friedmanna.

4.1. Imitowanie „przyspieszonej ekspansji” Wszechświata za pomocą niejednorodnego rozkładu materii

Hipoteza przyspieszonej ekspansji Wszechświata, napędzanej przez „ciemną energię”, jest próbą wyjaśnienia obserwacji supernowych typu Ia. Supernowa typu Ia jest końcowym momentem ewolucji węglowo-tlenowego białego karła w układzie podwójnym, na który spada materia z gwiazdy towarzyszącej. Z powodów, których nie będziemy tu dyskutować, uważa się, że maksymalna jasność absolutna wszystkich supernowych tej klasy jest taka sama.

Mierzając przesunięcie ku czerwieni w obserwowanym widmie supernowych, można obliczyć odległość do każdej z nich, zakładając, że jest ściśle spełnione prawo Hubble’a – a następnie obliczyć spodziewaną jasność obserwowaną na Ziemi. Okazało się, że faktycznie obserwowana maksymalna jasność jest mniejsza od oczekiwanej. Aby wyjaśnić tę rozbieżność, należało zmodyfikować używany wcześniej model Wszechświata. Supernowe są dalej od nas, niż myśleliśmy, Wszechświat musiał więc rozszerzać się z większą prędkością niż stary model przewidywał. Próby dopasowania różnych modeli R–W do obserwowanych wyników doprowadziły do wniosku, że najdokładniejszą zgodność daje model z $\Omega_k = 0$, $\Omega_m \approx 0,3$, $\Omega_\Lambda \approx 0,7$ – w którym gęstość energii związana z Λ jest ponad dwukrotnie większa od gęstości energii związanej ze zwykłą materią i tzw. ciemną materią w galaktykach.

Niezerowa wartość Ω_Λ oznacza, że Wszechświat musiał rozszerzać się ruchem przyspieszonym. To zrodziło zagadkę tzw. „ciemnej energii” i wielkie programy badawcze zmierzające do jej rozwiązania.

Proszę zwrócić uwagę na wyróżnione fragmenty zdań. W interpretację obserwacji jest uwikłana geometria modelu kosmologicznego. Wnioski, że stała kosmologiczna ma dużą wartość i że Wszechświat rozszerzał się

ruchem przyspieszonym wynikły z założenia, że musimy używać modeli R–W.

To ostatnie stwierdzenie musi być dokładnie zrozumiane – bo jego niewłaściwe rozumienie stworzyło fałszywą legendę. To, co należy wyjaśnić, to relacja między obserwowaną jasnością supernowych typu Ia a przesunięciem ku czerwieni w ich widmach. „Przyspieszona ekspansja” Wszechświata nie jest obserwowanym zjawiskiem wymagającym objaśnienia, tylko elementem interpretacji obserwacji, związanym z przyjętą a priori klasą modeli R–W. Jeśli potrafimy w modelu niejednorodnym odtworzyć obserwowany kształt funkcji $D_L(z)$, to „przyspieszona ekspansja” nie jest do niczego potrzebna.

Wiedząc to, zobaczmy, jak można objaśnić „pociemnienie supernowych” za pomocą modelu L–T (według pracy Iguchi, Nakamura i Nakao [18]).

Dla uproszczenia rachunku zakładamy, że obserwator znajduje się w środku symetrii. Tor promienia biegnącego do obserwatora jest dany równaniem

$$\frac{dt}{dr} = -\frac{R_r}{\sqrt{1+2E(r)}}. \quad (4.1)$$

Wielkości obliczone wzdłuż takiego promienia będziemy zaznaczali wskaźnikiem ps . Dla centralnego obserwatora przesunięcie ku czerwieni jest rozwiązaniem równania ([2]):

$$\frac{1}{1+z} \frac{dz}{dr} = \left[\frac{R_{tr}}{\sqrt{1+2E}} \right]_{ps}, \quad (4.2)$$

zaś odległość jasnościowa od źródła w punkcie odpowiadającym wartości z jest dana tym samym wzorem, co w modelach R–W:

$$D_L = (1+z)^2 R|_{ps}. \quad (4.3)$$

Rozwiązania $t_{ps}(r)$ i $z_{ps}(r)$ równań (4.1) i (4.2) mogą być znalezione numerycznie, gdy zadamy funkcje $M(r)$, $E(r)$ i $t_B(r)$. Jako jedną z nich wybieramy

$$M = M_0 r^3, \quad (4.4)$$

gdzie $M_0 = \text{const}$ (to jest definicja współrzędnej r).

Rozpatrujemy teraz dwa przypadki:

- Przypadek 1: $E = E_0 r^2$, gdzie $E_0 = \text{const}$ (takie samo E , jak w modelu Friedmanna), natomiast dla $t_B(r)$ wybieramy następującą definicję uwikłaną:

$$R|_{ps} = \frac{1}{H_0(1+z)} \int_0^z \frac{dz'}{\sqrt{\Omega_m(1+z')^3 + \Omega_\Lambda}}, \quad (4.5)$$

gdzie $\Omega_m = 0,3$ i $\Omega_\Lambda = 0,7$, zaś H_0 jest obecną wartością stałej Hubble’a. Równania (4.1), (4.2) i (4.5) tworzą układ, który można rozwiązać numerycznie.

- Przypadek 2: $t_B = 0$, tzn t_B ma taką samą formę jak w modelu Friedmanna. Tym razem używamy (4.5) jako uwikłanej definicji $E(r)$ i układ równań można znów rozwiązać numerycznie.

Równanie (4.5) jest tu definicją funkcji t_B albo funkcji E . Na pytanie, dlaczego wybraliśmy taką definicję, moglibyśmy odpowiedzieć „bo tak się nam podobało”. Ale,

oczywiście, wybór był nieprzypadkowy. Porównanie (4.5) z (3.8) pokazuje, że (4.5) definiuje taki sam związek między odległością jasnościową $D_L = (1+z)^2 R|_{ps}$ a przesunięciem ku czerwieni z , jaki występuje w „standardowym” modelu Friedmanna z $\Omega_k = 0$, $\Omega_m = 0,3$ i $\Omega_\Lambda = 0,7$. Tylko, że... osiągnęliśmy ten sam wynik z zerową stałą kosmologiczną, tzn. bez egzotycznej „ciemnej energii”, i przy ekspansji ruchem opóźnionym – jak nakazują znane z Ziemi prawa grawitacji. Oznacza to, że gdybyśmy do interpretacji obserwacji używali modelu L–T zamiast R–W, pomysł z „ciemną energią” i „przyspieszoną ekspansją” nie przyszedłby nikomu do głowy.

W ramach użytych tu dwu modeli można wyjaśnić „pociemnienie supernowych” w jasny fizycznie i elegancki sposób. W obydwu przypadkach podstawą objaśnienia jest fakt, że w modelu L–T każda kulista powłoka materii ewoluuje niezależnie od pozostałych, ze swoimi własnymi danymi początkowymi. Zatem każda powłoka ma swój własny moment wyjścia z wybuchu początkowego (czyli „wiek”) i swoją własną początkową prędkość. „Ekspansja z przyspieszeniem” oznacza, że im wcześniejszy moment w historii Wszechświata badamy, tym wolniejsza była ekspansja. Zatem, gdy dokonujemy obserwacji na stożku świetlnym, im dalsze zdarzenie widzimy, tym wolniejsza była ekspansja w jego sąsiedztwie. Ten efekt jest imitowany w obydwu omówionych wyżej modelach L–T.

W pierwszym przypadku powłoki dalsze od obserwatora wynurzyły się z wybuchu początkowego wcześniej niż bliższe. W momencie, gdy przecinają przeszły stożek świetlny obserwatora, są więc „starsze” i uciekają z mniejszą prędkością niż przewidywałby model Friedmanna bez stałej kosmologicznej.

W drugim przypadku dalsze powłoki zostały „wystrzelone” z osobliwości początkowej z mniejszą prędkością, niż w modelu Friedmanna z $\Lambda = 0$.

Fakt, że pociemnienie supernowych można wyjaśnić w modelu niejednorodnym bez wprowadzania stałej kosmologicznej i „ciemnej energii”, jest już częściowo znany w środowisku astronomicznym. Niestety, wiedza ta stwarzała się z pewną legendą, która jest, z punktu widzenia logiki, nieprawdziwa. Mianowicie, próbując intuicyjnie wytłumaczyć mechanizm pociemnienia w modelu L–T, kilku autorów zrobiło to, przed czym powyżej przestrzegaliśmy: zamiast objaśniać pociemnienie, objaśniali przyspieszenie. Rozumowali tak: Jeśli niejednorodny rozkład materii ma powodować przyspieszenie, czyli jeśli bliższe nas galaktyki muszą od nas uciekać szybciej niż dalsze, to nasza Galaktyka musi znajdować się w centrum wielkiej pustki, która rozszerza się szybciej niż gęstszy obszar daleko od nas. Promień tej pustki jest oceniany różnie przez różnych autorów, największa spotykana w literaturze wartość wynosi ok. 1,5 Gpc [10], najmniejsza 200 Mpc albo nawet mniej [19]. Nasza Galaktyka powinna znajdować się względem środka tej pustki w miejscu, które też jest różnie oceniane przez różnych autorów, od ok. 15 Mpc [10] do

ok. 150 Mpc [20]. Wniosek ten bywa stawiany jako argument przeciwko modelowi L–T z punktu widzenia „zasady kosmologicznej”: położenie w pobliżu środka tak wielkiej pustki jest wyróżnione².

Na ten zarzut można podać kilka odpowiedzi, przy czym niekoniecznie jest to kompletny zbiór argumentów:

1. Istnienie tej wielkiej pustki nie jest udowodnione – wiara w jej istnienie opiera się na „intuicyjnym rozumieniu Natury”. Jak widzieliśmy powyżej, relację $D_L(z)$ można odtworzyć w modelu L–T za pomocą tylko jednej funkcji dowolnej, druga funkcja z pary (E, t_B) była ustalona i niewykorzystana przy dopasowywaniu modelu do danych. Istnieje co najmniej jedna praca [20], w której pokazano, że w modelu L–T można równocześnie odtworzyć obserwacje supernowych i zapewnić jednorodny rozkład materii na powierzchni stałego (obecnego) czasu.

Tymczasem dla pewnej grupy astronomów alternatywa „model R–W albo wielka pustka” stała się nowym dogmatem. Zakładając, że rzeczywiście mamy tylko te dwie możliwości do wyboru, można w obronie wielkiej pustki podać jeszcze takie argumenty:

2. Co wolimy: zajmować umiarkowanie wyróżnione położenie we Wszechświecie wypełnionym materią dobrze znaną z laboratorium i zachowującą się zgodnie z ziemskimi prawami fizyki, czy „typowe” położenie we Wszechświecie, w którym 70% masy składa się z materii niewidzialnej, o nieznanych własnościach fizycznych i chemicznych – wymyślonej przez kosmologów tylko po to, aby mogli nadal używać ich ulubionej klasy modeli kosmologicznych i nie musieli uczyć się teorii względności ponad absolutne minimum?

3. Gdzie dokładnie, we Wszechświecie pełnym pustek i zagęszczeń najrozmaitszych kształtów, powinien znajdować się obserwator, aby jego położenie było „typowe” i nie „łamało zasady kosmologicznej”? Gdyby znajdował się w miejscu, gdzie gęstość materii jest równa średniej kosmicznej, byłoby to też „nietypowe” i „nieprawdopodobne”, bo takich miejsc praktycznie nie ma. W obserwowanym Wszechświecie nie ma miejsc, które nie znajdowałyby się wewnątrz jakiejś niejednorodnej struktury – i jeśli ta niejednorodność jest jedną z wielu, nie ma sensu mówienie o „wyróżnionym położeniu”. Odległość równa 1% promienia od środka pustki jest tak samo „kopernikańska” albo „antykopernikańska”, jak 50% czy 80%.

4. Położenie w pobliżu środka wielkiej pustki wynika z faktu, że do interpretowania wyników obserwacji użyliśmy modelu sferycznie symetrycznego. Gdybyśmy użyli do tego celu modelu Szekelesa, albo innego, który nie ma żadnej symetrii, pustka nie miałaby żadnego środka. Aby konsekwentnie interpretować obserwacje za pomocą modelu niesymetrycznego musielibyśmy jednak uwzględnić fakt, że „prawo Hubble’a” jest całkowicie złamane – „funkcja Hubble’a” w takich modelach nie tylko jest nieliniowa jako funkcja odległości, ale też zależna od kierunku. Pozorna sprzeczność z zasadą kosmologiczną jest

²Rozrzut liczb podawanych przez różnych autorów jest tak wielki, że trudno traktować te dywagacje z całkowitą powagą. Na potrzeby dalszej dyskusji założmy jednak, że mamy tu do czynienia z poważną nauką.

więc w tym przypadku przejściowa i zniknie, kiedy nauczymy się używać modeli bez symetrii.

5. Jeśli nawet znajdujemy się w samym środku wielkiej pustki, to najprawdopodobniej nie jest ona jedyna we Wszechświecie. Nie potrafimy tego powiedzieć z całą pewnością, bo nawet nie wiadomo, jak duża musi być „nasza” pustka. Jeśli jest ich więcej, nasze położenie staje się mniej wyjątkowe.

4.2. Powstawanie i ewolucja struktur

Tradycyjnie, badając ewolucję układu fizycznego zadaje się warunki początkowe w pewnej chwili t_1 (np. w hydrodynamice – początkowy rozkład gęstości i prędkości), za pomocą równań ewolucji oblicza się stan układu w późniejszej chwili $t_2 > t_1$, a następnie porównuje się wyniki obliczeń z pomiarami dla chwili t_2 . Jeśli występuje niezgodność między wynikami obliczeń a pomiarami, zmieniamy warunki początkowe i próbujemy lepiej „wstrzelić się” w znany stan końcowy.

Modele niejednorodne pozwalają zmodyfikować to podejście w sposób, który byłby niemożliwy za pomocą metod perturbacyjnych. Mianowicie, jako dane początkowe dla równań ewolucji można zadać niezależne informacje zebrane o stanach w obydwu chwilach – t_1 i t_2 , a następnie dobrać model, który przeprowadza zadany stan w t_1 w zadany stan w t_2 . Na przykład, można zadać następujące pary danych:

Dane w t_1	Dane w t_2
rozkład gęstości	rozkład gęstości
rozkład prędkości	rozkład gęstości
rozkład prędkości	rozkład prędkości

Pokażemy działanie tej metody dla pierwszego przypadku. Szczegóły i przykłady zastosowań są podane w pracach [21–23].

Zakładamy $\Lambda = 0$ i szukamy rozwiązania (3.10) w postaci $t(R)$. Trzeba osobno rozpatrzyć przypadki $E > 0$ i $E < 0$, zaś dla $E < 0$ także osobno podprzypadki, gdy stan końcowy jeszcze się rozszerza i gdy już się zapada. Tu zajmiemy się tylko przypadkiem $E > 0$.

Można łatwo (numerycznie) przeliczyć rozkład gęstości w dowolnej ustalonej chwili czasu na rozkład funkcji $R(r)$ w tej samej chwili (używamy $M(r)$ jako zmiennej radialnej):

$$\kappa\rho = \frac{2M_{,r}}{R^2R_{,r}} \equiv \frac{6}{(R^3)_{,M}} \iff R^3 = \int_0^M \frac{6}{\kappa\rho(\tilde{M})} d\tilde{M}. \quad (4.6)$$

Wtedy rozwiązania (3.10) dla $t = t_i$, $i = 1, 2$, są dane przez

$$t_B = t_i - \frac{M}{(2E)^{3/2}} \times \left[\sqrt{(1 + 2ER_i/M)^2 - 1} - \operatorname{arccosh}(1 + 2ER_i/M) \right], \quad (4.7)$$

gdzie $R_i \stackrel{\text{def}}{=} R(t_i, M)$. Odejmując (4.7) dla $t = t_1$ od tego samego równania dla $t = t_2$ dostajemy

$$\begin{aligned} & \sqrt{(1 + 2ER_2/M)^2 - 1} - \operatorname{arccosh}(1 + 2ER_2/M) \\ & - \sqrt{(1 + 2ER_1/M)^2 - 1} + \operatorname{arccosh}(1 + 2ER_1/M) \\ & = \frac{(2E)^{3/2}}{M} (t_2 - t_1). \end{aligned} \quad (4.8)$$

To równanie jest uwikłaną definicją funkcji $E(M, t_1, t_2, R_1, R_2)$ (można rozwiązać je numerycznie). Podstawiając takie E do (4.7) znajdujemy funkcję $t_B(M, t_1, t_2, R_1, R_2)$, a wtedy funkcje E i t_B jednoznacznie definiują model L–T przeprowadzający założony stan początkowy w założony stan końcowy.

Zanim użyjemy (4.8) jako numerycznej definicji funkcji $E(M)$, musimy sprawdzić, czy to równanie ma jakiegokolwiek rozwiązanie, a jeśli ma, to czy rozwiązanie jest jednoznaczne. To sprawdzenie jest prostym ćwiczeniem rachunkowym, a wynik jest następujący: rozwiązanie istnieje wtedy i tylko wtedy, gdy spełniona jest następująca nierówność:

$$t_2 - t_1 < \frac{\sqrt{2}}{3} (a_2^{3/2} - a_1^{3/2}) \stackrel{\text{def}}{=} \mathcal{L}, \quad (4.9)$$

gdzie $a_i \stackrel{\text{def}}{=} R_i/M^{1/3}$. Rozwiązanie jest wtedy jednoznaczne.

To jest warunek konieczny i dostateczny na istnienie ewolucji z $E > 0$ przeprowadzającej stan $R(t_1, M)$ w stan $R(t_2, M)$. Jest on równoważny stwierdzeniu, że pomiędzy t_1 i t_2 $R(t, M)$ rosło szybciej niż rosłoby w modelu L–T z $E = 0$ (przestrzenie płaskim).

W przeciwnym wypadku, gdy $t_2 - t_1 > \mathcal{L}$, można te same dwa stany powiązać ewolucją L–T z $E < 0$. Dodatkowe kryterium pozwala rozpoznać, czy stan końcowy będzie ekspandujący, czy zapadający się. Stan końcowy będzie ekspandujący, jeśli

$$t_2 - t_1 \leq (a_2/2)^{3/2} \times \left[\pi - \arccos(1 - 2a_1/a_2) + 2\sqrt{a_1/a_2 - (a_1/a_2)^2} \right] \quad (4.10)$$

i zapadający się w przeciwnym wypadku.

Miarą rozkładu prędkości dla ustalonej chwili czasu jest funkcja $R_{,t}/M^{1/3}$ (w granicy Friedmanna jest ona zależna tylko od t).

W serii prac Krasieńskiego i Hellaby [21–23] zastosowano to podejście do opisu powstawania obecnie obserwowanych struktur z małego zaburzenia gęstości albo prędkości nałożonego na jednorodne tło w chwili ostatniego rozproszenia promieniowania CMB. Założonym końcowym produktem ewolucji (w chwili dzisiejszej, $\approx 10^{15}$ lat po Wielkim Wybuchu) była gromada galaktyk, pustka albo pojedyncza galaktyka z czarną dziurą wokół centrum.

Najlepsze wyniki osiągnięto w przypadku gromady galaktyk. Założyliśmy, że początkowe zaburzenie ma masę równą 0,01 masy dzisiejszej gromady galaktyk, a pozostałe parametry w granicach dopuszczanych przez pomiary promieniowania tła (względna amplituda gęstości $\Delta\rho/\rho = 10^{-5}$, względna amplituda prędkości $\Delta v/v = 10^{-4}$). Jako reprezentanta obecnego stanu wybraliśmy gromadę

A199 z katalogu Abella, dla której w literaturze podany jest „Uniwersalny Profil Gęstości”.

Możliwym źródłem niepowodzenia jest następujący problem: jeśli zakładamy rozkłady gęstości dla t_1 i t_2 , to rozkłady prędkości dla tych dwu chwil możemy obliczyć ze znalezionej modelu. Obliczony rozkład prędkości może wtedy mieć amplitudę niezgodną z wynikami obserwacji. Okazało się jednak, że wszystko mieści się w dopuszczalnych granicach: założona amplituda gęstości dla $t = t_1$ rzędu 10^{-5} implikuje amplitudę prędkości tego samego rzędu i vice versa.

Ciekawym wynikiem ubocznym tego badania było stwierdzenie, że zaburzenia prędkości skuteczniej generują struktury niż zaburzenia gęstości. Wynika to z następujących faktów:

- Jeśli początkowa gęstość jest jednorodna, to obliczona początkowa amplituda prędkości potrzebna do wygenerowania gromady galaktyk wynosi $5,2 \cdot 10^{-4}$ – jeszcze w dopuszczalnych obserwacyjnie granicach.
- Jeśli początkowy rozkład prędkości jest jednorodny, to obliczona początkowa amplituda gęstości wynosi $1,2 \cdot 10^{-2}$ –1000 razy za dużo.

Zatem zaburzenie prędkości może prawie samodzielnie wygenerować gromadę galaktyk, podczas gdy zaburzenie gęstości jest o wiele mniej skuteczne.

Przy równoczesnym działaniu zaburzeń gęstości i zaburzeń prędkości jest możliwe odwrócenie profilu: początkowe zagęszczenie może przekształcić się w rozrzedzenie i vice versa.

Te dwie uwagi powinny być poważnie potraktowane przez specjalistów od tradycyjnej teorii powstawania struktur, według której struktury powstają z samych zaburzeń gęstości.

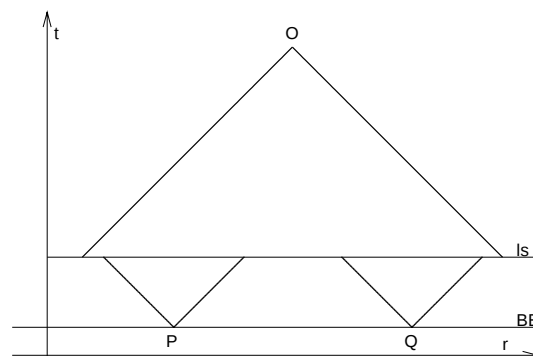
Mieliśmy w naszym podejściu trudność z wygenerowaniem pustki o realistycznych parametrach [24]. Aby odtworzyć profil gęstości w typowej obserwowanej pustce, początkowe zaburzenie gęstości musiało być o 2 rzędy wielkości większe od dopuszczalnego, a początkowe zaburzenie prędkości o 1 rząd za duże. Problem ten rozwiązał potem samodzielnie K. Bolejko [25] – rozbieżność znika, gdy uwzględnimy oddziaływanie materii z promieniowaniem podczas ostatniego rozproszenia.

4.3. Rozwiązanie „problemu horyzontu” bez użycia modeli inflacyjnych

Zakłada się zwykle, że emisja mikrofalowego promieniowania tła (CMB) była zdarzeniem natychmiastowym, a nie procesem rozciągającym w czasie. Zdarzenie to nazywa się ostatnim rozproszeniem promieniowania, które nastąpiło w chwili $t_{ls} \approx 3 \cdot 10^5$ lat po Wielkim Wybuchu (BB) [2]. Związany z tym zdarzeniem „problem horyzontu” jest jakościowo objaśniony na rys. 4.

Rozwiązanie tego problemu zaproponowane przez modele inflacyjne polegało na zastąpieniu zwykłej materii

(mieszaniny promieniowania i cząstek o niezerowej masie) przez „pole skalarne”, o właściwościach podobnych do stałej kosmologicznej. To „pole skalarne” miało istnieć w okresie 10^{-34} do 10^{-32} s po Wielkim Wybuchu i działać odpychająco, wskutek czego Wszechświat rozszerzał się z prędkością rosnącą wykładniczo (czyli „nadymał się” – stąd nazwa „inflacja”)³. Dzięki tej szybkiej ekspansji promień stożka świetlnego o wierzchołku na Wielkim Wybuchu w chwili t_{ls} stawał się znacznie większy niż promień widzialnej obecnie części Wszechświata.



Rys. 4. „Problem horyzontu” we Wszechświecie Robertsona–Walkera, schematyczna ilustracja w współrzędnych metryki (3.1). Linia BB to Wielki Wybuch, linia ls to hiperpowierzchnia ostatniego rozproszenia. O to dzisiejszy obserwator. Przeszły stożek świetlny obserwatora O obejmuje w chwili $t = t_{ls}$ dużo większą objętość niż objętość wewnątrz któregośkolwiek pojedynczego stożka świetlnego wyemitowanego podczas Wielkiego Wybuchu (dwa takie stożki pokazano). Zatem, różne obszary, które dziś widzimy, nie mogły być w przyczynowym kontakcie przed emisją promieniowania CMB, a mimo to temperatura promieniowania CMB docierającego od nich do nas jest prawie taka sama, jak gdyby źródła tego promieniowania oddziaływały przedtem i ustaliły równowagę między sobą. Jak to możliwe? To jest właśnie „problem horyzontu”.

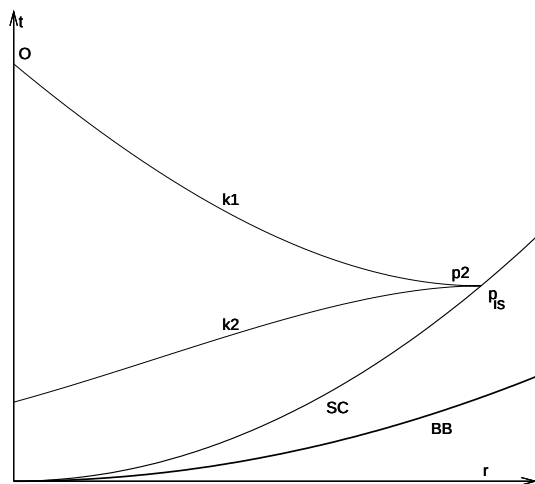
Modele inflacyjne zagnieździły się mocno w dzisiejszej kosmologii, zainspirowały różne badania obserwacyjne i teoretyczne i dzisiaj ich entuzjaści podają znacznie więcej argumentów za ich użytecznością [28,29]. Kosmologowie sceptycznie i krytycznie nastawieni do paradygmatu, metod i ideologii inflacji stanowią nieliczną mniejszość – patrz np. [30,2]. Tym niemniej podkreślić trzeba, że „problem horyzontu”, którego rozwiązanie rozreklamowano jako jeden z wielkich sukcesów inflacji, należy do gatunku tych, które stały się problemami dopiero po ich rozwiązaniu – nie były one przedtem uważane za zasadnicze ani ważne. Natomiast inflacja stworzyła kilka nowych problemów. Na przykład, do dziś nie wiadomo, jakie są fizyczne przyczyny początku „ery inflacji” i jak można zatrzymać rozszerzanie wykładnicze, aby przejść do normalnej ewolucji grawitacyjnej zwykłej materii [30].

³Pole skalarne $\phi(x)$ w modelach inflacyjnych spełnia równanie $g^{\alpha\beta}\phi_{;\alpha\beta} - \partial V/\partial\phi = 0$, gdzie potencjał $V(\phi)$ dobiera się tak, aby objaśnić badane zjawisko poprzez odpowiednią zmienność ϕ w czasie i przestrzeni.

Najprostsze rozwiązanie tych problemów polega na porzuceniu idei inflacji.

Ponadto, po prawie 20 latach okazało się, że „problem horyzontu” można rozwiązać w ramach modelu L–T bez wprowadzania inflacji ani nowych form materii – to rozwiązanie zostało zaproponowane przez M.N. C         [31]. Wystarczy za  o  y  ,   e funkcja $t_B(r)$ ma lokalne minimum w   rodku symetrii $r = 0$, co implikuje istnienie przecięcia pow  ok. Jako  ciowe obja  nienie jest podane w podpisie do rys. 5; szczeg  ły rachunkowe mo  na znaleźć w oryginalnej pracy [31] i w podr  czniku [2]. Dla zrozumienia rys. 5 nale  y wiedzie   dwie rzeczy:

1. Przeci  cie pow  ok znajduje si   w miejscu, gdzie $R_{,r} = 0$.
2. We wsp  rzednych (t, r) u  ytych w (3.9), promie     wiatlny przechodz  cy przez przeci  cie pow  ok ma styczn   poziom  .



Rys. 5. Rozwi  zanie „problemu horyzontu” w modelu L–T z funkcj   $E = 0$. Radialny promie     wiatlny k_1 wyslany w przesz  o   przez obserwatora O trafia w przeci  cie pow  ok SC w punkcie p_{is} . Z punktu p_2 na k_1 le  ącego nieco p   niej ni   p_{is} inny radialny promie   jest wyslany w przesz  o   w kierunku   rodku symetrii. Jak pokazuje rachunek [2], k_2 musi trafi   w   rodek w chwili p   niejszej od Wielkiego Wybuchu i od przeci  cia pow  ok, niezale  nie od po  o  enia obserwatora. Wskutek tego wszystkie obszary nieba widziane przez obserwatora w dowolnej chwili mog  y by   ze sob   przy  czynowo powi  zane poprzez wsp  lne   r  d  o w przesz  o  ci.

Zauwa  amy,   e jest to rozwi  zanie „problemu horyzontu” skuteczniejsze ni   za pomoc   inflacji. W modelach inflacyjnych, ich parametry s   dobierane tak, aby „problem horyzontu” znikn   dla chwili obecnej i dla pewnego okresu po niej. Po odpowiednio d  ugim czasie problem ten musi jednak wyst  pi   znowu. Rozwi  zanie za pomoc   modelu L–T jest permanentne [31,10].

4.4. Zaburzenia rosn  ce i malej  ce

W rachunkach perturbacyjnych zidentyfikowano dwa rodzaje zaburze   jednorodno  ci: rosn  ce z czasem i malej  ce z czasem. W   cis  tych modelach niejednorodnych

mo  na t   klasyfikacj   powi  za   z innymi geometrycznymi charakterystykami. Jako pierwsi zrobili to Silk w roku 1977 [26] dla modelu L–T oraz Goode i Wainwright w roku 1982 [27] dla modelu Szekeres (patrz tak  e [2]), w obydwu przypadkach za  o  ono $\Lambda = 0$.

Goode i Wainwright pokazali,   e rosn  ce z czasem zaburzenia metryki s   generowane przez niejednorodno  ci rozk  adu energii $((E/M^{2/3})_{,r} \neq 0)$, a malej  ce – przez nier  wnoczesno  c wybuchu pocz  tkowego ($t_{B,r} \neq 0$). Te pierwsze nazwiemy zaburzeniami typu E, te drugie – zaburzeniami typu BB.

Silk powi  za   te zaburzenia z fizycznymi charakterystykami modelu L–T: kontrastem g  sto  ci $D \stackrel{\text{def}}{=} \rho_{,r}/\rho$ i kontrastem krzywizny $K \stackrel{\text{def}}{=} {}^{(3)}R_{,r}/{}^{(3)}R$, gdzie ${}^{(3)}R$ jest krzywizn   przestrzeni $t = \text{const}$. Wyniki by  y nast  puj  ce:

1. W modelu z $E > 0$ (rozszerzaj  cym si   do niesko  czono  ci):

	typ E $t \rightarrow t_B$	typ E $t \rightarrow \infty$	typ BB $t \rightarrow t_B$	typ BB $t \rightarrow \infty$
D	0	$0 < D < \infty$	∞	0
K	$0 < K < \infty$	0	∞	$0 < K < \infty$

2. W modelu z $E < 0$ (zapadaj  cym si   do osobliwo  ci ko  cowej po sko  czonym czasie):

	typ E $t \rightarrow t_B$	typ E $t \rightarrow t_C$	typ BB $t \rightarrow t_B$	typ BB $t \rightarrow t_C$
D	0	∞	∞	∞
K	$0 < K < \infty$	0	∞	∞

4.5. Kr  tki prze  gl  d innych wyników uzyskanych w modelach L–T i Szekeres

- **Fluktuacje temperatury promieniowania t  a.** Zbadano w  lyw niejednorodno  ci napotykanym przez promienie   wiat  ne w drodze do obserwatora na fluktuacje temperatury promieniowania t  a (tzw. efekt Reesa–Sciamey, ale w   cis  ej teorii) [10]. Kilka niezale  znych bada   da  o zgodne wyniki: fluktuacje temperatury nie mog  y by   wi  ksze ni   10^{-6} – 10^{-5} . (Dla przypomnienia: „niezwykle dok  adn  ” izotropi   temperatury CMB podawano w latach 60. i 70. jako „dow  d” jednorodno  ci Wszech  wiata, gdy dok  adno  c pomiar  w wynosi  a 10^{-2}).
- **Wp  lyw ekspansji Wszech  wiata na orbity planet.** Ten w  lyw jest niemierzalnie ma  y, ale niezerowy (orbita Saturna powinna powi  ksza  c sw  j promie   o 1   rednic   protonu na 1000 lat [34]).
- **Fale g  sto  ci.** Przestrzenne maksima i minima g  sto  ci na og   l nie wsp   poruszaj   si   z materi   – w  druj  y niezale  nie od ruchu materii [21,32,2].
- **Proste zwi  zki mi  dzy znakiem krzywizny, topologi   przestrzeni i typem ewolucji, znane z modeli Friedmanna, w og  lno  ci nie istniej  .** Model z dodatni   krzywizn   przestrzenn   mo  e by  c przestrzennie

nieograniczony i mieć nieskończony czas życia, model z ujemną krzywizną może mieć skończony rozmiar przestrzenny [2].

5. Krótka konkluzja

Teoria względności ma do zaoferowania kosmologii znacznie więcej niż proste modele R–W, odkryte prawie 90 lat temu. Modele niejednorodne pozwalają wyjaśnić większość obserwowanych zjawisk bez uciekania się do „nowej fizyki”.

Historia fizyki powinna być nauczycielem nas większej ostrożności w postulowaniu istnienia materii o niezwykłych własnościach dla „wyjaśnienia” niezrozumiałych zjawisk. W szczególności, pouczająca jest tu historia idei eteru i historia „kosmologii stanu niezmiennego” (Bondiego–Golda i Hoyle’a). Eter miał być „naturalnym” układem odniesienia dla fal elektromagnetycznych, kosmologia Bondiego–Hoyle’a miała rozwiązać „problem wieku” Wszechświata. W obydwu przypadkach rozwiązanie problemu znalazło się dzięki konsekwentnemu analizowaniu i doświadczalnemu testowaniu „starej fizyki” i przesuwaniu jej zastosowań poza wcześniej znane granice. To jest dobry przykład do naśladowania dla dzisiejszej kosmologii.

6. Dodatek: krótka historia modeli kosmologicznych

1. 1917 – model Einsteina [35]: gęstość materii stała. Nierealistyczny z dzisiejszego punktu widzenia, ale odzwierciedlający ówczesną wiarę astronomów. Okazało się, że nie istnieje rozwiązanie oryginalnych równań Einsteina ze stałą gęstością. Aby usunąć tę sprzeczność Einstein wprowadził stałą kosmologiczną.
2. 1922 – pierwsza praca Friedmanna [36]: gęstość materii i krzywizna przestrzeni stałe w przestrzeni, ale zmienne w czasie, zerowe ciśnienie, stała kosmologiczna dowolna. Krzywizna przestrzeni dodatnia. Związek z kosmologią przez pierwsze 10 lat dla nikogo niewidoczny.
3. 1924 – druga praca Friedmanna [36]: wszystko j.w., ale krzywizna przestrzeni ujemna.
4. 1929 – praca Robertsona [37]: wyprowadzenie rozwiązań Friedmanna z jasno sformułowanych założeń geometrycznych. Dopiero w tej pracy pojawił się model z zerową krzywizną przestrzenną, przeoczony przez Friedmanna.
5. 1933 – praca Lemaître’a [6]: niejednorodne uogólnienie modeli Friedmanna – gęstość materii i krzywizna przestrzeni sferycznie symetryczne, ale zależne od położenia w przestrzeni poprzez dwie dowolne funkcje zmiennej radialnej. Czasoprzestrzeń Schwarzschilda jest innym przypadkiem granicznym tego rozwiązania. Lemaître przedstawił dyskusję własności

geometrycznych, fizycznych i astrofizycznych tego modelu wyprzedzającą o kilka dziesięcioleci możliwości percepcyjne ówczesnych czytelników (i niektórych dzisiejszych czytelników). Między innymi podjął pierwszą próbę opisu powstawania galaktyk (nazywanych wtedy „mgławicami”)⁴ i dowiódł, że powierzchnia $r = 2m$ w rozwiązaniu Schwarzschilda nie jest osobiwa.

Przez wiele lat potem model ten był cytowany jako „model Tolmana”. Dziś lansowana jest nazwa „model Lemaître’a–Tolmana (L–T)”, ale tylko dla odróżnienia od kilku innych modeli Lemaître’a. Pierwszeństwo Lemaître’a nie podlega żadnym wątpliwościom.

6. 1934 – praca Tolmana [7]: dowód, przy użyciu modelu Lemaître’a, że modele Friedmanna są niestabilne względem zaburzeń rozkładu gęstości: amplituda gęstości dla lokalnych zagęszczeń i rozrzedzeń będzie powiększać się z czasem. To była przepowiednia, że podczas ewolucji Wszechświata powinny tworzyć się pustki, ale nikt jej nie zrozumiał. Tolman cytował Lemaître’a. Kto dziś nazywa ten model „modelem Tolmana” ujawnia, że nie czytał pracy Tolmana.
7. 1934 – praca Sena [38]: komplementarne uzupełnienie pracy Tolmana; dowód, że modele Friedmanna są niestabilne względem zaburzeń początkowego rozkładu prędkości. Ta praca zawiera bardzo wyraźne stwierdzenie, że amplituda rozrzedzeń gęstości będzie rosła z czasem („the models are unstable for initial rarefaction”).
8. 1947 – praca Bondiego [39]: wszechstronna interpretacja modelu L–T z punktu widzenia fizyki i astrofizyki. Bondi podał interpretację fizyczną funkcji występujących w tym modelu i pokazał, że może on opisywać proces powstawania czarnej dziury. Różne wyniki tej pracy stały się w następnych dziesięcioleciach inspiracją dla całych programów badawczych. (Model L–T bywa często nazywany modelem Lemaître’a–Tolmana–Bondiego).
9. Po roku 1970 nastąpił, narastający z czasem, „wysyp” prac, w których używano modelu L–T do opisu różnych procesów astrofizycznych; jest ich zbyt wiele, aby dyskutować je w ramach „krótkiej historii”. Są one opisane w monografiach [33] i [10].
10. 1975 – praca Szekeresa [9]: uogólnienie modelu L–T. Jedno z kilku znanych dziś rozwiązań równań Einsteina bez symetrii. Zawiera 5 funkcji dowolnych jednej zmiennej i może opisać większe bogactwo struktur niż model L–T, który zawiera tylko dwie funkcje dowolne.
11. 2006–09 – seria prac K. Bolejko [14–17], w których zastosowano model Szekeresa do opisu ewolucji struktur we Wszechświecie i propagacji promie-

⁴Niestety, metoda Lemaître’a była nieprawidłowa. Założył on, że galaktyki powstają wewnątrz obszaru, na którego granicy odpychanie kosmologiczne związane ze stałą Λ zaczyna przeważać nad przyciąganiem grawitacyjnym. Założenie to prowadziło do całkiem niezgodnych z obserwacjami rozmiarów galaktyk.

niowania przez Wszechświat wypełniony lokalnymi strukturami.

Cennych informacji, wykorzystanych w artykule, dostarczyli mi współautorzy naszej najnowszej monografii [10], Marie-Noëlle Célérier, Charles Hellaby i Krzysztof Bolejko. Wdzięczny jestem też uczestnikom seminarium z teorii względności w Instytucie Fizyki Teoretycznej UW, których docieklive pytania zadawane podczas mojego referatu spowodowały uzupełnienie kilku luk.

Literatura

- [1] S.A. Gregory, L.A. Thompson, *Astrophys. J.* **222**, 784 (1978).
- [2] J. Plebański, A. Krasieński, *An introduction to general relativity and cosmology* (Cambridge University Press, 2006).
- [3] L. Bruna, J. Girbau, *J. Math. Phys.* **40**, 5117 (1999).
- [4] L. Bruna, J. Girbau, *J. Math. Phys.* **40**, 5131 (1999).
- [5] L. Bruna, J. Girbau, *J. Math. Phys.* **46**, 072501 (2005).
- [6] G. Lemaître, *Ann. Soc. Sci. Bruxelles* **A53**, 51 (1933); angielskie tłumaczenie z komentarzem historycznym: *Gen. Rel. Grav.* **29**, 637 (1997).
- [7] R.C. Tolman, *Proc. Nat. Acad. Sci. USA* **20**, 169 (1934); przedruk z komentarzem historycznym: *Gen. Rel. Grav.* **29**, 931 (1997).
- [8] C. Hellaby, K. Lake, *Astrophys. J.* **290**, 381 (1985); errata: *Astrophys. J.* **300**, 461 (1986).
- [9] P. Szekeres, *Commun. Math. Phys.* **41**, 55 (1975).
- [10] K. Bolejko, A. Krasieński, C. Hellaby, M.N. Célérier, *Structures in the Universe by exact methods – formation, evolution, interactions* (Cambridge University Press, w druku).
- [11] C. Hellaby, A. Krasieński, *Phys. Rev. D* **77**, 023529 (2008).
- [12] A. Krasieński, *Phys. Rev. D* **78**, 064038 (2008).
- [13] C. Hellaby, A. Krasieński, *Phys. Rev. D* **66**, 084011 (2002).
- [14] K. Bolejko, *Phys. Rev. D* **73**, 123508 (2006).
- [15] K. Bolejko, *Phys. Rev. D* **75**, 043508 (2007).
- [16] K. Bolejko, „Volume averaging in the quasispherical Szekeres model”, *Gen. Rel. Grav.* (w druku).
- [17] K. Bolejko, „The Szekeres Swiss Cheese model and the CMB observations”, *Gen. Rel. Grav.* (w druku).
- [18] H. Iguchi, T. Nakamura, K. Nakao, *Progr. Theor. Phys* **108**, 809 (2002).
- [19] S. Alexander, T. Biswas, A. Notari, D. Vaid, arXiv:0712.0370.
- [20] K. Enqvist, T. Mattsson, *J. Cosm. Astropart. Phys.* **02** (2007), 019 (2007).
- [21] A. Krasieński, C. Hellaby, *Phys. Rev. D* **65**, 023501 (2002).
- [22] A. Krasieński, C. Hellaby, *Phys. Rev. D* **69**, 023502 (2004).
- [23] A. Krasieński, C. Hellaby, *Phys. Rev. D* **69**, 043502 (2004).
- [24] K. Bolejko, A. Krasieński, C. Hellaby, *Mon. Not. R. Astron. Soc* **362**, 213 (2005).
- [25] K. Bolejko, *Mon. Not. R. Astron. Soc* **370**, 924 (2006).
- [26] J. Silk, *Astron. Astrophys.* **59**, 53 (1977).
- [27] S.W. Goode, J. Wainwright, *Phys. Rev. D* **26**, 3315 (1982).
- [28] T. Padmanabhan, *Structure Formation in the Universe* (Cambridge University Press, 1993).
- [29] P.J.E. Peebles, *Principles of Physical Cosmology* (Princeton University Press, 1993).
- [30] G.F.R. Ellis, T. Rothman, *Postępy Fizyki* **38**, 511 (1987).
- [31] M.N. Célérier, *Astron. Astrophys.* **362**, 840 (2000).
- [32] G.F.R. Ellis, C. Hellaby, D.R. Matravers, *Astrophys. J.* **364**, 400 (1990).
- [33] A. Krasieński, *Inhomogeneous Cosmological Models* (Cambridge University Press, 1997).
- [34] R. Gautreau, *Phys. Rev. D* **29**, 198 (1984).
- [35] A. Einstein, *Sitzungsber. Preuss. Akademie der Wissenschaften (Berlin)*, 142-152 (1917); angielskie tłumaczenie: H.A. Lorentz, A. Einstein, H. Minkowski, H. Weyl, *The Principle of Relativity* (Methuen, reprinted by Dover (1923)).
- [36] A.A. Friedmann, *Z. Physik* **10**, 377 (1922), *Z. Physik* **21**, 326 (1924); angielskie tłumaczenie obydwu prac z komentarzem historycznym: *Gen. Rel. Grav.* **31**, 1985 (1999); addendum: *Gen. Rel. Grav.* **32**, 1937 (2000).
- [37] H.P. Robertson, *Proc. Nat. Acad. Sci. USA* **15**, 822 (1929).
- [38] N.R. Sen, *Z. Astrophysik* **9**, 215 (1934); przedruk z komentarzem historycznym: *Gen. Rel. Grav.* **29**, 1473 (1997).
- [39] H. Bondi, *Mon. Not. Roy. Astr. Soc.* **107**, 410 (1947); przedruk z komentarzem historycznym: *Gen. Rel. Grav.* **31**, 1777 (1999).

Techniczne aspekty zderzacza LHC

Jan Kulka

Wydział Fizyki i Informatyki Stosowanej, Akademia Górniczo-Hutnicza, Kraków

Technical merit of Large Hadron Collider

Abstract: New accelerator LHC is commonly known due to the high energy physics experiments and expectations of hadron collision results. General description how the accelerator complex works and new technologies applied there are presented in this paper. Superconductive magnets and cooling cryogenic system on one side and energy extraction system for magnets protection on the other are the main subjects. Finally, technical problems which caused the project delay and future plans of LHC commissioning are shortly explained.

Zainteresowanie nowym eksperymentem fizyki wysokich energii skupiło się na jego przyszłych wynikach w doświadczeniach ze zderzeniami hadronów. Poniżej przedstawiono opis działania zespołów akceleratorów w CERN-ie oraz ważniejsze osiągnięcia techniki użyte przy budowie tego akceleratora. Nadprzewodnikowe magnesy i układy kriogeniczne do ich schładzania oraz systemy ekstrakcji energii i zabezpieczeń na wypadek zaniku stanu nadprzewodzącego stanowią główny cel opisu. Przedstawiono przyczyny awarii wpływających na opóźnienie projektu oraz plany ponownego rozruchu LHC.

Wielki Zderzacz Hadronów (LHC – Large Hadron Collider) jest urządzeniem badawczym fizyki wysokich energii zbudowanym w Europejskim Centrum Badań Jądrowych (CERN) pod Genewą. Zderzającymi się cząstkami przeciwbieżnymi będą w pierwszej fazie eksperymentu protony o energii nominalnej 7 TeV. W drugiej fazie eksperymentu badane będą zderzenia jąder ołowiu o energii 2,76 TeV/nukleon.

Zatwierdzenie projektu LHC nastąpiło na posiedzeniu Rady CERN-u w grudniu 1994 r. Natomiast harmonogram prac zatwierdzony został w grudniu 1996 r., kiedy uzyskano zapewnienia od krajów nie będącymi członkami CERN-u: Indii, Japonii, Kanady, Rosji i Stanów Zjednoczonych o ich wkładzie materiałowym w budowę LHC.

Zainteresowanie nowym eksperymentem fizyki wysokich energii skupia się na jego przyszłych wynikach w doświadczeniach ze zderzeniami hadronów. Interesującym aspektem tego największego (i najkosztowniejszego) przedsięwzięcia w dziejach nauki są też aspekty techniczne, wymagające najwyższego kunsztu w wielu dziedzinach jak: uzyskiwanie wysokiej próżni, wytwarzanie silnych pól magnetycznych przez solenoidy nadprzewodzące, techniki kriogeniczne, mikrofalowe systemy przyspieszania cząstek, sterowanie wiązkami cząstek wysokoenergetycznych. Zorganizowanie pracy około 1000 fizyków, inżynierów i techników jest wyzwaniem organizacyjnym najwyższej skali trudności. W artykule starano się tylko przybliżyć najważniejsze i najciekawsze z tych problemów, związanych z uzyskaniem wiązki cząstek o wysokiej energii, pomijając opis gigantycz-

nych detektorów znajdujących się w czterech punktach zderzeń.

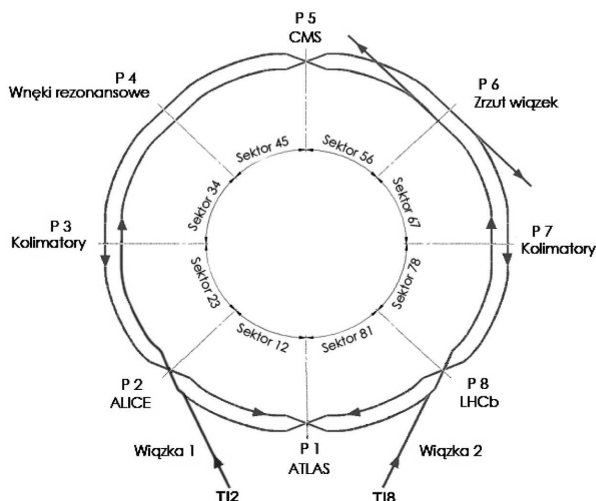
1. Tunel akceleratora

W tym nowym akceleratorze wykorzystuje się infrastrukturę pozostałą po poprzednim akceleratorze LEP, czyli zderzaczu elektronowo-pozytonowym, w którym wiązki e^- oraz e^+ rozpędzane były do energii 100 GeV.

W roku 2000 zakończono eksperymenty fizyczne przy zderzaczu LEP, a trzy lata później podziemny tunel tego akceleratora był opróżniony z jego elementów. Podziemne roboty górnicze związane były z budową dwu tuneli (łącznie 5,6 km) łączących LHC z SPS (Synergiczny Protonowy) – na rys. 1 oznaczone jako TI2 oraz TI8 – oraz budową dwu rozległych kawern dla eksperymentów ATLAS i CMS.

W punkcie 6 mają początek instalacje służące do bezpiecznego usuwania cząstek z akceleratora. Są to dwa ślepo zakończone tunele styczne do dwu trajektorii. Długość każdego z nich wynosi 750 m. Na ich końcu znajduje się grafitowy pochłaniacz cząstek usuwanych z LHC.

Na rysunku 1 przedstawiono rozmieszczenie eksperymentów fizyki wysokich energii (ATLAS, ALICE, CMS, LHCb) oraz szkic torów przyspieszanych protonów. Obwód akceleratora składa się z 8 łuków o długości 3 km każdy oraz 8 odcinków prostych łączących te łuki. Pośrodku każdego odcinka prostego znajdują się szyby wentylacyjne i transportowe łączące tunel akceleratora z powierzchnią. Jest osiem takich miejsc dostępu do LHC. Konwencja numeracji jest: eksperyment ATLAS punkt 1



Rys. 1. Rozmieszczenie eksperymentów fizyki wysokich energii na obwodzie LHC

i dalej zgodnie ze wskazówkami zegara. Pozostałe eksperymenty posadowione są w punkcie 2 – ALICE, w pkt. 5 – CMS, w pkt. 8 – LHCb. W punktach 3 i 7 znajdują się zespoły kolimatorów do usuwania z jonowodów cząstek zbytnio oddalających się od zadanej trajektorii.

Tunel akceleratora wydrążony jest w skałach osadowych masywu Jury, przeważnie na terytorium francuskim, na średniej głębokości 100 metrów. Płaszczyzna jego obwodu jest nachylona ze względów geologicznych pod kątem $1,4^\circ$ (co jest istotne dla przepływów kriogenicznych). Naziemne inwestycje strukturalne zrealizowano dla potrzeb: skraplarek helu, układów chłodniczych, systemów wentylacyjnych oraz podstacji energetycznych. Elementy te są rozproszone na całym obwodzie LHC (27 km). Należy nadmienić, że dla potrzeb transportu wielkogabarytowego dla LHC, jakim było opuszczanie magnesów dipolowych (długości 17 m z osłonami, opuszczanych w pozycji poziomej) dostępny był tylko jeden szyb.

2. System akceleratorów wstępnego przyspieszania protonów i jonów ołowiu

Protony krążą w LHC zgrupowane w pęczki o długości 7 cm i przekroju poprzecznym od 300 μm na łuku akceleratora do 16 μm w punktach zderzeń (ATLAS i CMS). Średnia liczba protonów w pęczku wynosi $1,15 \cdot 10^{11}$.

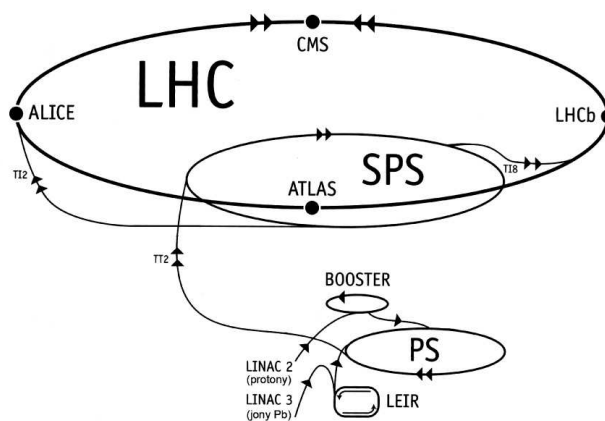
Maksymalna liczba pęczków protonów w LHC wynosi 2808, a to oznacza, że na jeden cykl przyspieszania potrzeba netto 0,6 ng wodoru.

Zanim pęczek cząstek znajdzie się w LHC przechodzi on przez ciąg akceleratorów pośrednich (rys. 2). Początek stanowi akcelerator liniowy LINAC 2. Źródłem protonów jest gazowy wodór. Z butli przepływa on do plazmotronu, gdzie po jonizacji wspomaganej działem elektronowym następuje separacja protonów od elektronów (napięcie 90 kV) i magnetyczne ogniskowanie protonów przed ich przelotem do akceleratora liniowego LINAC 2. W tym akceleratorze, o przemiennym polu elektrycznym z anten

o narastającej długości (akcelerator typu Alvareza), protony uzyskują energię 50 MeV.

Dalej protony będą przyspieszane w akceleratorach kołowych. Pierwszy z nich, czteropięścieniowy Booster, przyspiesza protony do energii 1,4 GeV. Z niego protony, już uformowane w pęczki, transferowane są synchronicznie do synchrotronu protonowego PS. W tym akceleratorze prędkość protonów osiąga 99,93% prędkości światła, co odpowiada energii 25 GeV. 3,6-sekundowy cykl przyspieszania protonów w PS kończy się przekierowaniem pęczków protonów poprzez tunel transferowy TT2 do SPS. W tym akceleratorze protony osiągają energię 450 GeV.

Przykładowy cykl napełniania LHC pęczkami protonów jest następujący: z Boostera w dwu cyklach (po 3 pęczki) napełnia się PS, potem przekierowuje się je do SPS, który wypełnia się w czasie 4 cykli PS. Po tym rozpoczyna się przyspieszanie w SPS, aż do energii 450 GeV. Osiągnięcie tej energii jest sygnałem do przemieszczenia tego pakietu 24 pęczków protonów do LHC.



Rys. 2. Kompleks akceleratorowy w CERN-ie

Napełnianie akceleratorów pęczkami cząstek wymaga nie tylko precyzji przestrzennej, ale przede wszystkim czasowej. W LHC odległość pomiędzy pęczkami wynosi 7,5 metra (25 ns). Oznacza to, że potencjalnych miejsc do umieszczenia pęczka protonów na trajektorii tego akceleratora jest 3554. Natomiast maksymalne projektowane wypełnienie wynosi 2808 pęczków. Różnica wynika z potrzeby zachowania części trajektorii wolnej od cząstek (przez około 1 μs), gdy załączane są magnesy sterujące napełnianiem jonowodu cząstkami i magnesy opróżniające jonowód z cząstek. Ta chwila czasu potrzebna jest do uzyskania nominalnej wartości indukcji pola przez te magnesy. Część miejsc na trajektorii jest nieobsadzona w celu symetryzacji pomiędzy eksperymentami CMS i ATLAS oraz dla kodowania położenia pęczka w wiązce.

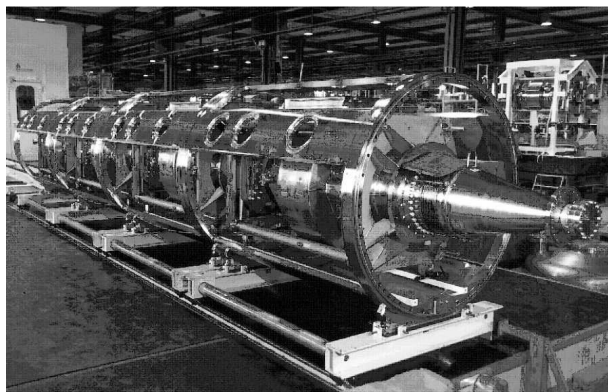
W eksperymencie, w którym zderzają się pęczki jąder ołowiu, proces wstępnego przyspieszania jest nieco odmienny od poprzedniego. W celu uzyskania jonów ołowiu podgrzewa się ten metal, aby uzyskać jego parowanie. Bombardując te pary strumieniem elektronów wytwarza się jony ołowiu. Po wstępnej separacji dominujących

frakcji Pb29+ są one przyspieszane w akceleratorze LINAC 3 do energii 4,2 MeV/nukleon, aby po przejściu przez folię węglową ulec dalszej jonizacji o przewodze jonów Pb54+. Niskoenergetyczny Pierścień Jonów (LEIR) nadaje im energię 72 MeV/nukleon. Stąd trafiają do PS uzyskując energię 5,9 GeV i dopiero opuszczając PS przechodzą przez drugą folię, gdzie obdzierane są z reszty elektronów (Pb82+).

Końcowa trajektoria przyspieszanych wiązek dla eksperymentu z protonami począwszy od PS aż do LHC jest identyczna jak dla eksperymentu z jonami ołowiu.

3. Mikrofalowy system przyspieszania

Jedynym miejscem w LHC, w którym dokonuje się przekazywanie energii przyspieszanym cząstkom, są wnęki rezonansowe umieszczone w punkcie 4. Dla każdej trajektorii są dwa moduły po 4 wnęki w jednym kriostacie, jak na rys. 3. Częstotliwość pracy 400 MHz narzuca rozmiary kriostatu, gdyż długość wnęki wynosi $\lambda/2$, a odstęp pomiędzy wnękami $3\lambda/2$. Rozmiary wnęki powodują, że na tym odcinku akceleratora odległość pomiędzy jonowodami należało zwiększyć z nominalnych 194 mm do 420 mm, aby druga wiązka (nie przyspieszana w tym module) przelatowała w jonowodzie znajdującym się w tym samym kriostacie, ale poza wnęką. Wnęki zasilane są poprzez 22-metrowe falowody z 16 klistronów o mocy 300 kW każdy umieszczone w osobnej kawernie przyległej do tunelu. Dla takiego przyspieszającego modułu potrzeba aż trzech systemów próżniowych dla: wnęki, jonowodu z cząstkami nie przyspieszanymi oraz kriostatu.



Rys. 3. Wnęki rezonansowe LHC w trakcie montażu

Aby utrzymać wnękę w stanie nadprzewodzącym na wewnętrzną powłokę wnęki rezonansowej naniesiono niob, a temperatura jej pracy wynosi 4,5 K. Technologia ta umożliwia uzyskanie wartości przyspieszającego pola elektrycznego 5,5 MV/m.

4. Cykl przyspieszania protonów, własności wiązki

Napełnianie LHC pęczkami protonów trwa około pół godziny (poprzez tunel TI2 na trajektorię zgodną z ru-

chem wskazówek zegara i poprzez TI8 na trajektorię przeciwną). W tym czasie poprzez wnęki rezonansowe kompensowane są straty energii krążących pęczków protonów. Umieszczenie ostatniego pakietu z SPS na trajektorii LHC oznacza możliwość rozpoczęcia przyspieszania w tym akceleratorze. O długości tego cyklu decydują dwa współzależne parametry: maksymalna szybkość narastania indukcji pola magnetycznego w magnesach nadprzewodnikowych (10 A/s) oraz natężenie pola elektrycznego we wnękach rezonansowych (5,5 MV/m). Protony osiągną swoją energię nominalną (7 TeV) po około pół godzinie od napełnienia LHC.

Położenie przestrzenne krążących pęczków protonów monitorowane jest co 56 metrów przez pojemnościowe detektory położenia wiązki. Odstępstwo od zakładanej trajektorii stanowi sygnał dla obwodów sterowania magnesami głównymi i korekcyjnymi. Ponadto pęczek, sam w sobie, jako zbiór identycznie naładowanych cząstek, ma tendencję do powiększania swoich rozmiarów. Z tego to powodu, co 56 metrów na trajektorii wiązki umieszczone są magnesy kwadrupolowe ogniskujące pionowo w płaszczyźnie prostopadłej do toru wiązki, a za następne 56 metrów w kierunku poziomym. Gradient 223 T/m tego ogniskującego pola w magnesach kwadrupolowych wymaga, by nominalne natężenie prądu w uzwojeniach tych magnesów wynosiło 11 870 A. Zastosowanie tego typu magnesów pozwala utrzymać poprzeczne rozmiary wiązki na poziomie 0,3 mm w obszarze łuku maszyny.

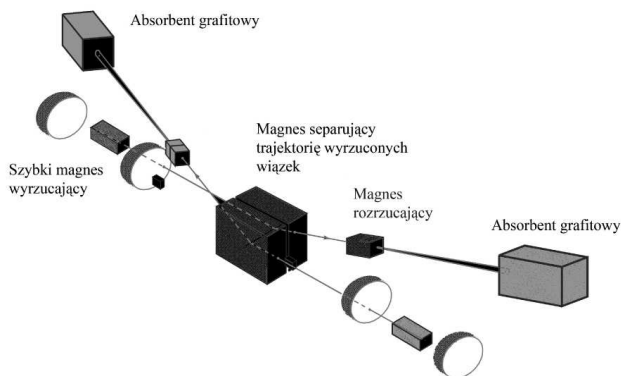
Protony opuszczające akcelerator SPS są już cząstkami relatywistycznymi (dla $E = 450$ GeV wartość $v/c = 0,999998$). Dalsze ich „przyspieszanie” oznacza, przy niewielkim wzroście prędkości, wzrost ich masy relatywistycznej, która przy energii 7 TeV osiąga wartość prawie 7500 razy większą od masy spoczynkowej. Utrzymanie coraz większej masy na torze o ustalonej krzywiźnie przy prawie stałej prędkości ładunku zbliżonej do prędkości światła wymaga, by indukcja pola magnetycznego generującego siłę dośrodkową wzrastała proporcjonalnie do energii cząstek. Dla protonów o energii 7 TeV, wielkość indukcji pola magnetycznego w głównych dipolach wynosi 8,3 tesli.

Dla zwiększenia prawdopodobieństwa zderzenia, tuż przed punktami zderzeń pęczek protonów jest dodatkowo ściskany do rozmiarów poprzecznych 16 μm (ATLAS, CMS). Osiągnięto to poprzez zastosowanie silnie ogniskujących magnesów kwadrupolowych, po 3 magnesy z każdej strony czterech eksperymentów. To ostatnie ogniskowanie pęczków odbywa się dopiero po osiągnięciu zamierzonej energii i uzyskaniu stabilnej trajektorii wiązki.

Spodziewany czas trwania zderzających się wiązek w LHC szacowany jest na około 10 godzin. Po tym czasie wiązka degeneruje się w takim stopniu, że przestaje być użyteczna dla eksperymentów. Główną przyczyną degeneracji jest oddziaływanie wzajemne pomiędzy pęczkami poruszającymi się w przeciwnych kierunkach w jednym jonowodzie na długości 524 m. W takiej sytuacji oraz w przypadku zaistnienia zagrożenia braku stabilności trajektorii, wiązka zrzucana jest do absorbentów w postaci

bloków grafitu o długości 9 metrów znajdujących się w zakończeniach tuneli styrczych w punkcie 6.

Energia jednej wiązki przy pełnym wypełnieniu jonowodu LHC wynosi 362 MJ (to wystarcza do stopienia 500 kg miedzi). Depozycja tej energii odbywa się w czasie 0,1 ms. Decyzja o usunięciu wiązek z akceleratora oznacza dla specjalnych niskiindukcyjnych magnesów jednozwojowych włączenie pól magnetycznych, które nieznacznie odchyłają wiązki z trajektorii nominalnej i skierują je do magnesu separującego.



Rys. 4. Schemat usuwania wiązki z jonowodów

Cząstki po przełocie przez ten separujący magnes odchylane są w kierunku pionowym. Około 150 metrów za tym magnesem znajduje się specjalny magnes rozrzucający pęczki w taki sposób aby tworzyły one w chwili uderzenia w powierzchnię grafitu spiralę o średnicy 30 cm. Dzięki małemu przekrojowi czynnemu grafitu na zderzenia, energia wiązki jest rozpraszana w całej jego objętości. W wyniku wyhamowania cząstek w graficie jego temperatura wzrasta do 850 °C. Po obniżeniu prądu w magnesach LHC do 763 A, można rozpocząć ponowne napełnianie LHC pęczkami cząstek z SPS.

5. Hybrydowa struktura LHC

Rozległość akceleratora LHC sprawia, że systemy kriogeniczne, energetyczne oraz chłodnicze i wentylacyjne muszą być podzielone tak, aby minimalizować straty transportowe. I tak skraplarki helowe zlokalizowano w punktach parzystych LHC. Hel tłoczony jest od takiego punktu do końca sektora (3,5 km), np. z punktu 4 do całego sektora 34 i sektora 45 (rys. 1). Chłodzenie jednego sektora od temperatury otoczenia do 20 K trwa 25 dni. Przez następne 3 dni odbywa się wypełnianie zbiornika ciśnieniowego magnesu nadprzewodnikowego helem, by osiągnąć temperaturę 4,5 K. Dodatkowe dwa dni są potrzebne dla schłodzenia tych zbiorników wraz z zawartością do 1,9 K.

Zasilanie energetyczne magnesów nadprzewodnikowych i konwencjonalnych jest szeregowo dla poszczególnych typów magnesów, a zespoły prostownikowe znajdują się w kawernach w pobliżu szybów. Tam też umieszczono sektorowe układy rozpraszania energii (w sytuacjach awa-

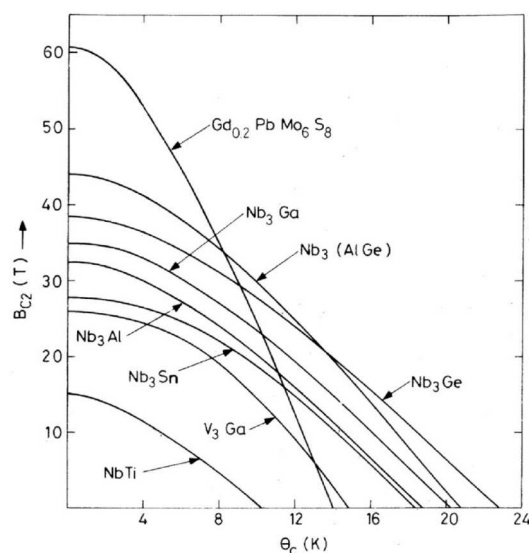
ryjnych) z magnesów nadprzewodnikowych. Zarówno systemy kriogeniczne, jak i zasilania energetycznego potrzebują wydajnego systemu chłodzenia i wentylacji. Układy wytwarzania wody lodowej oraz chłodnie kominowe znajdują się w punktach parzystych LHC (rys. 1).

Średnie zapotrzebowanie zespołu akceleratorów wraz z LHC na energię elektryczną wynosi 120 MW, przy czym oprócz wymienionych wyżej systemów dla LHC należy uwzględnić akceleratory wstępne, których magnesy są konwencjonalne, tzn. z uzwojeniem miedzianym chłodzonym wodą. Oczywiście należy uwzględnić pobór energii przez aparaturę do eksperymentów fizycznych również znajdującą się pod powierzchnią ziemi. Jako zasadę przyjęto niestosowanie ciekłego azotu w instalacjach podziemnych. Natomiast na powierzchni jest on stosowany jako podstawowe chłodziwo w wymiennikach ciepła skraplarek helowych.

6. Kable nadprzewodzące

Potrzeba około 10 tysięcy różnego typu magnesów nadprzewodnikowych wymusiła dopracowanie techniki wytwarzania drutów nadprzewodzących i w dalszym etapie produkcji nadprzewodzących kabli.

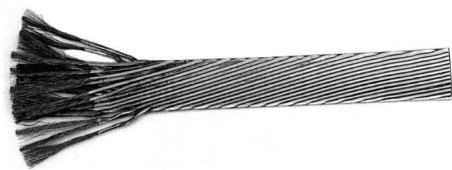
Spośród wielu znanych dzisiaj materiałów nadprzewodzących (rys. 5), ze względu na własności elektryczne, magnetyczne i koszty produkcji, tylko dwa z nich nadają się do produkcji drutów. Są to NbTi oraz Nb₃Sn. O wyborze stopu NbTi zadecydowały jego własności mechaniczne (ciągliwość, brak pękania kruche) mimo gorszych niż dla Nb₃Sn takich parametrów, jak temperatura krytyczna i krytyczna gęstość prądu.



Rys. 5. Charakterystyki: indukcja krytyczna – temperatura krytyczna potencjalnych materiałów na druty nadprzewodzące

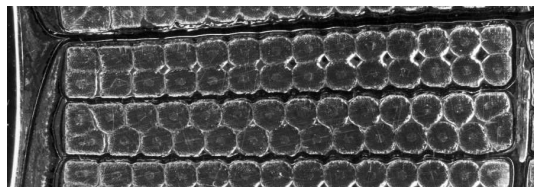
Najbardziej rozpowszechniony drut do produkcji kabli dla LHC posiada średnicę 1,06 mm i składa się z miedzianego rdzenia, wokół którego znajduje się 8900 włókien

ze stopu niob–tytan każde o średnicy 7 μm i zewnętrznej otuliny – stabilizatora miedzianego (rys. 6). Produkcja



Rys. 6. Kabel typu Rutherforda z usuniętą z końcówek drutów otuliną miedzianą. Szerokość 15 mm.

takich drutów polega na wielokrotnym przeciąganiu i termicznym usuwaniu naprężeń materiału wyjściowego którym jest rura miedziana o średnicy powyżej 0,5 m z ułożonymi w jej wnętrzu prętami stopu Nb–Ti. Optymalizacja formy splotu kabla pod względem mechanicznym (siły oddziaływań pomiędzy drutami po załączeniu prądu) i elektrycznym (druty nie są od siebie izolowane) wykazała, że śrubowy skok kabla sprasowanego trapezoidalnie powinien wynosić między 200 a 230 mm w zależności od liczby drutów (od 28 do 36) w kablu i maksymalnego natężenia prądu zasilania. Trapezowy przekrój kabla daje możliwość formowania cewki magnesu równoległe do osi rury jonowodu bez stosowania przekładek wypełniających. Na kabel cewki przed jej formowaniem nakładana jest izolacja kaptonowa przepuszczająca nadciekły hel.



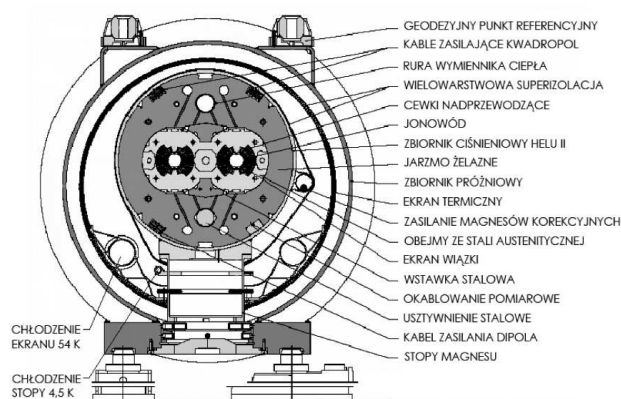
Rys. 7. Przekrój poprzeczny uzwojenia magnesu dipolowego. Widoczne dwa 28-drutowe kable typu Rutherforda.

7. Magnesy nadprzewodnikowe

Na 26,7-kilometrowym obwodzie LHC aż 21,4 km zajęte jest przez kriostaty magnesów nadprzewodnikowych (18,7 km przez dipolowe i 2,7 km kwadrupolowe). Nadprzewodnikowe magnesy korekcyjne: sektupolowe, oktapolowe i dekapolowe znajdują się w kriostatach wymienionych magnesów głównych. Na całym obwodzie LHC znajdują się 1232 nadprzewodnikowe magnesy dipolowe o długości ponad 15 metrów oraz 392 nadprzewodnikowe kwadrupole o długości ponad 6 m każdy. Magnesy nadprzewodnikowe stanowią przykład konstrukcji, w której spletają się technika wysokiej próżni, kriogenika i nadprzewodnictwo.

Rysunek 8 przedstawia przekrój przez magnes dipolowy w płaszczyźnie prostopadłej do trajektorii przyspie-

szanych cząstek. Wszystkie magnesy rozlokowane na łukach LHC są dwuaperturowe. Pole magnetyczne wytworzone przez prąd płynący w solenoidach zamyka się poprzez rdzeń żelazny (pionowe płyty o grubości 6 mm). Kriogeniczne warunki pracy magnesów nadprzewodnikowych wymagają zastosowania płaszcza próżniowej izolacji od otoczenia. Ciśnienie w objętości pomiędzy zbiornikiem próżniowym a ciśnieniowym zbiornikiem helu II (tzw. zimnej masy) utrzymuje się na poziomie 10^{-4} Pa. Ta próżnia izolacyjna jest ze względu na ułatwienie napraw podzielona na odcinki 214 m.



Rys. 8. Przekrój poprzeczny dwuaperturowego magnesu nadprzewodnikowego

Wymagania dotyczące wielkości próżni w jonowodzie są znacznie wyższe. Próżnia w tej strefie wynosi 10^{-8} Pa. Jej wielkość wynika z oczekiwanego czasu życia wiązki (100 godzin) ograniczonego przez rozpraszania cząstek na atomach gazów reszkowych nie usuniętych przez układy próżniowe. Wewnątrz jonowodu znajduje się miedziany perforowany ekran wiązki. Perforacja ma umożliwić uwięzienie atomów gazów reszkowych w przestrzeni pomiędzy rurą jonowodu pokrytą nieparującym getterem¹ a ekranem wiązki. Strumień cząstek odpowiada natężeniu prądu 0,5 A. Tworzący się w ekranie i przemieszczający się ładunek-obraz generuje wraz z absorpcją promieniowania synchrotronowego moc z ilości 0,4 W na metr długości jonowodu. W celu usunięcia tego ciepła, ekran wiązki wyposażono w kapilary do przepływu helu o temperaturze 4,5 K. Nie jest to podstawowy problem z usuwaniem ciepła.

Obudowa kriostatu znajduje się w temperaturze otoczenia, a zimna masa z rdzeniem magnesu w 2 K. W celu ograniczania dopływu promieniowania ciepłego do zimnej masy owinięta jest ona 10-warstwową izolacją z folii aluminiowej przełożonej siatką poliestrową. Dodatkowo na folię naniesiona jest warstwa refleksyjna, co sprawia, że dopływ ciepła poprzez tą izolację nie przewyższa 60 mW/m^2 . Ta izolacja chłodzona jest przez kontakt

¹Getter (z ang.) – substancja wprowadzana do wnętrza urządzeń próżniowych po ich odpompowaniu w celu usunięcia z nich resztek gazu poprzez jego pochłanianie – red.

cieplny z ciśnieniowym zbiornikiem helowym. Temperatura jej zewnętrznej warstwy nie przewyższa 20 K.

Następna warstwa izolacji termicznej nie ma kontaktu z zimną masą. Rysunek 9 pokazuje jej usytuowanie wokół już zaizolowanej zimnej masy. Blacha aluminiowa w formie łuku wsparta jest na łożu. Na tak przygotowanej konstrukcji nakłada się następną, tym razem 15-warstwową superizolację.



Rys. 9. Montaż ekranu termicznego wokół magnesu dipolowego

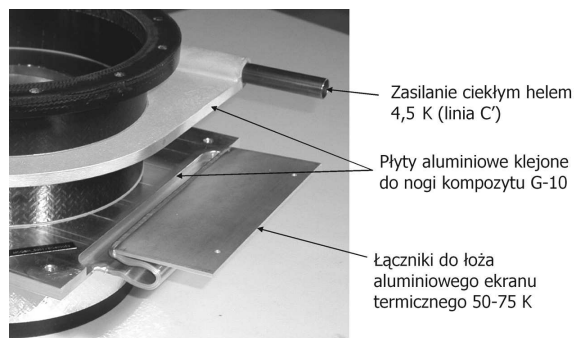


Rys. 10. Łoże magnesu dipolowego z przygotowanymi wspornikami (stopami) dla „zimnej masy”. Z lewej strony przepust dla chłodzącego gazowego helu.

W łożu z zamkniętego profilu aluminiowego znajduje się szczelny okrągły kanał dla medium chłodniczego. Jest nim hel gazowy o temperaturze 50 K i ciśnieniu 20 barów. Kanał ten jest łączony szeregowo przez wszystkie magnesy w sektorze. Takie rozwiązanie konstrukcyjne izolacji cieplnej minimalizuje dopływ ciepła do poziomu $1,1 \text{ W/m}^2$.

27-tonowa zimna masa magnesu dipolowego spoczywa wewnątrz kriostatu na trzech stopach. Ich konstrukcja ma zminimalizować dopływ ciepła od korpusu krio-

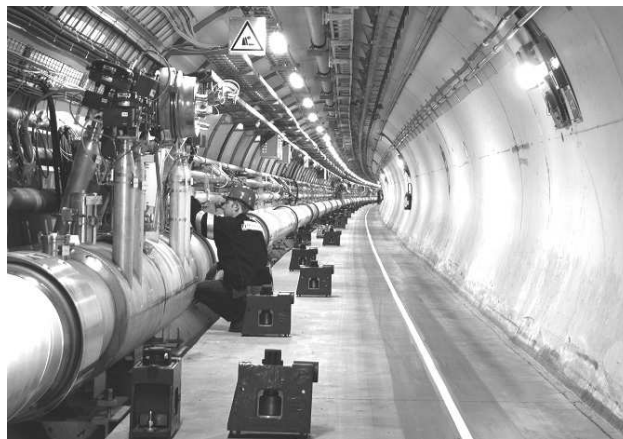
statu, a ponadto umożliwić stabilność mechaniczną kurczącego się o 40 mm w trakcie chłodzenia magnesu dipolowego. Kompozytowe nogi posiadają przekładki z płyt aluminiowych, z których dolna ma połączenie termiczne do łoża, natomiast górna jest zasilana osobnym obiegiem ciekłego helu o temperaturze 4,5 K.



Rys. 11. Stopa magnesu dipolowego z elementami układu chłodzenia

8. Obwód kriogeniczny. Nadciekły hel

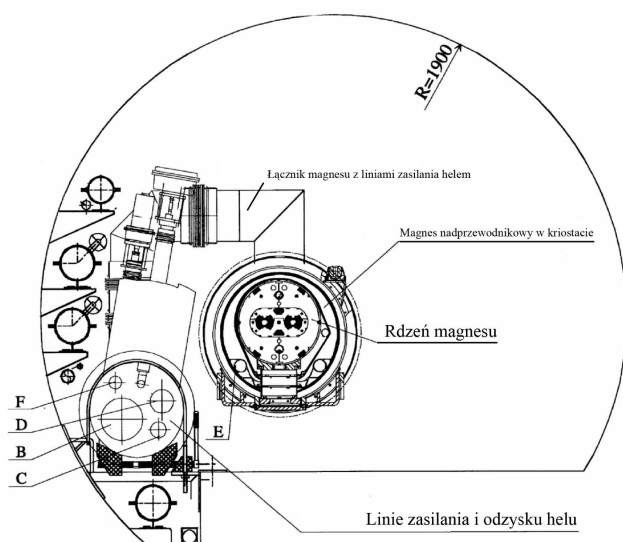
Na łuku akceleratora, co 107 metrów znajduje się połączenie linii dystrybucji helu do magnesów nadprzewodnikowych. Na rysunku 12 pokazano helociąg z wystającymi z niego pneumatycznymi zaworami regulacji przepływu helu do magnesów i króciec z rurami na ciekły hel. Linia ta zostaje w dalszej fazie instalacji zasłonięta magnesami.



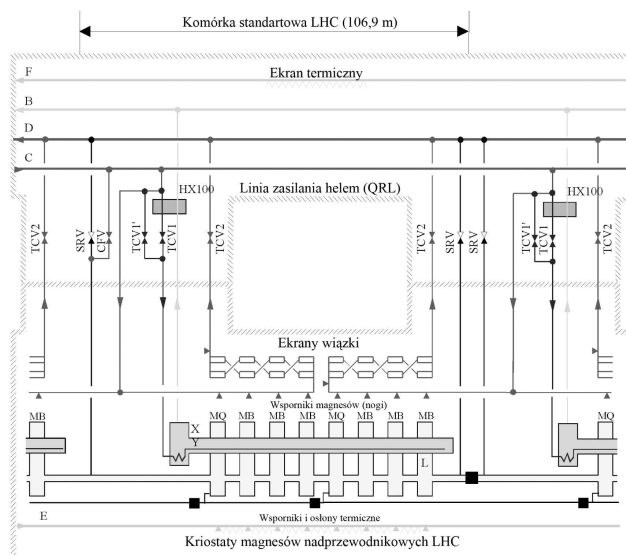
Rys. 12. Końcowa faza montażu linii zasilania kriogenicznego. Obok linii widoczne przygotowane do czopowania w podłodze wsporniki magnesów. Namalowana na podłodze biała linia służy do automatycznego kierowania ruchu pociągu rozwozowego magnesów i elementów konstrukcyjnych.

Druga strona łącznika do linii zasilania helem (QRL) znajduje się na magnecie kwadrupolowym. Wszystkie przepusty są wykonane ze stali nierdzewnej i spawane w ochronnej atmosferze argonowej.

Kriostat QRL zawiera 4 przewody z helem o różnej temperaturze i kierunku przepływu (rury B, C, D i F). Stanowi on równocześnie osobny obwód dla instalacji próżniowych z izolacją termiczną podobną do tej w magnesach. Poprzez jeden łącznik hel dopływa do ciągu (rys. 14): magnes kwadrupolowy (MQ), trzy magnesy dipolowe (MB), magnes kwadrupolowy, trzy magnesy dipolowe.



Rys. 13. Względne rozmieszczenie linii zasilania kriogenicznego i magnesów w tunelu LHC



Rys. 14. Struktura standardowej komórki układu kriogenicznego na łuku akceleratora LHC

W łączniku z linią C znajduje się zawór Joule'a–Thomsona (HX100) rozprężający hel od 1,3 bara do 16 milibarów. Uzyskany w ten sposób nadciekły hel prze-

plywa swobodnie rurą Y do końca ciągu magnesów i wypełnia rurę wymiennika ciepła X. Nadciekły hel jest medium chłodzącym dla każdego 27-tonowego rdzenia magnesu dipolowego i 6,5-tonowego kwadrupola. Ze względu na własności dielektryczne rozprężonego helu, nie nadaje się on do chłodzenia uzwojeń magnesu. Cewki magnesu i jego jarzmo zanurzone są całkowicie w statycznym helu o ciśnieniu 1,3 bara, zalewane wcześniej helem poprzez linię SRV. Kontakt cieplny pomiędzy tymi obydwooma obwodami helu zapewnia miedziana rura wymiennika ciepła. Po dwu dniach schładzania poprzez rozprężony hel temperatura cewek obniża się od 4,5 K do 1,9 K, a wtedy cewki można zasilić prądem o natężeniu nominalnym 11 850 A.

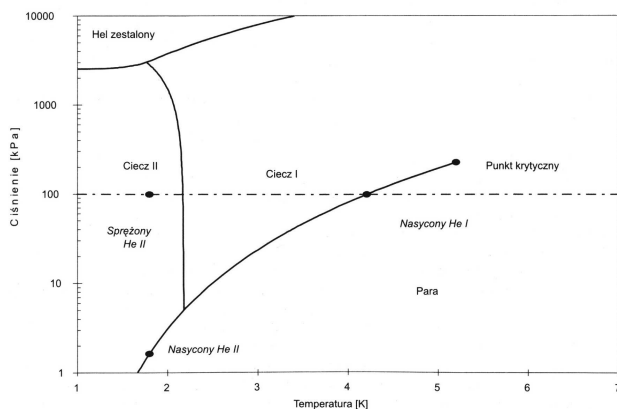
Niestety stan nadprzewodzący w kablach Nb–Ti jest wrażliwy na ustępujące naprężenia mechaniczne, radiację itp. prowadzące do lokalnego wzrostu temperatury i w efekcie zaniku nadprzewodnictwa. Takie przejście oznacza groźbę wydzielania się energii zgromadzonej w polu magnetycznym całego dipola (7 MJ) w miejscu powstania przewodnictwa oporowego. Powstanie takiego miejsca powoduje pojawienie się napięcia na kablach doprowadzających zasilanie do dipola. Jeśli po upływie 3 milisekund napięcie to pozostaje nadal, to jest to sygnał dla układów ochrony magnesu do załączenia oporowej taśmy grzewczej znajdującej się pomiędzy warstwami cewki nadprzewodnika. Akcja podgrzania cewki wyprowadza uzwojenie ze stanu nadprzewodzącego do stanu normalnego, a w jej wyniku energia pola magnetycznego wydzielana się w całej objętości magnesu, a prąd płynący szeregowo przez magnesy w sektorze omija ogrzane uzwojenie przez diodę. Wysokoprądowa dioda jest zintegrowana z korpusem zimnej masy. Brak takiego zabezpieczenia doprowadziłby do trwałego zniszczenia magnesu. Pod każdym magnesem nadprzewodnikowym znajduje się moduł ochrony magnesu na wypadek utraty nadprzewodnictwa. Jest to zespół stale naładowanych kondensatorów gromadzących energię niezbędną do ogrzania cewek magnesu powyżej temperatury krytycznej nadprzewodnika.

Po raz pierwszy zastosowano na tak ogromną skalę, 27 km, obwody kriogeniczne z nadciekłym helem do utrzymywania temperatury 1,9 K. Jest to temperatura pracy uzwojenia głównych magnesów nadprzewodnikowych: dipolowych i kwadrupolowych oraz magnesów korekcyjnych: sektupolowych, oktapolowych i dekapolowych. W temperaturze 2,17 K hel staje się cieczą nadciekłą (rys. 15), co polepsza jego własności penetrujące (brak lepkości).

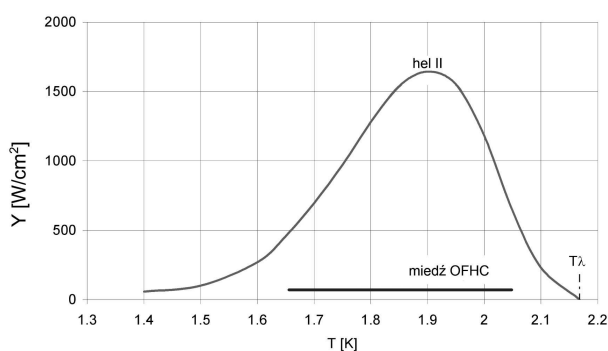
Dodatkowym powodem zastosowania helu II jako chłodziwa jest jego ogromna w stosunku do miedzi² przewodność cieplna, zwłaszcza w pobliżu temperatury 1,9 K (rys. 16).

We wszystkich instalacjach kriogenicznych LHC zgromadzone jest 120 ton helu, przy czym 30 ton przypada na skraplarki.

²Nawet miedzi OFHC (Oxygen-Free High Conductivity copper), czyli miedzi beztlenowej odznaczającej się wysoką przewodnością.



Rys. 15. Diagram fazowy helu



Rys. 16. Przewodność cieplna helu II poniżej 2,17 K

9. Nadprzewodniki wysokotemperaturowe w przepustach prądowych

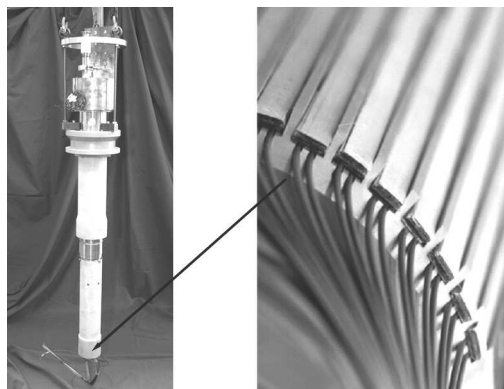
Po raz pierwszy na wielką skalę zastosowano materiały ceramiczne z nadprzewodnika wysokotemperaturowego na bazie bizmutu, a mianowicie $\text{Bi}_2\text{Sr}_2\text{Ca}_2\text{Cu}_3\text{O}_x$, którego temperatura krytyczna wynosi 110 K. Jego techniczna nazwa to BSCCO 2223. Osiągnięcia ostatnich lat w produkcji tych materiałów pozwoliły na ich zastosowanie jako przepustów prądu pomiędzy kablem zasilającym magnesy znajdującym się w temperaturze pokojowej a nadprzewodzącą linią będącą w temperaturze 4,5 K. Zastosowanie ceramiki w tym wrażliwym miejscu ograniczyło 10-krotnie w stosunku do przepustów miedzianych przepływ ciepła do wnętrza kriostatów z ciekłym heliem. Wymiernym efektem jest 30% redukcja zapotrzebowania na moc elektryczną dla układów schładzania. Pozytywne wyniki testów instalacji prototypowych doprowadziły do ich powszechnego zastosowania we wszystkich takich przepustach (1182) dla prądów o natężeniu od 600 do 13 000 A.

Jakkolwiek technologię wytwarzania tych materiałów opracowano w CERN-ie, to ze względu na ilość potrzebnych przepustów produkcję powierzono dwu firmom zewnętrznym. W sumie zużyto 31 km taśmy z tego nadprzewodnika o szerokości 4 mm i grubości 0,2 mm. Ze względów technologicznych pojedynczy odcinek składał

się z 7 warstw taśmy umieszczonej w srebrnej osnowie z dodatkiem 4% złota. Przepust prądowy o całkowitej długości 1,5 m (rys. 17) jest umieszczony w indywidualnym kriostacie wypełnionym w części ciekłym heliem. Nadprzewodnik BSCCO 2223 jest umieszczony w dolnej części przepustu.

Odcinek ten o długości 35 cm znajduje się pomiędzy strefą ciekłego helu 4,5 K a miedzianym końcem przepustu w temperaturze 50 K. Pojedynczy przepust nadprzewodzący utworzony jest z 36 połączonych równolegle tych 7-warstwowych odcinków.

Obróbka mechaniczna, lutowanie i spawanie elektro- nowo z nadprzewodnikiem Nb-Ti odbywa się w próżni. Dzięki dobranej technologii i konstrukcji wsporników dla taśm BSCCO przepływ chłodzącego medium (hel gazowy 20 K) wynosi 0,045 g/(s · kA). Zastosowanie tych nadprzewodników nie zapobiega konieczności ogrzewania ciepłej strony przepustu, jeżeli ochłodzone magnesy nie są zasilane prądem. Brak dopływu ciepła z kabli zasilających doprowadziłby do kondensacji pary wodnej na przepustach pogarszając własności izolacyjne przepustu.



Rys. 17. Przepusty prądowe do separacji cieplnej układów kriogenicznych od układów „ciepłych”. Po prawej – zespół równolegle połączonych taśm BSCCO zgrzanych z nadprzewodzącym drutem stopu Nb-Ti.



Rys. 18. Moduł połączeniowy energetycznych kabli zasilających poprzez nadprzewodniki wysokotemperaturowe z obwodami kriogenicznymi magnesów nadprzewodnikowych

10. Uruchomienie akceleratora, awaria, naprawy

Pierwotny harmonogram uruchamiania LHC zakładał testowanie kompletnych instalacji w sektorach. Tak więc jako pierwszy miał być uruchomiony sektor 78 z iniekcją cząstek z SPS poprzez tunel TI8. Niestety w trakcie testów ciśnieniowych (do 28 barów) linii zasilania kriogenicznego (QRL) doszło do deformacji linii. Negatywny wynik testu wymusił zmiany konstrukcyjne wnętrza kriostratu linii i sposobu mocowania linii do podłoża tunelu. Dla uniknięcia transportu długich segmentów QRL na powierzchnię, adaptowano kawernę punktu 6 na warsztaty naprawcze. Z powodu robót demontażowych w sektorze i transportu dla tych prac, zmieniono w tunelu trasę jazdy pociągów z magnesami. Wąskim gardłem jest tunel łączący z powierzchnią przeznaczony dla elementów wielkogabarytowych. Z uwagi na wrażliwość mechaniczną magnesów maksymalna prędkość transportowa w tunelu wynosiła 3 km/h. Oznaczało to, że dziennie można było umieścić na pozycji w tunelu średnio 3 magnesy. Należy nadmienić, że realizowana precyzja posadowienia geodezyjnego magnesu wynosiła 0,3 mm na 27 km obwodzie LHC. Ostatecznie ten negatywny wynik testu wymusił zmianę planu uruchamiania LHC. Od 2006 r. prace montażowe prowadzono praktycznie równocześnie na całej długości LHC z zachowaniem wyprzedzających prac przy QRL.

Podobny typ uszkodzenia wystąpił w 2007 r. w trakcie testów ciśnieniowych magnesów ogniskujących wiązkę przed wlotem w obszar zderzeniowy. Magnesy te (tzw. wewnętrznego trypletu) zostały dostarczone (24 sztuki) przez laboratoria Brookhaven (USA) i KEK (Japonia). W magnesach tych zastosowano inny niż w dipolach europejskich sposób mocowania w kriostracie i izolacji termicznej. Po tym teście magnesy wywieziono na powierzchnię i wymieniono sposób mocowania zimnej masy. Ta modyfikacja trwała prawie 3 kwartały.



Rys. 19. Przekrój łącznika ekranu wiązki z kompensacją termicznego skurczu magnesu

W tym czasie ogrzano wcześniej oziębione magnesy jednego sektora. Okazało się, że specjalny moduł łączący ekrany wiązki sąsiadujących magnesów zamknął złożonymi kontaktami ślizgowymi trasę przelotu wiązki. Ten moduł ma zapewnić ciągłość elektryczną pomiędzy magnesami dla ładunku-obrazu przemieszczającego się w ekranie wiązki oraz szczelność próżniową jonowodu,

gdy magnesy wskutek skurczu termicznego oddalają się od siebie o 40 mm. Defekt ten wykryto przez przypadek. Wykrycie miejsca jego występowania było trudne ze względu na brak nieinwazyjnej metody detekcji. Opracowanie metody polegającej na przepychaniu lekko sprężonym powietrzem nadajnika radiowego mieszczącego się w jonowodzie pozwoliło na wykrycie wszystkich takich miejsc blokady. Ciśnienie pchającego gazu dobrano tak, aby nie uszkodzić warstwy gettera w jonowodzie. Przelot nadajnika przez sektor 3,5 km trwa około 10 minut. Otwierano tylko te połączenia magnesów, w których nadajnik zatrzymał się.

Na koniec sierpnia 2008 r. wszystkie sektory LHC były sprawdzone pod względem kriogenicznym, próżniowym i elektrycznym dla prądów magnesów pozwalających utrzymać zamkniętą trajektorię cząstek. Pomyślnie przetestowano tunele łączące SPS i LHC dla protonów o energii 450 GeV.

10 września 2008 r. dokonano iniekcji protonów do LHC w celu sprawdzenia całej trajektorii. W ciągu godziny sprawdzono sektor po sektorze trajektorię lotu protonów i wykonano korekcję sterowania tak, że wstrzelona do LHC wiązka nr 1 (przeciwie do ruchu wskazówek zegara) pokonała cały obwód akceleratora i została wyrzucona do bloku grafitowego w punkcie 6. Test dla wiązki 2 był większym sukcesem, gdyż osiągnięto pełną synchronizację przelatujących cząstek z klistronami w punkcie 4.



Rys. 20. Uruchomienie LHC – 10 września 2008 r. Kierujący projektem Lyn Evans (z lewej) i dyrektor CERN-u Robert Aymar.

W cztery dni po uruchomieniu LHC z powodu awarii transformatora zasilającego jedną skraplarkę helu przebrano testy z wiązką w jonowodzie. W oczekiwaniu na usunięcie awarii energetycznej rozpoczęto testy zasilania magnesów, aby uzyskać pole dla protonów o energii 5 TeV w ostatnim nietestowanym sektorze 34. 19 września 2008 r. w trakcie testów energetycznych sektora 34 doszło do awarii w linii zasilającej dipole w strefie połączeń pomiędzy kwadrupolem a dipolem. W tym czasie w obwodach magnetycznych tego sektora płynął prąd o natężeniu 8,7 kA. Prawdopodobną przyczyną awarii było przejście

rezystywne w połączeniu tych dwu magnesów. W warunkach kriogenicznych oporność tego połączenia dla głównych kabli powinna wynosić 0,6 nanoohma. Przejście rezystywne jest jedną z hipotez, gdyż uwolniona energia około 700 MJ stopiła w łuku elektrycznym strefę połączenia. Parujący metal zniszczył mieszek kompensacyjny jonowodu zanieczyszczając na znacznej długości jonowód. Ponadto, energia uwolniona w strefie awarii spowodowała gwałtowne parowanie helu. Helowa fala uderzeniowa biegnąca w próżni izolacyjnej kriostatów strąciła z podpór magnes kwarupolowy znajdujący się w odległości 200 metrów od miejsca awarii. Magnes ten posiadał płytę zaporową sektorującą próżnię izolacyjną. Zniszczeniu uległa izolacja cieplna w obszarze próżni izolacyjnej. 39 dipoli i 14 kwadrupoli wymagało naprawy na powierzchni.

Wydarzenie to było przyczyną poszukiwań, w sektorach oziębionych, miejsc o podwyższonej oporności przejścia pomiędzy magnesami oraz opracowania metod wykrywania miejsc o rosnącym w czasie wydzielaniu ciepła.

W dniu 30 kwietnia 2009 r. zwieziono ostatni z naprawianych magnesów ponownie do tunelu LHC.

11. Plany ponownego uruchomienia LHC

Ponowny rozruch akceleratora przewidywany jest we wrześniu 2009 r. Początkowe testy będą prowadzone tylko z pilotowym pęczkiem protonów. Po ich zakończeniu uruchomione zostaną eksperymenty fizyczne na 43 pęczkach protonów, natomiast pod koniec roku ich liczba zwiększy się do 156. Energia zderzających się protonów będzie wynosić 5 TeV.

W pierwszej połowie roku 2010 zostaną zamontowane zawory upustowe na wszystkich dipolach. Zawory te 40-krotnie zwiększą przepustowość helu w przypadku jego gwałtownego parowania. Po tej modyfikacji akcelerator może pracować przy pełnej przewidywanej energii protonów 7 TeV.

12. Udział Polaków w budowie akceleratora

Struktura organizacyjna CERN-u sprawia, że wiele oddziałów tej instytucji zajmowało się tematyką LHC poprzez afiliację do tego projektu. Polacy nie będący pracownikami etatowymi CERN-u pracowali dla tego projektu poprzez firmy zakontraktowane do konkretnych zleceń, np. instalacje gazowe eksperymentu CMS – ZEW Wrocław, instalacje energetyczne w tunelu – Elektromontaż Południe. Inną formą uczestnictwa w budowie LHC były umowy o współpracy CERN – polskie instytucje naukowe i uczelnie. Instytut Fizyki Jądrowej PAN z Krakowa delegował średnio około 40 inżynierów fizyków i techników do prac przy kontroli jakości odbudowywanej linii zasilania kriogenicznego, połączeń magnesów nadprzewodzących, testów elektrycznych obwodów głównych i magnesów sterujących. Politechnika Wrocławska utrzymuje wieloletnią współpracę dotyczącą układów skraplarek helowych. Politechnika Krakowska współpracuje w projektowaniu systemów mechanicznych i badaniu niezawodności systemów w LHC. Akademia Górniczo-Hutnicza z Krakowa delegowała 40 osób do grup: uruchamiających sterowanie podziemnych instalacji kriogenicznych, budujących system zabezpieczeń magnesów na okoliczność przejścia rezystywnego i systemów ekstrakcji energii pola magnetycznego, dokumentacji technicznej układów chłodzenia i wentylacji, koordynacji technicznej prac w tunelu.

Alain Poicet, Jean-Philip Tock, Błażej Skoczeń, Andrzej Siemko, Tadeusz Kurtyka to osoby, którym autor dziękuje za wieloletnią współpracę przy budowie LHC. Dziekanom WFiIS AGH: Kazimierzowi Jeleniowi, Zbigniewowi Kąkolowi oraz Wojciechowi Łużnemu autor dziękuje za wspieranie idei współpracy AGH–CERN przy budowie LHC, a wicedyrektorowi IFJ PAN w Krakowie Grzegorzowi Polokowi za inspirację projektem i podpowiedzi organizacyjne.

W pracy wykorzystano zbiory archiwalne zdjęć z budowy LHC dostępne z witryny biblioteki CERN-u.

Wielki Program Wszechświata

Michał Heller

Centrum Kopernika Badań Interdyscyplinarnych, Kraków

Great Program of the Universe

Abstract: Has the Universe to obey our „macroscopic imaginations”? We should rely more on well founded physical theories than on our intuitive imaginations. All physical theories have mathematical structures, and these structures are usually dynamic structures. They do not only describe the physical world; but also should be thought of as our approximations to a Great Program being executed by the Universe.

Czy Wszechświat musi się podporządkowywać naszym „wyobrażeniom makroskopowym”? Do sprawdzonych teorii fizycznych winniśmy mieć większe zaufanie niż do naszych wyobrażeń. A wszystkie teorie fizyczne mają strukturę matematyczną i zwykle jest to struktura dynamiczna. Struktury matematyczne nie tylko opisują świat fizyczny, lecz raczej należy myśleć o nich jako o Wielkim Programie, który nasz Wszechświat wykonuje.

Ostatnio, nakładem wydawnictwa Iskry ukazała się książka *Zapis świata. Traktat metafizyczny* (Warszawa 2009). Autor, Piotr Wierzbicki – z wykształcenia polonista, pisarz, publicysta, parał się także polityką – podjął próbę opowiedzenia w języku literackim o dokonaniach współczesnej nauki i osnutych wokół nich swoich przemyśleniach metafizycznych. Rzecz bardzo charakterystyczna – fizycy bardziej pasjonują się szczegółami prowadzonych przez siebie eksperymentów lub własnościami badanych równań, natomiast humaniści, jeżeli w ogóle zainteresują się naukami ścisłymi, to natychmiast sięgają do ich prawdziwych lub rzekomych konsekwencji filozoficznych. Przyjrzyjmy się nieco bliżej próbie takich spekulacji, zaczerpniętej z książki Wierzbickiego.

Wierzbicki, rozważając problem czasu, chce zwrócić uwagę czytelnika na to, że niewykluczona jest sytuacja, w której czas nie istnieje i może się dopiero narodzić. W tym celu każe sobie wyobrazić „inny świat, stanowiący przez nieskończenie wielkich rozmiarów masywną bryłę nie złożoną z żadnych cząstek, tylko szczelnie wypełnioną materią tak gęsto utkaną, że nie ma tam miejsca ani na najmniejsze nawet oczko wolnej przestrzeni, ani na najmniejsze poruszenie”. Czy w takim wszechświecie może płynąć czas? Oczywiście nie! Bo „czas umie płynąć tylko między: między zdarzeniem a zdarzeniem”, a w takim wszechświecie nic się nie zdarza. Spróbujmy więc zaradzić sytuacji i wprowadzić do naszego hipotetycznego wszechświata zmianę, ale zmianę najdrobniejszą, która by jednak wystarczała, „by w tym wszechświecie ruszył czas”. Otrzy-

mamy w ten sposób „pigułkę czasu, elementarną, niepodzielną »drobinę«, która, zaszczepiona w tej martwej krainie, tchnie w nią życie”.

Doskonały chwyt dydaktyczny. Ale... Pomińmy pytanie, czy taka jednolita, nieskończona, doskonale ciągła masa mogłaby w ogóle istnieć. W naszym makroskopowym świecie masywne ciała istnieją tylko dlatego, że „u ich podłoża” działają prawa fizyki kwantowej. Pomińmy to pytanie. Popularyzacja ma swoje prawa i dla celów dydaktycznych można sobie coś takiego wyobrazić. Ale co robi Wierzbicki i wielu innych popularyzatorów (a także pisarzy *science fiction*)? Oni manipulują obrazem świata, ale nie są w stanie wyzwolić się z „ontologii” zakładanej przez ogląd makroskopowy, modyfikują tylko tę ontologię i przekomponowują stosownie do swoich potrzeb, lecz jest to ciągle ten sam nasz makroskopowy świat, jedynie „poukładany” w inny, może trochę bardziej dziwaczny sposób. Wcale ich za to nie ganimy. Nawet bardzo wyrafinowany fizyk-teoretyk, gdy myśli o swoich funkcjach falowych i macierzach gęstości, najczęściej przykłada do nich wyobrażenia rodem ze swego makroskopowego otoczenia. Może tylko jego wyobrażenia jest bardziej wytrenowana w manipulowaniu materiałem wyobrażeniowym. I na pewno swoje imaginacyjne obrazy traktuje on z większym dystansem i z bardziej ironicznym przymrużeniem oka.

Czy jednak świat rzeczywisty musi się podporządkowywać ontologii naszych wyobrażeń? Do sprawdzonych teorii fizycznych musimy jednak mieć większe zaufanie

niż do gry naszych wyobrażeń. A wszystkie teorie fizyczne mają strukturę matematyczną. Wyraz „struktura” wydaje się tu najbardziej trafny ze wszystkich, jakie przychodzą na myśl. Trzeba tylko do naszego potocznego rozumienia struktury dodać element zmienności dynamicznej. Katedra gotycka ma piękną strukturę architektoniczną, ale jest to struktura statyczna. Równania Lagrange’a lub Hamiltona także opisują pewną strukturę, ale ta struktura zmienia się w czasie, jest dynamiczna. Właśnie „struktura” a nie „obraz”, bo obraz jest wypełniony malarskimi detalami, podczas gdy struktura pozostaje bardziej abstrakcyjna (co nie znaczy mniej konkretna) i bardziej otwarta na oddziaływanie z innymi strukturami.

Szkielet matematycznej struktury nie zawsze daje się adekwatnie wypełnić wyobrażeniowym tworzywem. Spróbujmy na przykład opowiedzieć językiem, jakim posługujemy się na co dzień (wyrosłym przecież z doświadczenia potocznego), strukturę nieskończeniowymiarowej przestrzeni Hilberta i działające na niej operatory hermitowskie. Uwzględnijmy przy tym obraz Schrödingera, Heisenberga i Diraca-Tomonagi. I jeszcze to, że cała ta struktura jest izomorficzna z pewną abstrakcyjną C^* -algebrą. Wyobraźnia może kreślić tylko jakieś zamglone kontury, które niekiedy tylko pomagają – ale częściej przeszkadzają – w uchwyceniu właściwych związków strukturalnych. Dziwne i zaskakujące jest nie to, że sobie tego wszystkiego nie potrafimy wyobrazić, lecz to, że – dzięki posłuszeństwu wynikom pomiarów i dyscyplinie dedukcji matematycznych – potrafiliśmy tę formalną strukturę zrekonstruować. To właśnie, a nie co innego, dało nam wgląd do wewnętrznej architektury świata.

Często mówi się, że matematyka opisuje świat. Niewątpliwie tak jest, ale to tylko jej bardzo „zewnątrzna” funkcja. Rozpatrzmy przykład. Oto wykonujemy serię pomiarów jakiejś obserwabli kwantowej. Nie jesteśmy specjalnie zaskoczeni, że statystyka wyników jest zgodna z tym, co przewidują rachunki odwołujące się do spektrum badanej obserwabli. Takie jest przecież prawo fizyki: wartości własne obserwabli odpowiadają przewidywanym wynikom pomiarów. Operator hermitowski, działający na przestrzeni Hilberta, wraz ze swoim spektrum, poprawnie opisuje to, co dzieje się podczas odpowiednio długiej serii pomiarów pewnej własności układu kwantowego. Czy jest to naprawdę tylko opis? Oto istnieje jakiś układ kwantowy, który podczas wykonywanej na nim serii pomiarów, zachowuje się zgodnie z prawami fizyki i tak się szczęśliwie składa, że pewna struktura matematyczna „pasuje” do tego, co się tam dzieje, a my wykorzystujemy tę strukturę do opisywania całego procesu. Czy tak jest w istocie? Codzienna praktyka fizyków podpowiada, że jest z gruntu inaczej. Powiedziałbym – jest wręcz odwrotnie. Układ kwantowy wykonuje pewien program, który – jak wszystko na to wskazuje – jest częścią, lub aspektem, większego, kosmicznego programu i dzięki niezwykle zjawisku, jakim jest postęp nauki, udało nam się zrekonstruować fragment tego programu. Okazuje się, że teoria

operatorów hermitowskich na przestrzeniach Hilberta poprawnie rekonstruuje ten fragment kosmicznego programu, który odpowiada za oddziaływania układów kwantowych z większymi układami, jakie nazywamy aparatami pomiarowymi. A więc, ujmując to krótko, struktury matematyczne nie tylko opisują funkcjonowanie Wszechświata, ale je „zadają”.

Przykład z operatorami działającymi na przestrzeni Hilberta wybrałem celowo. Jest to bowiem struktura, która łączy wewnętrzne funkcjonowanie świata kwantów z wynikami wykonywanych przez nas eksperymentów: wartości własne operatora hermitowskiego określają możliwe wyniki pomiarów. Zadaje to kłam powszechnie uznawanej dychotomii: z jednej strony teoria, z drugiej eksperymenty. Teoria jest dobra, jeżeli dobrze przewiduje wyniki eksperymentów. Eksperymenty są ostateczną instancją, która weryfikuje teorię. Oczywiście, jest to prawda, ale należy ją swoiście rozumieć. Eksperymenty, razem z aparatem pomiarowym, są w pewnym sensie częścią teorii. Przecież operator hermitowski reprezentuje („zadaje”) i wielkość, którą mamy zmierzyć, i wyniki pomiarów (wartości własne), a operator hermitowski poza strukturą matematyczną, której jest częścią, w ogóle nie miałby sensu. Co więcej, nowoczesnego urządzenia eksperymentalnego (pomyślmy o akceleratorze w CERN-ie pod Genewą) nie dałoby się ani zaprojektować, ani zbudować bez istotnego udziału zmatematyzowanych teorii. Dzięki temu splątaniu teorii i doświadczenia współczesna fizyka bynajmniej nie jest mniej teoretyczna niż fizyka Galileusza, który swoje doświadczenia wykonywał za pomocą spadających kul, równi pochyłych i czasu odmierzanego własnym pulsem. Dziś doświadczenie nie jest czymś peryferyjnym w stosunku do teorii, lecz wchodzi do samego jej jądra. Zresztą i spadające kule Galileusza także wykonywały tylko instrukcje zapisane w Wielkim Programie Wszechświata. Tyle, że ludzkość uczyła się dopiero odcyfrowywać pierwsze sylaby strukturalne tego Wielkiego Programu.

Fizyka nowożytna stworzyła nową jakość w rozumieniu Wszechświata. Dziś nie można już rozumieć Wszechświata bez rozumienia fizyki.

★ ★ ★

Dzięki Nagrodzie Templetona zostało założone w Krakowie Centrum Kopernika Badań Interdyscyplinarnych. Jest to międzyuczelniana instytucja afiliowana do Uniwersytetu Jagiellońskiego i Papieskiej Akademii Teologicznej w Krakowie. Ma ona cele badawcze, edukacyjne i popularyzacyjne. Jej specyfiką jest studiowanie zagadnień powstających na styku nauk i filozofii, a także rozpatrywanie ich w kontekście ogólnokulturowym. Dzięki gościnności redakcji *Postępów Fizyki* rozpoczynamy publikowanie na łamach tego czasopisma cyklu esejów związanych z problematyką Centrum Kopernika, a dotyczących fizyki, tej najbardziej podstawowej ze wszystkich nauk empirycznych.

W stulecie Nagrody Nobla z fizyki za rok 1908 przyznanej Gabrielowi Lippmannowi za fotografię barwną*

Pierre Ranson, Robert Ouillon, Jean-Paul Pinan-Lucarré

Physique des milieux denses, IMPMC, Université P&M Curie, Paryż

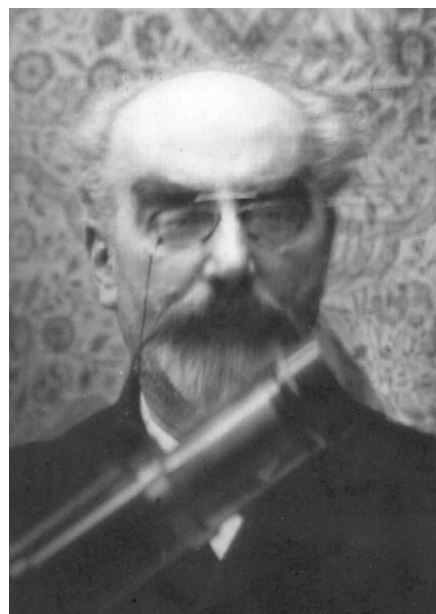
A centenary, G. Lippmann, Nobel prize of Physics 1908 for colour photography

Abstract: One century ago, on December 10 (1908), the Royal Swedish Academy of Sciences awarded the Nobel prize for physics to Professor Gabriel Lippmann of Sorbonne for his interference-based method to reproduce colours by photography.

Podstawy fotografii barwnej

Od niepamiętnych czasów człowiek chciał różnymi sposobami zachować obrazy otaczającego świata. Światło przychodzące od obserwowanych przedmiotów zostawia na siatkówce ich obraz, który następnie jest analizowany przez mózg. Stosownie do tego wyobrażenia człowiek próbował tworzyć wierne kopie tych obrazów, posługując się rzeźbą, rysunkiem, malarstwem, a od stu lat również sztuką fotograficzną. Oryginalne obrazy są barwne, ale receptory światłoczułe siatkówki (czopki) nie są zbyt dobrymi analizatorami częstości fal światła. Znane są 3 rodzaje tych receptorów, których maksimum czułości przypada w obszarze niebieskim (ok. 425 nm), zielonym (ok. 535 nm) i czerwonym (ok. 570 nm). Tak więc w fotografii wystarczyłoby w pierwszym przybliżeniu użyć płyt fotograficznych o właściwościach podobnych do siatkówki oka. W zwykłej fotografii barwnej osiągnęto to do niedawna, stosując głównie trójbarwną metodę subtraktywną, a obecnie, gdy dominuje technika kamer cyfrowych, używa się metod addytywnych. Celem tych metod jest naśladowanie odbierania kolorów przez oko ludzkie.

Metoda interferencyjna Gabriela Lippmanna [1] podąża zupełnie inną drogą. Jest to proces czysto fizyczny oparty na własnościach fal stojących, co z zasady daje wierne i niezmiennie powtórzenia prawdziwego składu widmowego promieniowania. W sztuce fotografii barwnej proces Lippmanna, mimo swojej elegancji, nie był jakimś tworzącym nową epokę wynalazkiem, był jednak przekonującą ilustracją interferencji harmonicznych fal płaskich w optyce, a sto lat później, gdy stały się dostępne źródła promieniowania laserowego, umożliwił nowoczesne zastosowanie efektu Lippmanna-Bragga.



Rys. 1. Obraz czarno-biały (płyta „autostereoskopowa” 13,6 × 18 cm) przedstawiający profesora Lippmanna w jego gabinecie, za teleskopem. Na zdjęciu widać pionowe prążki. Umieszczenie ich na powierzchni szklanej płyty pozwala obserwatorowi patrzącemu z odpowiedniej odległości widzieć przedmiot trójwymiarowo. Tę metodę opracował ostatecznie Eugène Estanave, który od 1908 r. pracował w laboratorium Lippmanna na Sorbonie.

Kim był Gabriel Lippmann?

Gabriel Lippmann (rys. 1), syn Francuzów zamieszkałych w Holerich w Luksemburgu, urodził się 16 sierpnia

*Artykuł opublikowany w *Europhysics News* 39, nr 6 (2008) został przetłumaczony za zgodą Autorów i Wydawcy. [Translated with permission, © 2008 European Physical Society and EDP Sciences]

1845 r. Po studiach w paryskiej Ecole Normale Supérieure zdecydował się obrać karierę uniwersytecką. Otrzymał na Sorbonie w 1883 r. katedrę fizyki matematycznej, a w 1886 r. – fizyki. Od 1886 r. aż do śmierci w 1921 r. był dyrektorem słynnego laboratorium L.R.P.S. (Laboratoire des Recherches Physiques de la Sorbonne), gdzie opracował swoją metodę interferencyjnej fotografii barwnej. Zdjęcia Lippmanna, które ilustrują ten artykuł, pochodzą z tego laboratorium, które zakończyło działalność w 1968 r. wraz z utworzeniem Université Pierre et Marie Curie (Paris VI).

Wśród najbardziej znanych jego studentów i współpracowników na Sorbonie byli laureaci Nagrody Nobla – Pierre Curie (1903 z fizyki) i Maria Skłodowska-Curie (1903 z fizyki i 1911 z chemii).

1810–1848: najwcześniejsze metody rejestracji barwnej [2]

W 1666 r. Izaak Newton wykazał, że światło białe utworzone jest z ciągłego zbioru składowych monochromatycznych, których barwy – od niebieskiej do czerwonej – można uzyskać przez rozszczepienie w szklanym pryzmacie. Pierwsze próby uzyskania i zapisania tych „jednorodnych” lub „czystych” barw na płaskiej trwałej powierzchni datują się na początek XIX w., tzn. dwadzieścia lat przed odkryciem fotografii przez jej pionierów Nicéphore’a Niépce’a (ok. 1826 r.) i Louisa Jacquesa Mandé Daguerre’a (1836 r.). Możemy tu wspomnieć, że Johann T. Seebeck, próbując zapisać różne kolory widma słonecznego, użył (1810) papieru nasyczonego roztworem wodnym chlorku srebra. Podobne badania prowadzili także sir John Herschel, Niépce de Saint Victor i Poitevin, ale ostatecznie najlepsze wyniki osiągnął Edmond Becquerel, który użył uczulonych powierzchniowo srebrem płyt typu Daguerre’a. Wykazał, że dobrze wypolerowana płyta srebrna pokryta cienką warstwą chlorku srebra zabarwia się pod wpływem światła, oddając przy tym barwy takie jak ma oryginał. Potem uzyskiwał słynne heliochromy (1838–48) Becquerela. Niektóre z nich przetrwały do dziś. Jednak nie znaleziono metody, która by utrzymała barwy na płycie, a co więcej, blakły one na świetle. Jednak metoda ta przez 20 lat nie wzbudzała zainteresowania, choć był to prawdopodobnie pierwszy sukces fotografii interferencyjnej wykorzystującej powstawanie fal stojących przez odbicie światła od warstwy srebra! To w każdym razie jest wyjaśnieniem pochodzenia barwnych obrazów Becquerela, jakie podał Wilhelm Zenker w 1868 r.

1861–1895: pierwsze próby uzyskania trójbarwnych obrazów fotograficznych [2,3]

Gdy James Clerk Maxwell ustalił w 1859 r., że wszystkie kolory, jakie widzi oko ludzkie, można wytworzyć, mieszając w odpowiednich proporcjach trzy zasadnicze barwy – czerwoną, zieloną i niebieską (RGB – red, green, blue), zainteresowanie fotografią kolorową nabrało nowego rozmachu. W 1861 r. Maxwell zrobił pierwsze

trwałe barwne fotografie stosując metodę addytywną. Zrobił 3 oddzielne czarno-białe fotografie przez filtry czerwony, zielony i niebieski (po wywołaniu płyt, negatywy zostały zamienione na pozytywy i każdy z nich został zrzucony przez taki filtr, jaki był użyty do jego rejestracji, a obrazy starannie nałożono na siebie na ekranie). Kolory można uzyskać również ze światła białego metodą subtraktywną przy użyciu odpowiedniego barwnika. W takim procesie zasadniczymi barwami są: cyjan, karmazyn, żółty i czarny (CMYK – cyan, magenta, yellow, black). Charles Cros i Louis Ducos du Hauron pierwsi przeprowadzili badania teoretyczne (1869 r.) i doświadczalne (1877 r.) uzyskania reprodukcji kolorów metodą subtraktywną.

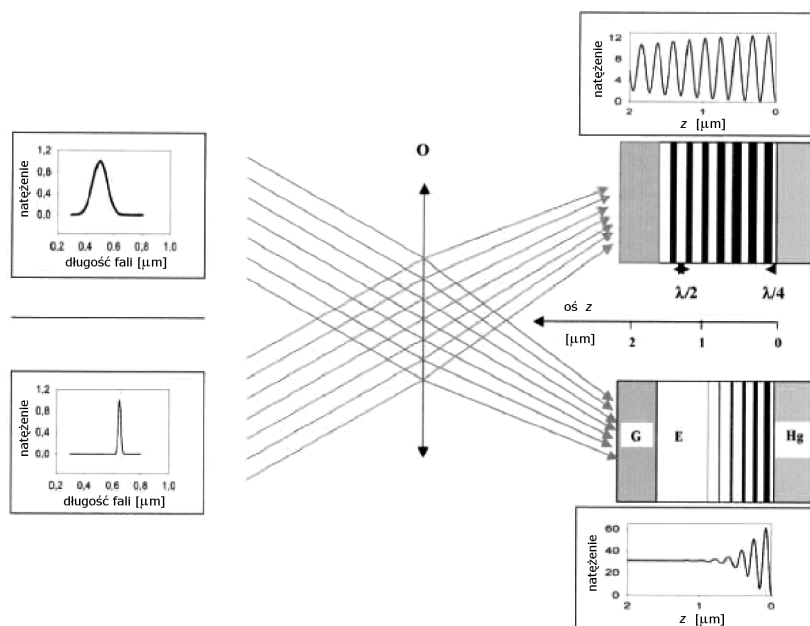
1891: metoda interferencyjna Gabriela Lippmanna – reprodukcja barw metodą fali stojącej [1–6]

Wytwarzanie kolorów we wszystkich metodach stosowanych w końcu XIX w. wymagało użycia barwników wprowadzanych do filtrów lub warstw światłoczułych w trakcie procesu fotograficznego. Jednak reakcja tych materiałów na światło i zachowanie barwy z upływem czasu nie podlegały dostatecznej kontroli. Ideą Lippmanna było wykorzystanie jedynie właściwości światła do wytworzenia barw. Uwagę skupił przede wszystkim na ścisłości związku między widmem częstości padającego światła a jego barwą. Wyzwaniem było zapisanie gęsto upakowanych fal harmoniczych zawierających się w naturalnej barwie jako trwałego znaku w warstwie światłoczułej, a potem reprodukcję z tego znaku odpowiedniego koloru (jest to analogia do fal akustycznych w fonografie – urządzeniu zapisującym i odtwarzającym dźwięki, którego teoretyczne wyobrażenie przedstawił Charles Cros, a w praktyce urzeczywistnił w 1877 r. Thomas Alva Edison).

Ten trwały znak może zostać stworzony w warstwie światłoczułej przez falę stojącą. Warstwa światłoczuła to żelatyna z rozpuszczonym w niej halogenkiem srebra. Falę stojącą uzyskuje się przez odbicie w zwierciadle i nałożenie promieni odbitych na padające (rys. 2). Genialny pomysł zastosowania zwierciadła rtęciowego zapewnił bliski kontakt między materiałem odbijającym i warstwą foto-czułą. Po procesach wywoływania, utrwalania i suszenia w płaszczyznach strzałek fali stojącej powstawały płaskie równoległe zwierciadła zrobione z cząstek srebra.

Dla padającego światła monochromatycznego o długości fali λ_0 (w próżni) odległość tych płaszczyzn d jest równa połowie długości fali w danej warstwie ($d = \lambda_0/2n$, gdzie n jest współczynnikiem załamania w warstwie światłoczułej). Ten obszar może rozciągać się na wiele mikrometrów.

Gdy światło nie jest monochromatyczne, rozkład natężeń w fali stojącej w małej odległości od zwierciadła rtęciowego staje się jednorodny skutkiem nakładania się poszczególnych „monochromatycznych” fal stojących. W wyniku tego rozkład cząstek srebra również staje się jednorodny.



Rys. 2. Proces Lippmanna fotografii barwnej. W fotografii barwnej każdy barwny przedmiot rozpatruje się jako bardzo wiele źródeł światła, z których każde ma widmo ciągłe w pewnym przedziale widmowym $\Delta\lambda$. W procesie Lippmanna promieniowanie każdego źródła jest ogniskowane przez soczewkę kamery (O), jak zwykle na płycie fotograficznej (G – podkład szklany, E – warstwa światłoczuła), ale ponadto promieniowanie jest odbijane przez zwierciadło rtęciowe (Hg) umieszczone na tylnej ścianie warstwy światłoczułej. W wyniku nakładania się promieniowania padającego i odbitego nad powierzchnią odbijającą powstaje fala stojąca. Po naświetleniu i wywołaniu tworzy się w tej warstwie układ płaskich metalowych zwierciadeł. Na rysunku pokazano dwie różne fale stojące: pierwsza pochodzi od czerwonego, wąskiego pasma widmowego ($\lambda = 0,65 \mu\text{m}$, $\Delta\lambda \approx 0,02 \mu\text{m}$), a druga – od szerokiego pasma niebieskiego ($\lambda = 0,5 \mu\text{m}$, $\Delta\lambda \approx 0,01 \mu\text{m}$). W pierwszym przypadku układ płytek wypełnia całą głębokość warstwy, a w drugim rozciąga się tylko na kilkaset nanometrów (E). Wytworzony układ metalowych płytek oświetlany światłem naturalnym (białym) wybiera z niego i odbija ku obserwatorowi tylko promieniowanie o składowych monochromatycznych, tych które wykorzystano w procesie zapisu.

Niemniej w każdym przypadku, jakkolwiek byłby rozkład przestrzenny, metaliczne płytki działają jak reflektor selektywny względem częstotliwości światła. Gdy ze źródła pada światło białe (o widmie ciągłym), układ płytek odbija ze znacznym natężeniem tylko te składowe monochromatyczne, które były obecne w procesie rejestracji (jest to konsekwencją interferencji konstruktywnej, którą również można analizować w kategoriach efektu Bragga). Wobec tego, metoda ta jest teoretycznie odpowiednia do reprodukcji naturalnych barw.

W języku nowoczesnej optyki możemy powiedzieć, że proces Lippmanna jest równoznaczny z podwójną transformacją Fouriera. W stadium rejestracji transformata Fouriera widma światła zapisuje się w profilu natężeń fali stojącej. W stadium reprodukcji zachodzi inna transformacja Fouriera przywracająca początkowy skład widmowy światła.

Praktyczne przeprowadzenie procesu Lippmanna okazuje się trudne. Warstwa światłoczuła musi spełniać wiele wymagań:

- musi być bardzo dobrze przezroczysta, gdyż przychodzące światło, zanim padnie na zwierciadło, przechodzi przez wszystkie warstwy (zwykle grubości kilku mikrometrów),

- musi być „bezziarnista”, tzn. ziarna halogenku srebra muszą mieć bardzo mały rozmiar $s \leq 10 \text{ nm}$. To jest przede wszystkim konieczne, aby uzyskać dużą rozdzielczość przestrzenną fali stojącej (dla światła niebieskiego o $\lambda_0 = 450 \text{ nm}$ odległość między prążkami interferencyjnymi jest $i = \lambda_0/2n = 150 \text{ nm}$), a także aby ograniczyć do minimum szumy powstałe przez dyfuzję lub rozproszenie.
- musi być izochromatyczna (uzyskuje się to przez dodanie odpowiednich substancji uczulających), aby dokładnie reprodukcja rozkładu amplitud w szerokich widmach częstotliwości, jakie dają naturalnie zabarwione przedmioty. Ten ostatni punkt jest kluczowym wymaganiem procesu Lippmanna.

Reprodukcja barwnych obrazów powstaje w procesie Lippmanna w wyniku odbicia światła białego na układzie metalicznych płytek srebrnych i jest funkcją sumy amplitud lokalnych odbić r wywołanych wewnątrz płytek w czasie procesu rejestracji przez każde drganie monochromatyczne. Lippmann definiuje izochromatyzm jako wartość wielkości r , która zależy tylko od amplitudy drgań, a nie zależy od ich częstotliwości. To stwierdzenie Lippmann szczegółowo wyjaśnił w teoretycznym artykule [7] opublikowanym w roku 1894. Po kilku latach technicznego ulepszania jakości emulsji (co było prowadzone z wynalazcami

płyty autochromowych Augustem i Louisem Lumière'ami) uzyskano wspaniałe, wierne obrazy. Sto lat później mają nadal żywe kolory (patrz II i III strona okładki), a niektóre z nich przedstawiające dobrze znane miejsca mają też wartość archiwalną.

Dlaczego więc proces Lippmanna nigdy nie zyskał wielkiej popularności?

W stadium rejestracji [8] proces Lippmanna miał pewne wady. Płyty, które wymagają odbijającego zwierciadła, nie są bardzo praktyczne. Ze względu na zastosowanie w emulsji bardzo drobnoziarnistego halogenku srebra czas ekspozycji (minuta w słońcu) jest zbyt długi, aby robić fotografie portretowe, a ponadto nie jest możliwe robienie kopii zdjęć.

W stadium oglądania obrazu płyta musi być oświetlona i oglądana pod właściwym kątem, aby oddawała prawdziwe kolory, a ich projekcja jest trudnym zadaniem. Ponadto, po roku 1910 wygodniejszy addytywny proces autochromiczny braci Lumière zdominował na prawie 20 lat fotografię barwną.

Jeśli chodzi o Lippmanna, to poświęcał czas poszukiwaniu urządzeń optycznych, które w fotografii umożliwiłyby uzyskiwanie obrazów trójwymiarowych. Zaproponował 3 marca 1908 r. umieszczenie na powierzchni obrazu układu małych wypukłych soczewek (układ „oka muchy”) zamiast nieprzezroczystych pasków jak w metodzie Estanave'a (patrz rys. 1). Jego pomysł został przedstawiony francuskiej Akademii Nauk pod tytułem „La Photographie Integrale”. Ostatecznie Estanave zrealizował w 1925 r. pomysły Lippmanna.

Odrodzenie idei Lippmanna

Wraz z rozwojem metody obrazowania stosującej rejestrację fal stojących w grubej (ok. 15 μm) warstwie światłoczułej (objętościowa holografia Deniszyuka [9], holografia barwna [10], przechowywanie danych optycznych w mikrowłóknach [11], ...) idee Lippmanna ożyły w 50 lat po przyznaniu mu Nagrody Nobla. W tych metodach dwie interferujące fale ze wspólnego źródła laserowego osiągają warstwę holograficzną z kierunków przeciwnych i tworzą hologram odbiciowy. W odwzorowaniu trójwymiarowym jedna z tych dwóch fal pada prostopadle na warstwę (wiązka referencyjna), przechodzi przez warstwę i jest odbijana przez badany przedmiot (wiązka przedmiotowa). Przedmiot zastępuje się płaskim zwierciadłem, aby stworzyć odbiciową siatkę pojedynczej linii.

Przeprowadzono dokładne badania „modelowania lippmannowskiego procesu barw” [12,13]. Chodziło o ulepszenie wysokorozdzielczej emulsji typu Lippmanna używanej w holografii. Ostatnio są w trakcie prace nad ulepszeniem warstwy światłoczułej o dużej przezroczystości i wzrostem czułości metody [14,15].

Bardzo prostym i praktycznym osiągnięciem są nowoczesne filtry holograficzne – karbowane filtry rowkowe (notch filters), które wytwarza się, rejestrując w ośrodku

światłoczułym wzór interferencyjny utworzony przez dwie wzajemnie spójne wiązki laserowe. Zastosowanie ich do eliminacji silnego sygnału laserowego i linii rayleighowskiej spowodowało uproszczenie i polepszenie działania zwykłej aparatury ramanowskiej [16].

Podziękowania

Zdjęcia płyt Lippmanna, które ilustrują ten artykuł (rys. 1 oraz II i III strona okładki), zrobił Alain Jeanne-Michaud, technik w IMPMC (Université Pierre et Marie Curie). Użył on kamery cyfrowej i oświetlenia dziennego. Portret na płycie Lippmanna pokazany jest z ramką ochronną przesłoniętą matowym czarnym papierem. Na innych zdjęciach ramek nie pokazano.

Tłumaczyła Barbara Wojtowicz
Warszawa

Literatura

- [1] G. Lippmann, *Comptes Rendus de l'Academie des Sciences* **112**, 274 (1891).
- [2] P. Connes, *J. Opt.* **18**, 147 (1987).
- [3] W. Alschuler, w: *31st Applied Imagery Pattern Recognition Workshop*, red. D.H. Schaefer, s. 3.
- [4] J.M. Fournier, *J. Opt.* **22**, 259 (1991).
- [5] J.M. Fournier, *J. Imaging Sci. Techn.* **38**, 507 (1994).
- [6] R. Ouillon, J.-P. Pinan-Lucarré, P. Ranson, *B.U.P.* **100**, 459 (2006).
- [7] G. Lippmann, *J. de Phys.* **3**, 97 (1894).
- [8] P. Hariharan, *J. Mod. Opt.* **45**, 1759 (1998).
- [9] Y.N. Deniszyuk, *Fundamentals of Holography* (State Optical S.I. Vavilov Institute, Leningrad 1978), s. 26.
- [10] H.I. Bjelkhagen, *Holography, Art and Design, 2002 London Conference*; www.holography.co.uk/conference/bjelkhagen/CON2HIB1.pdf.
- [11] A. Labeyrie, J.P. Huignard, B. Loiseaux, *Opt. Lett.* **23**, 301 (1998).
- [12] H. Nareid, H. Pedersen, *J. Opt. Soc. Am.* **8**, 257 (1989).
- [13] M. Jäger, H. Bjelkhagen, M. Turner, w: *Practical Holography XVIII: Materials and Applications*, red. T.H. Jeong, H. Bjelkhagen (*Proc. SPIE* **5290**, 306 (2004)).
- [14] Y. Gentet, P. Gentet, w: *Holography 2000*, red. T.H. Jeong, W.K. Sobotka (*Proc. SPIE* **4149**, 56 (2000)).
- [15] J. Belloni, *C. R. Physique* **3**, 1 (2002).
- [16] J. Barbillar, B. Roussel, E. Da Silva, *J. Raman Spectrosc.* **30**, 745 (1999).

PIERRE RANSON, ROBERT OUILLON I JEAN-PAUL PINAN-LUCARRÉ specjalizowali się w badaniach rozproszenia ramanowskiego, a głównym przedmiotem ich studiów była relaksacja fononowa w kryształach molekularnych. Swoją karierę uniwersytecką rozpoczęli w 1960 r. w LRPS na Sorbonie, w laboratorium, którego szefem przed siedemdziesięciu laty był profesor Lippmann. Ich promotor R. Dupeyrat poświęcił wiele czasu zachowaniu zbioru płyt Lippmanna. Stanowią one teraz kolekcję zbiorów UPMC. Pierre Ranson, Robert Ouillon i Jean-Paul Pinan-Lucarré, obecnie na emeryturze, nadal zajmują się tym zbiorem.

■ Grażyna Chełkowska

Grażyna Chełkowska (z domu Wnętrzak) urodziła się w 1954 r. w Żywcu. Doktorat z zakresu nauk fizycznych pod kierunkiem prof. Augusta Chełkowskiego uzyskała na Wydziale Matematyki, Fizyki i Chemii Uniwersytetu Śląskiego w Katowicach w 1982 r. Zatrudniona w Instytucie Fizyki UŚI realizowała swoje naukowe zainteresowania w zakresie struktury elektronowej oraz własności magnetycznych i transportowych związków międzymetalicznych na bazie ziem rzadkich.

Stopień doktora habilitowanego przyznała jej Rada Wydziału Matematyki, Fizyki i Chemii Uniwersytetu Śląskiego w Katowicach, na podstawie rozprawy habilitacyjnej „Magnetic and electronic structure properties of gadolinium based Laves phase intermetallic compounds”. Zatwierdzenie pracy habilitacyjnej nastąpiło w 1996 r. Tytuł naukowy profesora otrzymała 11 lutego 2009 r.



Opublikowała ponad osiemdziesiąt artykułów w czasopiśmie o zasięgu międzynarodowym, uczestniczyła w kilkudziesięciu konferencjach w kraju lub zagranicą. Jest autorką rozdziału „Compounds of rare earth elements and 4d or 5d elements” w jednym z tomów *Landolt-Börnstein, Magnetic Properties of Metals*, wydanego przez wydawnictwo Springer-Verlag w 2004 r.

Współpracuje z zagranicznymi ośrodkami m.in. z Instytutem Fizyki Uniwersytetu w Osnabrück, z Centre d'Elaboration de Materiaux et d'Etudes Structurales, CNRS w Tuluzie oraz z Forschungszentrum Jülich (Niemcy).

W minionej kadencji 2005–08 pełniła funkcję prodziekana do spraw kształcenia na Wydziale Mat.-Fiz.-Chem. UŚI, a od 1 września 2008 r. pełni tę funkcję w kolejnej kadencji.

Jest członkiem Polskiego Towarzystwa Fizycznego.

Prywatnie jest mamą Marysi (studentki psychologii) i Błażeja (licealisty). Lubi muzykę oraz podróże.

■ Andrzej Baran

Pochodzi z małej miejscowości Kaplonosy w powiecie włodawskim. Studiował fizykę na Uniwersytecie Marii Curie-Skłodowskiej w Lublinie. Po magisterium, we wrześniu 1971 r., podejmuje pracę w Zakładzie Agrofizyki Polskiej Akademii Nauk, a po dwóch miesiącach zostaje asystentem w Zakładzie Fizyki Teoretycznej UMCS. Idąc w ślady kolegów, podejmuje współpracę z prof. Adamem Sobieczewskim z ówczesnego Instytutu Badań Jądrowych. Pod jego kierunkiem pisze rozprawę doktorską „Dynamika jąder atomowych silnie zdeformowanych”.



Po doktoracie wyjeżdża do Monachium, gdzie na Uniwersytecie Technicznym pracuje z prof. Klaussem Dietrichem nad problemem gęstości poziomów jądrowych i przejść elektromagnetycznych w jądrach ziem rzadkich. W następnych latach współpracuje także z I.N. Michajłowem (Dubna), Jerzym Dudkiem i Johanem Bartelem (Strasbourg) i P. Quentinem (Bordeaux). W ostatnich latach intensywnie współpracuje z zespołem Witolda Nazarewicza z Oak Ridge National Laboratory. Tematem są stare problemy dynamiki rozszczepienia jąder atomowych w modelach Skyrme'a–Hartree'ego–Focka.

Poświęca wiele czasu dydaktyce i organizacji życia naukowego środowiska. Od roku 1994 organizuje coroczne Warsztaty Fizyki Jądrowej w Kazimierzu Dolnym nad Wisłą. Należą one do bardzo popularnych spotkań polskich i zagranicznych fizyków jądrowych. Redaguje również specjalny tom *International Journal of Modern Physics E* poświęcony Warsztatom. Wiele wysiłku włożył w organizację kierunku Informatyka na Wydziale Matematyki i Fizyki UMCS. Przez wiele lat pracował dodatkowo w WSZiA w Zamościu, prowadząc kilkadziesiąt prac licencjackich i kilkanaście magisterskich z informatyki stosowanej.

Tytuł naukowy otrzymał 3 kwietnia 2009 r.

Jego pasje to literatura, jazz, go, wędrówki, rower...

Ma rodzinę: żonę, dwóch synów i ukochanego wnuczka.

Katowicko-Krakowskie Seminarium Fazy Skondensowanej

W dniu 8 maja 2009 r. odbyło się w Krakowie w Instytucie Fizyki Uniwersytetu Jagiellońskiego Katowicko-Krakowskie Seminarium Fazy Skondensowanej zorganizowane przez prof. Andrzeja Szytułę – kierownika Zakładu Fizyki Ciała Stałego IF UJ i prof. Alicję Ratuszną – kierownika Zakładu Fizyki Ciała Stałego IF UŚ. Seminarium te powoli zaczynają stanowić tradycję obu uniwersytetów. Są organizowane na przemian w Katowicach i w Krakowie. Tegoroczne było kolejnym, ósmym już spotkaniem.

W seminarium wzięło udział około sześćdziesięciu osób.

Referat plenarny „Polski synchrotron – perspektywy” wygłosił prof. Edward Goerlich z UJ. Omówił możliwości eksperymentalne, jakie dla różnych dziedzin nauki i techniki stwarza synchrotron mający powstać w Krakowie. W porównaniu z indywidualnymi przyrządami laboratoryjnymi będzie to ogromny krok naprzód w nauce polskiej. Drogę do realizacji projektu budowy źródła promieniowania synchrotronowego otworzyło podpisanie przez Uniwersytet Jagielloński umowy wstępnej z MNiSW. Umowa określa skalę projektu, a jego ogólne założenia techniczne zostały już przyjęte. Realizację projektu przewidziano na 2014 rok.

Podczas trzech sesji Seminarium wygłoszono 21 krótkich, 15-minutowych referatów. Dotyczyły one najbardziej aktualnych badań z zakresu fazy skondensowanej prowadzonych w Polsce Południowej. Program wyglądał następująco:

Sesja pierwsza:

- Artur Chrobak – Wpływ relaksacji strukturalnej na magnetostrykcję w stopach amorficznych na bazie żelaza,
- Jerzy Goraus – Własności termodynamiczne i transportowe skutterudytów $\text{REPt}_4(\text{Sn,Sb})_{12}$ ($\text{RE} = \text{La, Ce, Pr, Nd, Eu, Yb, Sm}$) oraz $\text{MPd}_4(\text{Sn,Sb})_{12}$ ($\text{M} = \text{Ca, Sr, Ba, Eu, Yb}$),
- Władysław Borgiel – „Sputtering” na przygotowanej powierzchni – narzędzie nanotechnologii,
- Stanisław Baran – Własności magnetyczne ErCu_2Si_2 ,
- Łukasz Hawełek, Andrzej Burian – Transformacja nanodiament–nanocebulki węglowe,
- Paweł Starowicz – Badanie oddziaływań elektron–fonon w TaSe_2 przy pomocy kątownorozdzielczej spektroskopii fotoelektronów.

Sesja druga:

- Marzena Rams – Badania biologiczne nowego fotouczulacza pod kątem wykorzystania w PDT,
- Anna Mrozek, Agnieszka Szurko, Jarosław Polański, Alicja Ratuszna – Właściwości przeciwnowotworowe związków opartych na strukturze chinazolini,
- Zofia Drzazga – Autofluorescencja kości,
- Karina Maciejewska, Zofia Drzazga – Wpływ leków antyretrowirusowych na ontogenezę kości,

- Jerzy Kubacki – Przełączanie rezystywne w kryształach niobianu potasu,
 - Yuriy Natazon – Obliczenia kwantowo-mechanicznej struktury elektronowej MgH_2 ,
 - Sebastian Pawlus – Dynamika relaksacyjna w 1,4-poliizoprenach pod wysokim ciśnieniem – nowe wyniki doświadczalne,
 - Robert Pełka – Oznaki przejścia Berezinskiego–Kosterlitz–Thoulessa w warstwach magnetyku molekularnego.
- Sesja trzecia:
- Jakub Rysz – Wpływ struktury cienkich warstw plastikowych ogniw słonecznych na ich wydajność,
 - Justyna Jaczewska – Samoorganizacja według szablonu jako metoda tworzenia układów elektroniki plastikowej,
 - J. Czerwiec, R. Dąbrowski, J. Gasowska, M. Marzec, A. Mikułko, M. Wierzejska-Adamowicz, D. Ziobro, S. Wróbel – Właściwości dielektryczne nowych związków ciekłokrystalicznych z fazą ferroelektryczną i antyferroelektryczną,
 - Mirosław Gałązka – Sprzężenie parametrów porządku i jego wpływ na własności efektywnych wykładników krytycznych,
 - Andrzej Biborski – Termodynamika defektów punktowych w związkach międzymetalicznych: symulacje Monte Carlo układów otwartych,
 - Marzena Materska – Nanocząsteczki CoPt : struktura i własności magnetyczne,
 - Teresa Jaworska-Gołąb – Magnetyczne przejścia fazowe w pseudo-tetragonalnym międzymetalicznym związku $\text{NdMn}_{1,6}\text{Fe}_{0,4}\text{Ge}_2$ indukowane zmianami temperatury i ciśnienia.

Sesje prowadzili kolejno profesorowie: Andrzej Szytuła (UJ), Alicja Ratuszna (UŚ), Zofia Drzazga (UŚ) oraz Stanisław Urban (UJ). Poszczególni referenci rekrutowali się nie tylko z obu uniwersytetów ale też z Instytutu Fizyki Jądrowej PAN.

Na sesji plakatowej przedstawiono 18 prezentacji, wszystkie interesujące merytorycznie, w ciekawej szacie graficznej.

Seminarium stanowi znakomite forum dla doktorantów i nowych adiunktów, którzy występując przed „swoimi” nabierają szlifów koniecznych na konferencjach międzynarodowych. Ponadto daje okazję do nawiązania współpracy naukowej i zorientowania się w najbardziej aktualnych możliwościach doświadczalnych całego środowiska.

Należy podkreślić dużą staranność organizatorów, miłą atmosferę oraz wielki entuzjazm do dalszego rozwijania utrwalającej się tradycji tych seminariów. Może trochę szkoda, że czas na dyskusję był tak minimalny (praktycznie można ją było prowadzić jedynie podczas przerwy kawowej czy obiadowej). Niemniej bardzo bogaty program jednodniowego spotkania był niewątpliwie godny podziwu.

Małgorzata Nowina Konopka

Instytut Fizyki Jądrowej
im. H. Niewodniczańskiego PAN, Kraków

Jan Turkiewicz (1934–2009)

Jan Andrzej Turkiewicz był fizykiem jądrowym. W latach 1983–87 pełnił funkcję dyrektora Instytutu Problemów Jądrowych im. Andrzeja Sołtana w Świerku.

Urodził się 30 maja 1934 r. w Krzemieńcu na wschodnich kresach Polski. Pochodził z zasłużonych dla rozwoju kultury polskiej rodzin Turkiewiczów i Choroszczaków. Ojciec jego, z zawodu inżynier leśnictwa, był zarządcą dóbr słynnego Liceum Krzemienieckiego. We wczesnym dzieciństwie, w okresie okupacji sowieckiej stracił w tragicznych okolicznościach rodziców. Kilka miesięcy spędził w ukraińskim sierocińcu wraz z młodszym bratem, gdzie odszukała ich rodzina matki i zabrała do Warszawy. Podczas Powstania Warszawskiego ginie jego młodszy brat.

W 1952 r. ukończył Liceum im. Wojciecha Górskiego w Warszawie. Już w szkole średniej przejawiał wybitne uzdolnienia do matematyki i fizyki. Był także (wg relacji jego kolegów z Towarzystwa Wychowawców Szkoły Wojciecha Górskiego) zapalonym szachistą i świetnym bramkarzem w szkolnej drużynie piłki nożnej.

Studia wyższe w zakresie fizyki odbył w latach 1952–56 na Wydziale Matematyki, Fizyki i Chemii Uniwersytetu Warszawskiego. Tematem jego pracy magisterskiej, wykonanej pod kierunkiem prof. Mariana Danysza, było „Wyznaczenie współczynnika skurczu i zdolności hamującej dla emulsji jądrowej”.

Stopień doktora nauk matematyczno-fizycznych uzyskał w roku 1963 na Uniwersytecie Warszawskim na podstawie rozprawy pt. „Oddziaływanie bezpośrednie w reakcji $^{122}\text{Te}(n,p)^{122}\text{Sb}$ przy energii neutronów 14.1 MeV”. Promotorem pracy był prof. Zdzisław Wilhelmi. Habilitował się w roku 1969 w Instytucie Badań Jądrowych na podstawie rozprawy „Reakcje (n,α) i pewne zagadnienia oddziaływań bezpośrednich w obszarze niskich energii”. W roku 1975 uzyskał w Instytucie Badań Jądrowych tytuł profesora nadzwyczajnego nauk fizycznych, a w roku 1987, w Instytucie Problemów Jądrowych, tytuł profesora zwyczajnego.

Pracę zawodową rozpoczął na Uniwersytecie Warszawskim w katedrze kierowanej przez prof. Andrzeja Sołtana. Było to jeszcze w czasie ostatniego roku jego studiów. Po ukończeniu studiów związał się z Instytutem Badań Jądrowych (IBJ) przechodząc tam kolejne szczeble kariery zawodowej od asystenta (1956), adiunkta i kierownika pracowni (1963), docenta i kierownika zakładu (1970), do profesora nadzwyczajnego (1975). Po podziale IBJ w 1983 r. na trzy niezależne instytuty został dyrektorem jednego z nich (Instytutu Problemów Jądrowych (IPJ)). Funkcję tę pełnił do roku 1987. W roku 1997 przeszedł na emeryturę.

Jan Turkiewicz odbył dwa dłuższe staże zagraniczne. Jeden w latach 1957–58 w Centrum Badań Jądrowych (CEN) w Saclay (Francja), a drugi w latach 1969–70 w Zjednoczonym Instytucie Badań Jądrowych w Dubnej

(ZSRR). Jako ekspert Międzynarodowej Agencji Energii Atomowej w Wiedniu przebywał w roku 1984 w Mongolii, a w 1987 r. w Tunezji. Współpracował z wieloma ośrodkami, głównie z CEN w Bordeaux, IPN w Orsay, Instytutem Fizyki Uniwersytetu w Mediolanie i Instytutem Fizyki Jądrowej Ukraińskiej Akademii Nauk w Kijowie.



Jan Turkiewicz

Jako wieloletni kierownik Zakładu Reakcji Jądrowych IBJ, wiele wysiłku wkładał w organizację pracy podległego mu zespołu oraz w kształcenie młodej kadry. Był promotorem 11 prac doktorskich z zakresu fizyki jądrowej i jej zastosowań. Pełnił także funkcję recenzenta w licznych przewodach doktorskich, habilitacyjnych i profesorskich i był członkiem wielu rad naukowych i komisji. W latach 1964–68 prowadził na Uniwersytecie Warszawskim wykład „Modele jądrowe” dla studentów IV i V roku fizyki oraz wykład „Fizyka ogólna” dla studentów III roku matematyki studiów wieczorowych. W późniejszych latach 1971–77 prowadził wspólnie z prof. Zdzisławem Wilhelmi seminarium z fizyki jądrowej dla studentów starszych lat fizyki UW, uczestników Studium Doktoranckiego IBJ oraz pracowników Zakładu Reakcji Jądrowych IBJ i Zakładu Fizyki Jądra Atomowego UW. W latach 1983–86 pełnił funkcję kierownika dużego problemu węzłowego 04.3. „Badanie procesów jądrowych i ich zastosowanie w społeczno-gospodarczym rozwoju kraju”, obejmującego większość badań jądrowych w Polsce.

Pierwsze prace naukowe Jana Turkiewicza dotyczyły procesu trypartycji jąder ciężkich wywoływanej powolnymi i szybkimi neutronami. Prace te zyskały duże uznanie międzynarodowe i były referowane na drugiej Międzynarodowej Konferencji Pokojowego Wykorzystania Energii Jądrowej w Genewie w 1958 r. W kilkudziesięcioletnim okresie działalności naukowej Jana Turkiewicza centralne miejsce (blisko 70% wszystkich prac) zajmowały bada-

nia oddziaływań szybkich neutronów (stosujące reakcje (n,p) i przede wszystkim (n,α)), ważne w przyszłościowych reaktorach termojądrowych. Mierzac rozkłady kątowe i energetyczne emitowanych produktów reakcji oraz funkcje wzbudzenia, badał mechanizm tych reakcji, stwierdzając dominację oddziaływań wprost w przypadku jąder lekkich ($A < 20$) i ciężkich ($A > 140$) oraz tworzenia się jądra złożonego dla jąder o masach pośrednich ($30 \leq A \leq 70$). Z analizy funkcji wzbudzenia wyciągał wnioski dotyczące czasu życia jądra złożonego. Stwierdzał też klastrową (cząstki α) strukturę powierzchni jąder ciężkich. Badał również wpływ struktury powłokowej jądra tarczy na proces emisji cząstek α . Kilkanaście jego prac wykonanych w późniejszym okresie dotyczyło oddziaływań protonów, deuteronów, jonów helu i cięższych jonów (węgiel, tlen). Większość prac prof. Turkiewicza zyskała uznanie międzynarodowe. Cechowała je nie tylko sama wartość uzyskanych wyników eksperymentalnych, ale również interesująca interpretacja i ciekawe porównanie z wynikami teoretycznymi.

Profesor Turkiewicz jest autorem lub współautorem ponad 130 publikacji w międzynarodowych czasopismach oraz wielu referatów na prestiżowych konferencjach. Jego prace cieszyły się uznaniem w światowym środowisku fizyków i były często używane do testowania zaawansowanych modeli reakcji jądrowych, np. modelu typu „pre-equilibrium”. Warto zwrócić uwagę, że większość prac (a zwłaszcza te, które dotyczą badania reakcji (n,α) i (n,p)) wykonana została w kraju, skromnymi środkami,

którymi dysponowaliśmy tutaj. Prace te są nadal często cytowane w literaturze światowej. Udział prof. Turkiewicza w nich, wykonanych z reguły zespołowo, był bardzo istotny. Polegał na zaproponowaniu ich, wyborze metody badań i wkładzie w analizę i uogólnienie wyników. Dotyczy to zwłaszcza badania oddziaływania neutronów prędkich z jądrami, w której to dziedzinie stworzył on zespół liczący się w skali międzynarodowej (tzw. grupa warszawska).

Prace jego były często nagradzane, trzykrotnie były wyróżnione zespołową nagrodą Państwowej Rady ds. Pożojowego Wykorzystania Energii Jądrowej i raz nagrodą indywidualną I stopnia tej Rady.

W środowisku fizyków był ceniony za erudycję, dociekliwość i wnikliwość naukową. Przyszło mu jednak żyć i działać w bardzo trudnych warunkach i czasach. Nękały go również liczne choroby. Do końca jednak wykazywał silną wolę życia i walki z trudnościami.

Miał dwóch synów i czworo wnuków. Jego zamiłowania pozanaukowe to szachy, grzybobranie, wędkarstwo i wędrowki górskie po ukochanych Tatrach. Był również świetnym kierowcą. Lubił muzykę poważną, której słuchał jeszcze w ostatnich miesiącach i tygodniach swego życia spędzonych w domach opieki. Po wieloletnim zmaganiu się z nieuleczalną chorobą zmarł 28 lutego 2009 r. w domu opieki Tabita w Konstancinie-Jeziornie.

Marian Jaskóła, Zbigniew Moroz, Adam Sobiczewski

Instytut Problemów Jądrowych im. A. Sołtana, Świerk

KRONIKA

■ Tytuły profesorskie

Tytuł naukowy profesora nauk fizycznych nadany przez Prezydenta Rzeczypospolitej Polskiej otrzymali w dniu 3 kwietnia 2009 r.: Andrzej Kazimierz Baran (UMCS), Krzysztof Michał Pachucki (UW), Roman Józef Płaneta (UJ) i Piotr Jan Zieliński (IFJ PAN); 23 kwietnia 2009 r.: Jan Marek Antosiewicz (UW), David Bernhard Blaschke (UWr) i Zbigniew Włodarczyk (Uniwersytet Humanistyczno-Przyrodniczy Jana Kochanowskiego, Kielce).

<http://isip.sejm.gov.pl>

■ Nowi członkowie PAU

Na Walnym Zgromadzeniu Polskiej Akademii Umiejętności w dniu 20 czerwca 2009 r. zostali wybrani nowi członkowie czynni, korespondenci oraz zagraniczni. W Wydziale III Matematyczno-Fizyczno-Chemicznym są to m.in. fizycy: Ryszard Sosnowski, Tomasz Dietl i Arnold Wolfendale.

Nowy członek czynny PAU Ryszard Sosnowski (ur. 1932 r.) był członkiem korespondentem PAU od roku 2000. Jest profesorem emerytowanym Instytutu Problemów

Jądrowych i specjalistą w dziedzinie fizyki cząstek elementarnych i oddziaływań wysokich energii. Badając oddziaływanie proton-proton znalazł właściwą metodę wykrywania mezonów powabnych, a jego badania oddziaływań elektron-pozyton wniosły znaczący wkład do Modelu Standardowego. Jest autorem lub współautorem ponad 390 publikacji, które miały 13 tysięcy cytowań.

Odegrał zasadniczą rolę w doprowadzeniu Polski do członkostwa w CERN-ie. Pełnił wiele ważnych funkcji, m.in. był wiceprzewodniczącym Rady CERN-u, członkiem Komitetu Wykonawczego Europejskiego Towarzystwa Fizycznego, członkiem Centralnej Komisji Kwalifikacyjnej ds. tytułów i stopni naukowych. W 1973 r. otrzymał Nagrodę PAN im. Marii Curie-Skłodowskiej, w 1994 r. Medal Smoluchowskiego PTF. Jest członkiem rzeczywistym PAN.

Członkiem korespondentem PAU został Tomasz Dietl (ur. 1950 r.), profesor Instytutu Fizyki PAN. Jest specjalistą fizyki półprzewodników i fizyki niskich temperatur. Zajmuje się nanotechnologią, spintroniką półprzewodnikową, fizyką układów nieuporządkowanych i mezoskopowych. Jest autorem lub współautorem 230 prac oryginalnych i ponad 100 artykułów przeglądowych, cytowanych

ponad 7 tysięcy razy. Jest członkiem korespondentem PAN, członkiem Komisji Niskich Temperatur IUPAP, laureatem Nagrody PAN im. Marii Curie-Skłodowskiej (1997 r.), Nagrody Naukowej Fundacji Humboldta (2003 r.). W 2006 r. otrzymał najwyższe polskie wyróżnienie naukowe – Nagrodę Fundacji na rzecz Nauki Polskiej.

Członkiem zagranicznym PAU został Arnold Wolfendale (ur. 1925 r.), obywatel Wielkiej Brytanii, fizyk i astronom, profesor emerytowany Uniwersytetu w Durham. Prowadzi badania z dziedziny kosmologii i promieniowania kosmicznego, przede wszystkim pochodzenia promieniowania kosmicznego, fluktuacjami geomagnetycznymi i słonecznymi. Od wielu lat współpracuje z Zakładem Promieni Kosmicznych Uniwersytetu Łódzkiego, a wielu polskich fizyków prowadziło prace w jego instytucie w Durham. W latach 1991–95 był Astronomem Królewskim, a w roku 1992 otrzymał od Polskiego Towarzystwa Fizycznego Medal Smoluchowskiego.

www.pau.krakow.pl

Barbara Wojtowicz

■ Nagroda Templetona dla d’Espagnata

W roku 2009 Nagrodę Templetona za postępy w badaniach i odkryciach dotyczących realności duchowych przyznano Bernardowi d’Espagnat. Uczony ten (ur. 1921 r.) otrzymał tę Nagrodę (wartości 1 mln funtów) za prace nad filozoficznymi implikacjami mechaniki kwantowej i stworzenie teoretycznych podstaw badania naruszania nierówności Bella. Jest to już siódmy fizyk (poprzednim był ks. prof. Michał Heller) uzyskujący w ciągu ostatnich 10 lat tę Nagrodę, która została ufundowana w roku 1972 przez filantropa sir Johna Templetona.

D’Espagnat studiował fizykę w Ecole Polytechnique, doktoryzował się w Institute Henri Poincaré u Louisa de Broglie’a. W CERN-ie był jednym z twórców oddziału teorii. Pracował nad zagadnieniem, jak doświadczalnie można by sprawdzić słuszność nierówności Bella. „Musimy przeprowadzić próby, które wykazałyby, że mechanika kwantowa jest istotnie prawdziwa” – powiedział kiedyś. Taki dowód został uzyskany w 1981 r., gdy wyniki doświadczeń nad polaryzacją fotonów przeprowadzonych przez Alaina Aspecta, z którym d’Espagnat blisko współpracował, wykazały, że nierówność Bella jest istotnie naruszona.

Phys. World 22, nr 4 (2009)

Barbara Wojtowicz

■ Frank Steglich doktorem honoris causa Uniwersytetu Jagiellońskiego

Wniosek o nadanie profesorowi Frankowi Steglichowi godności doktora honoris causa Uniwersytetu Jagiellońskiego złożyła grupa sześciu profesorów z Wydziału Fizyki, Astronomii i Informatyki UJ. Recenzowali go profesorowie Henryk Szymczak z Instytutu Fizyki PAN i Andrzej Ślebarski z Instytutu Fizyki Uniwersytetu Śląskiego. Na podstawie obu recenzji popartych przez władze wymienionych placówek Senat UJ, uchwałą z dnia 29 października 2008 r. postanowił przyznać profesorowi Frankowi Steglichowi tę godność.

Uzasadnieniem uchwały były przede wszystkim ogromne osiągnięcia naukowe profesora Franka Steglicha, który:

- odkrył nadprzewodnictwo w układach ciężkich fermionów, a także współistnienie magnetyzmu i nadprzewodnictwa,
- udowodnił, że magnetyzm i nadprzewodnictwo w tych układach nie tylko się nie wykluczają, ale także, co było zaskakujące dla opinii środowiska fizyków, mogą być spowodowane przez jeden wspólny mechanizm mikroskopowy,
- odkrył szereg nowych układów wykazujących kwantowe przejście fazowe, które obecnie są w centrum światowej uwagi fizyków materii skondensowanej.

W uzasadnieniu decyzji Senat UJ podkreślił również zasługi profesora Steglicha w rozwijaniu kontaktów Niemiec z ich wschodnimi sąsiadami, a zwłaszcza z Polską. Profesor Steglich nawiązał współpracę z polskimi fizykami ponad 30 lat temu, kiedy nasze kraje dzieliła jeszcze żelazna kurtyna. Współpracy tej nie przerwał bardzo trudny okres stanu wojennego, podczas którego nawet łączność telefoniczna była niezmiernie utrudniona. Profesor wielokrotnie przyjeżdżał do Polski, uczestniczył w konferencjach, seminariach i szkołach organizowanych w naszym kraju, publikował swoje wyniki w polskim czasopiśmie *Acta Physica Polonica B* (15 prac!) i *Acta Physica Polonica A*, podnosząc w ten sposób jego prestiż.

Uroczystość nadania tytułu z wręczeniem dyplomu i epitafium odbyła się 16 grudnia 2008 r. w Collegium Maius. Laudację wygłosił prof. Józef Spałek. Wspominał kilka dat i faktów z życiorysu Franka Steglicha jak: lata studiów fizyki w Münster i w Getyndze (1960–66) czy 1969 – rok obrony pracy doktorskiej. Obecnie profesor Steglich jest dyrektorem Instytutu Chemii i Fizyki Materiałów im. Maxa Plancka w Dreźnie. Dzięki niemu pracuje tam wielu stypendystów zagranicznych, w tym także z ośrodków polskich.



Prof. Frankowi Steglichowi (z prawej) gratulacje składa prof. Józef Spałek (fot. Anna Wojnar)

Znacznie więcej uwagi poświęcił laudator omówieniu osiągnięć naukowych nowego doktora honoris causa. Były to prace doświadczalne, co dla teoretyka jest szczególnie imponujące. Największe odkrycie profesora Steglicha miało miejsce w 1979 roku i polegało na zaobserwowaniu nadprzewodnictwa w układzie ciężkich fermionów CeCu_2Si_2 o własnościach magnetycznych. Oznaczało to, że związkami wyjściowymi dla nadprzewodników wysokotemperaturowych są izolatory magnetyczne, a nie związki metaliczne i pociągnęło za sobą odkrycie całej klasy układów nowych nadprzewodników, które nie mogą być opisane teorią Bardeena, Coopera i Schrieffera.

Innym odkryciem Steglicha było wynalezienie szeregu nadprzewodników, jak UPd_2Al_3 , w którym występuje jednocześnie uporządkowanie antyferromagnetyczne i nadprzewodnictwo. Jest to jeden z wielu przykładów zaobserwowania bardzo bogatego diagramu różnych faz i stanów kolektywnych w sposób powtarzalny. Daje to znakomite świadectwo precyzji, ostrożności i kunsztu doświadczalnego profesora Steglicha.

Sukcesem Profesora jest również podjęcie badań kwantowych przemian fazowych i zjawisk krytycznych w pobliżu temperatury zera bezwzględnej. Badania własności tych ciekawych materiałów w ultraniskich temperaturach w skali makroskopowej ujawniają kwantową naturę tych układów, a obserwowane prawa skalowania w funkcji temperatury, wielkości pól zewnętrznych czy ciśnienia dostarczają podstawowej wiedzy na temat, jak rozumieć kolektywne ciecze kwantowe. Przemiany fazowe stanowią jedno z centralnych zagadnień współczesnej fizyki kwantowej. Występuje ono również w astrofizyce, fizyce atomowej czy układach jądrowych przy bardzo wysokich energiach. Jednak układy materii skondensowanej można badać łatwiej i przy znacznie niższych kosztach. Profesor Frank Steglich wnosi znaczny wkład w rozwój tych badań.

Dziękując za przyznanie godności prof. Frank Steglich opowiedział o swoich wieloletnich kontaktach z polskimi fizykami. Wymieniał nazwiska spotkanych osób z różnych ośrodków (Kraków: UJ i AGH, Wrocław: Instytut Niskich Temperatur i Badań Strukturalnych PAN, Katowice: Uniwersytet Śląski), nakreślał rozwój ich kariery naukowej, przypominał wspólne przedsięwzięcia.

Współpraca z Polakami była silnie uwarunkowana przez sytuację polityczną. Walcząca z komunistycznym reżimem Solidarność wzbudzała podziw i szacunek ludzi Zachodu. Prof. Frank Steglich z dumą przyjmował znaczek związku i wpinał go do klapy marynarki, uczestniczył w tajnych spotkaniach członków Solidarności. Na własne oczy widział represje ze strony władz komunistycznych wobec Polaków, ale także w pewien sposób sam ich doświadczył. Wracając z konferencji we Wrocławiu został przetrzymany wiele godzin przez straż graniczną w Zgorzelcu, przez co stracił bilet na samolot. Z opresji wyratował go pewien Duńczyk, który zabrał go samochodem. Rekompensatą mógł być przejazd przez Drezno – miasto, w którym się urodził, a którego od 34 lat nie widział! Ale i to się nie udało – przejeżdżali przez nie w nocy, gdy panowała tam głęboka ciemność komuny.

Prof. Steglich stwierdził, że w czasach Zjednoczonej Europy, kiedy współpraca naukowa obejmuje coraz szersze kręgi światowe, zarówno nauka jak i sztuka przyczyniają się do zbliżenia ludzi, a w szczególności tych, którzy przez długie lata żyli w skrajnie różnych społeczeństwach.

Na zakończenie wyraził swoją ogromną wdzięczność za przyjęcie go do Senatu Uniwersytetu Jagiellońskiego przez nadanie mu tytułu doktora honoris causa, co jest kolejnym dowodem na ścisłe związki pomiędzy Krakowem i Dreznem sięgające bardzo dalekiej historii. Przypomniał, że w 1697 r. elektor saski Fryderyk August I (August Mocny) został królem Polski Augustem II i kiedy zmarł w 1733 r. pochowano go na Wawelu, a jego serce przewieziono do Katedry Drezdeńskiej.

Prof. Frank Steglich wierzy, że więzy z Polakami nie tylko nie zanikną, ale umocnią się i rozkwitną w przyszłości, gdyż, jak napisał wielki poeta niemiecki Wolfgang von Goethe „Es geht nicht über die Freude, die uns das Studium der Natur gewährt” („nie ma większej radości jak badanie natury”). Ostatnie zdanie powiedział po polsku: Życzę Państwu powodzenia we wspólnej pracy naukowej, przynoszącej liczne korzyści społeczności naukowej zarówno Krakowa jak i Drezna.

Małgorzata Nowina Konopka

■ Ernst Bauer doktorem honoris causa Uniwersytetu Marii Curie-Skłodowskiej

W dniu 16 kwietnia 2009 r. odbyło się uroczyste nadanie tytułu doktora honoris causa Uniwersytetu Marii Curie-Skłodowskiej w Lublinie profesorowi Ernstowi Bauerowi, amerykańskiemu uczoneму, który swe życie naukowe podzielił między Niemcy i Stany Zjednoczone. Tę najwyższą godność Uniwersytet przyznał Profesorowi za jego fundamentalny wkład w badania wzrostu kryształów, stworzenie podstaw mikroskopii opartej na zjawisku dyfrakcji powolnych elektronów oraz trzydziestoletnią współpracę z fizykami lubelskimi. Podczas uroczystości sylwetkę profesora Ernsta Bauera przedstawił promotor doktoratu, prof. Mieczysław Jałochowski z Zakładu Fizyki Powierzchni i Nanostruktur Instytutu Fizyki UMCS.

Lata nauki profesora Bauera w liceum przypadły na czas II wojny światowej. Po maturze w 1949 r. Ernst Bauer studiował nauki fizyczne i matematyczne na uniwersytecie w Monachium. Stopień doktora uzyskał w dwa lata po ukończeniu uniwersytetu w 1955 r. Praca doktorska poświęcona była badaniom wzrostu cienkich warstw kryształów jonowych. Pierwsza z jego fundamentalnych publikacji naukowych dotyczyła termodynamicznej teorii epitaksjalnego wzrostu kryształów. Opublikowana w *Zeitschrift für Kristallographie* **110**, 372 (1958), wprowadzała nowe nazewnictwo typów wzrostu kryształów na powierzchni, powszechnie zaakceptowane i stosowane przez środowisko naukowe całego świata. Publikacja ta doczekała się blisko tysiąca cytowań. Poszukując lepszych warunków do prowadzenia badań, wkrótce po uzyskaniu stopnia doktora Ernst Bauer wyjechał do USA, gdzie pracując w Michelson Laboratory w China Lake w Kalifornii, w 1962 r. opracował nową



Profesorowi Ernstowi Bauerowi (z lewej) gratulacje składa Dziekan Wydziału Matematyki, Fizyki i Informatyki UMCS prof. Zdzisław Rychlik

metodę badania powierzchni, znaną jako mikroskopia niskoenergetycznych elektronów LEEM (Low Energy Electron Microscope) i zbudował prototyp tego urządzenia. Pierwszy w pełni działający mikroskop LEEM powstał w Niemczech, w kierowanym przez niego od 1969 r. Instytucie Fizyki Uniwersytetu Technicznego w Clausthal-Zellerfeld. Mikroskop pozwala obrazować na bieżąco zjawiska zachodzące na powierzchni ciała stałego (np. adsorpcję, przejścia fazowe, wzrost warstw, reakcje chemiczne itp.).

W latach osiemdziesiątych Instytut w Clausthal-Zellerfeld stał się szeroko w świecie znanym ośrodkiem badań powierzchni. Wówczas to na stażach naukowych w Clausthal-Zellerfeld przebywało ponad sześćdziesięciu pięciu fizyków z Europy, Ameryki Północnej i Południowej, Japonii, Chin, Indii i Korei, a wśród nich czternastoosobowa grupa polskich naukowców z Lublina i Wrocławia. W 1996 r. profesor Bauer rozpoczął, jako distinguished professor, pracę w Arizona State University, gdzie stworzył nowe wersje mikroskopu LEEM z nowymi możliwościami badawczymi powierzchni i nanostruktur. Powstały m.in. mikroskopy SPLEEM (Spin Polarized Low Energy Electron Microscope), w którym powolne elektrony o spolaryzowanych spinach wykorzystywane są do badań warstw ferromagnetycznych, czy SPELEEM (Spectroscopic Photo Emission and Low Energy Electron Microscopy) dostosowany do badań spektromikroskopowych.

Profesor Ernst Bauer jest autorem ponad czterystu prac naukowych cytowanych prawie dziesięć tysięcy razy, a także jednej książki i osiemdziesięciu dwóch artykułów przeglądowych. Od wielu lat jest członkiem komitetu naukowego czasopisma *Surface Science* i komitetu wydawniczego *Physica Status Solidi (a)*. Wśród wielu przyznanych mu nagród i wyróżnień należy wymienić Nagrodę Gaede Niemieckiego Towarzystwa Próźniowego, Nagrodę Davissona–Germera Amerykańskiego Towarzystwa Fizycznego czy Nagrodę Naukową kraju Dolnej Saksonii.

Profesor Jałochowski szczególnie ciepło ukazał silne i długotrwałe związki Ernsta Bauera z naukowcami Za-

kładu Fizyki Powierzchni i Nanostruktur Instytutu Fizyki UMCS, podkreślając takie cechy osobowości Profesora, jak: życzliwość, wyrozumiałość i partnerstwo w pracy naukowej. Wymienił też osobiste sukcesy we wprowadzaniu nowych technik badawczych i tematyki naukowej w Instytucie Fizyki UMCS, zainspirowane wielokrotnymi pobytami w Clausthal-Zellerfeld. Do dziś trwa współpraca naukowa pomiędzy Zakładem Fizyki Powierzchni i Nanostruktur IF UMCS a profesorem Bauerem w Arizonie, dzięki której powstają liczne wspólne publikacje.

Przyjmując profesora Ernsta Bauera do grona swoich doktorów honorowych, Uniwersytet Marii Curie-Skłodowskiej wyraził szczególne uznanie dla jego wybitnych osiągnięć naukowych w dziedzinie fizyki powierzchni oraz zasług dla środowiska fizyków lubelskich.

W dniu poprzedzającym uroczystość nadania tytułu doktora honoris causa UMCS profesorowi Ernstowi Bauerowi Zakład Fizyki Powierzchni i Nanostruktur IF UMCS zorganizował naukowe spotkanie pod nazwą International Workshop on Surface and Nanostructures Characterization, na które przybyli naukowcy z całego świata, współpracujący z profesorem, jego uczniowie i przyjaciele, w większości zajmujący się badaniami powierzchni z wykorzystaniem mikroskopii LEEM. W trzech sesjach wykłady wygłosili: prof. Takanori Koshikawa z Osaka Electro-Communication University, prof. Michael Altman z Hong Kong University of Science and Technology, prof. Lothar Fritsche z Technical University of Clausthal, dr Ludovic Douillard z Commissariat à l'Energie Atomique Saclay, prof. Marek Przybylski z Max Planck Institute of Microstructure Physics w Halle, prof. Adam Kiejna z Uniwersytetu Wrocławskiego, prof. Marek Szymoński z Uniwersytetu Jagiellońskiego, dr Ryszard Zdyb z UMCS oraz dr Krzysztof Grzelakowski z firmy Option z Wrocławia. Tematyka wystąpień dotyczyła m.in. takich zagadnień, jak: dynamika wzrostu cienkich warstw, zjawisko dyfuzji i transportu masy na powierzchni, obliczenia teoretyczne z pierwszych zasad wzrostu nanostruktur na powierzchni, teoria uporządkowania ferro- i antyferromagnetycznego, efekt rozmiarowy i studnie kwantowe w ultracienkich warstwach żelaza, technologia wytwarzania nanostruktur organicznych na powierzchni ciała stałego, ich charakteryzacja i manipulacja, oraz badanie plazmonów na zrekonstruowanych powierzchniach.

Elżbieta Jartych

■ Małgorzata Nowina Konopka w *Europhysics News*

Małgorzata Nowina Konopka, nasza krakowska korespondentka, wieloletnia wypróbowana współpracowniczka redakcji *PF*, której doniesienia z Polski Południowej można znaleźć w niemal każdym zeszycie *Postępów* (także w tym), wypłynęła na szerokie, europejskie wody – od tego roku jest członkiem rady redakcyjnej (Editorial Advisory Board) *Europhysics News*, magazynu Europejskiego Towarzystwa Fizycznego. Serdecznie gratulujemy!

Redakcja *Postępów Fizyki*

POSTĘPY FIZYKI W INTERNECIE

<http://postepy.fuw.edu.pl>

► ARCHIWUM

spisy treści wszystkich zeszytów

► ARTYKUŁY DO POBRANIA

m.in. przekłady wykładów noblowskich:

- Wolfgang Ketterle – Gdy atomy zachowują się jak fale: kondensacja Bosego–Einsteina i laser atomowy
- Raymond Davis Jr. – Pół wieku z neutrinami słonecznymi
- Masatoshi Koshihara – Narodziny astrofizyki neutrin
- Riccardo Giacconi – Narodziny astronomii rentgenowskiej
- Aleksiej A. Abrikosow – Nadprzewodniki drugiego rodzaju i sieć wirów
- Anthony J. Leggett – Początki historii nadciekłego helu-3 widziane oczami teoretyka
- Witalij Ł. Ginzburg – O nadprzewodnictwie i nadciekłości (co mi się udało zrobić, a czego nie) oraz o „kanonie fizyki” u zarania XXI wieku
- Frank Wilczek – Asymptotyczna swoboda: od paradoksu do paradygmatu
- David J. Gross – Odkrycie asymptotycznej swobody i narodziny QCD
- David Politzer – Dylemat nagradzania
- Roy J. Glauber – Stulecie kwantów światła
- Theodor W. Hänsch – Umiętkowanie dokładności
- John L. Hall – Pomiar częstości optycznych – szansa dla zegarów optycznych i nie tylko
- John C. Mather – Od Wielkiego Wybuchu do Nagrody Nobla i jeszcze dalej
- George F. Smoot III – Anizotropie kosmicznego mikrofalowego promieniowania tła: ich odkrycie i wykorzystanie
- Albert Fert – Geneza, rozwój i przyszłość spintroniki
- Peter A. Grünberg – Od fal spinowych do gigantycznego magnetooporu (GMR) i dalej

oraz wykłady z ostatnich Zjazdów Fizyków Polskich:

Białystok '99, Toruń '01, Gdańsk '03, Warszawa '05, Szczecin '07

► MATERIAŁY DODATKOWE

uzupełnienia niektórych artykułów

PRENUMERATA

Postępy Fizyki można zaprenumerować w jeden z następujących sposobów.

► PRZEZ ODDZIAŁY PTF (tylko prenumerata krajowa dla członków PTF i studentów):

Cena rocznej prenumeraty krajowej w 2009 r. wynosi 48 zł. Dostawa *Postępów* odbywa się za pośrednictwem Oddziałów.

► PRZEZ ZARZĄD GŁÓWNY PTF (tylko prenumerata krajowa):

Wpłaty należy dokonać na konto Zarządu Głównego PTF: 19 1020 1097 0000 7802 0001 3128 (PKO BP IX O/Warszawa) lub w Biurze Zarządu Głównego PTF.

Cena rocznej prenumeraty krajowej w 2009 r. wynosi 60 zł. Dostawa *Postępów Fizyki* następuje drogą pocztową pod wskazany adres.

► PRZEZ PRZEDSIĘBIORSTWA KOLPORTAŻU PRASY:

RUCH (<http://www.prenumerata.ruch.com.pl>)

KOLPORTER (<http://sa.kolporter.com.pl>)

GARMOND PRESS (<http://www.garmond.com.pl>)

Cena rocznej prenumeraty krajowej w 2009 r. wynosi 72 zł.

Prenumerata ze zleceniem dostawy za granicę – patrz <http://www.ruch.pol.pl>.

Dostępne są również zeszyty archiwalne – prosimy o kontakt z redakcją.

INFORMACJE DLA AUTORÓW

Artykuły powinny mieć charakter przeglądowy i być dostępne dla ogółu fizyków. Prace należy nadsyłać pod adresem redakcji. O przyjęciu pracy do druku decyduje komitet redakcyjny. Prac niezamówionych i niezakwalifikowanych do druku redakcja nie zwraca. Bardziej szczegółowe informacje na temat układu i sposobu przygotowania pracy znajdują się na stronie internetowej *Postępów Fizyki*.

REKLAMA W POSTĘPACH FIZYKI

Zapraszamy – szczególnie przedstawicieli producentów aparatury oraz sprzętu i oprogramowania komputerowego, wydawców podręczników i książek naukowych oraz popularnonaukowych – do zamieszczania ogłoszeń reklamowych w *Postępach Fizyki*. Nasze czasopismo dociera do większości polskich fizyków, z których wielu decyduje o bieżących zakupach uczelni, instytutów i szkół. Zainteresowanych prosimy o kontakt z redakcją pod adresem: postepy@fuw.edu.pl.

POSTĘPY FIZYKI (ADVANCES IN PHYSICS)

Founded in 1949, published bimonthly in Polish with titles in English by the Polish Physical Society with a support of the Ministry of Science and Higher Education and the Physics Faculty of the Warsaw University.

INFORMATION FOR SUBSCRIBERS

A subscription order can be sent through the local press distributor or directly to „RUCH” S.A. Oddział Krajowej Dystrybucji Prasy, ul. Jana Kazimierza 31/33, skrytka pocztowa 12, 00-958 Warszawa, Poland (for details see <http://www.ruch.pol.pl>).

WKRÓTCE W POSTĘPACH

- *Ludwika Lipińska o nanomateriałach zol-żel dla optoelektroniki i fotowoltaiki*
- *Ryszard Horodecki o kryptografii kwantowej*
- *Stanisław Bajtlik o Światowym Roku Astronomii 2009*
- *Marcin Kostur i Jerzy Łuczka – Motory molekularne, Część II*

Aleja Królewska (Zielony Dywan) w ogrodach Wersalu (płyta Lippmanna 9×12 cm). To jest obraz archiwalny! Aleja Królewska zaczyna się przy fontannie Latony i biegnie do fontanny Apollina. W głębi widzimy Wielki Kanał mający prawie milę długości.



Pałac Luksemburski w Paryżu (płyta Lippmanna 9×12 cm). Ta fotografia wiernie oddaje kolory pałacu i ogrodów luksemburskich, tak jak one wyglądały 100 lat temu.



Z prawej: Fotografia młodego człowieka (płyta Lippmanna 9×12 cm). Jedną z wad „bezziarnistej” emulsji Lippmanna w zastosowaniu do zwykłej fotografii jest jej bardzo mała czułość. Za pomocą tego procesu zrobiono bardzo niewiele fotografii ludzi lub zwierząt.

U dołu: Widmo słoneczne (płyta Lippmanna $6,5 \times 8,5$ cm, ekspozycja od 30 min do 2 h). Widmo słoneczne lub widmo elektrycznego łuku grafitowego to pierwszy zapis kolorów za pomocą interferencji. Wydaje się, że Lippmannowi zabrało dwanaście lat – od 1879 r. – uzyskanie emulsji o dostatecznie drobnym ziarnie, by móc zarejestrować fale stojące wytworzone w płycie. W grudniu 1891 r. została zakomunikowana paryskiej Akademii Nauk praca Lippmanna „Fotografia barwna” przedstawiająca jego proces uzyskiwania trwałych i znakomych kolorów.





XL ZJAZD FIZYKÓW POLSKICH

KRAKÓW, 6-11 WRZEŚNIA 2009

Współorganizatorzy

Oddział Krakowski PTF, Akademia Górniczo-Hutnicza, Instytut Fizyki Jądrowej PAN, Politechnika Krakowska,
Uniwersytet Jagielloński, Uniwersytet Pedagogiczny

Kontakt: XL ZJAZD FIZYKÓW POLSKICH, Oddział Krakowski PTF, Instytut Fizyki UJ, ul. Reymonta 4, PL-30059 Kraków;
telefon: +48 12 663 56 84; fax: +48 12 633 84 94;

Małgorzata Nowina-Konopka (IFJ) — tel. 012 662 82 63; e-mail: malgorzata.nowina-konopka@ifj.edu.pl

Rejestracja i informacja: <http://www.ptf.agh.edu.pl/XL-zjazd/>; Anna Tylek — tel. 012 663 38 58 lub 012 663 38 29;

Joanna Hoszko — tel. 012 663 38 58 lub 012 663 38 29

Komitet Honorowy:

Stanisław Kardynał Dziwisz
Arcybiskup Metropolita Krakowski
Prof. dr hab. Andrzej Białas
Prezes Polskiej Akademii Umiejętności
Prof. dr hab. Marek Jeżabek
Dyrektor IFJ-PAN w Krakowie
Prof. dr hab. Michał Kleiber
Prezes Polskiej Akademii Nauk
Prof. dr hab. Maciej Kolwas
President Elect of the European
Physical Society
Prof. dr hab. Karol Musiał
JM Rektor Uniwersytetu Jagiellońskiego

Komitet Programowy:

Marek Demiański UW
Robert Gałązka IFPAN — przewodniczący
Wojciech Gawlik UJ
Ryszard Horodecki Uniw. Gdański
Marek Kuś CFT PAN
Roman Mlepnas UAM Poznań
Stefan Pokorski UJ
Maria Rożarska IFJ Kraków
Józef Sznajd INTIBS PAN Wrocław
Józef Szudy UMK Toruń
Andrzej Twardowski UW
Karol Wysokiński UMCS Lublin