

Praca umieszczona jest na stronie internetowej poświęconej rozważaniom dotyczącym natury światła:

<http://fiertek.3miasto.net.pl/>

sprawdzone i poprawiane 16-17.09.2021

Światło! Czym naprawdę jest?

Paweł Fiertek

Akademia Marynarki Wojennej

„Pierwszy mój błąd polegał na tym, że uznałem drogi planet za doskonałe koła i błąd ten kosztował mnie tym więcej, że opierał się na autorytecie wszystkich filozofów oraz zgodny był z metafizyką” - 1609r, Johannes Kepler [1].

Historia, którą przyjdzie mi opowiedzieć rozpoczyna się w małym, dość słabo wyposażonym laboratorium studenckim, którego lokalizacja wydaje się być bardzo symboliczna. Związane jest to z tym, że znajduje się ono tuż nad dużym audytorium, w którym od wielu lat prowadzone są wykłady z fizyki dla studentów Politechniki Gdańskiej. Jakby jacyś nieznani planiści poprzez symbolikę tego miejsca chcieli nam przekazać, że nad wykładaną tam wiedzą nadal góruje empiria i to ona jest ostateczną instancją stwierdzającą co jest, a co nie jest prawdą? Początek roku akademickiego 2006 rozpoczął się bez jakichkolwiek godnych uwagi zdarzeń, a po kilku tygodniach przyszedł czas, by małe laboratorium zaczęło wypełniać się regularnie studentami. Wśród nich była również nieśmiała, z pozoru niczym niewyróżniająca się studentka o imieniu Kamila. Pewnego dnia zajęcia na które się udała były wyjątkowo prowadzone przez profesora, z którym miała wykłady. Usiadła niepewnie przy swoim stanowisku zatytułowanym: „Doświadczenie Younga” i szybko studiując pozostawioną tam instrukcję, nerwowo rozglądała się, czy profesor przypadkiem nie nadchodzi. Za długo czekać nie musiała. Po wyrzuceniu kilku nieprzygotowanych osób i przepytaniu kilku szczęśliwców, którzy sali opuszczać nie musieli, podszedł on do Kamili. Ponieważ była to ostatnia osoba jaka mu pozostała do przepytania, postanowił zrobić to nieco dokładniej. Usiadł więc obok:

- Dzień dobry.
- Dzień dobry panie profesorze.
- Z jakim doświadczeniem przyjdzie się pani zapoznać na tym stanowisku?
- Z tego co doczytałam, to będzie to doświadczenie Younga.
- Doczytałaś? No dobra, zobaczmy co sobie pani doczytała. Powiedz mi proszę, na czym polega istota tego doświadczenia?
- Na interferencji światła. {Pośpiesznie powiedziała Kamila.}
- No dobrze, ale przydałoby się to opisać trochę dokładniej. Nieprawdaż?
- Tak panie profesorze. Mamy tutaj siatkę dyfrakcyjną, przez którą przepuszczamy światło z lasera, a następnie oglądamy na ekranie obrazy interferencyjne. Później trzeba zmierzyć odległość siatki od ekranu oraz odległość między plamkami światła na ekranie, podstawić do wzoru i mamy już obliczoną stałą siatki dyfrakcyjnej.
- Proszę pani! W zasadzie to pytanie było o istotę doświadczenia, a nie jak go przeprowadzić i co w nim należy zmierzyć. Wobec czego zacznijmy od początku. Czy pamięta pani z wykładu, czym jest światło?
- Tak, ale tego nie bardzo rozumiem. Na wykładzie było coś o dualizmie falowo-korpuskularnym. W każdym razie pamiętam, że światło jest falą elektromagnetyczną w pewnym przedziale długości fali.
- No właśnie, a teraz proszę mi przypomnieć, jakie doświadczenie przeprowadził nieco ponad dwieście lat temu w roku 1801 Thomas Young i dlaczego było to tak ważne doświadczenie?
- Thomas Young przepuścił światło przez przesłonę, na której były dwa otwory. Następnie na ekranie umieszczonym w pewnej odległości od przesłony zaobserwował prążki interferencyjne. To znaczy odważnie interpretował je jako interferencyjne, bo w tamtym czasie dominowała hipoteza Newtona, która zakładała tylko korpuskularną interpretację zjawisk dotyczących światła. Mimo, że już za czasów Newtona, niejaki Hook postulował falową naturę światła. Zatem było to bardzo ważne doświadczenie, gdyż pozwoliło

przewyciężyć autorytet Newtona i zwróciło uwagę innych ludzi na falowy aspekt zjawisk dotyczących światła, to znaczy dyfrakcję i interferencję.

- Pięknie! Widzę, że jednak uważała pani na wykładach. Jak pani już wspomniała, Young w swoim doświadczeniu wykorzystał przesłonę z dwiema blisko leżącymi szczelinami. Co się zatem dzieje z światłem, kiedy przechodzi ono przez szczelinę?

- Skoro światło jest falą, to tak jak każda fala ulega ona załamaniu na krawędzi. Szczelina jest niczym innym jak dwoma blisko leżącymi krawędziami. Tak więc, światło ugina się w jedną i drugą stronę otworu, następnie rozchodzi się jakby ze źródła punktowego. To mi przypomina coś o zasadzie Huygensa, która zakłada, że każdy punkt czoła fali jest traktowany jak źródło nowej fali punktowej. Jeśli dobrze rozumiem, to faza drgań fali wychodzącej z szczeliny powinna być taka sama jak faza drgań fali wchodzącej?

- Zgadza się. {Z zadowoleniem oznajmił profesor widząc, że choć jedna z osób na sali wiedziała coś więcej niż to, co było w instrukcji laboratoryjnej. W pełni podbudowany uzyskaną odpowiedzią kontynuował rozmowę.}

- Jak już pani wspomniała wcześniej, na ekranie światło interferuje ze sobą i powstaje obraz, który nazywamy obrazem interferencyjnym. Czy mogłaby pani powiedzieć coś więcej na ten temat?

- Interferencji? Jak pamiętam z wykładów, jest to nakładanie się w jednym miejscu różnych fal tego samego rodzaju, w tym również dotyczy to fal elektromagnetycznych. Więc, muszą na ekran padać co najmniej dwie różne fale, aby mógł powstać obraz interferencyjny. Dlatego też, w oryginalnym doświadczeniu była przesłona z dwiema szczelinami. W naszym przypadku mamy siatkę dyfrakcyjną. Czy są na niej dwie szczeliny?

- Oj! Pani Kamilo, chyba nie bardzo pani wie, co to jest siatka dyfrakcyjna. Wspomniała pani na samym początku naszej rozmowy, że będzie pani liczyła stałą siatki dyfrakcyjnej. Okazuje się jednak, że nie wie pani, co to jest siatka dyfrakcyjna i co za tym idzie, czym jest jej stała.

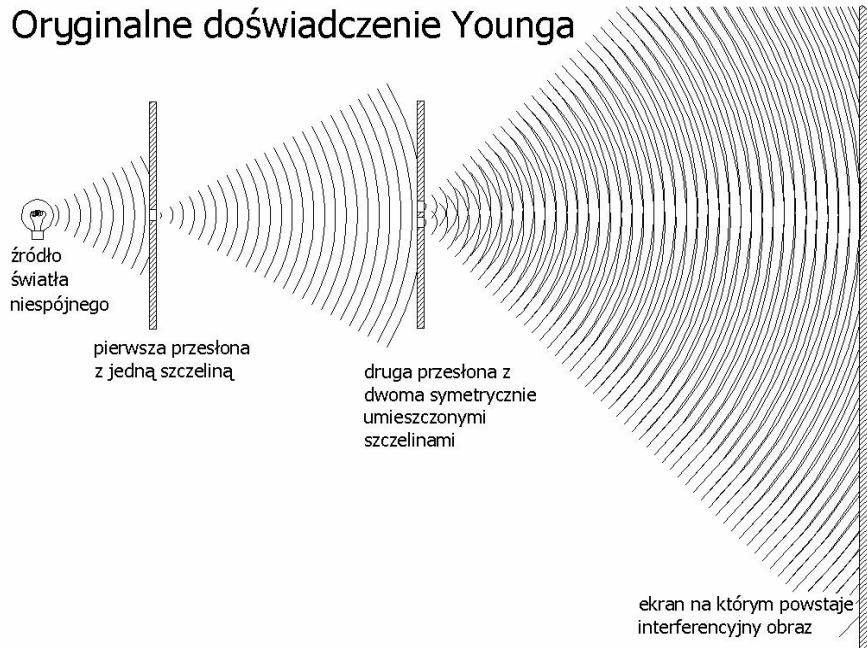
- W instrukcji nie było napisane; co to jest? {Cichym głosem pobrząkała pod nosem.}

- Pani Kamilo, siatka dyfrakcyjna to nic innego jak np. zwykłe szkiełko z naciętymi bardzo gęsto małymi rowkami, które pełnią funkcję przesłony. Jak również, może być to ...

- Już wiem! {Krzyknęła z entuzjazmem Kamila, nie przejmując się lekkim zażenowaniem profesora, który nie był przyzwyczajony do tego by mu przerywać}. Obszar nie zarysowany to szczeliny w naszej przesłonie, a odległość między nimi to stała siatki dyfrakcyjnej. Ale w takim razie, po co Young stosował dwie przesłony w swoim doświadczeniu?

- Widzi pani, najlepiej jak wyjaśnimy to na rysunku (Rys. 1).

Oryginalne doświadczenie Younga



Rys. 1 Schemat oryginalnego doświadczenia Younga.

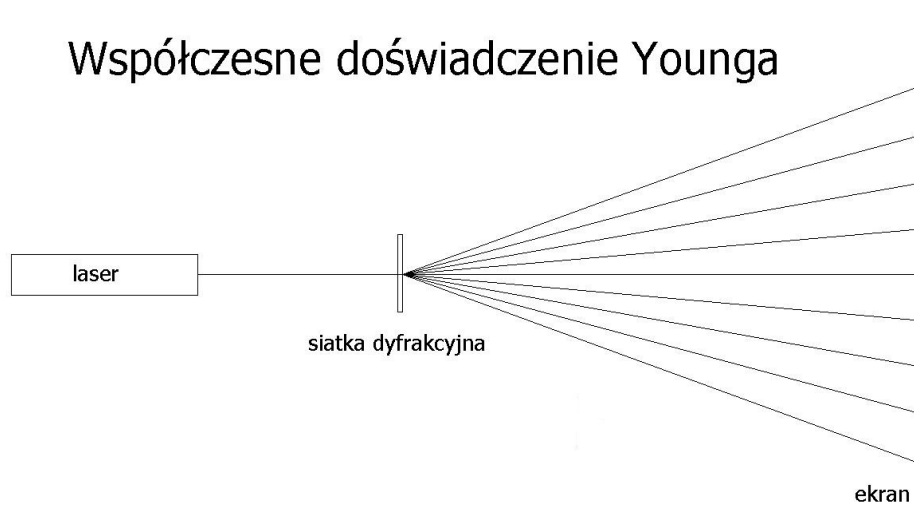
- Jak zapewne pani pamięta Young wykonał swoje doświadczenie ponad dwieście lat temu i to w czasach, gdy poglądy na temat światła były zdominowane rozważaniami Newtona, który zdecydowanie sprzeciwiał się, jakoby światło było jakąkolwiek falą. Tak więc, przedstawienie dowodu na falową naturę światła przez Younga musiało być dostatecznie przekonujące, by nie został on wyśmiany i posadzony o herezje. Wiedząc o tym, zadbał on o szczegóły istotne dla dalszych swoich rozważań tak, aby inni słuchacze nie byli w stanie mu cokolwiek zarzucić. Czy domyśla się pani z rysunku, o jaki szczegół może chodzić?

- Chyba tak. Jak już wspomniałam, faza fali wychodzącej z szczeliny musi być taka sama jak faza fali padającej na przesłonę. Z tego co widzę na rysunku, na drugiej przesłonie, gdzie mamy dwie szczeliny dociera do nich czoło fali o tej samej fazie. To znaczy, że faza fal wychodzących z obu szczelin musi być taka sama. Czy to jest aż tak ważne?

- Tak i to bardzo, ponieważ to, co się dzieje na ekranie bezpośrednio uzależnione jest od tego; jaka jest różnica faz między jedną a drugą wiązką światła padającego na ekran. Young tłumaczył powstawanie prążków na ekranie poprzez różnicę dróg jaką światło pokonuje dochodząc do danego punktu leżącego na ekranie od jednej i drugiej szczeliny. Jeśli ekran i przesłona są nieruchome to jesteśmy zmuszeni stwierdzić, że w każdym momencie czasu droga, jaką pokonuje światło z przesłony do punktu leżącego na ekranie nie ulega zmianie. Jeżeli obraz ma pozostać nieruchomy, to również nie może zmieniać się w czasie różnica faz fal świetlnych wychodzących z poszczególnych szczelin. Aby temu zaradzić Young musiał wykorzystać jeszcze jedną przesłonę, która była mu potrzebna tylko po to, by światło docierające do obu szczelin miało taką samą fazę (Przy założeniu, że nie mamy przypadkowych przesunięć fazy w czasie. Chodzi o to, że do ekranu dochodzą jednocześnie fale, które zostały wyemitowane w różnych momentach czasu.). Bez tej dodatkowej szczeliny, jego słuchacze słusznie zarzuciliby mu fałsz twierdząc, że światło dzienne jako fala jest w pełni niespójne. We współczesnej wersji tego doświadczenia nie korzysta się z pierwszej przesłony. Może pani pamiętać, jak ten problem jest rozwiązany obecnie?

- W instrukcji było coś wspomniane o tym, że obrazy interferencyjne nie mogą powstać dla światła niespójnego, stąd do doświadczenia wykorzystuje się światło lasera, które jest światłem spójnym. Myślę, że w ten właśnie sposób rozwiązano wspomnianą niedogodność (Rys. 2).

Współczesne doświadczenie Younga



Rys. 2 Schemat współczesnego doświadczenia Younga.

- Dokładnie. Wobec czego proszę powiedzieć mi, czym jest światło spójne?
- W świetle przytoczonych faktów jest to bardzo proste pytanie panie profesorze. Musi to być taka własność światła, że jeśli weźmiemy dowolny przekrój wiązki świetlnej prostopadły do kierunku jego rozchodzenia się, to w każdym punkcie tego przekroju, faza jaką posiada fala świetlna musi być taka sama. Wówczas, takie światło padając na przesłonę nawet z wieloma szczelinami, da na wyjściu wiele źródeł fal wychodzących z każdej szczeliny niezależnie, ale o tej samej fazie bądź stałym przesunięciu fazowym (np. skręcenie siatki względem kierunku padania światła laserowego).
- Widzę, że jesteś bardzo bystra! Dokładnie o to chodziło. Dwieście lat temu ludzie nie posiadali źródeł światła spójnego i słuchacze Younga bez problemu mogli mu zarzucić, że światło jakie używa do swoich doświadczeń nie posiada wspomnianej własności. Można powiedzieć nawet więcej. Światło termiczne pochodzące od mocno rozgrzanych obiektów takich jak żarówka czy żarzący się kawałek metalu, nie tylko posiada niejednakową fazę w przekroju wiązki świetlnej, ale faza danej fali zmienia się również przypadkowo w czasie w różnych punktach. Można więc powiedzieć, że takie źródło światła jest zupełnie niespójne zarówno czasowo jak i przestrzennie. Wynika to stąd, że każdy atom rozgrzanego obiektu emituje kwant światła niezależnie i w przypadkowym kierunku.
- Rozumiem więc, że dla światła niespójnego z każdej szczeliny siatki dyfrakcyjnej powinna wychodzić fala świetlna w zupełnie przypadkowej fazie w stosunku do innych i to jeszcze zmiennej przypadkowo w czasie? Gdyby na ekran dochodziło światło z wielu szczelin z przypadkowym przesunięciem fazowym, niemal wszystkie fale dodałyby się tak, że zdecydowana większość wygasłaby się dając ciemny obraz, czyli brak obrazu.
- Dokładnie o to chodziło. Dlatego nie może powstać obraz w doświadczeniu Younga dla światła niespójnego i dlatego w wielu skryptach opisujących wspomniane doświadczenie możemy się często spotkać ze stwierdzeniem, że dla otrzymania obrazu w tych zjawiskach potrzeba bezwzględnie światła spójnego. No dobra pani Kamilo czas ucieka, a tu trzeba zrobić następujące zadanie. Mamy tutaj laser wskaźnikowy jako źródło światła spójnego o długości fali 632,8nm, siatkę dyfrakcyjną zestaw cienkich drucików, regulowaną pojedynczą przesłonę i siatkę ochronną od monitora składającą się z wielu krzyżujących się pod kątem prostym cienkich drucików. Pani zadaniem będzie w pierwszej kolejności ustalić stałą siatki dyfrakcyjnej, później grubość drucika i odległość między krawędziami pojedynczej szczeliny. Co do drucika i szczeliny podane wzory są przybliżone. Później, jak starczy pani czasu, to można się pani jeszcze zapoznać z siatką „ochronną” od monitora, która pozwala uzyskać bardzo charakterystyczne dobrze widoczne obrazy interferencyjne w postaci siatki punktów.
- Dobrze panie profesorze.

Po skończonej rozmowie Kamila szybko zabrała się za ustawianie sprzętu laboratoryjnego i niebawem mogła już przeprowadzić pierwsze pomiary. Po przepuszczeniu światła laserowego przez siatkę dyfrakcyjną zmierzyła odległość siatki od ekranu jak również odległość pierwszej i drugiej plamki światła od środka obrazu. Po uwzględnieniu wszystkich zmierzonych wielkości określiła stałą siatki dyfrakcyjnej jako: $4,87 \pm 0,02 \mu\text{m}$ i zabrała się za pomiar pozostałych wielkości. Spore było jej zdziwienie, gdy przyszedł czas na doświadczenie z pojedynczą szczeliną. Ku swojemu zdziwieniu stwierdziła, że pojedyncza szczelina również daje takie same prążki jak siatka dyfrakcyjna, tylko o znacznie mniejszym rozsunięciu od środka obrazu i większej ilości punktów. Pamiętała niedawno skończoną rozmowę z profesorem, że do powstania obrazu na ekranie w doświadczeniu Younga potrzeba co najmniej dwóch źródeł światła, by można było mówić o interferencji.

- Przecież to stoi w jawnej sprzeczności z teorią! {Krzyknęła mimochodem. Usłyszawszy to profesor podszedł bliżej i zapytał.}

- Proszę tak nie wrzeszczeć! Co się stało pani Kamilo?

- Panie profesorze. Na początku zajęć stwierdziliśmy, że w doświadczeniu Younga muszą być co najmniej dwa źródła światła o tej samej fazie, by można było mówić o stabilnym obrazie interferencyjnym na ekranie. Tutaj mam tylko jedną szczelinę, a co za tym idzie tylko jedno źródło światła, a nadal widzimy prążki interferencyjne. Chyba coś nie tak z tą interpretacją pana Younga? {Powiedziała pewnym głosem dumnie podnosząc głowę.}

- Niekoniecznie. Rozwiązanie problemu prążków interferencyjnych powstałych po przejściu światła przez pojedynczą szczelinę wymaga bardziej złożonego wyjaśnienia. Pamięta pani zasadę Huygensa?

- Tak, każdy punkt czoła fali jest traktowany jako źródło nowej fali.

- Tak więc, czoło fali świetlnej padającej na pojedynczą szczelinę możemy podzielić na wiele źródeł nowych fal punktowych, które znajdują się między krawędziami szczeliny. Następnie, światło z tych źródeł dalej interferuje na ekranie. Jeśli odległość między tymi źródłami światła będziemy zmniejszać do zera, a ich ilość do nieskończoności, to po sumowaniu (czyli scałkowaniu) w zakresie od jednej do drugiej krawędzi dla tych fal uzyskamy obraz prążkowy na ekranie. Tak więc, teoria falowa nie widzi w tym żadnej sprzeczności.

- Rozumiem, że światło nadal musi być przy tym spójne?

- Myślę, że raczej na pewno.

- Wie pan, panie profesorze. Proszę mi wybaczyć, ale ta interpretacja wydaje mi się mocno naciągana.

- Dlaczego?

- Ponieważ zastanawiam się, po co Young zastosował dwie szczeliny i dlaczego do analizy obrazu zastosował tylko odległość między szczelinami? Jeśli obraz na ekranie powstaje również z pojedynczej szczeliny, to należy ten fakt uwzględnić również przy analizie dwóch szczelin, a on tego nie zrobił. Ponadto, obraz po przejściu przez dwie szczeliny według wyprowadzonego przez Younga wzoru zależy tylko od odległości pomiędzy szczelinami, a nie wielkości samej szczeliny. Ponieważ wielkość szczeliny i odległość między nimi nie musi być taka sama i z reguły nie jest, to być może powinniśmy oczekiwać dwóch różnych nakładających się na siebie obrazów, a tego nie obserwowałam. Ponadto, proszę zauważyć, że schemat doświadczenia Younga jest błędny (Rys. 1), gdyż zakłada on, że po przejściu przez pierwszą przesłonę światło zachowa się jak fala i równomiernie oświetli drugą przesłonę. Tymczasem, druga przesłona nie będzie równomiernie oświetlona w obszarze szczeliny, gdyż tam powinny pojawić się prążki od pierwszej szczeliny!

- Faktycznie, ma pani rację. Na razie proszę przyjąć to bez zastrzeżeń, później poszukam odpowiedzi na zaistniałe wątpliwości.

Profesor zmieszany kłopotliwym pytaniem pośpiesznie odszedł od stanowiska w kierunku innej grupy studentów, którzy akurat potrzebowali pomocy, pozostawiając przy tym Kamilę samą sobie z powstającymi w jej głowie nowymi wątpliwościami. Ta, nie bardzo wiedząc co z sobą robić, zaczęła bawić się siatką od monitora wsadzając i wyjmując ją z wiązki światła laserowego i obserwując przy tym pojawiające się i znikające na ekranie obrazy interferencyjne. W pewnym momencie uświadomiła sobie, że

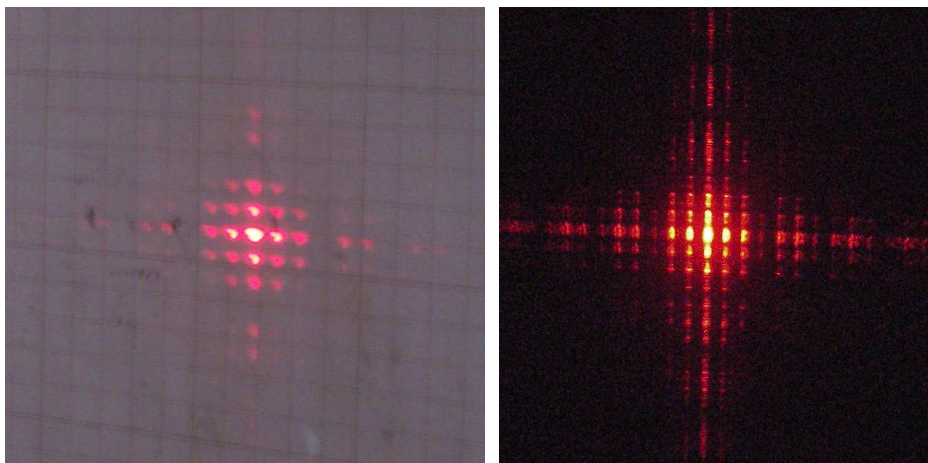
podobny problem i to daleko większy, gdyż bardziej widoczny, ma miejsce i w tym przypadku. Zawołała więc profesora.

- Panie profesorze, czy mogę pana prosić do siebie? Mam pewne wątpliwości z którymi nie umiem sobie poradzić.

- Tak proszę, w czym problem? {Spokojnym krokiem profesor podszedł do stanowiska mając nadzieję, że tym razem chodzi o coś znacznie prostszego.}

- Panie profesorze, zaczęłam bawić się tą dużą siatką od ekranu monitora. Jak rozumiem, są to równoległe cienkie druciki metalowe naciągnięte na ramkę.

- Owszem, w dwóch prostopadłych kierunkach, stąd powinna pani zobaczyć bardziej złożony obraz, niż to co wychodzi przy zastosowaniu siatki dyfrakcyjnej czy pojedynczej szczeliny (Rys. 3).



Rys. 3 Obraz prążków powstały przez przepuszczenie wiązki laserowej przez siatkę od monitora.

- Zgadza się obraz prążków układał się nie w jednym, ale w dwóch kierunkach. W każdym razie zastanawiałam się, co mogę obliczyć na podstawie tego obrazu? W przypadku doświadczenia z siatką dyfrakcyjną miałam do policzenia odległość między szczelinami, w tym wypadku byłyby to odległości między drucikami siatki. W przypadku doświadczenia z pojedynczymi drucikami powstawały prążki interferencyjne, które według skryptu zależą od grubości drucika. Natomiast doświadczenie z pojedynczą szczeliną miało umożliwić mi obliczenie odległości między krawędziami, co w tym przypadku odpowiadałoby odległości między krawędziami drucików w tejże siatce. Tak więc, zastanawiałam się panie profesorze, co właściwie jestem w stanie z tego obrazu obliczyć? Zwłaszcza, że każda z wymienionych wielkości posiada fizycznie inny wymiar i powinna dawać zupełnie różne obrazy na ekranie, a widzimy tylko jeden obraz.

- Wie pani co? {Po chwili drżącym głosem kontynuował.} - Nie wiem. To co pani mówi jest bardzo ciekawe. Pozwoli pani, że później głębiej się nad tym zastanowię. Proszę powiedzieć mi, ile wyszła stała siatki dyfrakcyjnej?

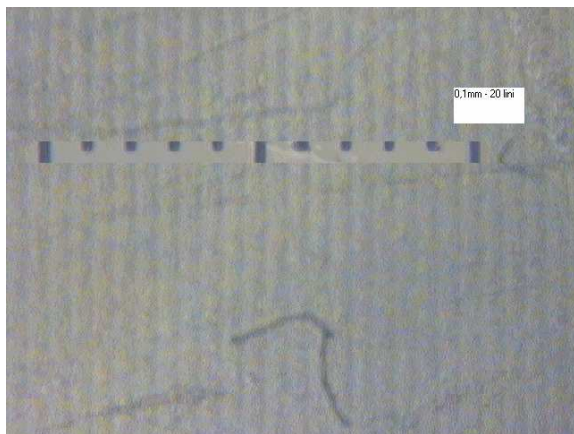
- $4,87 \pm 0,02 \mu\text{m}$, panie profesorze.

- Zapomniała pani o dwóch cyfrach znaczących! {Profesor skarcił nieznacznie Kamilę za zły zapis końcowy} Widzi pani, czas dobiega już końca i należałoby oddać sprawozdanie. Czy zdążyła już pani wszystko opisać?

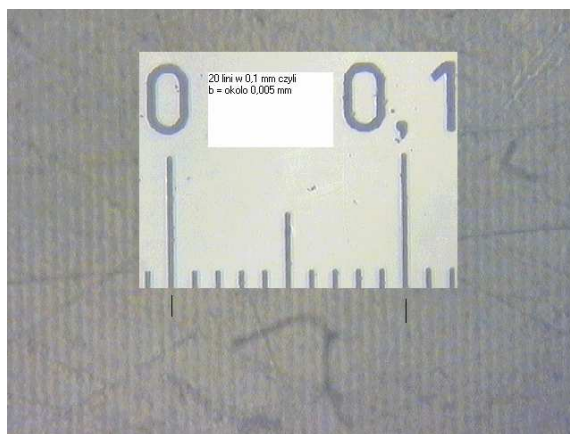
- Jeszcze nie panie profesorze, ale wszystko mam na brudno w notatkach.

Po przeglądnięciu notatek i sprawdzeniu poprawności pozostałych obliczeń profesor wystawił ocenę końcową i z ulgą uciekł przed dalszymi pytaniami na które na razie nie znał odpowiedzi. Skorzystał z tego, że po chwili przyszła kolejna grupa oznajmiając, że skończyła już swoją pracę i chce oddać sprawozdanie. Pod koniec dnia wrócił do laboratorium i powtórzył wszystkie doświadczenia dyfrakcyjno-interferencyjne

nie znajdując jednak żadnego wyjaśnienia dla pozostawionych przez Kamilę wątpliwości. Zwątpiwszy nieco w poprawność wykładanej przez siebie i innych teorii postanowił przynajmniej sprawdzić, czy wyprowadzony przez Younga wzór, wykorzystywany do obliczenia stałej siatki dyfrakcyjnej jest aby na pewno poprawny. W tym celu zabrał siatkę do laboratorium, gdzie mógł wykonać jej zdjęcia w mikroskopie optycznym dla różnych powiększeń, które następnie porównał ze zdjęciami wzorca wykonanymi na tym samym sprzęcie (Rys. 4, Rys. 5).



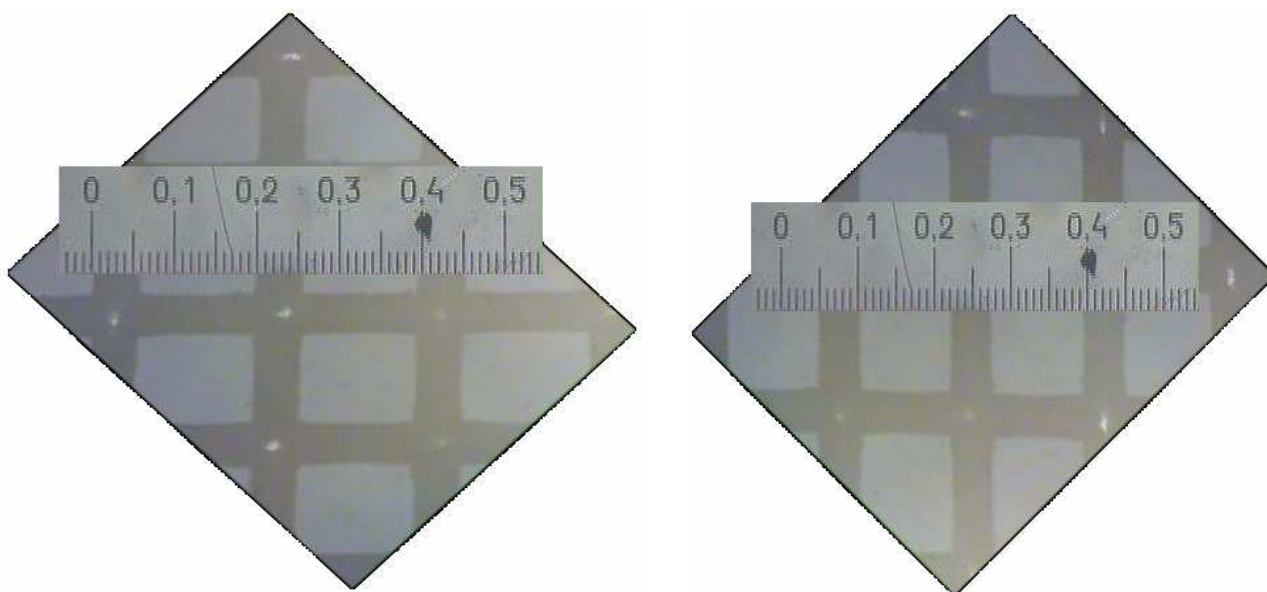
Rys. 4 Zdjęcie siatki dyfrakcyjnej pod mikroskopem optycznym z zaznaczoną skalą pochodzącą od wzorca.



Rys. 5 Zdjęcie siatki dyfrakcyjnej pod mikroskopem optycznym z zaznaczoną skalą pochodzącą od wzorca.

Na podstawie zrobionych zdjęć profesor określił, że na 0,1mm wzorca przypada 20 linii siatki dyfrakcyjnej zarówno na jednym jak i drugim zdjęciu. Stwierdził więc, że odległość między szczelinami dla laboratoryjnej siatki dyfrakcyjnej wynosi około $5\mu\text{m}$. Pamiętając, że wynik obliczeń z doświadczenia studentki wyszedł poniżej tej wartości, wrócił do laboratorium i samodzielnie przeprowadził własne pomiary i obliczenia raz jeszcze. Uzyskał nieznacznie zaniżony wynik, ale bliski temu co wynikało z pomiarów pod mikroskopem. Zaistniała, stosunkowo niewielką rozbieżność wyników spokojnie można było zwać na karby małej dokładności sprzętu pomiarowego znajdującego się w laboratorium studenckim. Tym samym uznał on uzyskany wynik jako zgodny z odległością szczelin w siatce dyfrakcyjnej!!!

Następnie profesor zainteresował się siatką od monitora. W końcu, nie wiedział co odpowiedzieć na pytanie Kamili. Jaką wielkość fizyczną możemy określić na podstawie rozważań Younga w odniesieniu do siatki od monitora? Wykonał więc zdjęcie siatki oraz wzorca jak poprzednio, tylko przy użyciu innego obiektywu (Rys. 6). Wykonane zdjęcie wykazało, że siatka składa się z włókien o średnicy wynoszącej $\sim 50\mu\text{m}$ w obu kierunkach. Natomiast odległość między włóknami na zdjęciu była inna w zależności od mierzonego kierunku. W pierwszym wypadku wynosiła $\sim 200\mu\text{m}$, a w drugim $\sim 170\mu\text{m}$. Stąd, odległość między krawędziami włókien również była inna i zależała od kierunku pomiaru, odpowiednio: $\sim 140\mu\text{m}$ i $\sim 120\mu\text{m}$.



Rys. 6 Zdjęcie siatki od monitora dla dwóch różnych ustawień siatki. Wzorzec przedstawia skalę w [mm].

Po określeniu wymiarów siatki przy pomocy mikroskopu, przyszedł czas na pomiary położenia prążków interferencyjnych. Profesor zmierzył odległość siatki dyfrakcyjnej od ekranu oraz wykonał pomiary położenia pięciu najbliższych prążków (rzędów od $k=1$ do $k=5$). Następnie obliczył odległości między szczelinami ze wzoru wyprowadzonego na bazie założeń falowej interpretacji Younga i oszacował niepewności uzyskanych wyników (Tab. 1). Przy czym długość fali światła lasera została określona na podstawie pomiarów położenia prążków interferencyjnych dla siatki dyfrakcyjnej o znanej odległości między szczelinami (600 linii na 1mm).

Tab. 1 Obliczenia prążków dla siatki od monitora dla dwóch prostopadły kierunków.

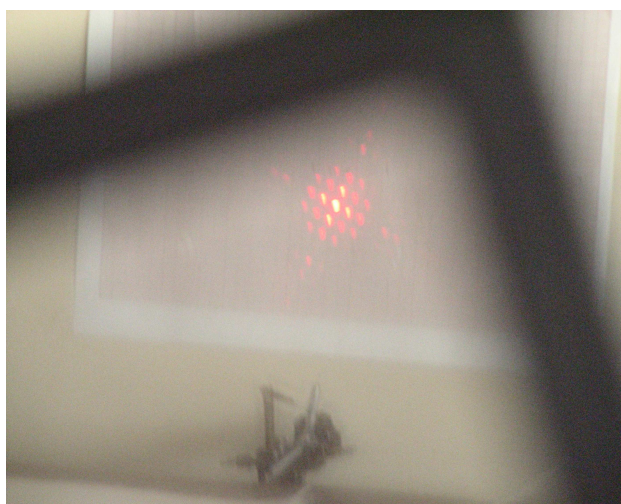
$\lambda = 639,6 \pm 1,1\text{nm}$					$L = 1575 \pm 2\text{mm}$		$S_x = 2\text{mm}$		
odpowiedź dla położenia poziomego				błąd	odpowiedź dla położenia pionowego				błąd
k	x	d	Sd	względny	k	x	d	Sd	względny
	[mm]	[mm]	[mm]	[%]		[mm]	[mm]	[mm]	[%]
1	13	0,155	0,024	15,5	1	11,5	0,175	0,034	19,4
2	25	0,1612	0,0067	4,2	2	22,5	0,1791	0,0092	5,1
3	37	0,1634	0,0033	2,0	3	33	0,1832	0,0045	2,5
4	50	0,1612	0,0023	1,4	4	45	0,1791	0,0028	1,6
5	62	0,1625	0,0023	1,4	5	55,5	0,1815	0,0025	1,4

Uzyskane wyniki mimo istniejących stosunkowo niewielkich rozbieżności (wyniki zaniżone w stosunku do tych uzyskanych na podstawie zdjęcia) najbardziej odpowiadają odległościom między drucikami w siatce. Przy czym odległość ta jest tym samym, co odległość między szczelinami (otworami w siatce). Wynik ten bardzo zadowolił profesora, gdyż był bliski temu, czego należało się spodziewać z obowiązującej falowej interpretacji. Niemniej jednak, obraz interferencyjny na ekranie wykazuje również zmianę amplitudy poszczególnych prążków, która ujawnia istnienie jakby dodatkowych „większych” prążków (Rys. 3). Położenie środków tych „prążków” jest nieco trudniej ustalić, gdyż dotyczy położenia maksimum amplitudy zmian natężenia normalnych prążków, a nie położenia prążka interferencyjnego jako takiego. Dlatego dokładność położenia tak określonej wielkości profesor przyjął jako $S_x = \pm 5\text{mm}$ i określił tą wielkość jako: $\sim 55\text{mm}$ i $\sim 100\text{mm}$ w obu kierunkach (w pionie i poziomie), odpowiednio dla pierwszego i drugiego rzędu ($k=1$ i $k=2$). Dla tak określonych parametrów z falowych rozważań Younga uzyskujemy następujące

odległości między „szczelinami”: odpowiednio $d_1 = 0,03664 \pm 0,00088\text{mm}$ $d_2 = 0,04032 \pm 0,00084\text{mm}$, odpowiednio dla $k = 1$ i $k = 2$. Tym samym, uzyskana w ten sposób z obliczeń wielkość najbardziej odpowiada grubości drucika, która jest taka sama dla pomiaru w obu kierunkach. Wątpliwości profesora wzbudziły w tym momencie dwie kwestie: Po pierwsze, „duże prążki” wcale nie są prążkami, tylko określają położenie maksimum zmiennej amplitudy zwykłych prążków. W tym momencie rozważania w myśl interpretacji Younga niekoniecznie muszą być słuszne dla wyjaśnienia tych dodatkowych efektów w obrazie interferencyjnym. Po drugie, uzyskany wynik $\sim 40\mu\text{m}$, a tym bardziej $\sim 37\mu\text{m}$ znacząco odbiega od średnicy drucika, który wynosi $\sim 50\mu\text{m}$ (z pomiaru uzyskanego za pomocą mikroskopu).

Nie będąc do końca pewnym słuszności zastosowania rozważań falowych Younga do wyciągania wniosków dotyczących położenia „dużych prążków”, profesor uznał, że uzyskana rozbieżność nie stanowi jeszcze problemu, gdyż położenie „dużych prążków” wymaga być może innego wytłumaczenia i niekoniecznie musi zależeć od średnicy drucików w siatce. Natomiast pozostałe wyniki, mimo że nieznacznie rozbieżne z wynikami uzyskanymi za pomocą mikroskopu, odpowiadają rozważaniom teoretycznym Younga. W tym przypadku położenie prążków uzależnione jest od odległości między szczelinami w siatce. Zarówno tej od monitora, jak i dyfrakcyjnej. Następnie słusznie uznając, że uzyskane niewielkie rozbieżności w pomiarach odległości między szczelinami mogą być wynikiem stosunkowo niedużej dokładności laboratoryjnego sprzętu, spokojnie udał się do domu na zasłużony odpoczynek, nie zajmując się dalej tą kwestią.

Minął tydzień, podczas którego Kamila zarzucona innymi obowiązkami nie miała czasu głębiej zastanowić się nad wątpliwościami jakie trapiły ją od ostatnich zajęć. Tym razem trafiła na wyjątkowo proste i nudne ćwiczenie, które wykonała bardzo szybko. Chcąc już wychodzić z sali jeszcze raz zerknęła na doświadczenie Younga i zauważyła, że nikogo tam nie ma. Najwyraźniej byli nieprzygotowani, pomyślała sobie. Usiadła przy stanowisku i włączwszy laser zaczęła bawić się siatką od monitora. Tak jak poprzednio obraz pojawiał się i znikał kiedy wkładała i odsuwała siatkę od wiązki światła laserowego. W pewnym momencie, kiedy obserwowała pojedynczą plamkę światła laserowego rozproszonego na gołej ścianie, machnęła siatką przed oczami i zobaczyła obraz, który ją wielce zdumiał i zaskoczył. Było to coś, czego nie było w instrukcji, jak również profesor tydzień temu o tym nie wspominał. Oglądając pojedynczą plamkę na ekranie przez siatkę, można zobaczyć taki sam obraz jak w normalnym doświadczeniu (Rys. 7).



Rys. 7 Obraz stołu laboratoryjnego z szyną optyczną na której zamocowano laser widziany przez siatkę „ochroną” od komputera. Po prawej stronie powiększenie wykonanego zdjęcia.

Przypadek ten sprawił, że postanowiła znacznie głębiej poznać sprawę i zaczęła sprawdzać powstałe obrazy nie tylko dla siatki od monitora, ale i siatki dyfrakcyjnej. Dość szybko przekonała się, że siatka

dyfrakcyjna daje dokładnie taki sam efekt. Ponieważ oko ludzkie to nie to samo, co duży ekran bądź ściana na której wcześniej obserwowała takie same obrazy. Zastanawiała się, czy aby na pewno ma do czynienia z tym samym zjawiskiem. Czy może zbieżność widzianych obrazów jest tylko przypadkowa? Dalsza zabawa z siatkami doprowadziła ją do odkrycia, że warunek spójności światła jednak nie jest konieczny do powstania tzw. „obrazu interferencyjnego”! O ile w podstawowym doświadczeniu na siatkę padało światło wyemitowane bezpośrednio z lasera, a dopiero później na ekran. Warunek ze spójnością światła dał się łatwo obronić. Kamila rozumowała; o ile światło padające na ścianę jest spójne (skoro wychodzi bezpośrednio z lasera), o tyle rozpraszając się na chropowatej ścianie o chropowatości znacznie przekraczającej długość fali światła, nie może ono dalej zachować swojej spójności. Ponieważ obserwowane przez siatkę obrazy niczym nie ustępowały od tych widzianych w doświadczeniu podstawowym, przeto uznała to za dowód, że spójność światła nie jest jednak potrzebna do powstania obserwowanych zjawisk.

Tego dnia Kamila nie była w stanie zasnąć. W jej głowie kłębiło się wiele pytań i myśli. Postanowiła następnego dnia odwiedzić profesora i opowiedzieć mu o swoim nowym odkryciu. Przecież Young nie bez powodu zastosował pierwszą przesłonę z pojedynczą szczeliną {Która i tak nie daje takiego rozkładu światła (dyfrakcji) jak zakładał, bo przecież pojedyncza szczelina też daje na ekranie prążki!}. Również, nie bez powodu w książkach jest napisane, że należy bezwzględnie zastosować źródło światła spójnego (np. laser), aby uzyskać „obraz interferencyjny”. Zresztą, profesor na zajęciach wyraźnie jej wytłumaczył, dlaczego jest to tak ważne założenie. Tymczasem obserwuje obrazy, które według jej własnego rozumienia powstają dla światła niespójnego.

Profesor wysłuchał Kamilę ze sporym zainteresowaniem, choć nie bardzo wierzył w jej słowa. Ostatecznie zaproponował spacer do pracowni laboratoryjnej, gdzie mogli spokojnie wszystko obejrzeć i wspólnie przemyśleć. Po dotarciu na miejsce Kamila pośpiesznie włączyła laser i na ścianie ukazała się świecąca czerwona plama.

- Panie profesorze, proszę spojrzeć. Jak przetniemy wiązkę laserową siatką od monitora to otrzymujemy obraz interferencyjny (Rys. 3). W myśl tego co pan mi ostatnio tłumaczył możemy powiedzieć, że na siatkę pada spójna wiązka światła laserowego, gdzie na szczelinach siatki światło ulega dyfrakcji, a na ekranie interferencji. Tak więc, jeśli pominiemy wcześniejsze zastrzeżenia odnośnie tego, że mamy pojedynczy obraz, to wszystko wydaje się w porządku odnośnie falowej interpretacji Younga. Kluczowe założenie o spójności światła jest spełnione poprzez zastosowanie w doświadczeniu światła laserowego.

Po odsunięciu siatki od wiązki laserowej, na ekranie ponownie pojawił się pojedynczy punkt od rozproszonego światła. Kamila kontynuowała rozmowę.

- Powinniśmy sobie zadać wpierw następujące pytanie. Co możemy powiedzieć o świetle padającym na ścianę?

- Jest monochromatyczne (jedna barwa - dana długość fali bądź częstotliwość - inaczej kolor) oraz spójne.

- A co możemy powiedzieć o świetle, które dociera do nas do oka?

- Nadal jest monochromatyczne, ale czy jest spójne? Z pewnością nie.

- No właśnie, W takim razie proszę spojrzeć na świecąca plamkę światła przez siatkę od monitora.

Profesor wziął do ręki siatkę i przytrzymał ją przed swoją twarzą. W tym momencie jego oczom ukazał się taki sam widok, jak dla doświadczenia podstawowego (Rys. 7). Obraz był mniej intensywny, ale to akurat nie było zaskoczeniem, gdyż mniej światła przy większej odległości mogło dotrzeć do jego oczu. Istotne było raczej pytanie; Dlaczego w ogóle widzi obraz „interferencyjny”, skoro światło przechodzące przez siatkę nie było spójne? Czyżby obserwowane zjawisko jednak nie wymagało spójności światła? Ale na to przecież nie mógł się zgodzić. Wiedział doskonale jakie to może mieć konsekwencje dla obowiązującej obecnie teorii światła i nie tylko.

- A może światło rozproszone na ścianie jest przynajmniej częściowo spójne? {Oznajmił i pomyślał, że to co widział można by nadal tłumaczyć przy pomocy falowej interpretacji.}

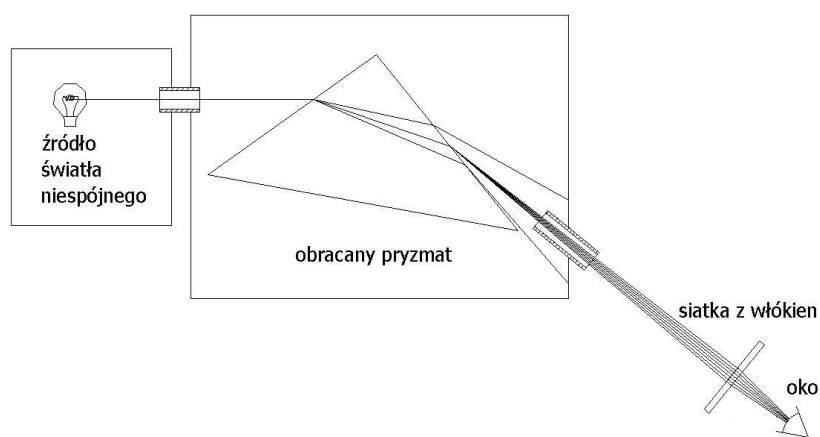
- Nie pomyślałam o tym, że światło rozproszone na ścianie mogło jednak zachować część pierwotnej spójności. Wówczas, faktycznie problemu by nie było. A gdyby tak zamiast plamki od lasera obejrzeć przez siatkę inne źródło światła, które na pewno nie byłoby spójne. {zaproponowała Kamila}

- To nie jest głupi pomysł. Myślę, że moglibyśmy użyć zwykłej żarówki. Światło termiczne jakie emituje żarówka jest ewidentnie niespójnym źródłem światła. W tym przypadku poszczególne atomy rozżarzonego włókna niezależnie wypromieniowują kwanty światła w różnych kierunkach i w różnym momencie czasu. O czym już wcześniej wspominałem.

- Już wiem! {krzyknęła Kamila} Wczoraj przerabiałam na zajęciach doświadczenie polegające na badaniu fotorezystora i fotodiody. W tym doświadczeniu wykorzystuje się żarówkę jako źródło światła białego, które następnie jest przepuszczane przez obracany pryzmat. W ten sposób zmienialiśmy długość fali światła, które padało na badane fotooporniki itd.. Ponieważ w tym doświadczeniu używaliśmy żarówki o sporej intensywności świecenia, po przepuszczeniu przez pryzmat niemal monochromatyczne światło powinno być nadal dość intensywne.

- To do dzieła, sprawdźmy co wyjdzie!

Przygotowali więc stanowisko, tak jak to przedstawiono na rysunku (Rys. 8).



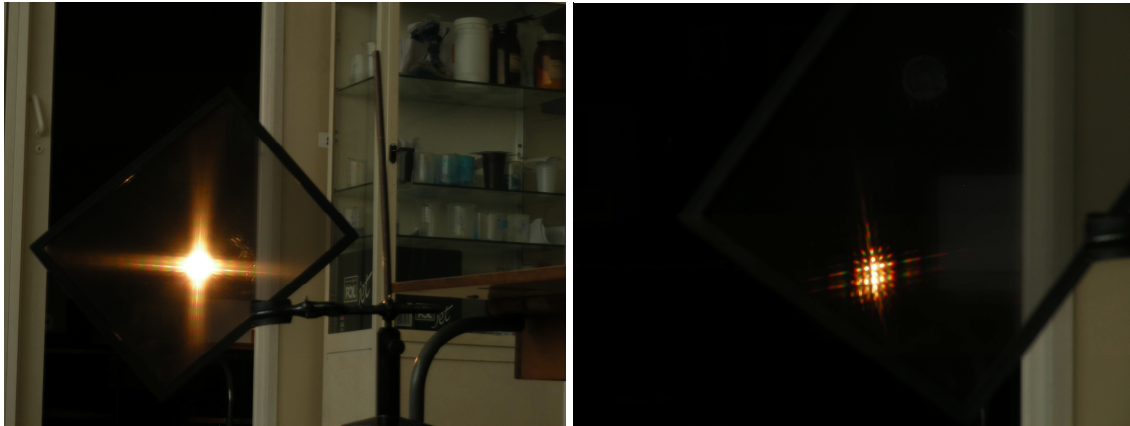
Rys. 8 Schemat doświadczenia ukazującego powstawanie prążkowych obrazów dla światła całkowicie niespójnego (żarówka).

Ponieważ przygotowanie stanowiska wymagało praktycznie jedynie włączenia zasilania żarówki i zdjęcia fotorezystora, zajęło im to zaledwie chwilę. Kamila szybko ustawiła barwę światła na czerwoną, by na ekranie wyglądała jak plamka pochodząca z rozproszonego światła laserowego. Profesor wziął siatkę i przyłożył ją sobie przed twarzą. Nie krył wielkiego zdumienia, kiedy okazało się, że widzi to samo co wcześniej i to znacznie wyraźniej, gdyż światło docierające do niego było znacznie intensywniejsze. Teraz nie dało się już ukryć, że nawet, jeśli wcześniejsze przypuszczenia o częściowej spójności byłyby słuszne. To w tym momencie i tak nie miało to już żadnego znaczenia, gdyż widziany przez niego obraz pochodził od światła ewidentnie niespójnego (żarówki). Oznaczać to mogło tylko jedno, że warunek spójności dla obserwowanego przez nich zjawiska (doświadczenie Younga) nie jest konieczny!

Kamila widząc wielce zaskoczoną minę profesora szybko domyśliła się, że widzi on coś ciekawego. Szybko podeszła do niego i poprosiła o siatkę, by też zobaczyć to co i on. Oboje szybko zrozumieli wagę tego odkrycia, ale nie wiedzieli jeszcze co o tym myśleć. Tymczasem Kamila podbiegła do pryzmatu i zaczęła nim obracać prosząc profesora by raz jeszcze spojrzął przez siatkę. Tym razem zmieniały się kolory, ale obraz był ciągle taki sam. Jedynie nieznacznie zmieniała się jego wielkość, co związane jest z tym, że światło o różnej barwie jest rozpraszane na siatce pod nieco innym kątem.

- A co będzie jeśli spojrzymy przez siatkę bezpośrednio na żarówkę? {zapytała Kamila}

- Zobaczmy (Rys. 9). {odparł profesor}.



Rys. 9 Obraz żarówki widzianej przez siatkę od ekranu.

- Ale piękne obrazy! {krzyknęła z zachwytu Kamila} Teraz wszystkie kolory tworzą wspólny niezwykle barwny obraz. {Który dla światła o dużej intensywności daje przepiękny wizualnie efekt}. Ale i dla małej żarówki wyraźnie widoczne są poszczególne kolorowe prążki.

- Nie ma wątpliwości. {Oznajmił przyciszonym głosem profesor} Zjawisko opisane przez Younga nie wymaga światła spójnego i będzie to miało dalekosiężne konsekwencje dla wytłumaczenia tego zjawiska.

- Czy nie da się tego zjawiska wyjaśnić, nie wykorzystując założenia o spójności światła?

- Obawiam się, że nie za pomocą falowej interpretacji. Stosowana obecnie interpretacja tego zjawiska jest w zasadzie czysto geometryczna i sprowadza się do liczenia różnicy dróg jakie pokonują poszczególne wiązki światła. Jeśli obserwowane obrazy są stabilne, a układ nie ulega jakimkolwiek zmianom, to zmuszeni jesteśmy przyjąć, że różnice dróg również nie ulegają zmianie w czasie. Wówczas, warunki początkowe dla światła muszą być dobrze zdefiniowane. W praktyce oznacza to konieczność założenia, że światło jest spójne. Gdyby było inaczej i fazy wiązek światła wychodzących z poszczególnych szczelin zmieniałyby się losowo, to wówczas na ekranie też mielibyśmy przypadkowe przesunięcia w fazie, niezależnie od rzeczywistej różnicy dróg. W taki sposób nie da się już wytłumaczyć istnienia minimów i maksimów natężenia światła na ekranie w oparciu o interferencję. Odkrycie, że w doświadczeniu Younga spójność światła nie jest konieczna, oczywiście nie zaprzecza istnienia tego zjawiska, ale oznacza odrzucenie interpretacji falowej bazującej na interferencji. A przynajmniej na takiej interpretacji, która bazuje na liczeniu dróg między „źródłami światła” a punktem na ekranie (punktem obserwacji światła). Obawiam się, że nie jest możliwe wytłumaczenie tego zjawiska w obecny sposób, nie odwołując się do założenia o spójności światła. Niemniej jednak, jeszcze zagłębę do kilku książek i upewnię się, co dokładnie na ten temat można przeczytać w podręcznikach?

- Póki co, to rozumiem, że można poprosić kogoś np. jakiegoś profesora na wykładzie z fizyki o wytłumaczenie tego zjawiska, a wykładowca i tak będzie musiał wcześniej czy później skorzystać z założenia o spójności światła, jeśli będzie chciał wykorzystać tłumaczenie bazujące na interpretacji falowej. Następnie wystarczy konsekwentnie pilnować, by nie mógł z tego założenia skorzystać, gdyż jest ono sprzeczne z doświadczeniem (zjawisko ma miejsce również dla światła absolutnie nie spójnego) i ostatecznie wykładowca nie będzie umiał wytłumaczyć tego zjawiska w ten sposób.

- Jeśli będzie trzeba w założeniach przyjąć, że początkowa faza w pojedynczej szczelinie bądź dla wielu szczelin (bądź rozkład fazy, gdy faza zmienia się o określoną wielkość przy przejściu od jednej do kolejnej szczeliny) jest dobrze ustalona, to owszem, może być problem z wyjaśnieniem tego zjawiska na gruncie teorii falowej. Gdyż faktycznie będzie to sprzeczne z założeniem o spójności światła, jeśli doświadczenie pokazuje, że światło spójne być nie musi. Czyli, nie wolno nam zakładać, że znamy rozkład faz fali świetlnej na wejściu do szczelin. Co więcej, mamy wręcz obowiązek w takim przypadku założyć, że faza wejściowa jest przypadkowa w przestrzeni i zmienna przypadkowo w czasie. Niemniej jednak należy zauważyć, że światło może być spójne w jakimś obszarze. Np. dla wiązki światła niespójnego wychodzącego z

nieśpójnego źródła możemy założyć, że pewien mały wycinek takiej wiązki będziemy mogli już traktować jako spójny. Wówczas, nawet jeśli założymy, że źródło światła jest nieśpójne jako całość, to możemy założyć, że przynajmniej dla poszczególnych szczelin mamy światło spójne i na tej podstawie nadal moglibyśmy próbować korzystać z rozważań Younga.

- Chyba nie bardzo mogę się z panem zgodzić, panie profesorze. {Po dłuższym namyśle Kamila kontynuowała}

Po pierwsze, źródło światła pochodzenia termicznego jest emitowane w postaci fotonów, które są między sobą zupełnie nieśpójne. Wobec czego, lokalna spójność wychodzącej z takiego źródła wiązki światła dotyczyłaby tylko pojedynczego fotonu. Skoro tak, to niech mi pan powie profesorze, jaki jest przekrój poprzeczny pojedynczego fotonu i czy aby na pewno zmienia się on wraz z odległością od źródła? W co absolutnie nie wierzę, gdyż byłoby to zupełnie niedorzeczne i sprzeczne z doświadczeniami w których odbiera się pojedyncze fotony niezależnie od odległości źródła światła od detektora i wielkości detektora. Poszczególne fotony nie są w tej samej fazie i powinny się częściowo znosić, gdyby ich przekroje poprzeczne były znaczne i się nakładały na siebie. A w przypadku nałożenia się wielu fotonów w ogóle nie powinniśmy nic widzieć dla nieśpójnych źródeł światła, co jest sprzeczne z doświadczeniem, gdyż światło pochodzące od źródła nieśpójnego jednak widzimy i to nawet wówczas gdy jest bardzo intensywne. Czyli wtedy, gdy powinno nakładać się na siebie wiele nieśpójnych fotonów.

Po drugie, rozmiar szczelin dla siatki od monitora jest dość znaczny w stosunku do długości fali (Długość fali to nie to samo co wielkość poprzeczna fotonu !!!). W związku z czym, na pewno należałoby założyć, że dla poszczególnych szczelin w przypadku siatki od monitora mamy inną fazę. W zasadzie, to wielkość szczelin jest tak duża, że pomiędzy krawędziami szczeliny siatki zmieści się wiele fotonów. Stąd, nawet dla pojedynczej szczeliny należy przyjąć, że między jej krawędziami mamy wiele źródeł fal o przypadkowej fazie. Tak przypadkowej, jak przypadkowa jest faza emitowanych fotonów na poszczególnych atomach nieśpójnego źródła światła. Dlatego panie profesorze, nie ma mowy, aby zgodzić się z pana rozumowaniem jakoby nadal można było traktować ten przypadek, tak jakby światło było spójne. Jeśli źródło światła jest w pełni nieśpójne, to jest w pełni nieśpójne ze wszystkimi tego konsekwencjami. W innym przypadku możemy się spierać w nieskończoność, jaki jest przekrój poprzeczny pojedynczego fotonu. Nawet jeśli założymy, że wielkość szczeliny jest tak mała, że może przez nią przejść tylko jeden foton. Wówczas taki foton jeśli dotrze do ekranu i spotka inny foton pochodzący od innej szczeliny, to będzie miał względem niego przypadkowe przesunięcie w fazie. Wynik sumowania fal takich fotonów będzie miał przypadkową wartość niezależnie od miejsca padania tych fotonów na ekran. Co więcej, takie fotony musiałyby zostać wyemitowane w innym czasie. Szczególnie dla prążków bardziej oddalonych od środka obrazu i bardziej oddalonych szczelin.

- Faktycznie, można by długo dyskutować dla jakiej części przekroju wiązki nieśpójnej mamy wiązką spójną. Niedawno sprawdzałem pod mikroskopem wielkość siatki użytej w tym doświadczeniu. Uzyskałem następujące wyniki: odległość między krawędziami wyniosła $120\mu\text{m}$ i $140\mu\text{m}$ w zależności od kierunku pomiaru. Natomiast odległość między szczelinami wyniosła: $170\mu\text{m}$ i $200\mu\text{m}$. Jeśli założymy, że wielkość przekroju pojedynczego fotonu jest związana z jego długością fali, to dla światła użytego w tym doświadczeniu zmieściłoby się między krawędziami $130\mu\text{m}/640\text{nm}(\text{Tab.1}) \approx 200$ szerokości fotonów. Czyli wystarczająco dużo, by nie wolno było zakładać, że w obszarze jednej szczeliny mamy jedną fazę! Tym bardziej nie wolno zakładać, że mamy taką samą fazę dla poszczególnych szczelin. A więc rzeczywiście mamy problem, bo nie możemy korzystać z założenia o spójności światła!

- Wydaje mi się panie profesorze, że założenie iż przekrój poprzeczny fotonu jest rzędu długości fali jest dość naiwny. Już stwierdziliśmy wcześniej, że światło emitowane w żarówce jest źródłem termicznym dla którego zakładamy, że każdy atom emituje światło w przypadkowym kierunku i fazie. To sugeruje, że przekrój poprzeczny fotonu powinien być raczej rzędu wielkości atomu. Szczególnie jest to rozsądne, gdy uwzględnimy absorpcję fotonów przez pojedyncze atomy, bądź złożone cząsteczki, np. barwniki.

- Istotnie, nie pomyślałem o tym. Absorpcja fotonów przez barwniki bądź inne centra absorpcji sugerują, że foton musiałby być znacznie bardziej zlokalizowany przestrzennie niż to, co sugeruje taki parametr jak długość fali. W takim przypadku należałoby przyjąć, że wielkość poprzeczna fotonu jest co najwyżej tego rzędu wielkości co duża pojedyncza cząsteczka barwnika, jeśli nie mniejszy. A to oznacza, że tym bardziej

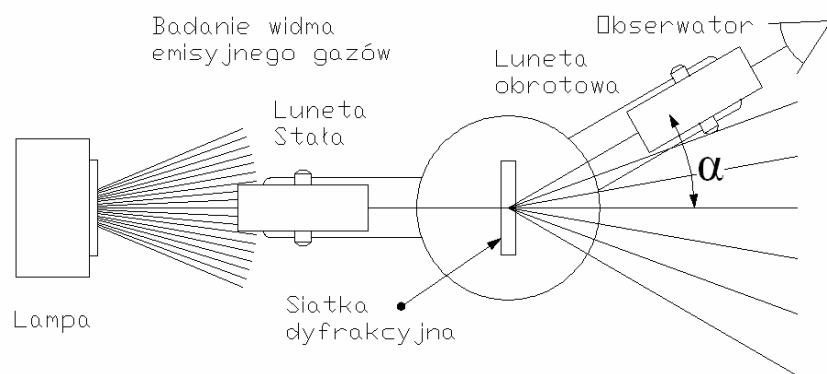
nie możemy zakładać, że uzyskaliśmy wiązkę spójną po przepuszczeniu światła przez pojedynczą szczelinę, a co dopiero dla wielu szczelin. Gdyż ilość nie nakładających się na siebie niespójnych fotonów jaka mieściłaby się między krawędziami szczeliny byłaby znacznie większa niż to co oszacowałem wcześniej. Szczególnie, że nie bardzo potrafię sobie wyobrazić mechanizmu pełnej absorpcji fotonu przez atom bądź związek chemiczny, dla którego przekrój poprzeczny absorbowanego fotonu byłby wielokrotnie większy niż wielkość cząsteczki, która go absorbuje. Oczywiście jest to problem w przypadku czysto mechanicznego podejścia do omawianych zagadnień i dalszego trzymania się falowych ustaleń dotyczących natury światła. Mechanika kwantowa na to pozwala, ale za to bez podania jakiegokolwiek sensownego mechanizmu. Oczywiście mechanizmu rozumianego w sensie klasycznym. Czyli takiego, który można sobie w pełni wyobrazić.

- {Po dłuższej chwili milczenia, Kamila spoglądając nerwowo na zegarek postanawia zakończyć rozmowę} Czy mogę po zajęciach popracować w laboratorium? To zagadnienie jest niesamowicie frapującą zagadką, może uda mi się jeszcze znaleźć jakieś inne nieścisłości. Bardzo proszę profesora.

- Dobrze pani Kamilo, proszę sprawdzić kiedy pracownia jest wolna. Klucze będzie mogła pani wziąć z portierni. Jak by były jakieś problemy, to proszę się powołać na moją osobę. Wówczas dadzą pani klucze do laboratorium. Ja tymczasem poszperam w literaturze i spróbuję znaleźć jakieś wyjaśnienie. Może ktoś już coś podobnego wytłumaczył. Na razie muszę dobrze przemyśleć to co dziś się wydarzyło.

- Dobrze panie profesorze. Jak tylko będę mogła, pojawię się w pracowni i spróbuję jeszcze raz przeprowadzić te doświadczenia. Może znajdzie coś nowego.

Przez najbliższe dni Kamila nie miała zbyt wiele czasu na rozważania. Niemniej jednak przygotowując się do kolejnych zajęć laboratoryjnych dokonała ciekawego odkrycia. Gdyby nie ostatnia wizyta w laboratorium pewnie nie zwróciłaby na to uwagi. Ale ostatnie wydarzenia wyostrzyły jej czujność odnośnie doświadczeń dotyczących światła. Tym razem miała do wykonania doświadczenie polegające na badaniu widma emisyjnego gazów. Ku swojemu zdziwieniu zauważyła, że widnieje tam wzór na stałą siatki dyfrakcyjnej, który jest identyczny z używanym w doświadczeniu Younga. Tyle, że tym razem nie było w tekście mowy o jakimkolwiek założeniu, że światło ma być bezwzględnie spójne. Tymczasem, schemat stanowiska laboratoryjnego umieszczony w skrypcie laboratoryjnym utwierdził ją w przekonaniu, że faktycznie ma do czynienia dokładnie z tym samym doświadczeniem (Rys. 10).

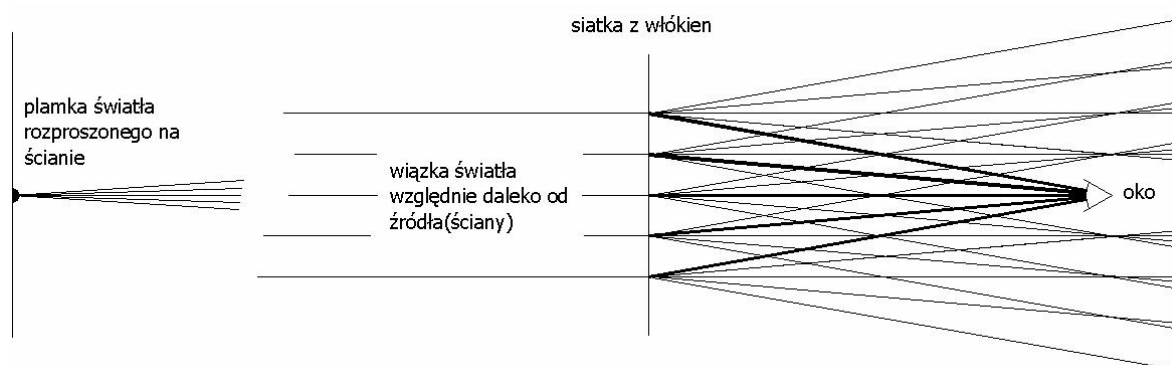


Rys. 10 Schemat doświadczenia do badania widma emisyjnego gazów przy zastosowaniu siatki dyfrakcyjnej.

Schemat doświadczenia nie zawierał lasera jako źródła światła spójnego, a jedynie lampę jarzeniową, która emitowała światło zdecydowanie nie spójne. Stąd, nie zdziwiło ją to, że nie wspomniano o koniecznym założeniu (spójności światła) dla wyprowadzenia zastosowanego wzoru. Zresztą, w skrypcie nie było żadnego wyprowadzenia dla wzoru na stałą siatki dyfrakcyjnej. Byłoby to przecież w jawnej

sprzeczności z przeprowadzonym doświadczeniem. Jednocześnie zrozumiała, że doświadczenie jakie przyjdzie jej wykonać, będzie w zasadzie potwierdzeniem jej wcześniejszej obserwacji, że w doświadczeniu Younga spójność światła nie jest potrzebna! Zaczęły ją jednak nachodzić wątpliwości; na ile poprawny jest wzór, który przyjdzie jej w tym doświadczeniu zastosować? Skoro jego wyprowadzenie wymaga założenia, które w tym przypadku nie jest spełnione. Oznaczać to może tylko dwa przypadki: Wzór jest błędny i tylko w przybliżeniu oddaje relacje geometryczne mierzone w tym doświadczeniu albo jest przypadkowo poprawny, ale jego wyprowadzenie wymaga zupełnie innej interpretacji, dla której założenie o świetle spójnym nie będzie wymagane. Z wcześniejszej rozmowy z profesorem wiedziała, że uzyskane pod mikroskopem obrazy dają podobne wyniki na stałą siatki, jak te uzyskane z pomiarów położenia prążków na ekranie i oparte obliczeniami na podstawie rozważań Younga. Ale czy jest to poprawna formuła opisująca rzeczywiste relacje geometryczne występujące w tym doświadczeniu, czy tylko przypadkowa zbieżność z wynikami doświadczalnymi? Na to Kamila odpowiedzi jednak nie знаła.

Następnego dnia korzystając z kilkugodzinnego okienka postanowiła spędzić nieco czasu w laboratorium w poszukiwaniu kolejnych niespodzianek. Powtórzyła wszystkie wcześniej przeprowadzone doświadczenia i zaglądnęła również do ciemni, gdzie przeprowadza się badanie widma emisyjnego gazów. Przeprowadzone doświadczenia nie wykazały jej niczego nowego, poza tym, że potwierdziły jej wcześniejsze wnioski. W pewnym momencie, kiedy po raz kolejny oglądała prążki „interferencyjne” przez siatkę od monitora, zadała sobie jeszcze raz pytanie. Dlaczego obraz, który widzę na siatce od monitora jest taki sam jak na ścianie? Przecież moje oczy to nie ściana, na której widzę rozproszone światło w kilku miejscach. Natomiast widzę taki sam obraz jaki zobaczyłabym na ekranie. Tyle, że widzę go na siatce. Zastanawiając się nad tym zrozumiała, że to co widzi to światło, które zostało rozproszone na różnych częściach siatki, które dochodzi do jej oczu z różnych jej obszarów. Znając z wcześniejszego schematu (Rys. 10) zachowanie się światła przechodzącego przez siatkę szybko zrozumiała jak wyjaśnić obserwowane obrazy (Rys. 11).



Rys. 11 Schemat przedstawiający mechanizm powstawania obrazu na siatce.

Z tych rozważań wynika, że powstawaniu tego obrazu towarzyszy dokładnie takie samo zjawisko, jakie ma miejsce podczas podstawowego doświadczenia Younga (dyskretne rozszczepienie wiązki świetlnej padającej na siatkę). Gdy siatka jest blisko światła rozproszonego na ścianie, wówczas światło pada na siatkę pod różnymi kątami. Z doświadczenia Younga z obracaną siatką Kamila wiedziała już, że jeśli światło pada na siatkę pod innym kątem niż kąt prosty, to jest rozproszone pod większym kątem (prążki się rozsuwają na boki). Ten efekt może tłumaczyć to, że dla małej odległości plamka-siatka obraz zaczyna się zlewać w jeden punkt, tak samo jak ma to miejsce podczas zmniejszania odległości między siatką a ekranem w oryginalnym doświadczeniu Younga. Przyjęta interpretacja tłumaczy również, dlaczego widziany obraz zmienia się podczas zmiany odległości między okiem a siatką. Zwiększenie tej odległości prowadzi do zwiększenia się obrazu.

By upewnić się co do swoich wniosków Kamila podeszła do stołu laboratoryjnego i wzięła do ręki siatkę dyfrakcyjną, którą przytrzymała na odległość wyciągniętej ręki, a nie tak jak wcześniej przy oku. Patrząc przez nią na źródło światła punktowego widziała tylko jedną plamkę. Zrozumiała, że jeśli siatka dyfrakcyjna rozprasza światło pod znacznie większym kątem niż siatka od monitora oraz jest od niej znacznie mniejsza, to w tej sytuacji nie widzi pozostałych prążków. Aby je zobaczyć należałoby mieć większą siatkę dyfrakcyjną lub przesunąć tą trzymaną w rękę nieco na bok. Kiedy to zrobiła, zobaczyła pierwszy prążek. Kolejne wymagały dalszego przesunięcia siatki w bok. Co ciekawe, widziała światło na siatce dyfrakcyjnej mimo, że nie patrzyła bezpośrednio na plamkę światła na ścianie. Zrozumiała, że jej wcześniejsze wnioski były poprawne, a przeprowadzony eksperyment tylko je potwierdza. Światło jest rozpraszane na siatce w konkretnych kierunkach i tam leży przyczyna obserwowanego zjawiska. Zabawa z siatką dyfrakcyjną w zasadzie niczego nowego nie wniosła w odniesieniu do siatki od komputera. Mogła się o tym przekonać przesuwając dużą siatkę tak, by widzieć centralny prążek nie na jej środku, ale blisko jej krawędzi. Wówczas zobaczyła tylko część obrazu. Nie widziała tej części, gdzie nie było siatki. W tym obszarze nie było krawędzi, które zakrzywiłyby tor lotu fotonów.

Zaobserwowane tego dnia zjawiska wzmocniły jej przekonanie o braku potrzeby użycia światła spójnego, ale zasadniczo niczego nowego nie wносиły. Zatem ponownie zmontowała podstawowe stanowisko do badania zjawiska Younga i zaczęła się bawić siatką dyfrakcyjną przemieszczając ją to bliżej to dalej od ekranu. Obserwując na ekranie dobrze widoczne, zbliżające się bądź oddalające od siebie plamki rozproszonego na ekranie światła zauważyła, że ich jasność praktycznie się nie zmienia. Zgodnie z powszechnie przyjętą falową interpretacją, jasność plamki świetlnej należy przecież interpretować jako natężenie fali świetlnej. Założyła więc, że w małej odległości od szczelin istnieje jakiś obszar, gdzie nakładają się fale świetlne pochodzące od poszczególnych szczelin w taki sposób, że sumowanie jest konstruktywne i widzimy jasną plamkę. Tyle, że końcowa amplituda tego sumowania zależy bezpośrednio od amplitud (natężenia) poszczególnych fal. Ponieważ prążki powstają pod dużym kątem, wobec czego Kamila założyła, że fala pochodząca z każdej szczeliny również będzie rozchodzić się pod dużym kątem. Co za tym idzie jej energia musi rozkładać się na coraz większy obszar wraz ze zwiększaniem odległości od szczelin. Taki efekt jest nierozdzielnie związany ze spadkiem natężenia fali wraz z odległością i jest ono tym większe im większy jest kąt, pod jakim rozchodzi się fala.

Opisane zjawisko powinno prowadzić do bardzo szybkiego spadku jasności powstałych prążków „interferencyjnych”, ale tak się nie dzieje. Magiczne przeniesienie energii z jednego obszaru ekranu tam, gdzie jest wygaszenie do innego, gdzie mamy wzmocnienie nie wchodzi w rachubę! Podekscytowana nowym odkryciem postanowiła niebawem odwiedzić profesora.

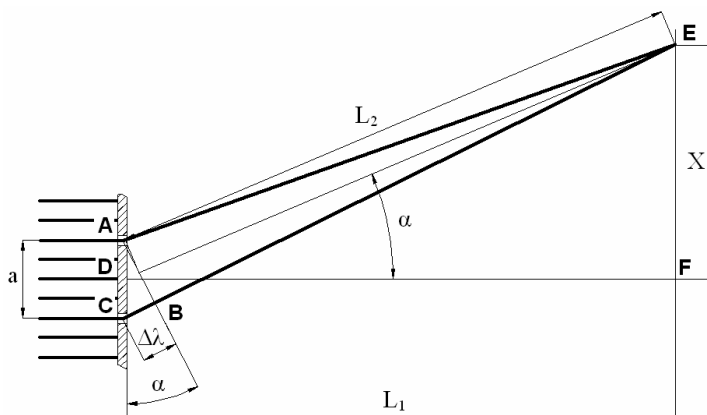
Tymczasem, profesor zaglądnął do podręcznika Mechaniki Kwantowej i przypomniał sobie przytoczone tam rozważania. Ponadto, mając na względzie wcześniejsze uwagi Kamili co do rozważań nad pojedynczą szczeliną postanowił sprawdzić obliczenia falowe dla przypadku, gdy na ekran pada światło z bardzo wielu źródeł (nieskończona ilość szczelin oddalonych na odległość dążącą do zera). Najprościej, byłoby jednak rozpocząć obliczenia dla skończonej ilości szczelin tak, by można było jako parametr ustawiać ich ilość oraz odległość między nimi i zobaczyć jak będzie się zmieniać rozwiązanie (obraz na ekranie) wraz ze zmianą tych parametrów. W efekcie zasymulował działanie siatki dyfrakcyjnej, dla której można przyjąć, że ilość szczelin na które pada światło jest znaczna, a ich odległość niewielka. W książce jak i skrypcie do laboratorium znalazł wyprowadzony już wcześniej wzór na stałą siatki dyfrakcyjnej (odległość między szczelinami).

$$a = \frac{k \cdot \lambda}{\sin \alpha} \quad , \quad (1)$$

gdzie: k jest numerem kolejnego prążka, λ - długość fali, a α - kąt pod jakim ugina się na siatce dyfrakcyjnej obserwowany strumień światła (Rys. 10).

Niemniej jednak zastosowany wzór nie uwzględnia liczby szczelin, a jedynie odległość między nimi. Co więcej, z wzoru wynika, że kiedy odległość między szczelinami dąży do zera, to $\sin(\alpha)$ musiałaby dążyć do nieskończoności. Tak więc, tłumaczenie jakie przytoczył podczas rozmowy z Kamilą w myśl tego wzoru nie miało sensu. Po chwili namysłu przypomniał sobie, że przecież zastosowany wzór jest przybliżony i

wyprowadzony tylko dla dwóch szczelin, a nie ich dużej ilości. Narysował na kartce schemat wyprowadzania stosowanej powszechnie formuły, by bliżej się temu przyjrzeć (Rys. 12).



Rys. 12 Dyfrakcja światła na dwóch szczelinach.

Z rysunku jasno wynika, że trójkąt ABC jest tylko w przybliżeniu trójkątem prostokątnym. By faktycznie można było tak założyć, należałoby przyjąć, że odległość $L_1 \gg a$. Wówczas proste AE i CE są prawie równoległe. Dla takiego założenia można określić $\sin(\alpha)$ jako:

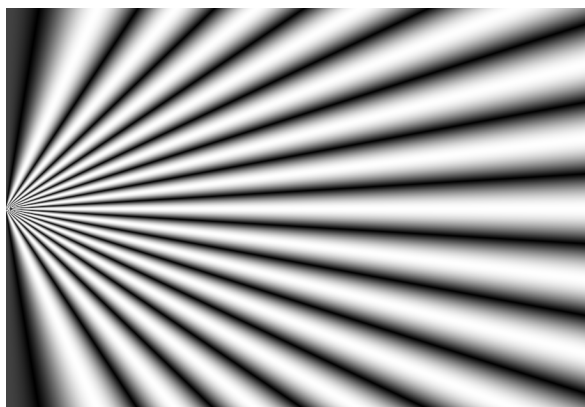
$$\sin \alpha = \frac{\Delta \lambda}{a} \quad (2)$$

$$\sin \alpha = \frac{X}{L_2} \quad (3)$$

Następnie za $\Delta \lambda$ profesor podstawiał $k\lambda$ jako warunek dla których przesunięcie fazowe daje interferencję konstruktywną (wzmocnienie) i uzyskał z wzoru 2 wzór 1. Po chwili przypomniał sobie, że przecież jest to wzór z którego korzysta się przy wyznaczeniu stałej siatki dyfrakcyjnej podczas badania widma emisyjnego gazów. Pamiętając ostatnią rozmowę z Kamilą również zauważył, że przecież dla tego doświadczenia nie jest spełniony warunek spójności światła i zastosowany wzór w zasadzie nie posiada już teoretycznego oparcia. Niemniej jednak miał pewne wątpliwości, gdyż we wspomnianym doświadczeniu korzysta się jeszcze ze szczeliny. Tak, że światło zanim dojdzie do siatki dyfrakcyjnej przechodzi wpierw przez szczelinę, gdzie zgodnie z przyjętą interpretacją falową możemy założyć, że wychodzi już jako światło spójne. Po tym spostrzeżeniu postanowił wykonać dokładniejsze obliczenia dla modelu falowego, nie przejmując się warunkiem o spójności światła.

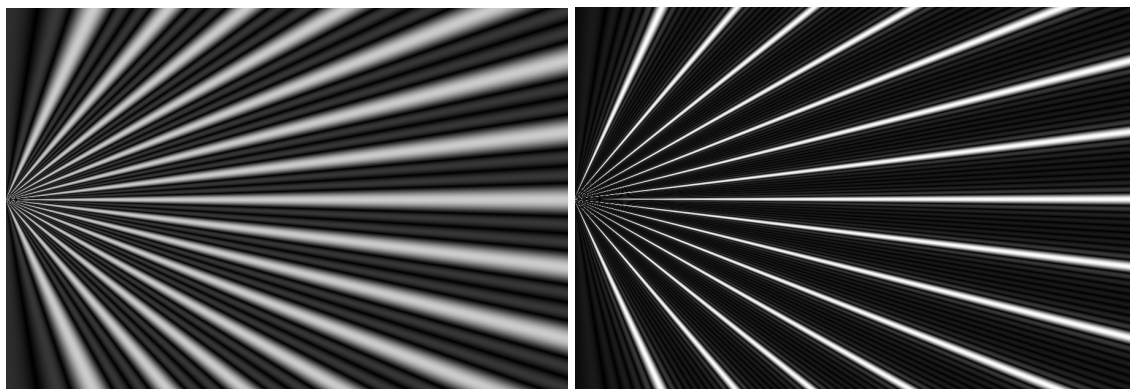
Owszem, mógł jeszcze połączyć wzór 2-gi z 3-cim i przy zastosowaniu prawa Pitagorasa wyprowadzić znaną mu zależność na stałą siatki dyfrakcyjnej jaką stosuje się powszechnie na zajęciach laboratoryjnych poświęconych temu tematowi. Ale, zdając sobie sprawę z przyjętych przybliżeń, postanowił przeprowadzić obliczenia numeryczne. Komputer może przecież policzyć dokładnie przebytą drogę i to dla wielu szczelin {pomyślał}. Stąd, można dokładnie zsumować wszystkie dochodzące na ekran fale, przy uwzględnieniu przesunięcia fazowego liczonego dla każdej szczeliny osobno bez uciekania się do przybliżeń.

Napisanie stosownego programu nie zajęło zbyt wiele czasu. Po wprowadzeniu danych: $L_1 = 1,5\text{m}$, $a = 5\mu\text{m}$ oraz $\lambda = 632,8\text{nm}$ (odpowiadających doświadczeniu z siatką dyfrakcyjną) uzyskał następujący wynik (Rys. 13).



Rys. 13 Rozkład światła dla symulacji falowej doświadczenia Younga dla dwóch szczelin (widziany z góry). W osi poziomej jest odległość od siatki dyfrakcyjnej.

Z przeprowadzonej symulacji od razu zauważył, że wiązki światła wychodzące od siatki są o wiele za szerokie w stosunku do doświadczenia. Pamiętając jednak o tym, że symulacja została wykonana dla zaledwie dwóch szczelin, profesor nie zdziwił się tym wynikiem. W końcu rzeczywisty obraz dla dwóch szczelin zawiera prążki leżące blisko siebie. Natomiast, nadmiar prążków bocznych w symulacji wyjaśnił tym, że w rzeczywistej dyfrakcji natężenie światła nie rozkłada się tak samo we wszystkich kierunkach oraz zmienia się wraz z odległością. W swojej symulacji nie uwzględnił tego faktu. Gdyby to zrobił, {tak rozumował}, powinien spodziewać się coraz mniejszego natężenia jasności dalej położonych prążków. Zadowolony z uzyskanych wyników wykonał kolejne symulacje dla coraz to większej ilości szczelin (Rys. 14).



Rys. 14 Rozkład światła dla symulacji falowej doświadczenia Younga dla pięciu i dziesięciu szczelin (widziany z góry). W osi poziomej jest odległość od siatki dyfrakcyjnej.

Z przeprowadzanych symulacji profesor stwierdził, że położenie poszczególnych prążków nie zależy od ilości szczelin. Natomiast, szerokość widzianych na ekranie prążków interferencyjnych powinna znacząco od tego parametru zależeć. Powinny być one tym węższe, im światło przechodzi przez większą ilość szczelin.

Dalsze rozważania profesora zostały zakłócone pukaniem. Kamila odważnie weszła do środka komunikując, że była w laboratorium i znalazła kolejną niespójność teorii z doświadczeniem. Profesor nieco się uśmiechnął i przywitał.

- Dzień dobry pani Kamilo.

- Dzień dobry panie profesorze. O! to przypomina rysunek ze skryptu. Czy to przedstawia doświadczenie Younga? {Spojrzała na ekran monitora}

- Tak, przygotowałem symulację i jak na razie to niemal zgadza się z doświadczeniem. Co panią dziś sprowadza pani Kamilo?

- Byłam niedawno w laboratorium i zrobiłam kilka doświadczeń. Przyszłam się nieco w ich sprawie poradzić. Ale wpierw chciałam dodać, że przygotowując się do kolejnego zadania zatytułowanego: „*Badanie widma emisyjnego gazów*” zauważyłam, że jest to dokładnie takie samo doświadczenie jak doświadczenie Younga. Tyle, że w tym doświadczeniu mamy lampę jarzeniową zamiast lasera jako źródła światła. Zastosowany w tym zadaniu wzór jest taki sam jak w doświadczeniu Younga, ale nie jest spełnione założenie o spójności światła. Więc nie jestem pewna na ile mogę poważnie traktować podaną w tym zadaniu zależność między stałą siatki, a kątem pod jakim widzimy dany kolor. W każdym razie potwierdza to wcześniejszy wniosek, że dla doświadczenia Younga światło spójne nie jest konieczne.

- Proszę pani. Możliwe, że pani nie zauważyła, że w schemacie tego doświadczenia jest umieszczona przesłona. Oczywiście, że źródło światła zastosowane w tym doświadczeniu jest całkowicie niespójne z wszystkimi tego konsekwencjami o jakich mówiliśmy już wcześniej, ale zanim dojdzie ono do siatki dyfrakcyjnej przechodzi wpierw przez szczelinę. Stąd, formalnie problemu nie ma, gdyż światło w myśl rozważań Younga po przejściu przez szczelinę jest światłem spójnym.

- Widzi pan panie profesorze, ale ja jednak o tym pamiętałam i już to przemyślałam i sprawdziłam. Po pierwsze, w myśl rozważań Younga, światło po przejściu przez szczelinę daje falę kulistą a nie prążki, co jest sprzeczne z doświadczeniem, o czym już mówiliśmy wcześniej. Po drugie, aby było to słuszne (uzyskanie spójnej wiązki światła po przejściu przez pojedynczą szczelinę) wielkość szczeliny pod względem formalnym musiałaby być przynajmniej porównywalna do długości fali. Właśnie w tym miejscu zgodzić się z pana rozumowaniem nie mogę, gdyż będąc niedawno w laboratorium wykonałam wspomniane doświadczenie zmieniając wielkość szczeliny umieszczonej za źródłem światła. Oczywiście, dla wąskiej szczeliny mamy różnokolorowe prążki widziane pod różnymi kątami. Tyle, że te prążki są bardzo wąskie. Kiedy kręciłam pokrętkę regulującą szczelinę, prążki robiły się szersze. Ostatecznie odkręciłam szczelinę maksymalnie jak się dało. Wyszło w przybliżeniu 5mm! Widziane przez „lunetę” prążki poszerzyły się również na około 5-6mm. Czyli widziałam prążki w przybliżeniu tak szerokie jak szeroko była odkręcona szczelina! Bez tej szczeliny też występują prążki, ale są one odpowiednio szerokie, gdyż światło padające na siatkę ma jeszcze większy przekrój (symulacja komputerowa dyfrakcji i interferencji dla wielu szczelin pokazuje odwrotną zależność, Rys. 13 i 14). Tak więc doszłam do wniosku, że szerokość widzianego prążka we wspomnianym doświadczeniu zależy od szerokości padającej na siatkę wiązki światła. Zatem prążki i to dobrze widoczne można zaobserwować również wtedy, gdy szczelina jest tak szeroka, że nie jest już szczeliną i nie może być traktowana jako źródło fali spójnej. Szczególnie, że „prążki” są widoczne również wtedy, gdy i takiej formalnej szczeliny nie ma! Zresztą, zapomniał pan o naszej wcześniejszej rozmowie na temat ilości niespójnych fotonów pomiędzy krawędziami pojedynczej szczeliny, gdy pada na nie światło niespójne. A w tym przypadku mówię o szczelinie odkręconej na szerokość nawet około 5mm! Po trzecie, nawet jeśli przyjmiemy, że światło po przejściu przez szczelinę jest światłem spójnym. To jest ono spójne tylko przestrzennie, ale nie czasowo! A przecież, do danego obszaru na ekranie (oka) dochodzą jednocześnie fotony, które zostały wyemitowane w innym czasie bo wyszły z innych szczelin i miały inną drogę do pokonania.

- Jeśli faktycznie jest tak jak pani mówi, to doświadczenie związane z badaniem widma emisyjnego gazów jest istotnie kolejnym dowodem na brak konieczności użycia światła spójnego w doświadczeniu Younga. Wówczas, używany w tym doświadczeniu wzór będzie trzeba wyprowadzić dla innych założeń teoretycznych. Dodatkowo, założenie przynajmniej lokalnej spójności światła dla niespójnej wiązki padającej nie wytrzyma argumentu odnośnie niespójności czasowej. W tym przypadku wymagana jest zarówno spójność czasowa i przestrzenna padającej na siatkę dyfrakcyjną wiązki światła.

- Też tak myślałam. Oczywiście nie mam na razie pojęcia jak to wytłumaczyć inaczej, a tym samym jak wyprowadzić zależność jaką podano w podręcznikach. Na razie przyjmę roboczo, że wspomniana zależność (1) jest poprawna pod względem ilościowym. Co nie jest takie oczywiste przy braku teoretycznych podstaw

do wyprowadzenia tej zależności i nie zmienia to faktu, że kiedyś będzie ją trzeba wyprowadzić dla założeń, które nie będzie można tak łatwo podważyć.

Kamila rozglądając się po pokoju zauważyła otwartą książkę z fizyki, która leżała obok na stole

- O! Tutaj jest opisane doświadczenie podobne do tego co robiłam na zajęciach. Co to jest za książka panie profesorze?

- Pani Kamilo, to jest podręcznik do Mechaniki Kwantowej. Wiem, że jeszcze nie miała pani wprowadzenia do tego zagadnienia. W kolejnych semestrach zapozna się pani z tym przedmiotem dokładniej, gdyż jest to podstawa dla zrozumienia zagadnień z fizyki współczesnej.

Usiedli więc przy stole, gdzie leżała otwarta książka, którą profesor wcześniej przeglądał.

- Czy będę mogła pożyczyć tą książkę?

- Myślę, że tak. Chociaż dobrze by było wpierw opanować nieco podstawy. Zwłaszcza matematykę. Bez tego trudno jest przejść przez ten kurs fizyki.

- Rozumiem.

- Akurat wstęp, gdzie jest opisane doświadczenie Younga, nie zawiera zaawansowanej matematyki. Rozdział „Cząstki i fale” poświęcony jest zagadnieniu tzw. dualizmu korpuskularno falowemu. W mechanice kwantowej przyjmuje się, że każdej cząstce materialnej można przypisać fale, tzw. falę materii. Zapewne dobrze pani wie, że w fizyce klasycznej przyjmuje się, że materia składa się z dobrze zdefiniowanych cząstek materialnych. To znaczy posiadających określoną masę, położenie, pęd i energię.

- Tak, pamiętam. Na zajęciach z mechaniki określaliśmy prędkość, położenie oraz tor ruchu pocisku, który był reprezentowany przez punkt materialny. Prowadzący zajęcia wspominał, że w fizyce klasycznej te wszystkie parametry są zawsze dobrze określone i teoretycznie można je określić z dowolną dokładnością. Co to jest fala materii? Jeszcze nie spotkałam się z tym pojęciem.

- Owszem, w fizyce klasycznej tak, ale nie w mechanice kwantowej. Zanim przejdziemy do omawiania pojęcia fali materii, należy wpierw przypomnieć sobie nieco historii. Jak pani zapewne dobrze pamięta, światło zostało opisane przez Jamesa Maxwella jako fala elektromagnetyczna. Wcześniej o świetle jako fali rozważali też inni badacze, tacy jak Young czy Fresnel, ale pierwszymi to byli Christiaan Huygens i Robert Hooke, który mocno wspierał się z Newtonem o naturę światła. Autorytet Newtona jednak był tak wielki, że przez długi czas obowiązywał jego punkt widzenia, który zakładał, że światło jest pewnego rodzaju cząstką podobnie rozumianą jak inne obiekty materialne. Wówczas, wiele prostych zjawisk potrafiono zinterpretować zarówno poprzez odwołanie się do falowej jak i korpuskularnej interpretacji. Dopiero pomiary prędkości światła przechyliły poglądy uczonych w kierunku interpretacji falowej. Innym kluczowym doświadczeniem przemawiającym za przyjęciem falowego poglądu na naturę światła jest właśnie doświadczenie Thomasa Younga. Związane jest to z tym, że Young zinterpretował swoje doświadczenie poprzez odwołanie się do zjawisk dyfrakcji i interferencji, które są spotykane w zjawiskach falowych oraz za pomocą swojego rozumowania powiązał różne wielkości geometryczne występujące w tym doświadczeniu. Nikt ówczesny nie miał pomysłu, jak wyjaśnić interferencję „cząstek” Newtona. Dlatego, na długi czas zapomniano o jego korpuskularnej hipotezie, aż do czasu Max Plancka. Planck poszukując formuły, która poprawnie tłumaczyłaby widmo emisyjne ciała doskonale czarnego, musiał w swoich rozważaniach dokonać bardzo niewygodnego jak na ówczesne poglądy posunięcia. Mianowicie założył, że światło emitowane jest w porcjach o określonej energii. Wszystko byłoby w porządku gdyby nie fakt, że przyjęta przez niego hipoteza implikowała istnienie światła jako niezależnych porcji energii (co było bliższe korpuskularnej hipotezie), ale jednocześnie prowadziła do otrzymania formuły, która poprawnie opisywała dane doświadczalne. Był to więc pierwszy wyłom w dominacji falowej interpretacji światła. Hipoteza Plancka była przez ówczesnych ignorowana i traktowana jedynie jako niewygodna sztuczka matematyczna. Potrzebna tylko do uzyskania poprawnej formuły, ale nie zmieniająca istoty światopoglądu ówczesnych naukowców na temat światła jako fali. Falowa interpretacja Maxwella miała się nadal dobrze. Jednak niewiele później bo już w 1905 roku Albert Einstein wykazał poprzez wytłumaczenie zjawiska fotoelektrycznego, że Planck jednak miał rację i światło faktycznie składa się z ściśle określonych porcji energii. Tym razem, nie była to już tylko matematyczna sztuczka, ale konkretne bezpośrednie wyjaśnienie wspomnianego doświadczenia, w którym Einstein musiał założyć, że światło ma korpuskularną naturę by wspomniane zjawisko wyjaśnić, a które nie udawało się w żaden sposób wytłumaczyć na bazie falowej

teorii Maxwella. Nieco później, niemniejsze zamieszanie w teorii dokonał Arthur Holly Compton, który badając rozpraszane pod różnym kątem promienie rentgenowskie doszedł do wniosku, że światło nie tylko ma określoną porcję energii jak chciał Planck i Einstein, ale również i pęd tak jak cząstki materialne. Ostatecznie przyjmuje się, że światło składa się z tzw. fotonów (cząstek światła).

- Dlaczego w takim razie nadal uczymy się w szkołach, że światło jest falą?

- Doświadczenia wykonane przez wspomnianych badaczy mogłyby całkowicie pograżyć falową interpretację światła już na początku XX wieku, gdyby nie doświadczenie Younga i pokrewne mu: np. Fresnela, które nadal tłumaczone jest jako doświadczenie dyfrakcyjno-interferencyjne. Ponadto, problemem jest również istnienie tzw. fal elektromagnetycznych (radiowych, mikrofal itd.), które wydają się być jednak tym samym co światło widzialne, a które jest znacznie trudniej tłumaczyć na gruncie pojęcia cząstki.

- To fakt, raczej nie da się twierdzić, że coś takiego jak fale radiowe nie istnieją. Przecież radio, telewizja (nie kablowa), telefony komórkowe itd. co by nie mówić działają ;-).

- No właśnie. Ponadto, większość rozważań dotyczących światła wychodzi z założenia, że światło nie posiada masy, co również doskonale pasuje do falowego światopoglądu. Dopiero teoria względności Einsteina pozwoliła nadać fotonom masę relatywistyczną powiązaną z energią fotonów. Co może doskonale tłumaczyć np. efekt Comptona, tzn. to, że foton posiada pęd (czyli powinien posiadać masę!).

- Rozumiem więc, że dualizm falowo-korpuskularny dotyczący światła oznacza tyle, że światło jest falą i nie jest nią jednocześnie. Nic z tego nie rozumiem!

- Nie tylko ty Kamilo. Tak naprawdę to bardzo trudno jest pogodzić opis falowy z korpuskularnym tak, by stał się on pojęciowo i logicznie spójny. W zasadzie to niewiedomo; czy w ogóle jest to możliwe? Obecnie podchodzi się do tego w ten sposób, że raz korzysta się z falowej a innym razem z korpuskularnej interpretacji zależnie od tego, która jest w danym momencie wygodniejsza do zastosowania. Stosuje się przy tym formalizm matematyczny, który pozwala przejść z jednego opisu w drugi.

- Ale jakoś nadal nie potrafię sobie tego wyobrazić. Pamiętam z zajęć, że fala to przecież nic innego jak przemieszczające się w przestrzeni zaburzenie ośrodka. Natomiast opis korpuskularny to nic innego jak poruszające się w przestrzeni „światłne” pociski. Niech będzie fotony. Tego nie da się sensownie pogodzić.

- Dokładnie. Jak już pani wspomniała, fala jest przemieszczającym się w przestrzeni zaburzeniem. W przypadku dobrze znanego zjawiska, takiego jak rozchodzenie się dźwięku, zaburzeniem ośrodka są przemieszczone od położenia równowagi atomy gazu, czy ciała stałego. W przypadku światła, które niezwykle łatwo przemieszcza się w próżni zdefiniowanie ośrodka stało się praktycznie niemożliwe. Dlatego Maxwell i wcześniejsi badacze założyli, że istnieje hipotetyczny ośrodek, który nazwali eterem. Teoria Maxwella powstała jako próba określenia przez niego ściśłości tego hipotetycznego ośrodka i zakłada jego istnienie [5]! Następnie w ramach badań nad eterem Albert Michelson i Edward Morley wykonali serię doświadczeń, które pozbawiły nas złudzeń, co do istnienia tego hipotetycznego ośrodka. Obecnie nie naucza się o eterze w szkołach, gdyż hipotezę o jego istnieniu traktuje się jako błędną.

- Przepraszam, ale skoro teoria Maxwella zakładała istnienie eteru dla rozchodzenia się fal elektromagnetycznych, to czemu nadal jest traktowana jako prawdziwa?

- Owszem, pierwotnie teoria Maxwella była zależna od tego założenia i tak szczerze mówiąc, to nadal jest od niego zależna. Obecnie zakłada się, że w przestrzeni może istnieć niezależnie od materii pole elektryczne i magnetyczne, które następnie oddziałuje na materię. Stąd, możliwość przemieszczania się tzw. fali elektromagnetycznej przez próżnię. W teorii Maxwella były to różne stany (mechaniczne) eteru reprezentowane poprzez naprężenia i ruch tego ośrodka. Czy my to nazwiemy polem magnetycznym i elektrycznym opisującym stan elektro magnetyczny danego punktu przestrzeni (teoria pola), czy mechanicznym stanem eteru to matematycznie jest to nadal to samo, zmieniły się tylko nazwy pojęć. Teoria Maxwella przetrwała głównie za sprawą użyteczności wzorów, które są z nią związane. Notabene, większość wzorów w niej zawartych wcale nie jest autorstwa Maxwella. Jego zasługa polega głównie na zebraniu w jedną dość spójną całość dokonań innych badaczy.

- To wszystko jest bardzo ciekawe, ale na początku rozmowy wspomniał pan coś o fali materii. Nadal nie rozumiem co to jest?

- No tak, zapomniałem już, od czego zacząłem. Jak zapewne zdążyła się już pani przekonać. Zagadka, czym jest światło? Nie jest w zasadzie rozwiązana, gdyż trudno pogodzić w logiczną spójną całość to, że coś może

być zarówno falą i materią, a do tego nie posiada masy (obecnie rozumianej jako masy spoczynkowej). Ale żeby było ciekawiej, niejaki George Thomson dokonał zdumiewającego odkrycia, że prążki interferencyjne powstają również dla takich obiektów jak elektrony.

- Jak to możliwe? Elektrony mogą interferować? {Kamila ze zdziwienia drgnęła w fotelu} Na czym polegały te doświadczenia?

- Jak już pani wie, w doświadczeniu Younga wiązkę światła przepuszcza się przez przesłonę na której są dwie szczeliny. Wówczas na ekranie powstają charakterystyczne prążkowe obrazy, które są interpretowane jako obrazy interferencyjne. Kiedy przepuszczono nie fotony, ale wiązki molekularne (rozpędzone w jednym kierunku: jony, elektrony, protony itd.) przez kryształ to na ekranie składającym się z odpowiednich detektorów również zaobserwowano struktury przypominające prążki interferencyjne. Zaobserwowany efekt oczywiście wielce zdumiał badaczy, którzy nie umieli znaleźć racjonalnego wyjaśnienia obserwowanych efektów dla obiektów, które w klasyczny sposób nie sposób traktować jako falę.

- Chyba zaczynam rozumieć. Skoro przyjęto, że światło jest falą, którą należy czasami traktować jak cząstki, to korpuskularną materię zaczęto czasami traktować jako falę. Stąd, materia traktowana jako fala nazywa się falą materii.

- Dokładnie. Ponadto, taka interpretacja stała się jednocześnie niezwykle wygodna, wręcz niezbędna dla świeżo powstającej wówczas mechaniki kwantowej. Niels Bohr postulował, że moment pędu elektronu w atomie jest kwantowany. Takie założenie pozwala całkiem nieźle odtworzyć widmo emisyjne gazów, które pani niedawno badała w laboratorium.

- Jak to?

- Wynika to stąd, że elektron można traktować jako falę i na poszczególnych orbitach musi ona tworzyć falę stojącą, by się nie wygasić jako fala do zera. Zauważono, że taki warunek nie może być spełniony dla dowolnych odległości elektronu od dodatniego jądra oraz przy zastosowaniu jeszcze kilku dodatkowych postulatów, udało się fizykom całkiem nieźle opisać poziomy energetyczny jakie elektron może przyjmować w różnych atomach. Zakładając, że energia emitowanych jak i absorbowanych fotonów zależy od różnicy obliczonych poziomów energii, można określić długości fal (barwy) jakie powinien emitować np. atom wodoru. Model Bohra z czasem się nie przyjął jako zbyt uproszczony, ale zasadniczy postulat, by materię traktować jako falę materii stał się fundamentem mechaniki kwantowej.

- To jest niesamowicie fascynujące! {Z podnieceniem krzyknęła Kamila} Już nie mogę się doczekać przeczytania tej książki. Czy będę mogła ją od pana pożyczyć? Bardzo proszę.

- Pani Kamilo, niech się pani tak bardzo nie śpieszy. Wpierw proponowałbym poczytać co nieco na temat historii fizyki. Będzie to na pewno równie interesująca lektura. Jak już pani przejdzie przez bardziej zaawansowany kurs matematyki, wówczas może pani śmiało sięgnąć po ambitniejsze tytuły.

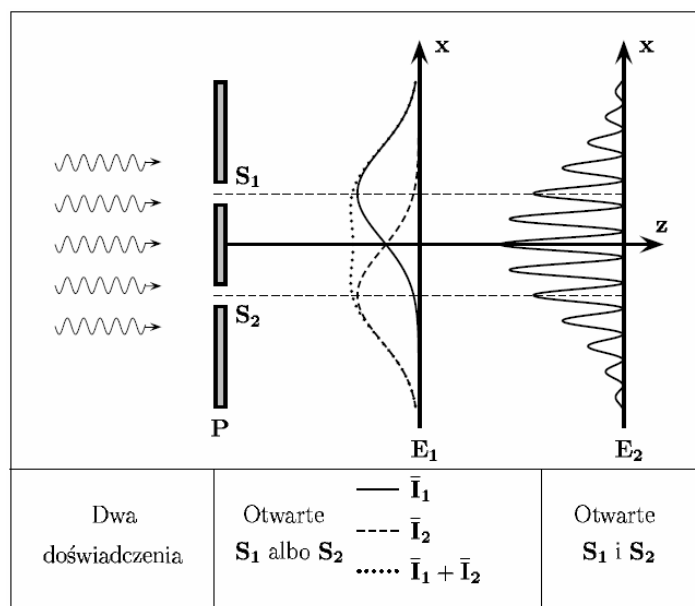
- Na początku wspominał pan profesor, że wyjaśni mi; jak jest wytłumaczone doświadczenie Younga w ramach mechaniki kwantowej.

- Ach, no tak. Jak zwykle się rozgadałem i zapomniałem. O czym to ja chciałem powiedzieć? Więc tak. W rozdziale Cząstki i fale opisane jest doświadczenie, które upoważnia fizyków do twierdzenia, że dualizm korpuskularno falowy faktycznie istnieje. Chodzi o to pani Kamilo, że skoro doświadczenie Younga udaje się dla cząstek materii, a co za tym idzie dla bytów którym łatwiej przyjąć korpuskularną a nie falową naturę, to równie dobrze można by próbować odrzucić interpretację falową jako taką dla oryginalnego doświadczenia ze światłem. Historycznie zdarzyło się tak, że pierwsze prążkowe obrazy zaobserwowano dla doświadczenia ze światłem. Więc łatwiej było wymyśleć interpretację falową. Dopiero później odkryto, że dla cząstek materii występuje podobne zjawisko. Tak więc, skorzystano z wcześniejszej interpretacji by wytłumaczyć nowo odkryte doświadczenie. A co by było, gdyby kolejność odkryć była inna?

- Ciekawe, nie pomyślałam o tym. Zapewne dla doświadczenia z wiązką elektronów rozproszoną na szczelinach poszukano by wyjaśnienia czysto korpuskularnego. Wówczas, doświadczenie Younga z światłem potraktowano by jako dowód, że światło również posiada tylko własności korpuskularne.

- Dokładnie, widzi pani jak ważna jest kolejność odkryć i następujących po nich idei. Zmienia to postrzeganie tego co widzimy i ukierunkowuje sposób w jaki interpretujemy nowe odkrycia. Tak więc, aby usankcjonować ideę, że obiekt materialny może być zarówno cząstką i falą przedstawia się w książkach następujące rozumowanie.

Profesor wyjął kartkę oraz ołówek i przerysował rysunek z podręcznika mechaniki kwantowej na którym przedstawione były wyniki przejścia światła przez przesłonę z dwoma szczelinami (Rys. 15).



Rys. 15 Schemat dwóch doświadczeń dyfrakcyjno-interferencyjnych na dwóch wąskich szczelinach mający wykazać słuszność koncepcji dualizmu korpuskularno-falowego omawianego w podręczniku mechaniki kwantowej [2].

- Co to za krzywe panie profesorze? Ta sinusoida o zmiennej amplitudzie zapewne przedstawia prążki widziane na ekranie w doświadczeniu Younga.
- Zgadza się. Ale to co jest najciekawsze w tym całym rozumowaniu, to obraz powstały po przejściu przez pojedynczą szczelinę. Środkowy obraz przedstawia krzywe natężenia światła, bądź padającej na ekran wiązki molekularnej (w doświadczeniu z elektronami itd.) pochodzącej nie od dwóch ale pojedynczej szczeliny. Krzywa ciągła przedstawia rozkład światła pochodzący od górnej szczeliny, a przerywana od dolnej.
- Rozumiem, że taki rozkład tłumaczony jest jako konsekwencja korpuskularnej natury światła.
- Owszem. Zakłada się, że na szczelinie fotony uderzają nieco w ścianę krawędzi. Wówczas, tak przypadkowo rozproszone fotony podążają dalej pod nieco innym kątem tworząc na ekranie jeden szeroki obraz szczeliny. Oczywiście obraz dla drugiej szczeliny jest nieco przesunięty względem obrazu powstającego dla pierwszej szczeliny.
- A co będzie jak obie szczeliny będą otwarte? Rozumiem, że otrzymamy klasyczne doświadczenie Younga i prążkowy obraz na ścianie, który przedstawiony jest na prawym wykresie.
- No właśnie. Jeden szeroki rozmyty obraz tłumaczy się korpuskularną naturą użytej w doświadczeniu wiązki, np. światła. Gdyby podobnie rozumować przy obu otwartych szczelinach (korpuskularnie) to powinniśmy spodziewać się, że uzyskamy taki obraz, jak kropkowana krzywa na środku rysunku. Natomiast doświadczenie pokazuje, że powstaje obraz prążkowy taki jak ten przedstawiony na prawym wykresie. Taki wynik z kolei wygodnie jest tłumaczyć jako aspekt falowej natury użytych wiązek: światła czy materii, np. elektronów. Wynika to stąd, że dla wytłumaczenia uzyskanego wyniku odwołujemy się do zjawisk: dyfrakcji i interferencji.
- Czyli uzyskamy to samo, co widziałam w laboratorium robiąc doświadczenia z dwoma szczelinami i laserem?
- Dokładnie tak. Ponieważ badając światło raz musimy odwołać się do falowej, a innym razem do korpuskularnej natury światła. Stąd, jest to dowód na prawdziwość hipotezy o istnieniu dualizmu

korpuskularno-falowego. Ponadto, jest to również prawdziwe dla cząstek materialnych, które czasami traktujemy jak fale.

- To się wydaje zbyt zawile. Nadal nie bardzo mogę sobie tego wyobrazić, by cząstki zdecydowanie materialne traktować jako falę. Mogę jeszcze raz przyjrzeć się rysunkowi?

- Oczywiście, proszę bardzo.

Kamila przyglądając się narysowanemu przez profesora rysunkowi, próbuje sobie przypomnieć co widziała w laboratorium, kiedy samodzielnie wykonywała podobne doświadczenia. W pewnym momencie ją olśniło i krzyknęła mimowolnie.

- To jest totalna lipa!

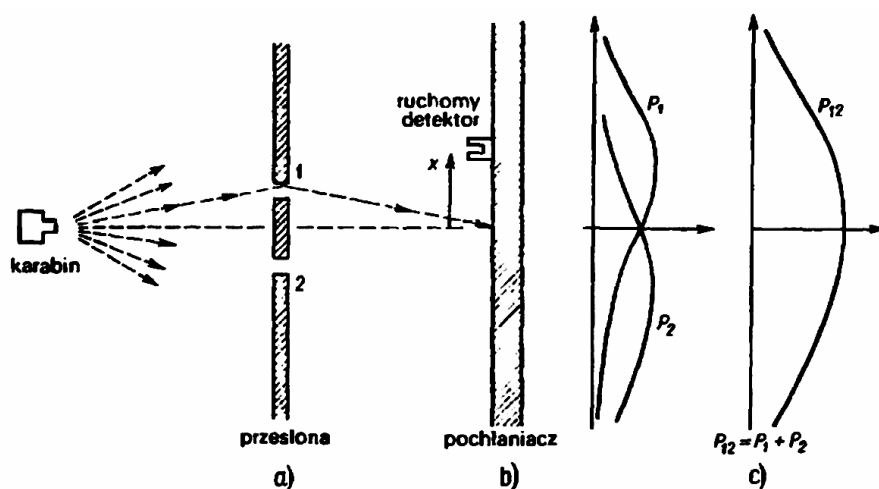
- Co proszę! {Odparł zdeglustowany profesor}

- O przepraszam. Proszę mi wybaczyć, ale czegoś tu nie rozumiem. Proszę spojrzeć na środkową część rysunku. Przecież światło przechodząc przez pojedynczą szczelinę nie tworzy takiego obrazu, jak jest tu przedstawione. Dobrze pamiętam, że jak przepuszczałam wiązkę światła laserowego przez pojedynczą szczelinę, to uzyskałam obraz prążkowy. Jeszcze mi pan profesor próbował to wytłumaczyć w jakiś mętny sposób, poprzez odwoływanie się do nieskończonej ilości szczelin oddalonych na zerową odległość.

- Hm. No tak, faktycznie tak było. Nawet zrobiłem dzisiaj symulację by sprawdzić jak się zachowa obraz przy zmianie takich parametrów jak: ilość szczelin i odległość między nimi. Czemu tego wcześniej nie zauważyłem?! {krzyknął z wrażenia}. Jak mogłem być tak ślepy, żeby tego wcześniej nie zauważyć! Faktycznie, światło przechodząc przez pojedynczą szczelinę tworzy obraz prążkowy, a nie tak, jak jest to przedstawione w podręczniku [2]! Na zajęciach laboratoryjnych pokazuję studentom, że światło przechodząc przez pojedynczą szczelinę tworzy obraz prążkowy. A tutaj, jednocześnie zaprzeczam sam sobie (na wykładach z mechaniki kwantowej) twierdząc coś zupełnie innego!

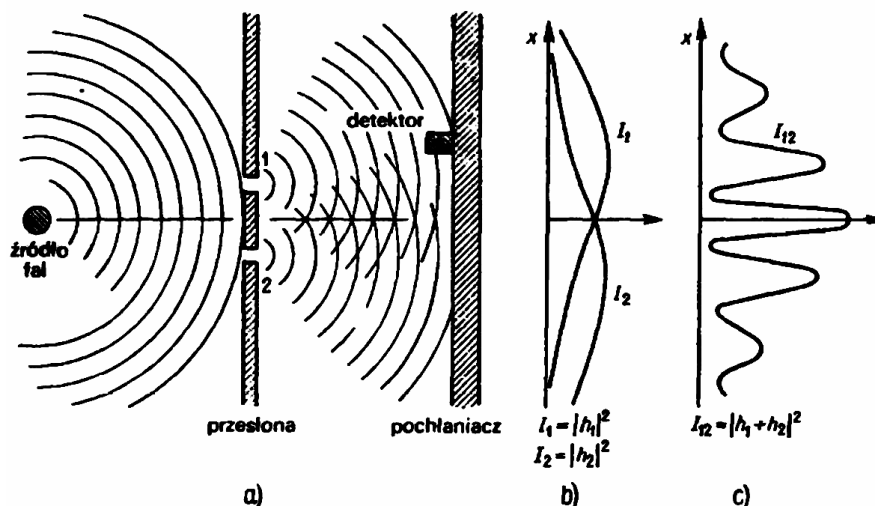
- Skoro światło nie zachowuje się tak jak pan przedstawił to na rysunku, to również wnioski wyciągnięte z tych rozważań są nieprawdziwe. Czyż nie?

- Oczywiście, ale dlaczego w bądź co bądź poważnym podręczniku już na wstępie pojawia się tak duży błąd? Czy jest możliwe uzyskanie dokładnie takiego obrazu na ekranie jak przedstawione jest to na rysunku? Przecież mechaniki kwantowej nie tworzą głupi ludzie, by nie zauważyć tak oczywistego błędu! W takim razie zajrzyjmy do jeszcze jednego podręcznika z fizyki: „Feynmana Wykłady z fizyki” Tom 1.2. W rozdziale zatytułowanym „Efekty kwantowe” mamy opis trzech podobnych doświadczeń: doświadczenie z pociskami; doświadczenie z falami oraz doświadczenie z elektronami. W pierwszym doświadczeniu Feynman opisuje eksperyment (według autora nigdy nie przeprowadzony) z karabinem maszynowym z którego wystrzeliwane są pociski w kierunku pancерnej przeszkody z dwiema szczelinami. Na ekranie za przeszkodą zliczane są pociski jakie dochodzą w danym czasie do danego punktu na ekranie. Obok schematu doświadczenia, Feynman przedstawia hipotetyczny wynik takiego eksperymentu (Rys. 16) [3].



Rys. 16 Doświadczenie z pociskami przechodzącymi przez pancerną przeszkodę na której zrobione są dwie szczeliny, przez które mogą przechodzić pociski [3].

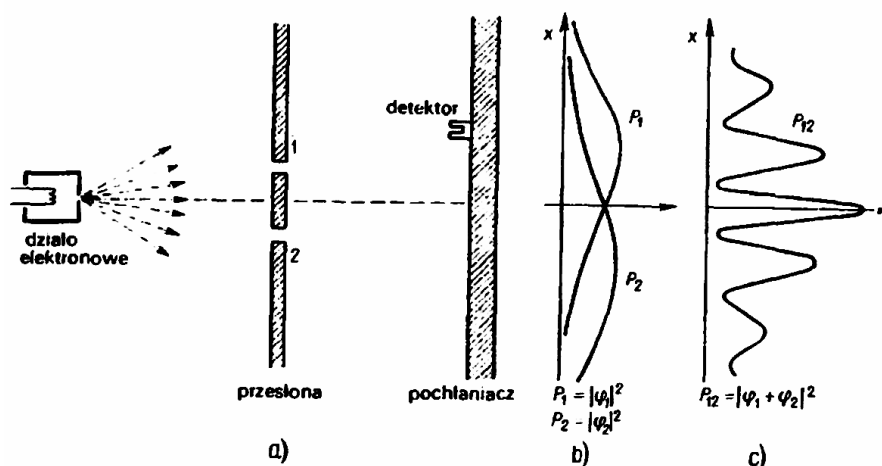
- Rozumiem, że takiego zachowania mamy się spodziewać dla obiektów czysto materialnych (korpuskularnych)?
- Oczywiście. Następnie Feynman przedstawia drugie doświadczenie z falami (Rys. 17).



Rys. 17 Doświadczenie z falami wodnymi przechodzącymi przez przesłonę z dwoma szczelinami [3].

- No tak, ale w tym przypadku autor założył, że źródło fali jest tak dobrane, by faktycznie mieć na obu szczelinach takie same fazy. Zresztą pamiętam z wykładów, na których było przeprowadzane podobne doświadczenie, że fala mechaniczna faktycznie na szczelinach tak się zachowuje, jak jest to przedstawione w podręcznikach. Stąd, przedstawiony w podręczniku [3] wynik doświadczenia z falą mechaniczną jest zgodny z doświadczeniem. Oczywiście dla różnych i przypadkowo zmieniających się faz na obu szczelinach nie uzyskamy już takiego obrazu jak to przedstawiono na prawej stronie rysunku. Zresztą, dla światła przedstawiony obraz i tak jest niezgodny z doświadczeniem, gdyż pojedyncza szczelina wystarczy do uzyskania prążków na ekranie!

- Owszem, światło zachowuje się inaczej. Nie mniej jednak w tym podręczniku Feynman opisuje przypadek przejścia światła przez pojedynczą szczelinę. Później sprawdzimy dokładniej, w jaki sposób jest to przez niego wytłumaczone. Na razie przyjrzyjmy się kolejnemu doświadczeniu. Tym razem z elektronami (Rys. 18).



Rys. 18 Doświadczenie z elektronami przechodzącymi przez przesłonę z dwoma szczelinami [3].

- W zasadzie, jest to zgodne z tym, co przedstawiał w swoim podręczniku pan Kryszewski [2]. Dla pojedynczej szczeliny płaski szeroki obraz. Natomiast, dla obu szczelin uzyskujemy obraz prążkowany. Tyle, że w tym przypadku Feynman nie wspomina o świetle, ale od razu opisuje zachowanie się elektronów. Przykładowo:

„Pierwszą rzecz, jaką stwierdzamy w naszym doświadczeniu z elektronami, to wyraźne pojedyncze dźwięki, jakie słyszymy z detektora (tzn. z głośnika). Wszystkie one są takie same. Nie ma „połówek” takich dźwięków.

Stwierdzamy również, że są one bardzo nieregularne. Coś w rodzaju: tik. tik-tik.....tik-tik..., itd., tak jak to już pewnie słyszeliście w czasie pracy licznika Geigera.”

- Stąd wynika, że dla elektronów mamy pomiar tylko całych elektronów, ale w takim razie jak mogą one interferować na ekranie. Przyczyna musi być inna i ma zapewne swoje miejsce na krawędziach.

- Niekoniecznie pani Kamilo. Feynman opisuje wynik doświadczenia zarówno dla przejścia elektronów przez pojedynczą szczelinę, co jest podobne do tego co widzimy dla cząstek materialnych i nie budzi większych wątpliwości oraz wynik dla przejścia elektronów przez dwie szczeliny. W tym drugim przypadku mamy obraz prążkowy, co odpowiada obrazowi falowemu i wskazuje na konieczność interpretacji za pomocą interferencji fal:

„Spróbujmy obecnie zanalizować krzywą z rys. 37.3 (Rys. 18) i przekonać się, czy potrafimy zrozumieć zachowanie się elektronów. Przede wszystkim chciałoby się powiedzieć, że ...”

Stąd właśnie koncepcja dualizmu korpuskularno falowego o czym już wspominałem wcześniej. Szczególnie, że dla światła również odbieramy „całe” fotony dla danej barwy, a nie pół fotonu. Czyli podobnie jak dla elektronu pełne tik-tik, a nie pół tik.

- Przepraszam panie profesorze, ale jakie tik-tik on opisuje. O jakim zachowaniu się elektronów w tym doświadczeniu pisze Feynman, skoro tutaj obok we własnej książce stwierdza on, że:

„Musimy od razu powiedzieć, że nie należy próbować przeprowadzenia takiego doświadczenia (w odróżnieniu od dwóch doświadczeń poprzednich, które można było wykonać). Doświadczenie to nie zostało nigdy wykonane. Rzecz w tym, że aby uzyskać efekt, który nas interesuje, urządzenie musiałoby być wykonane w niezmiernie małej skali. Przeprowadzamy więc „eksperyment myślowy”, który wybraliśmy z tych względów, że łatwo o nim mówić i myśleć. Wiemy, jaki wynik otrzymamy, ponieważ przeprowadzono już wiele eksperymentów, w których zachowano odpowiednią skalę i proporcję, co pozwoliło uzyskać efekty, które opiszemy.”

{Kamila oburzona tym co przeczytała, kontynuowała dalej.}

- Skoro doświadczenie według autora nigdy nie zostało wykonane, to jakim prawem autor opisuje wyniki takiego doświadczenia i jeszcze twierdzi, że ma ono taki a nie inny przebieg. Tak, jakby je wykonywał. Wyciąga z przytoczonych przez siebie wyników tegoż doświadczenia kontrowersyjne wnioski dotyczące

natury elektronów, ale przecież: „*Doświadczenie to nie zostało nigdy wykonane*”. Żeby było jeszcze śmieszniej, to autor stwierdza, że: „*nie należy próbować przeprowadzenia takiego doświadczenia*”. Czyli, według Feynmana, nie tylko nie należy wykonywać tego doświadczenia, ale nawet nie należy tego próbować! Ale sam narzuca swoją jedynie słuszną interpretację niewykonanego doświadczenia. Taką, która jest mu wygodna dla własnych rozważań. Owszem, autor odwołuje się do jakiś innych eksperymentów (nie wspomina jakich), ale te są pewnie bardziej pośrednie i tym samym mniej wiarygodne dla udowodnienia tezy o falowej naturze elektronów i swojego punktu widzenia.

- Pani Kamilo, proszę się tak nie unosić! Pan Feynman pisał swoją książkę jakiś czas temu. Kiedy ta książka powstawała takie doświadczenie mogło być jeszcze niezrealizowane. Od wydania tej książki minęło trochę czasu i możliwe, że ta część tekstu jest już nieaktualna. Choć należy to jeszcze sprawdzić dla lepszego rozeznania sprawy. Niemniej jednak muszę uczciwie przyznać, że rozumowanie autora jest nie merytoryczne i źle o nim świadczy. Opisywanie niewykonanego doświadczenia tak jakby było przez niego wykonane oraz przytaczanie wyników z niewykonanego doświadczenia dla potwierdzenia kontrowersyjnych wniosków jest bardzo naganne i niedopuszczalne. Tym bardziej dziwi, że autor sam przyznaje się do tego we własnej książce. Dziwne, że w tym miejscu zabrakło autorowi samokontroli nad tym co pisze. Szczególnie, że mówimy tu o bardzo inteligentnym człowieku. Możliwe, że tak bardzo wierzył w słusność swoich przekonań, że zaryzykował swoje dobre imię.

- Faktycznie to dziwne, bo przecież ta część jego książki, ze względu na wspomniany fragment, mogłaby słuszenie przypisać nie jednego o uśmiech politowania, przynajmniej nad intencją tego człowieka. Możliwe, że Feynman głęboko wierzył w prawdziwość tego co robił i nauczał. Choć duża wiara autorytetu we własne przekonania może pociągnąć wielu w kierunku tych przekonań. Sama wiara we własne racje nie czyni z tych przekonań prawdy. A odwoływanie się do „nigdy niewykonanych” doświadczeń powinno w nas uruchomić sygnał alarmowy. Byśmy tych przekonań nie brali bez oporu, wierząc tylko w sam autorytet autora.

- Słusznie pani Kamilo. Niemniej jednak wróćmy do początku. Wspomniane doświadczenie miało pokazać zachowanie się nie tylko cząstek (w tym przypadku elektronów) ale i fotonów po przejściu przez pojedynczą szczelinę i przez dwie szczeliny, tak jakby elektron lub foton mógł jednocześnie przejść przez obie szczeliny. Pani już wcześniej zauważyła, że dla światła nie uzyskujemy takiego obrazu jak przedstawiono w podręczniku. Dlaczego w takim razie w książkach przedstawia się coś innego, niż to, co możemy zaobserwować w laboratorium? Czy po przejściu światła przez pojedynczą szczelinę można uzyskać szeroki pojedynczy obraz bez prążków bocznych?

- Może wpierw zapytajmy się; Czym może różnić się szczelina z naszego „szkolnego” doświadczenia od tej uzyskanej w profesjonalnym doświadczeniu, skoro uzyskujemy tak różne jakościowo wyniki?

- Szczelina jest szczeliną i z punktu widzenia falowych rozważań w ogóle nie odwołujemy się do kształtu szczeliny i materiału z którego jest zrobiona a nawet jej temperatury itd.. Czyli niczym, czy słusznie?

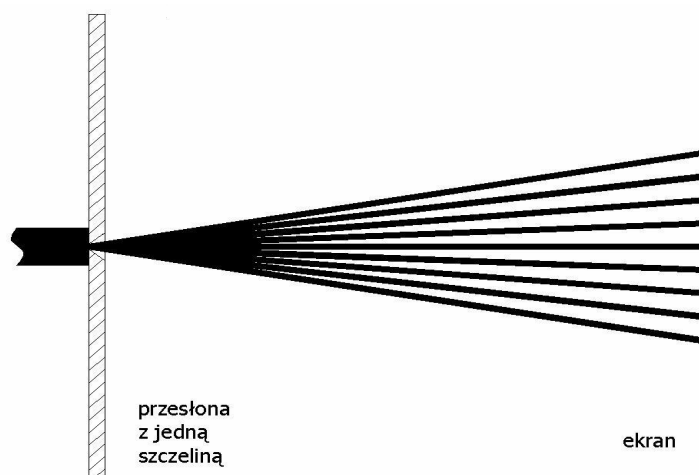
- Pewnie nie. {Niepewnie odparła Kamila}. Skoro nie uwzględnia się kształtu szczeliny w innych rozważaniach, to przyczyny rozbieżności należy poszukać gdzie indziej. Czym jeszcze może się różnić nasza szczelina od innych? Zapewne niczym specjalnym się nie różni. {Po chwili namysłu oznajmiła} Ale otoczenie szczeliny owszem!

- Jak to? {Ze zdziwieniem odparł profesor}

- W opisanym przez pana doświadczeniu mamy jedną przesłonę na której widnieją dwie szczeliny. By uzyskać dobry obraz interferencyjny należy takie szczeliny umieścić możliwie blisko siebie. Tak więc, aby wykonać doświadczenie z pojedynczą szczeliną wpierw trzeba zasłonić drugą. Stąd, w bliskiej okolicy nie zasłoniętej szczeliny znajduje się dodatkowa krawędź pochodząca od przesłony zasłaniającej drugą szczelinę, która może mieć wpływ na powstały obraz. W naszym podstawowym „szkolnym” doświadczeniu nie mamy dodatkowej krawędzi w okolicy szczeliny.

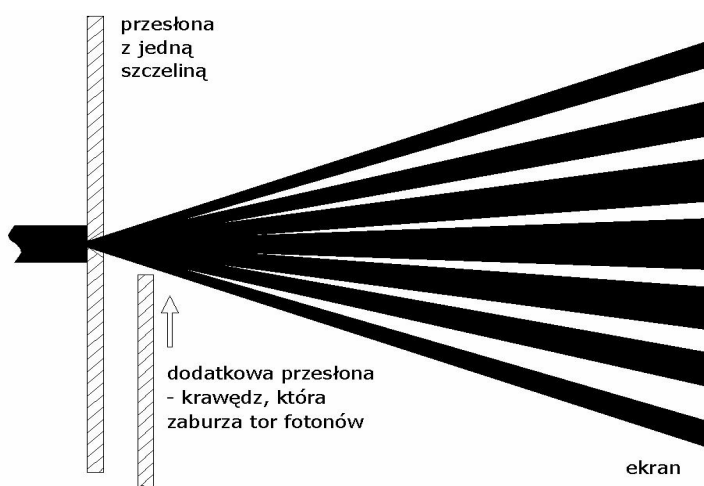
- Genialne! {Krzyknął entuzjastycznie profesor} To nic prostszego! Idźmy do laboratorium i sprawdźmy pani pomysł.

Przygotowanie stanowiska nie zajęło im wiele czasu i szybko uzyskali wynik, który niczym ich nie zaskoczył (Rys. 19).



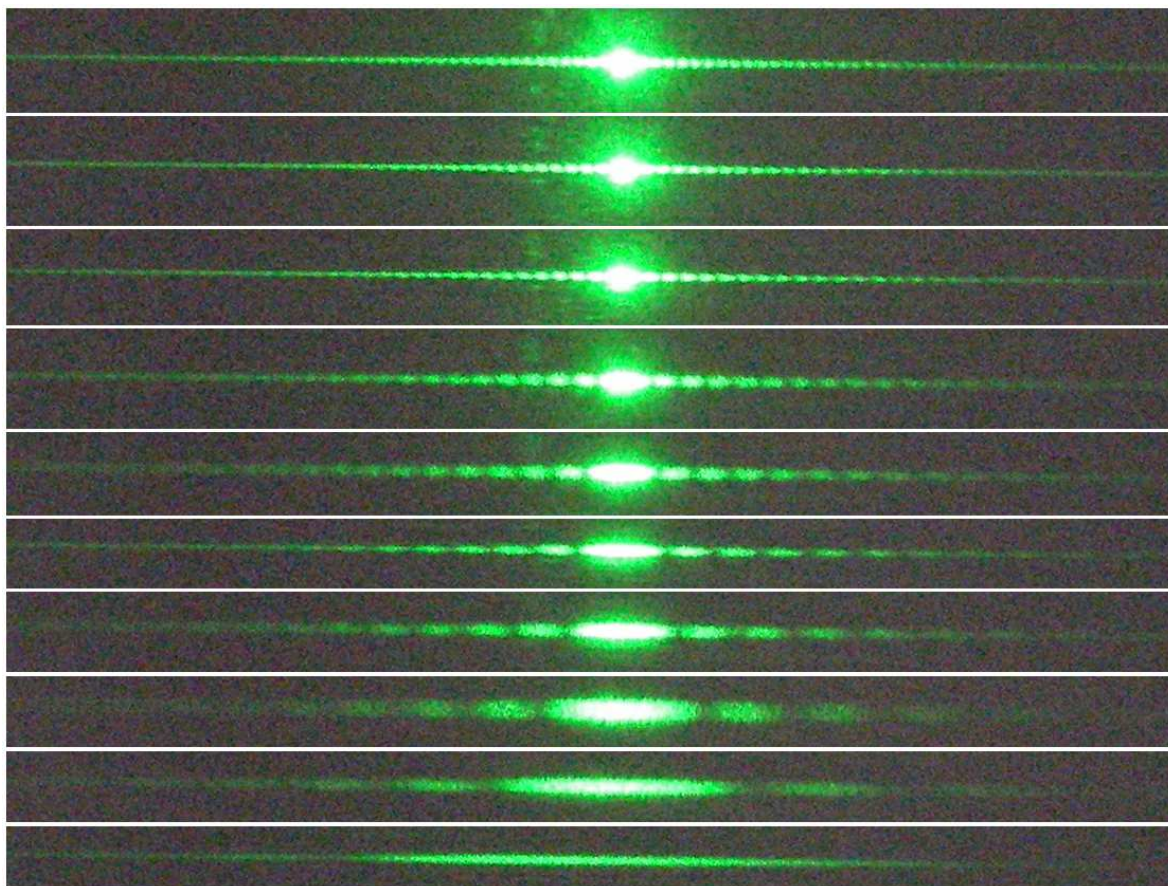
Rys. 19 Rozchodzenie się światła po przejściu przez pojedynczą szczelinę.

W ten sposób upewnili się, że w książce faktycznie jest poważny błąd. Pojedyncza szczelina daje prążkowy obraz. Następnie, Kamila zbliżyła grubą kartkę w kierunku pojedynczej szczeliny symulując zasłonięcie „drugiej” szczeliny przez dodatkową przesłonę (Rys. 20).



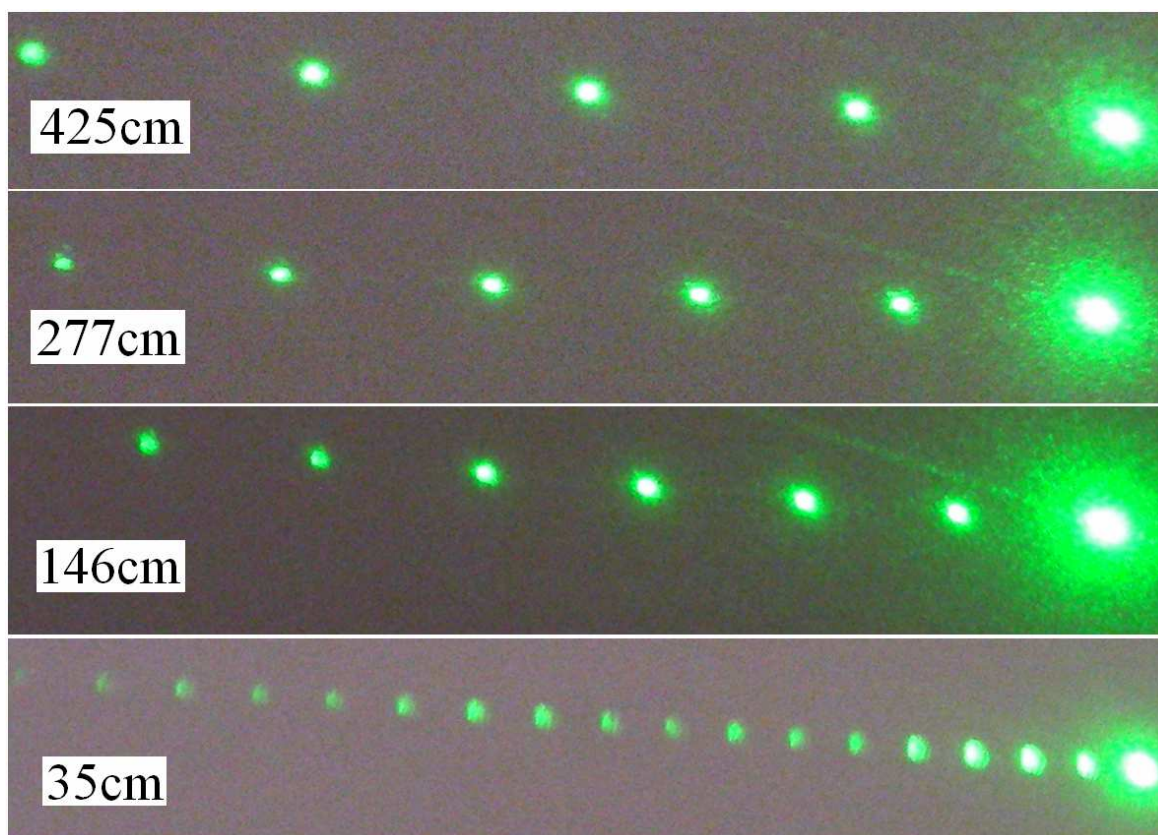
Rys. 20 Doświadczenie Younga z pojedynczą szczeliną i zaburzającą tor fotonów dodatkową krawędzią. W tym przypadku prążki ulegają spłaszczeniu i przesunięciu na boki.

Ich oczom ukazało się nieznanie im wcześniej zjawisko. Kiedy krawędź kartki zbliżała się do szczeliny, prążki na ekranie zaczęły rozsuwać się na boki i jednocześnie rozmywać (zwiększać szerokość poszczególnych prążków, Rys. 21).



Rys. 21 Obrazy powstałe z pojedynczej szczeliny, do której zbliżano dodatkową krawędź.

- No tak, i wszystko jasne! {krzyknęła Kamila}
- Kiedy krawędź jest dostatecznie blisko nie zasłoniętej jeszcze szczeliny, boczne prążki są na tyle rozsunięte na boki i rozmyte, że detektory ich nie wykryły. Natomiast, najjaśniejszy środkowy prążek został również rozmyty i utworzył taki rozkład światła jaki widnieje na rysunku w książce. {zauważył profesor}
- Czyli obraz widniejący w książce jest „prawdziwy” w tym sensie, że badacze mogli faktycznie uzyskać taki spłaszczony pojedynczy obraz jaki przedstawia się w podręcznikach.
- Owszem, ale wnioski stąd wyciągnięte są całkowicie błędne, gdyż pojedynczy rozmyty obraz na ekranie nie jest rezultatem przejścia światła przez pojedynczą szczelinę, gdzie otrzymujemy wiele prążków, ale jest efektem związanym z pojawieniem się w układzie dodatkowej krawędzi od obiektu zasłaniającego drugą szczelinę. Stąd, opisane w podręcznikach doświadczenie nie jest dowodem na istnienie dualizmu falowo-korpuskularnego, a wyciągnięte z niego wnioski i całe rozumowanie są po prostu błędne.
- Czyli, że źle wykonanego doświadczenia wyciągnięto nieprawdziwe i zupełnie absurdałne wnioski.
- Obawiam się, że dokładnie tak pani Kamilo.
- Korzystając z okazji, że jesteśmy już w laboratorium, chciałabym pokazać panu inną nieścisłość, którą zauważyłam wcześniej.
- Jeszcze jest coś nie tak? No tak, teraz sobie przypominam. Gdy pani weszła do pokoju to wspominała pani o jakimś nowym odkryciu. Proszę zacząć.
- Kiedy byłam tutaj ostatnim razem, wykonałam wcześniejsze doświadczenia, które wspólnie ostatnio wykonaliśmy. Upewniłam się co do poprawności naszych wcześniejszych wniosków odnośnie tego, że spójność wiązki światła w doświadczeniu Younga nie jest potrzebna. Następnie wróciłam do podstawowego doświadczenia z siatką dyfrakcyjną i zauważyłam, że jasność bocznych prążków „interferencyjnych” praktycznie nie zmienia się wraz z odległością od siatki (Rys. 22).



Rys. 22 Obrazy otrzymane w doświadczeniu Younga dla różnych odległości ekranu od siatki dyfrakcyjnej.

Początkowo nie zwróciłam na to uwagi, ale kiedy uświadomiłam sobie, że ostatni widoczny prążek jest pod stosunkowo dużym kątem, to zdałam sobie sprawę, że coś jest nie tak. Przecież w falowej interpretacji prążki są interpretowane jako interferencja konstruktywna fal pochodzących od wielu szczelin. Skoro tak, to fale od poszczególnych szczelin muszą rozchodzić się pod stosunkowo dużym kątem. Czyż nie?

- Oczywiście.

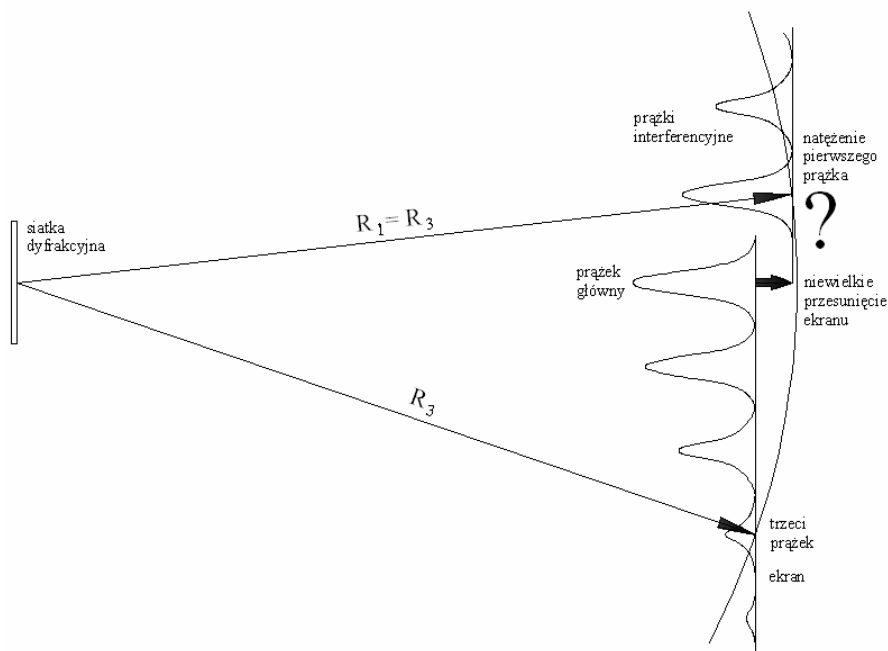
- Ale z mechaniki fal w ośrodku ciągłym pamiętałam, że tak rozchodząca się fala musi znacząco (stosunkowo szeroki kąt rozchodzenia się światła) zmniejszać swoje natężenie wraz z odległością, gdyż jej energia rozkłada się na coraz to większą przestrzeń (powierzchnię). Oczywiście nie posiadam dokładnego przyrządu na pomiar natężenia światła. Niemniej jednak przyłożyłam rękę blisko siatki dyfrakcyjnej około pięciu centymetrów, a następnie znacznie dalej na około dwa metry. Podczas tak dużego przemieszczenia nie stwierdziłam znaczących zmian jasności plamki na mojej ręce.

- Hmm. Sprawdźmy jak ta kwestia jest opisana w książce? Według podręcznika z mechaniki kwantowej pana Kryszewskiego, natężenie poszczególnych prążków bocznych w doświadczeniu Younga uzależnione jest od odległości prążka od szczelin. A dokładniej: „ A_i jest zależna od x , bo energia fali kulistej maleje wraz z kwadratem odległości od źródła (w tym przypadku szczeliny)”, gdzie A_i jest amplitudą fali dochodzącej do ekranu od i -tej szczeliny, natomiast „ x ” jest osią na ekranie przyłożoną wzdłuż obrazu „interferencyjnego” (Rys. 15) [2]. Boczne prążki leżą nieco dalej od szczelin i tym samym mają mniejsze natężenie światła, ale w pani opisie odległość zmieniła się około czterdziestokrotnie! To niemożliwe, by na taką odległość nie nastąpiła zauważalna zmiana jasności plamki światła dla poszczególnych prążków, skoro światło rozchodzi się po łuku ze szczeliny i to pod stosunkowo dużym kątem. Czyżby nasze założenia, co do rozchodzenia się światła ze szczeliny było nieprawdziwe?

Zaskoczony tym wnioskiem profesor umieścił siatkę dyfrakcyjną na osi optycznej i powtórzył doświadczenia o którym wspomniała Kamila. Ku swojemu zdziwieniu stwierdził, że było dokładnie tak jak

mówiła. Owszem plamka światła nieznacznie zwiększała swoją średnicę wraz z odległością, ale jasność zasadniczo nie zmieniała się w zauważalny sposób.

- Jak to możliwe, że jasność plamki światła praktycznie się nie zmienia? Zwłaszcza, że zgodnie z obowiązującym falowym opisem [2], prążki leżące bliżej osi optycznej powinny dość szybko zmniejszać swoje natężenie już przy niewielkiej zmianie odległości ekranu od szczelin, jeśli chcielibyśmy tłumaczyć spadek natężenia kolejnych prążków zmianą odległości (Rys. 23). Czy myślała już pani, jak wytłumaczyć zaistniałą niezgodność teorii z doświadczeniem?



Rys. 23 Schemat zmiany jasności prążków leżących bliżej środka obrazu przy niewielkiej zmianie odległości ekranu od siatki dyfrakcyjnej. Schemat opisany zgodnie z obowiązującą interpretacją falową [2].

- No cóż, myślałam o tym profesorze. Ale doszłam do dość absurdalnego wniosku, że energia musiałaby w jakiś magiczny sposób przenosić się z obszarów gdzie jest interferencja destruktywna do obszarów gdzie występuje interferencja konstruktywna.

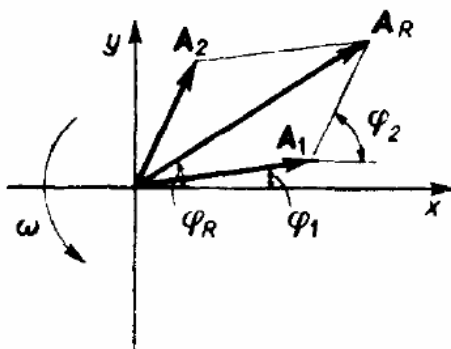
- No nie wiem proszę pani, przecież ekran równie dobrze może być kilometr dalej. I co wtedy? Energia przenosi się w mgnieniu oka z ogromnego obszaru do zaledwie kilku punktów. Nie! To jakiś absurd! Musi być bardziej racjonalne wyjaśnienie.

- No właśnie, być musi. Później doszłam do wniosku, że być może za zmiany natężenia światła dla poszczególnych prążków nie odpowiada różnica odległości od szczelin tak jak to pan przytoczył z książki, bo to jest akurat ewidentnie sprzeczne z doświadczeniem (Rys. 23). Aczkolwiek, to co pan przytoczył z podręcznika mechaniki kwantowej jest zgodne z prawdą. Tzn. natężenie światła dla każdej wiązki (widzianej jako pojedynczy prążek na ekranie) faktycznie maleje z kwadratem odległości licząc odległość od źródła. Ponieważ zakłada się, że światło jest falą, to mówimy o zmianie amplitudy fali wraz z odległością. Ale takie stwierdzenie jest prawdziwe również dla koncepcji czysto korpuskularnej, dla której natężenie światła byłoby tożsame z ilością fotonów przypadającą na jednostkę przekroju wiązki światła. Jeśli światło przemieszczałoby się w formie stożka, to wraz z kwadratem odległości rośnie przekrój wiązki i tym samym dana ilość fotonów rozkłada się na większą powierzchnię dając w efekcie spadek natężenia światła. Niemniej jednak, dotychczas przyjęliśmy założenie, że ilość fotonów przypadająca na daną jednostkę przekroju nie zależy od kąta wyjścia fotonów ze szczeliny. Dla takiego założenia nie byłoby innej możliwości, jak błędne przyjęcie, że za spadek natężenia poszczególnych prążków w doświadczeniu Younga

odpowiada różna odległości między prążkiem na ekranie a siatką dyfrakcyjną. Tymczasem, bardziej rozsądne jest przyjęcie, że szczelina nie zachowuje się dosłownie jak źródło punktowe, tak jak chce tego zasada Huygensa. Wówczas, możemy przyjąć, że światło wychodzące ze szczeliny ma różne natężenie w zależności od kąta na jakie odchyła się wiązka światła od kierunku osi optycznej. A co na ten temat pisze Feynman?

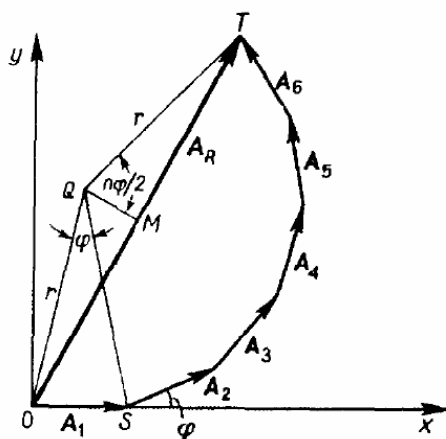
- Myślę, że Feynman myślał podobnie. Zaglądnijmy więc do jego książki. O! już mam.

W rozdziale 29-5 autor podchodzi do zagadnienia interferencji w ujęciu matematycznym tak samo jak my do zagadnienia opisu doświadczenia Younga. Czyli zakłada interferencję fal sinusoidalnych przesuniętych o pewną fazę w pewnym punkcie przestrzeni. Następnie pokazuje różne podejścia matematyczne do zagadnienia szukania fali końcowej będącej wynikiem nałożenia się wspomnianych fal sinusoidalnych [3]. Jedną z ciekawszych jest metoda geometryczna. W tej metodzie funkcje: $R_1 = A_1 \cos(\omega t + \varphi_1)$, oraz $R_2 = A_2 \cos(\omega t + \varphi_2)$ można przedstawić na osi współrzędnych (Rys 24).



Rys. 24 Geometryczna metoda składania dwóch fal cosinusowych. Należy sobie wyobrazić, że cały wykres obraca się z częstotliwością kątową ω przeciwnie do ruchu wskazówek zegara [3].

Dla takiego przedstawienia geometrycznego funkcje R_1 i R_2 należy rozważać jako wynik rzutowania wektorów z rysunku (Rys. 24) na oś x. W takim ujęciu względne położenie między wektorami nie zmienia się, gdyż oba wektory wirują z tą samą prędkością kątową. Ostatecznie wynik końcowy można uzyskać poprzez wektorowe dodanie obu wektorów R_1 i R_2 , a następnie wykonanie rzutu wektora wypadkowego na oś x. Takie podejście jest wygodne w przypadku rozważania większej ilości oscylatorów harmoniczných (Rys. 25) [3].



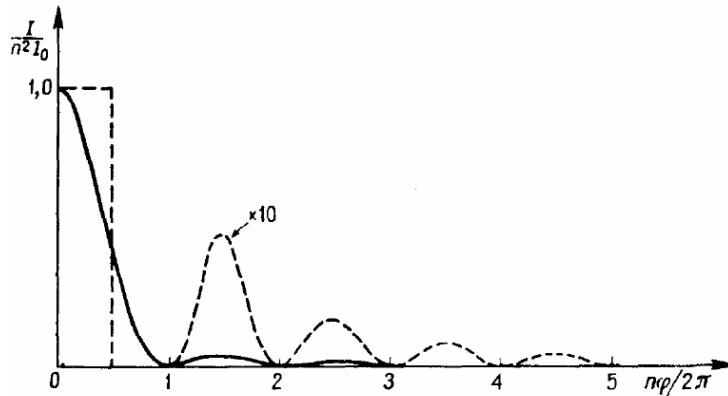
Rys. 25 Wypadkowa amplituda $n=6$ równoodległych źródeł, dla których różnica między sąsiednimi fazami wynosi φ [3].

- Ciekawe podejście matematyczne.

- Owszem, w ten sposób stopniowo dojdziemy do istoty rozumowania Feynmana. Pamięta pani jak wspomniałem, że dla pojedynczej szczeliny możemy założyć, że mamy nieskończenie dużą ilość źródeł fal oddalonych od siebie o nieskończenie małą odległość. Wówczas napotkamy problem dodawania nieskończenie wielu funkcji typu cosinus przesuniętych nieznacznie w fazie względem siebie. Będzie to dokładnie to samo co na rysunku (Rys. 25), ale przy długości wektora A_i dążącego do zera, różnicy faz między sąsiednimi źródłami dążącej do zera (Suma wszystkich przesunięć fazowych do zera nie dąży) oraz ilości wektorów (oscylatorów harmonicznnych) dążących do nieskończoności. W praktyce chodzi o takie sumowanie tych wektorów, dla których wyobrażamy sobie ogromną ilość n wektorów z rysunku (Rys. 25). I tutaj dochodzimy do sedna idei dyfrakcji według Feynmana. Mianowicie, dyfrakcja i interferencja jest tym samym zjawiskiem, ale o interferencji mówimy, gdy mamy do czynienia z niewielką ilością źródeł fal świetlnych. Natomiast dyfrakcja związana jest z wynikiem nakładania się fal pochodzących od wielu źródeł.

- W takim razie, jaki ostatecznie wynik uzyskał Feynman w wyniku takiego dodawania się fal pochodzących od wielu źródeł pochodzących z pojedynczej szczeliny?

- Autor nie opisuje dosłownie pojedynczej szczeliny, ale sytuację wynikającą z nałożenia się fal od wielu liniowo ustawionych oscylatorów harmonicznnych. Zanim przedstawię wynik końcowy jaki uzyskał Feynman. Przyjrzyjmy się bliżej rysunkowi (Rys. 25). Chodzi o to, że wielkość wektora A_R zakreślonego przez wektory A_i leżącego na łuku zależy generalnie od przesunięcia fazy φ , a ta zależy od kąta obserwacji naszych liniowo rozmieszczonych oscylatorów. Całkowity kąt o jaki się obróci ostatni wektor A_n wynosił $n\varphi$. Jest to przesunięcie fazowe dla ostatniego oscylatora. Pełen obrót to 2π . Zatem miarą kierunku obserwacji będzie wielkość $n\varphi/2\pi$. Również natężenie fali wypadkowej można przedstawić jako wielkość odniesioną do całkowitej wartości natężenia będącej sumą natężeń wszystkich oscylatorów bez przesunięcia fazowego. Feynman uzyskał w ten sposób zależność na natężenie fali wychodzącej od pojedynczej szczeliny (dyfrakcja) w funkcji kierunku rozchodzenia się fali (Rys. 26).



Rys. 26 Natężenie jako funkcja kąta fazowego dla wielkiej liczby oscylatorów o równym natężeniu [3].

- No nieźle, to udało mu się uzyskać dwie rzeczy za jednym zamachem. Mianowicie, występowanie prążków dla pojedynczej szczeliny oraz zmienne natężenie światła dla poszczególnych prążków. Tyle, że jeszcze nie wyjaśnił mi pan, jak Feynman wyznaczył poszczególne maksyma?

- To nie jest takie trudne pani Kamilo. Proszę spojrzeć na rysunek (Rys. 25). Jeśli założymy, że wszystkie wektory A_i są sobie równe, to możemy przyjąć, że n takich wektorów zakreśla kawałek okręgu. Wówczas łatwo wyprowadzić zależność na długość wektora A_R . Wystarczy wykorzystać trójkąt QOS do wyznaczenia długości r , a następnie trójkąt QMT do wyznaczenia połowy długości A_R . Stąd:

$$A_R = A \frac{\sin(n\phi/2)}{\sin(\phi/2)} \quad \text{oraz} \quad (4)$$

$$I = I_0 \frac{\sin^2(n\phi/2)}{\sin^2(\phi/2)} \quad (5)$$

- Teraz rozumiem! Dla różnych prążków mamy inne przesunięcie fazowe ϕ . Ponieważ $\sin(\phi/2)$ zmienia się wolniej niż $\sin(n\phi/2)$ dla dużych n . Dla kolejnego ϕ takiego, że $\sin^2(n\phi/2)=1$, mianownik $\sin(\phi/2)$ zmienił się niewiele. Stąd, kolejna amplituda będzie w okolicy $\phi=3\pi/n$. Gdy ten kąt podstawimy do mianownika, to uzyskamy $\sin^2(3\pi/2n)$. Dla dużego n jest to sinus bardzo małego kąta. Dlatego za mianownik możemy podstawić $(3\pi/2n)^2$. Wówczas, natężenie w przybliżeniu będzie wynosić $I=I_0(4n^2/9\pi^2)$. Ponieważ $I_{\max}=n^2I_0$, to kolejny pik na wykresie (Rys. 26) będzie wynosił w przybliżeniu $4/9\pi^2$. Czyli około 0,047. Kolejne prążki będą jeszcze mniejsze [3].

- No widzi pani. Feynman poprawnie wyjaśnił zachowanie się światła dla pojedynczej szczeliny. Wyjaśnił nie tylko występowanie prążków, ale również spadek ich amplitudy.

- Owszem, wydaje się, że poprawnie wyjaśnił ten problem. Czy spadek amplitudy poszczególnych prążków jest dokładnie taki jak przewiduje to Feynman, tego nie wiem. Ale sprawdźmy jeszcze jedną istotną kwestię. Rozważania Younga również wydawały się poprawnie przewidywać położenia prążków. Niestety wymagają założenia, które nie daje się utrzymać w świetle przeprowadzonych doświadczeń (spójność światła). A jakie założenia Feynman wykorzystał w swoich rozważaniach? Sprawdźmy to dokładniej.

- Dobrze, przyjrzyjmy się temu bliżej: [3].

„Rozważmy zatem układ n równoległych oscylatorów o takich samych amplitudach, ale różniących się między sobą fazami albo dlatego, że ich fazy zostały inaczej ustawione, albo dlatego, że obserwujemy je pod kątem powodującym różnicę w opóźnieniu czasowym. Niezależnie od przyczyny tej różnicy faz musimy obliczyć następującą sumę cosinusów:

$$R = A\{\cos(\omega t) + \cos(\omega t + \phi) + \cos(\omega t + 2\phi) + \dots + \cos[\omega t + (n-1)\phi]\}, \quad (30.1)$$

gdzie ϕ jest różnicą faz między kolejnymi oscylatorami obserwowanymi z ustalonego kierunku. Dokładnie mówiąc $\phi = \alpha + (2\pi d \sin(\theta/\lambda))$. Musimy teraz pododawać wszystkie te wyrazy. Zrobimy to sposobem geometrycznym. Pierwszy wyraz ma moduł A i fazę zero. Następny ma także moduł A , ale fazę równą ϕ . Dalszy ma znowu moduł A , ale fazę równą 2ϕ i tak dalej.”

- Tego się właśnie obawiałam. Widzi pan panie profesorze. Feynman zrobił założenie, że fazy wyjściowe poszczególnych oscylatorów muszą być dobrze zdefiniowane na początku. Założył on, że mamy do czynienia ze źródłem fali spójnej. Zarówno przestrzennie jak i czasowo. Oznacza to tylko tyle, że Feynman poprawnie wyprowadził zjawisko dyfrakcji na pojedynczej szczelinie, ale tylko dla przypadku, kiedy mamy bardzo dobrze dobrane warunki początkowe, czyli dysponujemy źródłem w pełni spójnej fali. Jego tłumaczenie będzie zupełnie bezwartościowe w przypadku rzeczywistego doświadczenia, które ma miejsce również dla światła niespójnego. Stąd, tak naprawdę on niczego nie wytłumaczył.

- Faktycznie. To będzie istotny problem, bo każdy eksperymentator w tym momencie odwołując się do pani doświadczeń z prążkami dla światła niespójnego, będzie mógł bez problemów podważyć założenia na których oparł swoje rozważania Feynman czy Young i inni mu podobni.

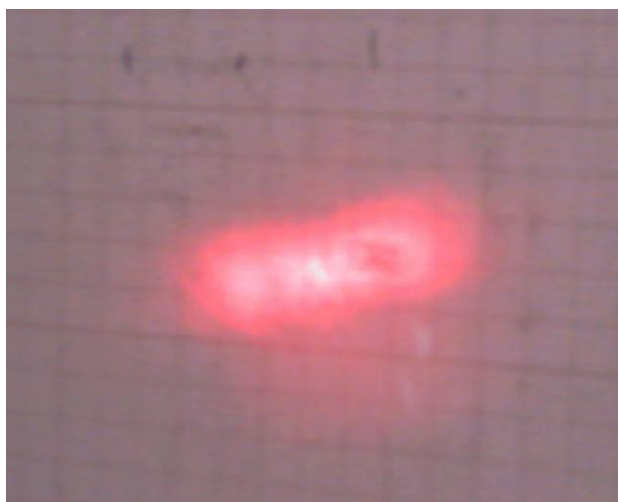
- Ostatecznie uważam więc, że całe zjawisko ma swoje miejsce w obszarze szczeliny i tylko tam. W tym obszarze dzieje się coś ciekawego, co rozdziela strumień światła na kilka niezależnych strumieni pod różnymi kątami. Dalej rozdzielone światło biegnie już niezakłócone prosto, aż dotrze do ekranu, detektora bądź innego obiektu. Wówczas, niezależnie gdzie postawi się ekran, zawsze będzie padać na niego tyle samo fotonów, co wyszło ze szczeliny pod danym kątem. Natomiast, natężenie poszczególnych prążków zależy tylko od tego ile fotonów zostanie odchylonych na krawędzi bądź szczelinie w danym kierunku.

Kamila nie wiedząc co dalej myśleć na temat światła, zaczęła chodzić po laboratorium. Zastanawiając się nad tym jak wyjaśnić rozdzielanie się światła na szczelinach. W głębi serca wierzyła, że musi istnieć jakieś racjonalne wyjaśnienie. Przechodząc obok innego stanowiska laboratoryjnego zabrała ze sobą soczewkę skupiającą i podeszła do profesora.

- Może zobaczymy, co się stanie z światłem, gdy przepuścimy go jeszcze przez soczewkę?

- Nic wielkiego nie będzie. {odparł profesor} Po prostu światło zostanie skupione w jednym miejscu i poleci dalej, dając na ścianie dużą plamę.

Kamila włączyła laser i umieściła soczewkę w osi optycznej. Na ekranie zamiast plamki światła pojawił się duży prostokąt (Rys. 27).



Rys. 27 Obraz rozproszonej przez soczewkę wiązki światła laserowego.

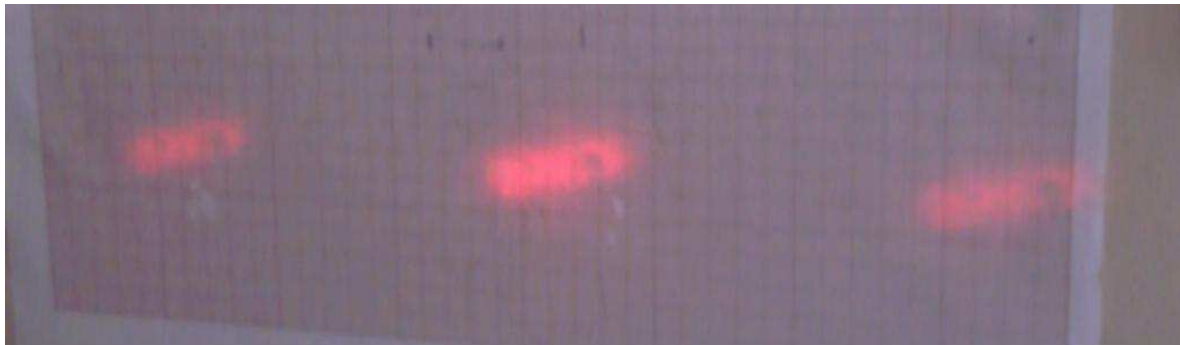
- Myślałam, że otrzymamy na ekranie dużą okrągłą plamę. Dlaczego kształt przypomina prostokąt?

- Ponieważ jest to światło pochodzące od diody laserowej. Jej kształt jest prostopadłościenny, a światło wychodzi z jednego boku, który ma akurat taki a nie inny kształt.

- Rozumiem. Czyli kształt wiązki świetlnej na całej długości musi mieć taki sam kształt. Zaś, soczewka po rozproszeniu światła ukazuje nam kształt tej wiązki, czyli prostokąt.

- Dokładnie. Skoro już bawimy się soczewką to sprawdźmy, co się stanie jak dodamy do układu siatkę dyfrakcyjną.

Kamila wzięła do ręki siatkę i postawiła ją tuż za soczewką. Ich oczom ukazał się powielony obraz przekroju wiązki, który różnił się od oryginału jedynie jasnością (Rys. 28).



Rys. 28 Obraz na ekranie przedstawia rozproszoną przez soczewkę i siatkę dyfrakcyjną wiązkę laserową.

- No tak. Obraz się rozmnożył. {zażartowała Kamila} Również jasność bocznych obrazów jest mniejsza, tak jak w oryginalnym doświadczeniu. Ciekawe, że mamy dokładnie takie same obrazy jak oryginalny.

- Czemu miałyby być inne?

- Nie wiem. Czym może się różnić takie doświadczenie od oryginalnego? Może ilością szczelin przez jakie przechodzi światło. No tak, przecież ilość szczelin przez które przechodzi światło to szerokość wiązki światła podzielona przez stałą siatki. Jeśli mamy taki przekrój wiązki jaki widzimy na ekranie to znaczy, że w różnym miejscu siatki dyfrakcyjnej światło musi przechodzić przez różną ilość szczelin. Czy obraz na ekranie zależy od ilości szczelin przez jaką przechodzi światło?

- Tak. Im szczelin przez jakie przechodzi światło jest więcej, tym powstające na ekranie prążki są węższe i na odwrót (wynik symulacji komputerowej Rys. 13 i 14).

- Rozumiem, że dotyczy to interpretacji falowej tego doświadczenia?

- Oczywiście, dlaczego pani pyta?

- Bo coś mi się w tym obrazie nie podoba. Skoro szerokość widzianego obrazu według obowiązującej teorii zależy od ilości szczelin uwzględnionych w obliczeniach, to niema możliwości byśmy uzyskali dokładnie taki sam kształt obrazu po przepuszczeniu przez siatkę dyfrakcyjną. A przecież bez siatki widzieliśmy dokładnie to samo!

- Hmm. Faktycznie, gdyby było jak mówiłem, to środek wiązki światła przechodziłby przez bardzo dużą ilość szczelin i obraz powinien być bardzo wąski. Natomiast na górze i dole przekroju wiązki światła, ilość szczelin przez które przechodzi światło maleje i obraz powinien stawać się coraz szerszy.

- Ale tak się nie dzieje. Ponadto, w środku obrazu jest dziura powstała być może z powodu zabrudzenia soczewki w laserze, która nie zmienia swojego położenia i kształtu, ale ma wpływ na ilość szczelin przez jaką przechodzi wiązka światła w środku obrazu.

- Dokładnie.

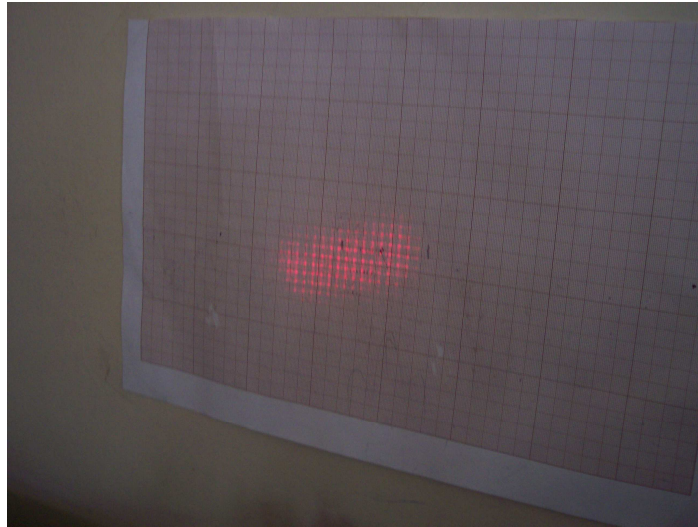
- Czyli po raz kolejny przewidywania teorii falowej są sprzeczne z doświadczeniem.

- Na to wygląda. Rano jak robiłem symulację to wydawało mi się, że wszystko się zgadza, ale się myliłem.

- Skoro zastosowanie soczewki i siatki dyfrakcyjnej pozwala uzyskać nowe nieścisłości z teorią, to może warto byłoby sprawdzić jeszcze siatkę od monitora.

- Dobry pomysł.

Po zdjęciu siatki dyfrakcyjnej Kamila przetrzymała siatkę od monitora tak jak poprzednio umieszczając ją za soczewką. Ich oczom ukazał się obraz o niezmienionej wielkości, ale składający się z wielu punktów (Rys. 29).

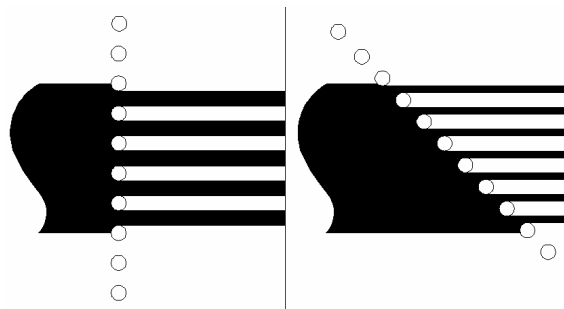


Rys. 29 Obraz na ekranie powstały z przejścia wiązki przez soczewkę i siatkę od monitora.

- Super! Coś nowego. {krzyknęła Kamila}
- Tego jeszcze nie widziałem. Zastanówmy się, co teraz widzimy? Kształt i wielkość obwiedni została niezmieniona w stosunku do obrazu uzyskanego bez siatki (Rys. 27). Co mogą przedstawiać punkty na ekranie?
- Myślę, że może otwory siatki. W końcu siatka to nic innego jak skrzyżowane prostopadłe włókna. Światło może przechodzić tylko między włóknami. I pewnie to teraz obserwujemy.
- Jak się o tym przekonać? Może pani poruszyć siatkę do przodu i do tyłu?
- Tak oczywiście.

Kiedy Kamila przesuwiała siatkę, ilość widzianych punktów ulegała zmianie, ale dla pewnego szczególnego położenia ich ilość była najmniejsza. Następnie, zmniejszając bądź zwiększając odległość siatki od soczewki ilość punktów rosła. Przy czym zakropkowany obraz dalej odtwarzał niezmieniony kształt prostokąta, który związany jest z kształtem przekroju wiązki laserowej.

- Już wiem! {krzyknęła Kamila} Kiedy mamy najmniej punktów świetlnych siatka znajduje się w ogniskowej soczewki. Oddalając bądź zbliżając ją do soczewki zwiększamy obszar przez który światło przechodzi przez siatkę. Stąd, obserwowana zmiana ilości punktów na ekranie.
- Możliwe, że ma pani rację. Sprawdźmy to jeszcze poprzez obrót siatki. Jeśli punkty na ekranie są rzeczywiście odzwierciedleniem otworów siatki, to podczas obracania jej w osi pionowej, ilość punktów na ekranie w kierunku poziomym będzie się zwiększała, a szerokość plamek światła w tym kierunku zmniejszać (Rys. 30).



Rys. 30 Schemat zmian rozchodzenia się światła podczas obrotu siatki w jednej osi.

Kamila obróciła siatkę zgodnie z poleceniem profesora. Obraz na ekranie zmieniał się zgodnie z przewidywaniami.

- Miała pani rację. Obraz faktycznie przedstawia otwory siatki w powiększeniu.
- Zobaczmy co się stanie jak damy soczewkę z drugiej strony.

Kamila przestawiła siatkę tak, by znajdowała się tuż przed soczewką (pomiędzy źródłem światła laserowego a soczewką). Dzięki temu światło z lasera wpierw przechodziło przez siatkę, a dopiero później przez soczewkę (Rys. 31).



Rys. 31 Stół laboratoryjny z laserem, siatką od monitora i soczewką tuż za siatką.

Początkowo obraz na ekranie nie różnił się niczym od poprzedniego. Obrót siatki również powodował obrót punktów na ekranie. Natomiast obrót w osi pionowej zmieniał ilość punktów widzianych w kierunku poziomym, tak jak poprzednio. Zmiana nastąpiła dopiero wówczas, gdy Kamila zaczęła zmieniać odległość soczewki od siatki. Niezależnie jak daleko była umieszczona soczewka od siatki, zawsze widzieli taką samą ilość punktów na ekranie.

- Ciekawe! {Zdziwiła się Kamila i po chwili dodała} No tak, ilość punktów na ekranie się nie zmienia mimo, że zmieniam odległość. Ale to jest akurat łatwo wyjaśnić. Jeśli wcześniej doszliśmy do wniosku, że ilość punktów jest zależna od ilości otworów przez które przelatuje światło, to oznacza to tylko tyle, że tym razem ilość oświetlonych przez laser otworów siatki nie ulegała znaczącej zmianie.

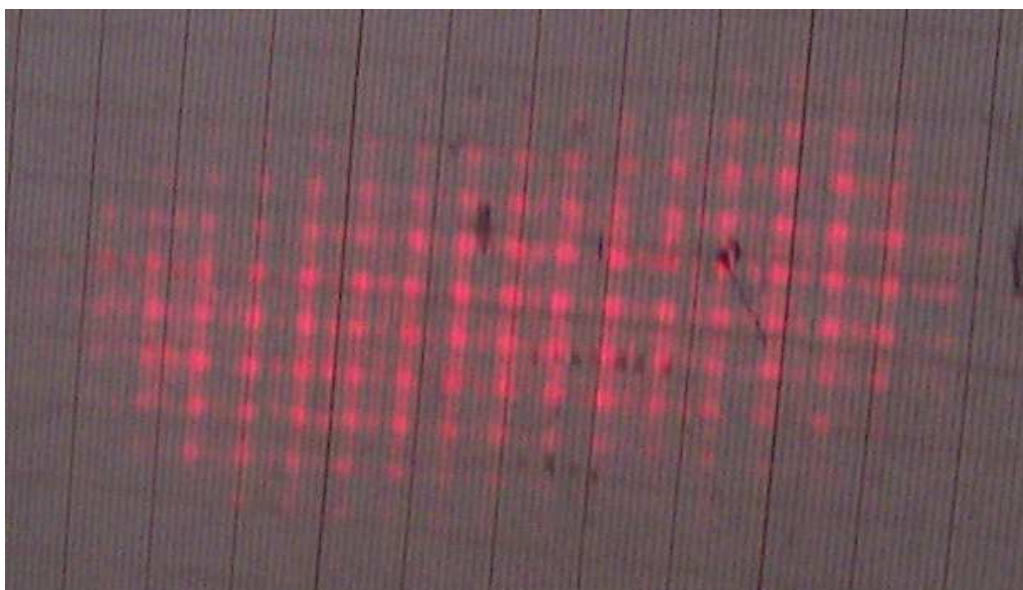
- No jasne! Przecież tym razem światło wpierw pada na siatkę i ma przy tym określony przekrój. Ponieważ światło lasera leci niemal równolegle, to podczas stosunkowo niewielkiego przemieszczenia siatki wzdłuż osi optycznej, oświetlana przez laser powierzchnia siatki niemal się nie zmienia. Za siatką światło jest już rozdzielone na poszczególne wiązki wychodzące z każdej dziury osobno i następnie odchylone przez soczewkę. Dzięki temu na ekranie widzimy w powiększeniu przekrój jaki tworzy wiązka światła za siatką.

- Super, ale gdzie są prążki "interferencyjne"? {Spytała Kamila}

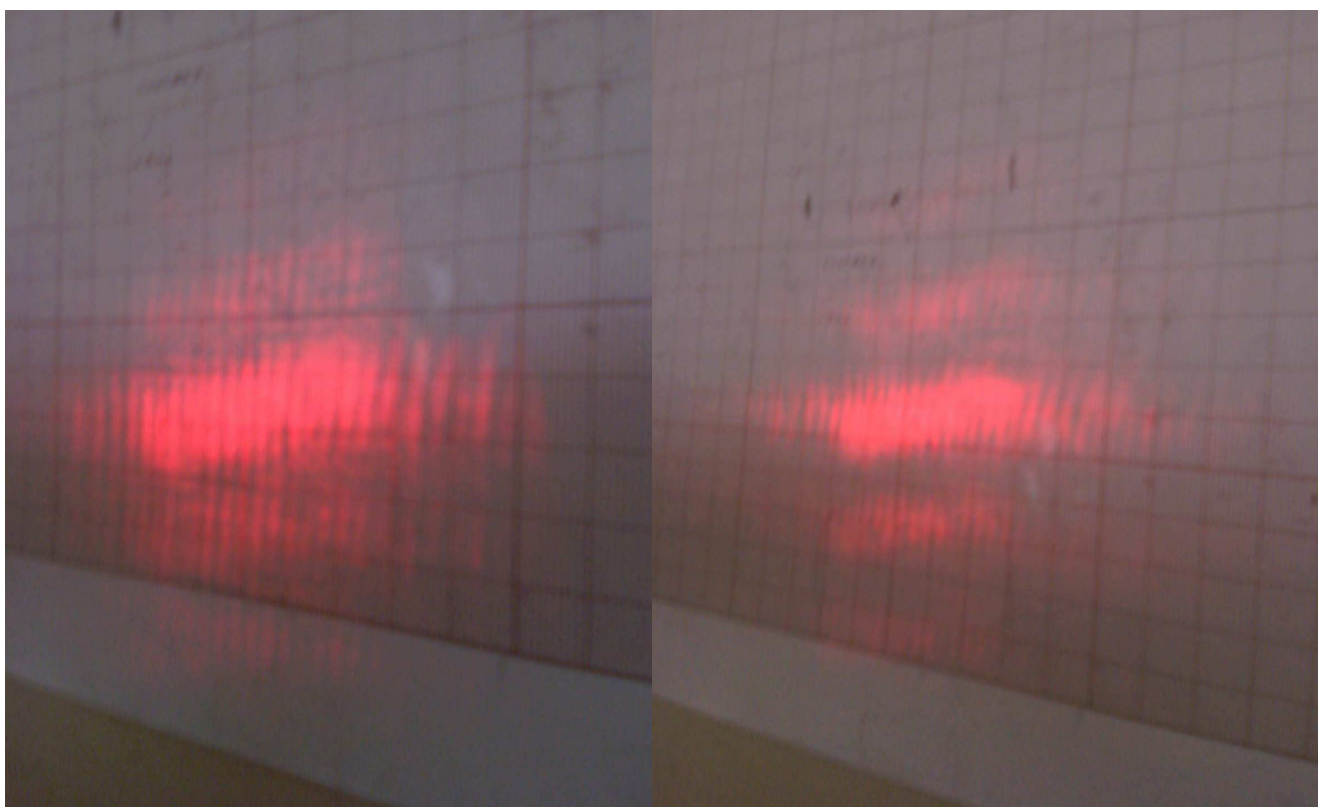
Profesor przerzucił siatkę na drugą stronę soczewki i ustawił ją w taki sposób, by było widać jak najmniejszą ilość otworów siatki.

- Są tutaj. {kontynuował} Każda mocniejsza plamka odpowiada otworowi siatki jak to już wcześniej zauważyliśmy. Kiedy się dobrze przyjrzymy, a dla takiego ustawienia widać to najlepiej (siatka w ogniskowej soczewki), każdy obraz otworu siatki posiada swój własny obraz interferencyjny.

- Faktycznie, już widzę. Przy wcześniejszym ustawieniu przyrządów (siatka ustawiona pomiędzy źródłem światła a soczewką) też się da zobaczyć dodatkowe prążki (Rys. 32). Może jak przesuniemy soczewkę dalej w kierunku ekranu to będzie to lepiej widać (Rys. 33).



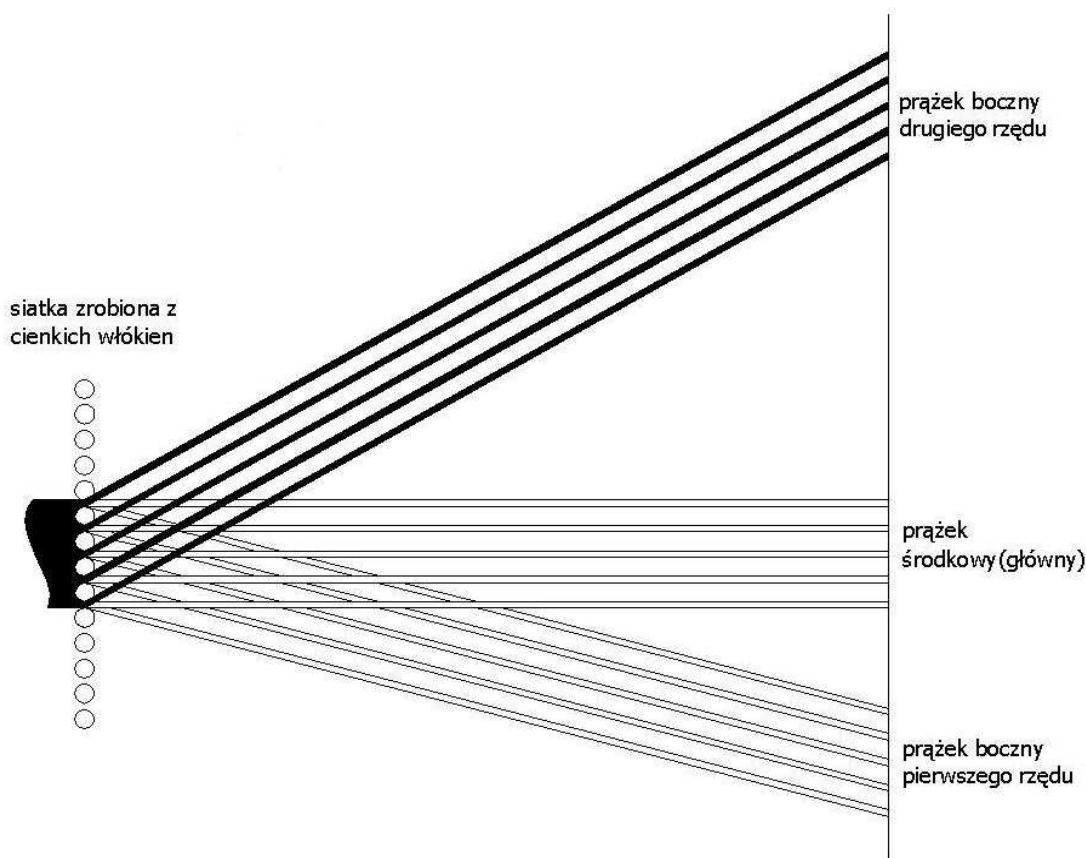
Rys. 32 Obraz na ekranie powstały z przejścia wiązki przez siatkę i soczewkę. Jaśniejsze punkty określają położenie otworów siatki. Słabsze punkty po obu stronach są prążkami „interferencyjnymi”.



Rys. 33 Obraz na ekranie powstały po przepuszczeniu wiązki laserowej przez siatkę i soczewkę rozpraszającą oddaloną nieco od siatki. Na prawym rysunku widzimy dalsze rozsuniecie się prążków bocznych podczas większego oddalania soczewki od siatki.

Podczas odsuwania soczewki od siatki obraz zaczął rozdzielać się na oddzielne obszary przedstawiające pozostałe prążki. Dla dostatecznie dużej odległości przypominał już ten z podstawowego doświadczenia (bez soczewki, Rys. 3). Niemniej jednak, zanim obraz prążków zlał się w jedną całość (coraz mniejszy obraz) można było zauważyć, że on również składał się z punktów, które łatwo można przypisać otworom siatki. Wielce to zdumiało zarówno Kamilę jak i profesora.

- Teraz już widzę, gdzie się ukryły nasze prążki. {zażartowała sobie}
- No tak, ale interesujące jest to, że one również składają się z punktów odpowiadających otworom w siatce oraz wyznaczają prostokątny kształt padającej na siatkę wiązki światła tak jak dla obrazu głównego.
- Myślę, że można to wyjaśnić tylko w ten sposób, że z każdego otworu musi niezależnie wychodzić kilka wiązek światła, a wiązki wychodzące pod danym kątem biegną równoległe (Rys. 34). {Z dumą dodała Kamila wiedząc, że jest to absolutnie sprzeczne z falową interpretacją}



Rys. 34 Schematyczne przedstawienie mechanizmu powstawania obrazu prążka bocznego składającego się z punktów świetlnych odwzorowujących otwory siatki i kształt wiązki światła padającego na siatkę.

- Ale to wymaga założenia, że światło jest rozpraszane dyskretnie pod określonymi kątami na każdym otworze (krawędziach) niezależnie dla każdego fotonu! {Jeśli oddziaływanie krawędzi na światło ma większy zasięg, niż odległość między szczelinami. To należy się spodziewać, że odchylenie światła na poszczególnych szczelinach nie będzie tak całkiem niezależne. Doświadczenie z pojedynczą szczeliną i przesłoną (Rys. 21) pokazało, że takie oddziaływanie jest dość znaczne w stosunku do długości fali światła. O czym można się samemu przekonać, wykonując wspomniane doświadczenie samodzielnie.}

- Zgadza się. Ale takie założenie tłumaczy, dlaczego widziany obraz prążków i głównej wiązki światła przedstawia otwory siatki przez który światło przechodzi. Soczewka tylko powiększa nam obraz przekroju wiązki, który biegnie już za siatką. Dzięki czemu możemy zobaczyć jak on wygląda naprawdę, a nie jako

pojedyncze rozmyte małe plamy. Tłumaczy to również, dlaczego dla siatki dyfrakcyjnej kształt poszczególnych prążków był taki sam, jak kształt wiązki głównej (Rys. 28).

- No tak, teoria falowa nie radzi sobie z odtworzeniem poprawnego kształtu prążka bocznego w tym doświadczeniu.

- Również w ten sposób można łatwo wyjaśnić kolejną niedogodność dla falowej interpretacji.

- Jaką?

- Zachowanie jasności prążków bocznych, które niemal nie zmieniają się wraz z odległością od siatki dyfrakcyjnej. Dla takiej interpretacji nieważne jest, gdzie znajduje się ekran. I tak powinniśmy uzyskać prążek o niemal niezmiętej jasności.

- Trudno się z tym nie zgodzić, jeśli przyjmujemy, że światło odchylone na szczelinach biegnie dalej niezakłócone w linii prostej. A o istnieniu prążków i ich jasności miałyby decydować tylko to, co się z nim dzieje w okolicy szczeliny.

- Przy takim założeniu, znika również problem tłumaczenia się: Dlaczego powstaje obraz prążkowy dla pojedynczej szczeliny, bądź krawędzi? Nie trzeba kombinować, co z czym interferuje i wymyślać absurdalnych koncepcji.

- A co jest nie tak z obrazem pochodzącym od pojedynczej krawędzi, pani Kamilo?

- Jak to co? Obraz prążków jest symetryczny i powstaje również po stronie cienia. Ponadto, prążki leżą szerzej niż szerokość padającej na krawędź wiązki światła. To niby co i z czym tam interferuje? A przecież nic z niczym w ogóle interferować nie musi. Wystarczy tylko znaleźć przyczynę, dla której światło (bądź cząstki) ulegają na krawędzi dyskretnemu podziałowi na kilka wiązek i sprawa rozwiązana.

- Niby tak, ale dla wielu fizyków to i tak będzie trudne do przełknięcia. Ale w świetle przeprowadzonych przez nas tutaj doświadczeń nie ma wątpliwości, że będziemy musieli taką ideę przemyśleć głębiej. Być może należy poszukać jeszcze innych doświadczeń, które nas naprowadzą do rozwiązania. O! Proszę zauważyć, że z pani rysunku (Rys. 34) wynika jeszcze jedna konsekwencja takiego, a nie innego wyjaśnienia obserwowanych w doświadczeniu efektów. Jeśli prążki boczne składają się z punktów, które przedstawiają otwory siatki (co pokazuje doświadczenie!). Wówczas oczywiste jest, że światło z każdego otworu musi biec równolegle w stosunku do innych strumieni światła. Tyle, że wówczas światło pochodzące z poszczególnych szczelin, teoretycznie nigdy nie spotyka się na ekranie!

- A to całkowicie wyklucza ideę interferencji w tym doświadczeniu.

- Zresztą, koncepcja odwołująca się do rozproszenia się wiązki na kilka innych, na każdej krawędzi niezależnie, jest spójna z koncepcją tłumaczącą pojawienie się prążkowego obrazu na siatce (Rys. 11).

- Czyli, zaczyna się nam rysować dość spójny obraz tego, co dzieje się w doświadczeniu Younga. Fotony są rozpraszane dyskretnie na każdej krawędzi niezależnie. Jeśli mamy wiele takich samych krawędzi blisko siebie (takie same warunki rozpraszania fotonów), wówczas w pewnych kierunkach znajdzie się wystarczająco dużo odchylonych w ten sposób fotonów byśmy mogli zaobserwować je na ekranie w postaci oddzielnego prążka. Takie tłumaczenie nie wymaga założenia spójności światła, gdyż na każdej szczelinie fotony są rozpraszane niezależnie. Wówczas, przyczyna obserwowanego zjawiska ma miejsce w obszarze szczeliny, a nie na ekranie i założenie o spójności "fali świetlnej" nie jest istotna. Należy przy tym zbadać, czy zachowanie fotonów w obszarze szczeliny nie zależy również od: geometrii szczeliny (nie tylko odległości między szczelinami), materiału z którego zrobiona jest szczelina oraz temperatury szczeliny.

- Słuszna uwaga pani Kamilo, teraz trzeba znaleźć dokładniejszy opis tego, co się dzieje w obszarze szczeliny oraz wyjaśnić, dlaczego światło zachowuje się tam tak, a nie inaczej? Ponadto, interesujące jest to, w jakimś stopniu bliskie sąsiedztwo innych szczelin lub krawędzi wpływa na rozproszenie poszczególnych fotonów. Działanie każdej z krawędzi wcale nie musi być takie całkiem niezależne, ale to wymaga głębszych studiów nad tym zagadnieniem. Przykładowo, światło przechodzące przez pojedynczą szczelinę oddziałuje jednocześnie z dwoma krawędziami itd.. Zastanawiam się, czy zrobić wprawdzie symulację opartą o rozważania falowe, by sprawdzić, czy uda się uzyskać teoretycznie na ekranie obraz odzwierciedlający rozmieszczenie poszczególnych otworów siatki w ramach prążka „interferencyjnego”. Tak jak to widzimy w doświadczeniu (Rys. 33).

- Myślę, że to jest ciekawy pomysł, ale mam obawy panie profesorze, czy nie będzie tutaj takiego samego problemu jak w rozważaniach pana Feynmana? Chodzi o to, że rozważania falowe oparte na interferencji jak

już stwierdziliśmy, są wrażliwe na warunki początkowe. Dokładnie chodzi o liczenie różnicy faz wynikających z różnicy dróg przebytych przez poszczególne wiązki światła. Takie podejście wymaga niestety dobrego zdefiniowania warunków początkowych, czyli założenie spójności światła użytego w doświadczeniu. A już udowodniliśmy wcześniej doświadczalnie, że zjawisko zachodzi również dla światła niespójnego. Gdyby nie to ograniczenie, to mogłabym powiedzieć, że Feynman już pokazał, że dla pojedynczej szczeliny mamy prążki. Jeśli założę, że każda szczelina działa zupełnie niezależnie, to rozpatrując każdą szczelinę z osobna mamy dokładnie taką sytuację jak tą przedstawioną tutaj na rysunku (Rys. 34). Z tym zastrzeżeniem, że dla rozpatrywania pojedynczej szczeliny jako tworów niezależnych od innych szczelin, powinniśmy uzyskać położenie prążków uzależnione od wielkości szczeliny, a nie odległości między szczelinami. Tymczasem, porównanie wymiarów siatki obliczonych na podstawie położenia prążków na ekranie ze zdjęciem siatki pod mikroskopem (Rys. 6 i Tab. 1) pokazało, że wynik jest bliższy odległości między szczelinami, a nie odległości między krawędziami.

Jeśli natomiast zrezygnujemy z założenia o spójności światła, to będzie miał pan zawsze problem z uczciwym określeniem warunków początkowych dla swojej symulacji. Przy czym, za „uczciwe” przyjmuję takie ustalenie warunków początkowych, które zakłada przypadkową i zmienną w czasie fazę dla poszczególnych wiązek światła wchodzących do symulowanego układu dla każdej szczeliny niezależnie. Obawiam się profesorze, że w taki sposób niczego sensownego w symulacji opartej na interferencji różnych fal pan nie uzyska.

- Niestety, ale ma pani chyba rację. Jak zwykle wszystko znów zamyka się na podstawowym założeniu dotyczącym spójności światła. Teraz to już chyba wszystko na dzisiaj. Zrobiło się już późno i czas się zbierać. Jak dojdzie pani jeszcze do innych ciekawych uwag dotyczących omawianych zagadnień, to proszę się nimi podzielić.

- Oczywiście panie profesorze. {Z dumą odpowiedziała Kamila i wyszła}

Przez kilka kolejnych dni profesor chodził notorycznie niewyspany. W jego głowie przewijało się tak wiele myśli, że miał trudności z zaśnięciem. Musiał przecież sobie wszystko na nowo w głowie poukładać. Ogromny bagaż wiedzy jaki nabrał w ciągu swojego życia nie ułatwiał mu tego zadania. Wręcz przeciwnie, trudno mu było się pogodzić z tym, że tak wiele zagadnień, które uważał za oczywiste i dobrze ugruntowane, będą wymagały weryfikacji w myśl nowych doświadczeń. W pewnym momencie zaświtała mu myśl, by rozważyć eksperyment myślowy, który polegałby na zamknięciu w skrzynce z małą dziurką nieznanego źródła światła. Wychodziłoby ono na zewnątrz tylko przez tę dziurkę w postaci wąskiego strumienia. Następnie zadał sobie pytanie; Co można powiedzieć na temat tego światła i jakich doświadczeń należy w tym celu użyć? Pierwsze co mu przyszło do głowy dotyczyło barwy światła. By to określić wystarczy przeciąć wiązkę światła np. białą kartką i zobaczyć jaka będzie barwa plamki świetlnej.

- Czy na pewno? {pomyślał} - No tak, oczywiście, że nie. Przecież już Newton zauważył, że światło może mieć różne barwy (widziane) w zależności od mieszaniny innych barw z których może się składać. Należy więc, zrobić to, co wieki temu wykonał angielski uczoney. Przeciąć wiązkę światła pryzmatem. Ponieważ światło o różnych barwach różnie się załamuje na szkło, światło rozdzieli się na wszystkie barwy z jakich się składa. Obecnie taki rozkład nazywa się rozkładem spektralnym światła i pozwala dokładnie wyznaczyć rozkład długości fal badanego źródła. W tym celu niekiedy używa się siatki dyfrakcyjnej.

Czy można coś więcej określić w kwestii takiej hipotetycznej wiązki światła? Owszem, przecież światło może być spolaryzowane. W tym celu wystarczy wziąć do ręki polaryzator i umieścić go na osi optycznej. {nic prostszego - pomyślał} Następnie światło skierować na ścianę i obracając polaryzator sprawdzić, czy jasność plamki się zmienia. Jeśli się okaże, że tylko niektóre barwy są spolaryzowane, to wówczas nie tylko zmieniałaby się jasność, ale i barwa. Tak więc, polaryzacja poszczególnych barw wiązki światła jest tą własnością, którą również można łatwo wyznaczyć doświadczalnie. W przypadku bardziej skomplikowanego rozkładu polaryzacji poszczególnych barw, można przecież posłużyć się również pryzmatem, a następnie sprawdzić polaryzatorem każdą barwę z osobna.

Jak na razie nie znalazł trudności doświadczalnych z określeniem takich cech jak rozkład barw i ich polaryzacja. Jaką jeszcze cechę może posiadać taka wiązka światła? Oczywiście może być spójna, bądź nie. Często przyjmuje się, że światło jest spójne, gdy mamy do czynienia ze źródłem światła pochodzącego od

lasera, ale w doświadczeniu przyjął przecież, że nic nie wie o źródle rozważanego przez siebie światła. Zastanowił się więc:

- Jakże należy wykonać doświadczenie by wykazać, że dane światło jest bądź nie jest spójne? Pierwsze co mu przyszło do głowy to doświadczenie Younga. Przecież niemal w każdej instrukcji laboratoryjnej poświęconej temu doświadczeniu jest wspomniane, że jest to warunek niezbędny do tego, by doświadczenie się udało. Niestety, nie! Przecież dobrze pamiętał jak oglądał obrazy „interferencyjne” dla światła niespójnego. Tak więc, niezależnie od przyczyny powstawania tych obrazów, wspomniane doświadczenie nie może być już zastosowane do zweryfikowania spójności wiązki światła.

Jakże jeszcze inne doświadczenie nadaje się do weryfikacji tej własności? {zastanawiał się Profesor}

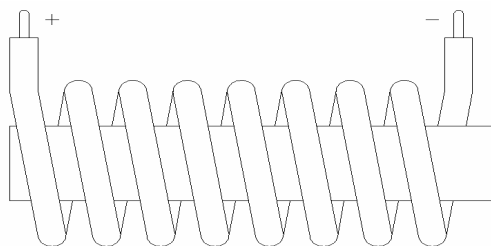
Szukając odpowiedzi na zadane przez siebie pytanie, profesor zadał sobie inne pomocnicze; W jakich doświadczeniach ważna jest spójność światła? Dość szybko doszedł do wniosku, że muszą to być te doświadczenia, które odwołują się bezpośrednio do interferencji. Wynika to dokładnie z tych samych powodów, co w doświadczeniu Younga. Postanowił przemyśleć kilka z nich i sprawdzić, czy faktycznie nadadzą się do roli weryfikatora cechy spójności światła. Jako pierwsze na warsztat wziął pierścienie Newtona oraz barwy powstające na cienkich warstwach. Niestety, nie był to dobry wybór. Szybko sobie przypomniał, że dla obu doświadczeń wystarczy zastosować światło słoneczne, które zdecydowanie nie jest spójne. Barwy powstające na cienkich warstwach widział chyba każdy, kto oglądał rozlaną ciekłą warstwę oleju na asfalcie bądź kałuży. Kolorowe bańki mydlane też są interpretowane jako interferencja światła na cienkich warstwach, a przecież chyba każdy już jako dziecko podziwiał ich piękne barwy w świetle dziennym. Dla tych wszystkich zjawisk wystarczyło niespójne światło słoneczne. (Czytelnikowi zostawię szczegóły udowodnienia, że dla wspomnianych doświadczeń w myśl obowiązujących interpretacji spójność światła jest niezbędna, o czym często się nie wspomina! Wynika to stąd, że światło odbite od powierzchni górnej „interferuje” ze światłem pochodzącym z innej części przekroju padającej wiązki światła, które to odbiło się na dolnej powierzchni cienkiej warstwy. Stąd, w całym przekroju wiązki padającej powinna być taka sama faza - światło spójne!).

Z pierścieniami Newtona też lepiej nie było, gdyż przypomniał sobie jak sam na wykładach pokazywał kolorowe koncentryczne pierścienie, które powstają po przejściu światła dziennego przez płytkę i soczewkę o bardzo dużym promieniu krzywizny. Na końcu przyszedł mu jeszcze do głowy pomysł z interferometrem. Początkowo ucieszył się, gdy przypomniał sobie, że zasada działania interferometru opiera się na nakładaniu na siebie dwóch fal spójnych, co prowadzi do powstania obszarów wygaszania oraz wzmacniania drgań. Tyle, że po chwili przypomniał sobie również, że i w tym temacie też będzie krucho, gdyż niejaki pan Michelson, o którym wspominał już wcześniej Kamili, ponad sto lat temu wynalazł i zbudował swój interferometr. Czyli znacząco przed wymyśleniem i zbudowaniem laserów. Stąd, interferometr Michelsona musiał pracować siłą rzeczy na świetle niespójnym.

Profesor zwątpił w to, że znajdzie jakiekolwiek doświadczenie, które nadawałoby się do wykazania spójności światła. Przypomniał sobie słowa Kamili: „Interpretacja falowa odwołująca się do interferencji jest interpretacją czysto geometryczną i jako taka zawsze odwołuje się do różnicy dróg, by w ten sposób wyjaśnić pojawienie się interferencji konstruktywnej (wzmocnienie) bądź destruktywnej (wygaszenie).” Co za tym idzie, zawsze dla takiej interpretacji będzie wymagane dobre zdefiniowanie warunków początkowych (początkowej fazy światła). W praktyce oznacza to konieczność wprowadzenia założenia, że użyta w doświadczeniu wiązka światła musi być spójna: przestrzennie, czasowo bądź jedno i drugie jednocześnie. W tym momencie profesor zrozumiał, że w praktyce wystarczy dla każdego zjawiska, w którym odwołujemy się do interferencji światła bądź fali materii wykazanie, że występuje ono również dla światła niespójnego by obalić falową interpretację. Ponieważ nie potrafił znaleźć w pamięci żadnego doświadczenia, które faktycznie (praktycznie) wymagałoby takiego założenia, to wielce się tym zmartwił. Czyżby nie potrafił wykazać doświadczalnie; Czy dane światło jest, bądź nie jest spójne? Jeśli tak, to pojęcie spójności jest pojęciem zupełnie niepotrzebnym, czysto teoretycznym i niemożliwym do weryfikacji doświadczalnie. Być może w ogóle nie istnieje coś takiego jak światło spójne, skoro nie potrafię tego sprawdzić doświadczalnie!

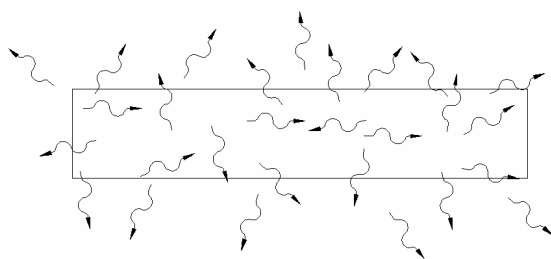
Po tych wnioskach postanowił nieco bliżej przyjrzeć się działaniu lasera, skoro powszechnie twierdzi się, że jest ono źródłem światła spójnego. Sercem każdego lasera jest ośrodek czynny, który umożliwia wzmocnienie natężenia światła pod wpływem jego wzbudzenia. W uproszczeniu zakłada się, że istnieją w

takim ośrodku miejsca (domieszki sieci krystalicznej, bądź cząsteczki np. CO_2 w laserze gazowym), które są w stanie przechwycić część energii z zewnątrz (przejsć do stanu wzbudzonego), a następnie oddać tą energię w postaci wyemitowanego fotonu. Energia potrzebna na wzbudzenie tych szczególnych miejsc pochodzi od światła o większej energii, wyładowania elektrycznego bądź innych przyczyn (np. reakcji chemicznych) (Rys. 35).



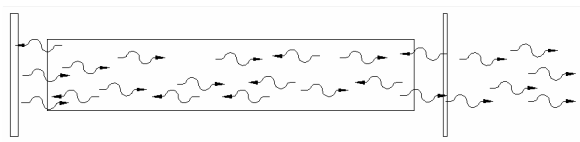
Rys. 35 Rdzeń lasera na bazie ciała stałego z nawiniętą lampą ultrafioletową.

By mogło być emitowane światło, rdzeń laserowy musi być dodatkowo intensywnie chłodzony. Powstaje wówczas tak zwana inwersja obsadzeń, która zgodnie z obowiązującym rozumieniem zjawisk jakie tam zachodzą prowadzi do promienistego pozbywania się energii w postaci światła o określonej barwie - świecenie (Rys. 36). W innym przypadku, ośrodek jedynie się ogrzewa.



Rys. 36 Promieniste wyświecanie energii zgromadzonej w rdzeniu laserowym.

Takie podejście nie spowoduje pojawienia się wiązki laserowej, a jedynie świetlistej poświaty kryształu, którą łatwiej przyrównać do świecenia się lampy z filtrem (jedna barwa) bądź zjawiska luminescencji. Atomy w stanie wzbudzonym po pewnym czasie wypromieniują swoją energię, ale w przypadkowym kierunku i fazie. Aby powstała wiązka laserowa należy do układu wstawić jeszcze dwa lustra, w tym jedno półprzepuszczalne (Rys. 37).



Rys. 37 Rdzeń laserowy otoczony dwoma lustrami, które umożliwiają wyświecanie energii promienistej w jednym kierunku.

Wówczas okazuje się, że wzbudzony ośrodek chętnie wypromieniuje swoją energię przed czasem wtedy, kiedy pada na niego foton o energii takiej samej jaka została wyemitowana. Przyjmuje się, że

wyemitowany dodatkowy foton posiada taki sam kierunek i fazę jak ten, który zainicjował jego emisję. Ponieważ lustro umieszcza się tak, by fotony biegnące w danym kierunku mogły zawrócić i ponownie wejść w obszar rdzenia laserowego. Wzbudzają one tym samym kolejne fotony, których kierunek wyznaczony jest przez lustro i staje się on dominującym kierunkiem emisji światła. Po krótkim czasie intensywność światła biegnącego prostopadle do luster wzrasta na tyle, że światło emitowane w innych kierunkach jest niemal pomijalne.

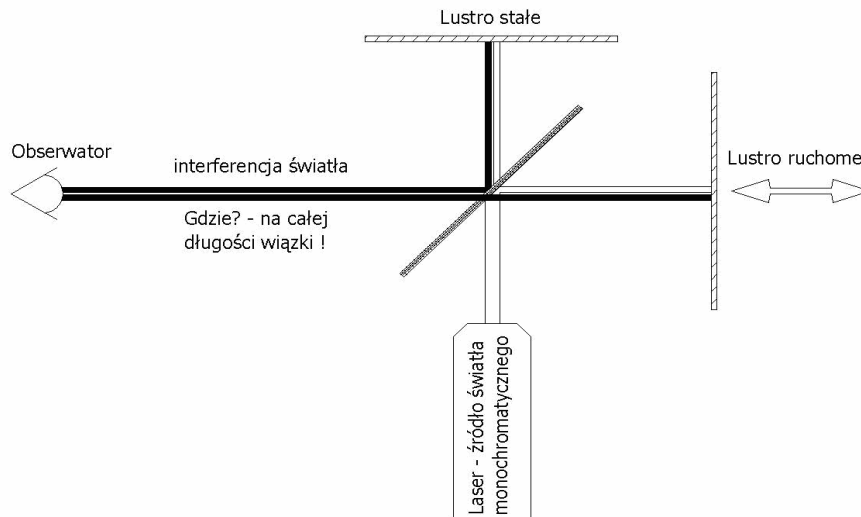
To właśnie przyjęcie, że wyemitowany foton posiada kierunek i fazę fotonu inicjującego emisję kolejnego fotonu jest głównym argumentem za tym, by światło lasera traktować jako spójne. Czy aby na pewno tak jest? {zastanowił się profesor} A co będzie, jeśli atomy emitują światło o tej samej energii (barwie) i kierunku co foton wyzwalaający świecenie, ale niekoniecznie dokładnie w tym samym momencie czasu? Np. z małym przypadkowym opóźnieniem. Wówczas twierdzenie, że jest to światło spójne nie jest uzasadnione. Nie posiadając jednak przekonujących argumentów, że tak mogłoby być, profesor przyjął oficjalną interpretację jako prawidłową i zaczął szukać innych możliwych nieścisłości.

W pewnym momencie zdał sobie sprawę, że każdy laser wytwarza światło niespójne w momencie rozruchu. Pierwszy etap polega przecież na wzbudzeniu ośrodka czynnego, który początkowo emituje światło w przypadkowym kierunku i różnych miejscach rdzenia laserowego niezależnie (niespójnie). Dopiero wzmocnienie tych fotonów, które zostały wyemitowane prostopadle do luster, powoduje wzmocnienie światła w żądanym kierunku i praktycznie zanik emisji w pozostałym. Tak się jednak składa, że w początkowym etapie w kierunku luster mogło zostać wyemitowanych zupełnie niezależnie wiele fotonów, które pochodziłyby z różnych części rdzenia i utworzyłyby wspólną wiązkę wyjściową. Nie ma przy tym żadnych podstaw by twierdzić, że którykolwiek z tych fotonów jest lepszy od innych. Są one równoprawne i teoretycznie każdy z nich ma taką samą szansę na kolejne powielenie. Oznaczałoby to tyle, że końcowa wiązka wychodząca z lasera składa się z dużej ilości tak powielonych fotonów pierwotnych, które nie były i nie są spójne między sobą. Ostatecznie, nie ma podstaw teoretycznych do twierdzenia, że wiązka laserowa jest wiązką spójną. Skoro pierwotnie spójną nie była, po wzmocnieniu również nią nie będzie, gdyż wiele „spójnie” wzmocnionych wiązek niespójnych nadal tworzy wiązkę niespójną. Uzyskany rezultat rozważań teoretycznych niespecjalnie zmartwił profesora, skoro już wcześniej doszedł do wniosku, że i tak nie zna żadnego doświadczenia, które umożliwiłoby mu praktyczne sprawdzenie spójności jakiegokolwiek badanej wiązki światła. Stąd, wyciągnięte przez niego wnioski miały charakter czysto akademicki bez specjalnego wpływu na dalsze rozważania.

Następnego dnia profesor spotkał Kamilę na holu prowadzącym do audytorium.

- Dzień dobry panie profesorze! {Krzyknęła i podbiegła uradowana ze spotkania}
- Dzień dobry, jak się spało? Wygląda pani na strasznie niewyspaną.
- Owszem. Niemal całą noc nie spałam szukając w internecie jakiś nowych materiałów o świetle i interferencji. Zastanawiałam się; Jakie to jeszcze inne ciekawe doświadczenia można byłoby wykonać?
- Znalazła pani coś nowego?
- Chyba tak. Przeglądałam materiały na temat interferometrów i znalazłam o nich informację, że wymagają koniecznie światła spójnego. Czy moglibyśmy sprawdzić konieczność zastosowania światła spójnego w tym doświadczeniu?
- Nie, nie ma takiej potrzeby pani Kamilo. Ja również nie próżnowałam wczorajszego wieczoru i już wiem, że dla interferometru spójność światła potrzebna nie jest.
- Spodziewałam się tego, ale tak naprawdę to chciałam sprawdzić coś jeszcze.
- Co takiego?

Kamila wyjęła zeszyt, gdzie miała narysowany schemat interferometru Michelsona, który składał się z lasera jako źródła fali świetlnej (interpretacja falowa), dwóch luster w tym jedno przesuwne oraz przesłony półprzepuszczalnej (Rys. 38).



Rys. 38. Schemat interferometru Michelsona. Światło przechodzi przez półprzepuszczalną powierzchnię, gdzie zostaje rozdzielone na dwie wiązki o tych samych natężeniach. Następnie po odbiciu od luster powraca, by po ponownym złożeniu interferować ze sobą.

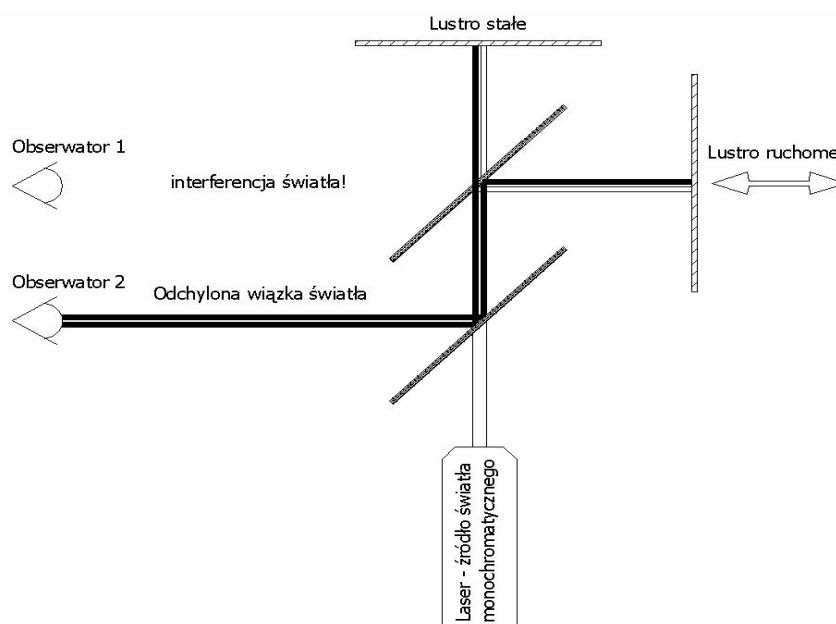
- Jak wiadomo. {kontynuowała rozmowę} Płytką półprzepuszczalną z założenia ma rozdzielać wiązkę światła na dwie wiązki o tej samej amplitudzie fali. Następnie rozdzielone wiązki przechodzą różną drogę optyczną i po ponownym połączeniu interferują ze sobą. W zależności od przesunięcia fazowego między ponownie połączonymi wiązkami zależy, czy obserwator zobaczy światło czy nie? Jeśli względna faza nie ulegnie zmianie, to mamy interferencję konstruktywną i obserwator widzi plamkę. Jeśli przesunięcie fazowe między wiązkami jest w przeciwfazie, to wiązki wygaszają się i obserwator nie widzi światła.
- Aby to miało sens, to światło musi być spójne.
- Oczywiście. Na razie zakładam poprawność interpretacji falowej.
- Rozumiem, proszę więc kontynuować pani Kamilo.
- W przypadku pośrednim mamy do czynienia z częściowym wygaszeniem światła. Jak będziemy zmieniać położenie jednego z ramion interferometru, to zmianie ulegnie droga jednej z wiązek światła i można wówczas obserwować zmianę jasności plamki na ekranie, aż do całkowitego wygaszenia. Przynajmniej tyle w teorii.
- Dokładnie, tak zakłada interpretacja falowa.
- Jak widzimy na schemacie (Rys. 38), od luster powracają odbite wiązki światła w kierunku płytki półprzepuszczalnej. Każda z tych wiązek światła niesie ze sobą energię, która może posłużyć np. do wypalenia dziury w kartce. Jeśli mamy interferencję konstruktywną to sprawa jest prosta. Cała energia światła idzie w kierunku obserwatora i obserwator może również wypalić dziurę (przypuśćmy, że dwa razy szybciej niż dla pojedynczej wiązki). Gdzie się jednak podziela energia światła w przypadku interferencji destruktywnej? Wówczas po stronie obserwatora kartka pozostałaby nietknięta.
- Istotnie, to ważne pytanie pani Kamilo.
- Czy w takim razie, zasada zachowania energii nie jest zachowana w tym doświadczeniu?
- Nie, to jakiś absurd! Zasada zachowania energii zawsze musi obowiązywać. Trzeba raczej poszukać jakiegoś wyjaśnienia tej trudności. W przeszłości wielokrotnie zdawało się różnym naukowcom, że obserwują niewyjaśniony ubytek, bądź nadmiar energii. Niemniej jednak, po dokładniejszych badaniach i przemyśleniu problemu, zasada zachowania energii dawała się zawsze obronić. Dlatego pani Kamilo, należy raczej poszukać tej brakującej energii. Czy wymyśliła już pani, gdzie mogłaby się ona podzielić?
- Owszem. Wpierw zadałam sobie pytanie; Czy wydzielą się ona w miejscu detekcji, w całej przestrzeni wzdłuż interferujących wiązek, czy na płytce półprzepuszczalnej? Następnie, zauważyłam, że z punktu widzenia formalnego nie jest ważne, jak daleko znajduje się obserwator. Jeśli założymy, że obserwator może

być gdziekolwiek, to może być również tuż za przesłoną półprzepuszczalną. Jeżeli mamy do czynienia z interpretacją destruktywną, obserwator nie powinien widzieć promienia światła. Postawiłam więc kolejne pytanie, czy podczas interferencji destruktywnej płytkę półprzepuszczalną nagrzewa się mocniej niż podczas interferencji konstruktywnej, gdzie jej nagrzewanie związane jest tylko z normalną niewielką absorpcją światła przez ośrodek materialny?

- Nie będzie łatwo sprawdzić to doświadczalnie. {wtrącił profesor}

- Wiem. Wymagałoby to czegoś więcej niż tylko studenckiego laboratorium. Ponieważ hipoteza o nagrzewaniu się płytki półprzepuszczalnej wydawała mi się mało realna, a przynajmniej trudna do weryfikacji doświadczalnie. Poszukałam innego wyjaśnienia.

Kamila dorysowała do wcześniejszego rysunku drugą płytkę półprzepuszczalną i dodatkowego obserwatora (Rys. 39).



Rys. 39 Schemat interferometru Michelsona z dodatkową płytką półprzepuszczalną. Światło przechodzi przez półprzepuszczalną powierzchnię, gdzie zostaje rozdzielone na dwie wiązki o tych samych natężeniach. Następnie po odbiciu od lusterek powraca, by po ponownym dotarciu do półprzepuszczalnej płytki skierować się w kierunku źródła światła. Dodatkowa płytkę półprzepuszczalną ujawnia powracającą w kierunku lasera wiązkę światła.

- Teraz rozumiem, co pani chce zrobić. Niestety w laboratorium, w którym ma pani zajęcia nie ma interferometru. Będzie trzeba poprosić innych, by pozwolili nam wykonać to doświadczenie, bądź wykonali je sami. Przyznaję pani Kamilo, że pani pomysł jest genialnie prosty. Wystarczy zaledwie jedna dodatkowa półprzepuszczalna przesłona by sprawdzić, czy podczas „interferencji destruktywnej” fotony po prostu nie powracają w kierunku źródła światła. Wówczas, nie byłoby problemu ze znikającą energią fotonów.

- Proszę jednak zauważyć, że z takiego doświadczenia można wyciągnąć znacznie więcej. Jeśli chodzi o przesunięcie fazowe, warunek na otrzymanie prążka w zgodzie z interpretacją falową jest taki sam dla wiązki biegnącej do obserwatora górnego jak i dolnego. I to niezależnie, czy jest to interferencja konstruktywna czy destruktywna.

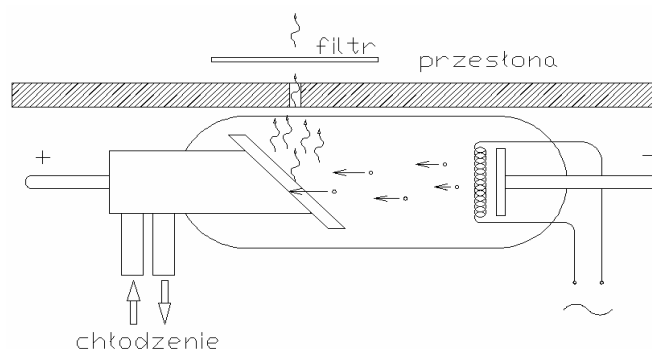
- Zgadza się, ale co z tego ma wynikać?

- Jeśli interpretacja falowa jest prawdziwa, to obaj obserwatorzy powinni widzieć jednocześnie wzmocnienie bądź wygaszenie.

- Ale to oznaczałoby, że problem znikającej energii nadal niebyły rozwiązany.

- Dokładnie! Dlatego przypuszczam, że światło w zależności od warunków na wejściu do płytki półprzepuszczalnej, raz wędruje do obserwatora pierwszego, a innym razem do obserwatora drugiego. Wówczas tak jak już pan zauważył, nie będzie problemów z zasadą zachowania energii.
- Tyle, że wymagałoby to poszukania zupełnie innego wyjaśnienia niż interferencja dla wytłumaczenia zjawisk tam zachodzących.
- Po tych wszystkich doświadczeniach chyba nie powinno to pana już dziwić? {zażartowała sobie}
- To nie było śmieszne pani Kamilo. Zupełnie nie mam pojęcia jak można byłoby to inaczej wyjaśnić.
- Proszę się nie martwić, ja też nie.
- Czy coś jeszcze pani znalazła w sieci, oprócz nieścisłości w kwestii interferometru?
- W sieci już nie, ale ostatnio na zajęciach z krystalografii mieliśmy wyprowadzenie wzoru Bragga dla wytłumaczenia pików rentgenowskich. Pan wykładowca odwoływał się do interferencji wiązek rentgenowskich. Czy w tym doświadczeniu, wiązka rentgenowska nie powinna być przypadkiem spójna?
- Hm. W zasadzie to interpretacja Bragga jest taką samą interpretacją jak wzmacnianie i wygaszanie się kolorów przy odbiciu światła na cienkich warstwach. Polega na wyznaczeniu różnicy dróg przebytych przez falę rentgenowską pomiędzy wiązkami odbitymi na różnych warstwach krystalograficznych. W zależności od kąta pod jakim pada wiązka rentgenowska, zmienia się również różnica dróg. Stąd, dla niektórych kątów padania mamy piki, a dla niektórych nie.
- No właśnie, ale wiązki rentgenowskie interferujące po wyjściu z kryształu pochodzą z różnych głębokości i tym samym z różnych miejsc przekroju wiązki padającej. Stąd, w całym przekroju wiązki padającej powinna być taka sama faza. Musi być spełniony warunek zarówno na spójność przestrzenną jak i czasową.
- Zgadza się. Dokładnie taką samą sytuację mamy w doświadczeniu z cienkimi warstwami.
- Czyli Bragg musiał założyć, że wiązka rentgenowska jest bezwzględnie spójna.
- Nie ma wątpliwości, musiał tak zrobić.
- W jaki sposób otrzymuje się wiązkę rentgenowską? Czy mógłby mi pan opowiedzieć? Proszę.
- Oczywiście mógłbym. Jeszcze nie miała pani tego na wykładzie?
- Nie, niestety nie.
- W takim razie, proszę podać mi swój zeszyt.

Profesor wziął zeszyt Kamili i starannie rozrysował schemat lampy rentgenowskiej, by na jego podstawie wyjaśnić jej działanie (Rys. 40).



Rys. 40 Schemat działania lampy rentgenowskiej.

- Lampa rentgenowska to nic innego jak bańka próżniowa w której umieszczono dwie elektrody. Przy elektrodzie ujemnej znajduje się spiralka podobna do włókna w żarówce Edisona, przez którą przepuszcza się prąd. Kiedy temperatura włókna podniesie się do znacznych wartości, w wyniku termo-emisji pojawia się intensywna chmura swobodnych elektronów, które następnie są intensywnie przyspieszane w kierunku dodatniej katody pod wpływem przyłożonego napięcia. Tak rozpędzone elektrony oddają część swojej energii elektronom znajdującym się w atomach katody (zwykle miedzi). Prowadzi to do lokalnego

wzbudzenia elektronów w atomach i emisji fotonów, kiedy ich stan wraca do poziomu podstawowego. Dokładnie tak samo jak ma to miejsce w laserach i stąd potrzeba intensywnego chłodzenia katody. Różnica polega na tym, że elektrony w atomach katody wzbudzane są do znacznie wyższych energii i stąd emitują fotony o większej energii niż ma to miejsce w typowych laserach, gdy wracają do stanu podstawowego.

- Rozumiem, że atomy katody nie wyświecają swojej energii od razu po tym jak zostały wzbudzone.

- Nie, ale czas między wzbudzeniem a emisją światła jest bardzo krótki.

- Niemniej jednak elektrony w takiej lampie uderzają w powierzchnię katody w przypadkowych momentach czasu i miejscu.

- Oczywiście, że tak. Elektrony w obszarze anody pojawiają się spontanicznie w wyniku termoemisji i tym samym rozpoczynają swoje przyspieszenie w różnych przypadkowych momentach czasu oraz z różnymi prędkościami początkowymi. Tak więc, dolatują do katody w sposób całkowicie niesynchronizowany.

- Czy chmura elektronowa oddziałuje na siebie podczas wędrówki od anody do katody?

- Tak. Elektrony odpychają się wzajemnie. Wobec czego ich prędkości i położenie podczas przelotu pomiędzy elektrodami będą się zmieniać w trudny do przewidzenia sposób, zależny od położenia i prędkości wzajemnej elektronów.

- Wobec czego przypadkowość wzbudzania atomów katody będzie jeszcze większa.

- Niekoniecznie, bo szybsze elektrony doganiając te wolniejsze są przez nie hamowane i na odwrót. Szybsze elektrony powodują przyspieszenie wolniejszych. Stąd, prędkości elektronów podczas przelotu powinny się raczej wyrównywać. Niemniej jednak, daleko jest tutaj do sytuacji, w której elektrony dochodziłyby do katody w sposób w pełni uporządkowany. Nawet jeśli oddziaływanie między elektronami nieco wyrównywałoby ich prędkości, to i tak poszczególne uderzenia elektronów w katodę są dość przypadkowe.

- Rozumiem więc, że nie możemy twierdzić iż wiązka rentgenowska jest wiązką fali spójnej.

- Jeśli chodzi o lampę rentgenowską, to zdecydowanie nie możemy twierdzić, że mamy do czynienia ze źródłem fali spójnej. Szczególnie, że nie tylko wzbudzenie atomów katody przez elektrony jest w przypadkowych momentach czasu ale i czas między wzbudzeniem elektronu a emisją fotonu promieniowania rentgenowskiego może być różny dla poszczególnych atomów.

- Skoro tak, to dlaczego w książkach warunek Bragga jest wyprowadzony przy zastosowaniu interpretacji falowej czyli interferencji. Skoro dla tej interpretacji niezbędny jest warunek o spójności padającej fali, a w praktyce warunek ten spełniony być nie może.

- Nie wiem co powiedzieć. Widocznie autorzy książek nauczeni ze szkoły, że to co piszą jest prawdziwe i poprawne przyzwyczaili się do tego, że tak jest ale nie zadali sobie trudu, by się zastanowić, czy aby na pewno?

- Czyli rozumiem, że to oznacza tylko jedno. Wzór wyprowadzony przez Bragga jest błędny, gdyż nie posiada teoretycznego uzasadnienia w tym sensie, że jego obecne teoretyczne wyprowadzenie wymaga założenia, które w praktyce nie jest spełnione.

- Niekoniecznie. Wzór może być poprawny, ale co najwyżej przypadkowo. W przeszłości fizyki już się zdarzało, że z błędnych hipotez wyciągano poprawne i zgodne z doświadczeniem wnioski dla których później znajdowano lepsze wyjaśnienia. Trudno powiedzieć, co jest obecnie prawdziwe pani Kamilo. Myślę, że wielu naukowców nie zgodzi się z pani wnioskiem, że warunek Bragga mógłby być błędny. Na założeniu poprawności tego warunku oparta jest cała krystalografia, fizyka ciała stałego i wiele, wiele innych zagadnień, które wydaje się, że poprawnie radzą sobie z opisem tego wszystkiego co nas otacza. Odrzucenie poprawności warunku Bragga byłoby dość kłopotliwe dla bardzo wielu naukowców.

- Rozumiem, ale to nie zmienia faktu, że będzie trzeba to porządnie przemyśleć i poszukać takiego wyjaśnienia obserwowanych zjawisk, które nie będzie można tak łatwo zakwestionować jak obecne. A w przyszłości może się okazać, że formuła Bragga faktycznie jest poprawna. Czas pokaże, czy aby na pewno tak jest?

- Wiem, że to nie będzie takie proste jak się być może pani wydaje, ale ma pani rację, że trzeba poszukać innego wyjaśnienia.

Z dumą Kamila podniosła głowę, że udało się jej znaleźć kolejną nieścisłość w teorii i wtem zauważyła na zegarze wiszącym nad wejściem do sali, że jest już mocno spóźniona na kolejne zajęcia.

- O przepraszam profesorze, ale za bardzo się rozgadałam i muszę już lecieć! Czy mogę przyjść na konsultację jak będę miała jeszcze jakieś pytania?
- Tak, proszę. {I obaj odeszli w pośpiechu w przeciwne strony. Profesorowi też śpieszyło się niezmiernie, ale to już z zupełnie innego powodu.}

Tymczasem Kamila pobiegła na zajęcia rachunkowe z fizyki, gdzie omawiane były zagadnienia dotyczące ruchu naładowanych cząstek w polu elektrycznym i magnetycznym oraz indukcja magnetyczna Faradaya. W ramach omawianych zagadnień przytoczony został wzór na siłę elektromotoryczną Faradaya ε , powstałą w wyniku zmiany w czasie strumienia pola magnetycznego Φ :

$$\varepsilon = - \frac{d\varphi}{dt} \quad (6)$$

$$\varphi = \oint_S \vec{B} \cdot \vec{S} \quad (7)$$

Na zajęciach omawiane były zadania w których zamknięty obwód elektryczny o powierzchni S znajdował się w jednorodnym polu magnetycznym o indukcji B skierowanym prostopadle do rozważanego obwodu. Wówczas wzór 7 znacząco się upraszczał do postaci:

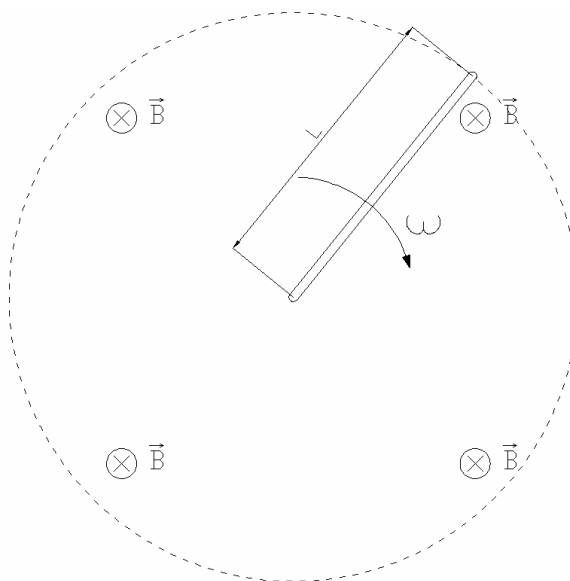
$$\varphi = B \cdot S \quad (8)$$

Zmiana strumienia magnetycznego następowała w zależności od rozważanego zadania poprzez jednostajną zmianę: wartości indukcji magnetycznej $B(t)$ lub pola przekroju obwodu elektrycznego $S(t)$, który był we wszystkich omawianych na zajęciach zadaniach łatwo wyznaczalny.

$$\varphi(t) = B(t) \cdot S(t) \quad (9)$$

Dla tak określonego strumienia pola magnetycznego wystarczyło policzyć jej pochodną od czasu i w ten łatwy sposób określić wartość siły elektromotorycznej pochodzącej od indukcji magnetycznej.

Kamila nie miała problemu ze zrozumieniem omawianych na zajęciach zagadnień, dlatego postanowiła w domu samodzielnie przeliczyć te zadania, których nie omówiono na zajęciach. Kiedy znalazła nieco wolnego czasu usiadła i przeczytała pierwsze zadanie: „Metalowy pręt o długości L wiruje w jednorodnym polu magnetycznym o indukcji B prostopadłym do płaszczyzny wirowania pręta. Oblicz różnicę potencjałów między końcami pręta, jeśli obraca się on jednostajnie względem swego końca z prędkością kątową ω .” Od razu pomyślała, że to przecież jakaś pomyłka. Na zajęciach omawiane były zadania w których można było łatwo wyznaczyć zamknięty obwód elektryczny i dla takiego obwodu liczona była zmiana strumienia magnetycznego. Tym razem nie ma żadnego zamkniętego obwodu elektrycznego (Rys. 41).



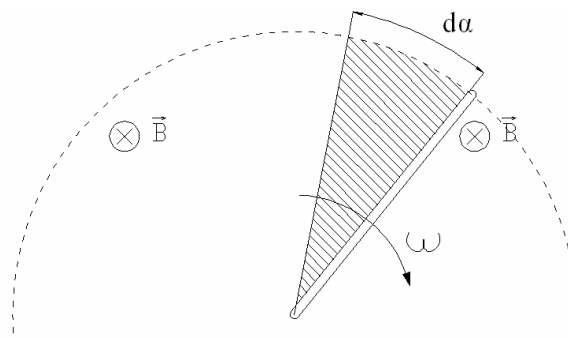
Rys. 41 Rysunek do zadania z wirującym w polu magnetycznym przewodzącym prętem.

Wiedziała przy tym, że szukana różnica potencjału jest niczym innym jak napięciem. Dlatego spodziewała się, że należy wyznaczyć siłę elektromotoryczną ε (wzór 6), którą można traktować jak zwykłe napięcie w prawie Ohma. Stąd:

$$\Delta V = U = \varepsilon = -\frac{d\varphi}{dt} \quad (10)$$

Niestety, w jaki sposób miała wyznaczyć zmianę strumienia magnetycznego, skoro nie wiedziała w jaki sposób wyznaczyć ma strumień magnetyczny? Przecież we wzorze 7 występuje wartość powierzchni obwodu zamkniętego, a w tym przypadku nie ma mowy o jakimkolwiek obwodzie zamkniętym. Tak więc, postanowiła następnego dnia odwiedzić profesora i dowiedzieć się jak należy rozwiązać ten problem.

Profesor zdziwił się, że Kamila przyszła do niego z tak prostym zadaniem. Podszedł więc do biurka, gdzie leżała czysta kartka i naskicował prosty rysunek (Rys. 42).



Rys. 42 Przedstawienie pola zakreślonego przez obracający się pręt w czasie \$dt\$.

- Proszę spojrzeć. W małym czasie \$dt\$ możemy przyjąć, że pręt przemieści się o kąt \$d\alpha\$. Wówczas pole jakie zakreśli pręt w tym czasie będzie takie jak to zakreskowane na rysunku. Czyli:

$$dS = \pi \cdot L^2 \cdot \frac{d\alpha}{2\pi} = \frac{L^2 \cdot d\alpha}{2} \quad (11)$$

Jednocześnie wiemy, że:

$$\omega = \frac{d\alpha}{dt} \quad (12)$$

Stąd: \$d\alpha = \omega \cdot dt\$. Podstawiając \$d\alpha\$ do wzoru 11 otrzymujemy:

$$dS = \frac{L^2 \cdot \omega \cdot dt}{2} \quad (13)$$

- Teraz rozumiem! {krzyknęła Kamila} Strumień magnetyczny będzie zależał tylko od zmiany przekroju \$dS\$, gdyż indukcja magnetyczna \$B\$ jest w zadaniu stała, Stąd:

$$\varphi(t) = B \cdot S(t) \quad (14)$$

Kiedy policzymy pochodną strumienia po czasie to uzyskamy siłę elektromotoryczną \$\varepsilon\$:

$$\varepsilon = -\frac{d\varphi}{dt} = -B \cdot \frac{dS}{dt} = -B \cdot \frac{d}{dt} \left(\frac{L^2 \cdot \omega \cdot dt}{2} \right) = -\frac{1}{2} B \cdot L^2 \cdot \omega \quad (15)$$

- Dokładnie o to chodziło. I jak pani Kamilo, nie było to trudne zadanie?

- Nie, nie było. Ale i tak mam pewne wątpliwości co do takiego rozwiązania. Przecież nadal to nie jest obwód zamknięty. To mi wygląda trochę jak sztuczka matematyczna, która pozwala coś tam policzyć, ale nie bardzo wiadomo o co dokładnie chodzi?

- Oj pani Kamilo. Co pani opowiada?

- No tak panie profesorze. Bo widzi pan, jak takie rozwiązanie ma się do twierdzenia, że zmiana strumienia pola magnetycznego w obwodzie zamkniętym wytwarza siłę elektromotoryczną. Dopiero teraz widzę na przykładzie tego zadania, że może się również pojawić problem w innych przypadkach. Kiedy obwód elektryczny nie jest zamknięty jak w tym zadaniu, ale zmianie ulega tylko wartość indukcji magnetycznej B , bądź układ elektryczny nie leży na płaskiej powierzchni. Jak wówczas określić przekrój S we wzorze 7? Ponadto, widzę również pewien problem z określeniem, gdzie dokładnie pojawia się napięcie i jakie? Prawo Faradaya (wzór 6) określa tylko wartość pojawiającej się siły elektromotorycznej tak, jak by indukowała się ona w całym obwodzie elektrycznym w jakiś magiczny sposób. Nie podaje wartości napięcia jaka pojawia się w dowolnej części obwodu. Mając właśnie takie wątpliwości, nie bardzo wiedziałam jak się zabrać za rozwiązanie tego zadania. Dlatego też przyszedłem z tym do pana. Niestety pańskie rozwiązanie nie daje odpowiedzi na moje wątpliwości.

- Czy zastanawiała się pani jak rozwiązać to zadanie inaczej?

- Tak, zastanawiałam się i dlatego przyszedłem, by sprawdzić, czy mój tok rozumowania jest poprawny i co z tego by wynikało.

- W takim razie słucham panią.

- Zaczniemy od założenia, że pręt składa się z elektronów swobodnych, które mogą się łatwo przemieszczać w obrębie pręta i pozostałych składników atomów, które będę dalej traktowała jako nieruchome dodatnie rdzenie. Czy jest to poprawne założenie?

- Może pani spokojnie takie założenie przyjąć. Jest ono zgodne z przyjętym obecnie wyobrażeniem na temat przewodników. Proszę kontynuować.

- Tak więc przyjąłem, że skoro pręt przemieszcza się w polu magnetycznym, to również tak samo przemieszczają się w nim elektrony swobodne jak i dodatnie rdzenie atomowe. Przy czym zrobiłam założenie, że ich prędkość względem pola magnetycznego jest taka sama jak prędkość styczna pręta. Czyli, czym dalej jest ona od osi obrotu tym jest większa.

$$v_s = \omega \cdot r$$

(16)

- To nie jest do końca prawda. Zakłada się w teorii, że atomy drgają wokół swoich średnich położenia i drgania te są tym większe im ciało posiada wyższą temperaturę. Ruch ten będzie się nakładał na prędkość styczną rdzeni atomowych. Natomiast prędkość elektronów swobodnych jest jeszcze większa. Możemy raczej założyć, że średnia prędkość elektronów w danym punkcie będzie równa prędkości stycznej pręta w tym punkcie.

- Rozumiem. {Kamila zamyśliła się}

- Proszę kontynuować pani Kamilo. Pani założenie co do prędkości elektronów swobodnych jak i rdzeni atomowych jest jak najbardziej do zaakceptowania.

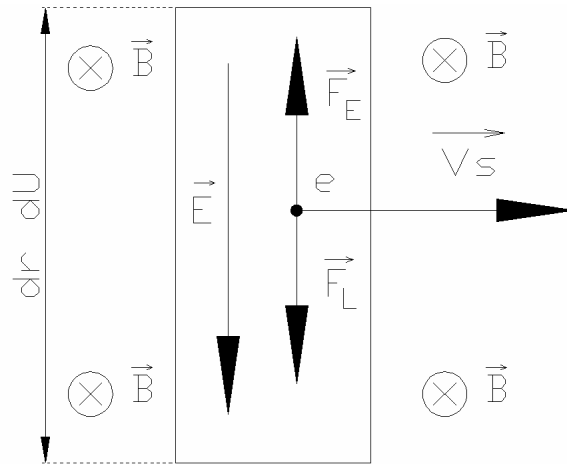
- Skoro elektrony swobodne i dodatnie rdzenie poruszają się w polu magnetycznym to założyłam, że działa na nie siła Lorentza. Przy czym przyjąłem, że rozważać będę tylko ruch elektronów, gdyż ruch rdzeni oznaczałby zmianę kształtu pręta.

- Słuszna uwaga.

- Kiedy na elektrony działa siła Lorentza w jednym kierunku, to próbują się one grupować na jednym końcu pręta. Po drugiej stronie tworzy się ich niedobór. W ten sposób powstaje różnica potencjałów czyli napięcie.

- Gdyby uwzględnić tylko siłę Lorentza to wszystkie elektrony przewodnictwa powędrowałyby na jeden koniec. Nieprawdaż?

- Owszem. Ale przemieszczone elektrony wytwarzają w przeciwnym kierunku pole elektryczne, które powoduje pojawienie się drugiej siły, która odpycha elektrony (Rys. 43).



Rys. 43 Schemat sił działających na elektrony swobodne w rozwiązaniu Kamili.

Kamila kontynuowała dalej:

- Tak więc, elektrony przemieszczają się tak długo, aż siła wypadkowa wyniesie zero. Kiedy tak się stanie powstaje stan równowagi i elektrony statystycznie dalej się już nie przemieszczają. Przedstawiony rysunek (Rys. 43) przedstawia kawałek pręta o długości dr oddalonego od osi obrotu na odległość r . Na tej odległości elektrony poruszają się względem pola magnetycznego z prędkością v_s (wzór 16) i działa na nie siła Lorentza:

$$F_L = e \cdot v_s \cdot B \quad (17)$$

Jednocześnie na odległości dr powstało w wyniku niejednorodnego rozkładu elektronów napięcie dU , które powoduje pojawienie się w tym obszarze pola elektrycznego o wartości:

$$E = \frac{dU}{dr} \quad (18)$$

Powstałe pole elektryczne powoduje pojawienie się siły elektrycznej:

$$F_E = E \cdot e = e \cdot \frac{dU}{dr} \quad (19)$$

- Sprytne pani Kamilo. Teraz rozumiem w pełni pani punkt widzenia. Chce pani przyrównać obie siły i wyciągnąć stąd wzór na dU :

$$dU = v_s \cdot B \cdot dr = \omega \cdot r \cdot B \cdot dr \quad (20)$$

- Dokładnie, elektrony będą się tak długo przemieszczać, aż siła wypadkowa działająca na nie będzie równała się zero. Wówczas, każdemu odcinkowi dr odpowiada pewna wartość napięcia dU . Jeśli dodamy wszystkie dU do siebie, to otrzymamy całkowite napięcie U . Tak więc, w stanie ustalonym mamy:

$$U = \sum dU = \int dU = \int_0^L \omega r B dr = \omega B \int_0^L r dr = \omega B \cdot \frac{1}{2} r^2 \Big|_0^L = \frac{1}{2} \omega B (L^2 - 0^2) = \frac{1}{2} \omega B L^2 \quad (21)$$

- Widzi pani! Wyszło to samo co w moim rozwiązaniu (wzór 15). Czyli pani rozumowanie wygląda na równie poprawne.

- Niekoniecznie równie poprawne, bo można zauważyć coś więcej.

- Co ma pani ma myśli?

- Teraz sobie właśnie uświadomiłam, że powstał jednak pewien problem. Oba rozwiązania dają taki sam wynik, ale pod względem fizycznym nie są sobie równoważne. Najpierw jednak chciałabym zauważyć, że moje rozwiązanie radzi sobie doskonale z wątpliwościami o których wspomniałam wcześniej. Dla każdego, nawet najmniejszego kawałka obwodu elektrycznego lub tylko części obwodu np. kawałka przewodnika tak jak wirujący w polu magnetycznym pojedynczy metalowy pręt z omawianego zadania, można określić ruch względny między polem magnetycznym a przewodnikiem i na tej podstawie określić siłę Lorentza

działającą na elektrony swobodne. Następnie, znając tą siłę obliczyć lokalny przyrost napięcia dU jaki powinien mieć miejsce w stanie ustalonym.

- Można więc dokładnie określić, gdzie i jakie napięcie będzie się indukować. Nie tylko ogólnie określoną siłę elektromotoryczną. Sprytnie pani Kamilo. Muszę przyznać, że jestem pełen uznania. Faktycznie wyższość pani rozwiązania nad moim jest znaczna.

- Tu nie chodzi o wyższość rozwiązania panie profesorze. Tylko o sens fizyczny. Pana rozwiązanie odwołuje się do ogólnej zasady określonej przez Faradaya poprzez określanie zmiany strumienia pola magnetycznego, który indukuje napięcie. Tak jakby, zmiana strumienia pola magnetycznego powodowała pojawienie się napięcia, które następnie powoduje ruch ładunków.

- Zgadza się, tak przecież jest. Zmienne pole magnetyczne powoduje pojawienie się pola elektrycznego, które następnie powoduje przepływ prądu jeśli obwód elektryczny jest zamknięty.

- I tak właśnie rozumując korzystał pan z prawa Faradaya (wzór 6). Tylko, że takie podejście wymagało przeprowadzenia sztuczki myślowej z określeniem zmiany przekroju dS . Tymczasem, w tym zadaniu nie mamy żadnej zmiany przekroju odvodu elektrycznego. Zresztą w ogóle nie ma jakiegokolwiek zamkniętego obwodu elektrycznego. Tak więc, pod względem sensu fizycznego pańskie rozwiązanie jest więcej niż naciągane i fizycznie wątpliwe. Natomiast w moim rozwiązaniu nie trzeba odwoływać się do rzeczy nieistniejących (brak zamkniętego obwodu elektrycznego).

- Hm, pod względem poprawności fizycznej faktycznie pani rozwiązanie wydaje się znacznie lepsze i wchodzi chyba znacznie głębiej w istotę zagadnienia.

- Cieszę się, że pan to docenia. Jednak uważam, że problem interpretacyjny związany z przyjęciem poprawności mojego rozumowania jest jednak większy. Czy zauważył już pan w czym leży problem?

Profesor dokładniej przyjrzał się rozwiązaniu Kamili i po dłuższym namyśle stwierdził.

- Faktycznie, w pani rozumowaniu jest inna kolejność przyczyny i skutku.

- Dokładnie! Względny ruch między ładunkiem elektrycznym a polem magnetycznym wywołany poprzez przemieszczenie przewodnika w polu magnetycznym bądź przemieszczeniem się linii sił pola magnetycznego * powoduje pojawienie się siły Lorentza działającej na ładunki elektryczne (* Zmiana wartości indukcji magnetycznej B może być interpretowana jako lokalne przemieszczenie się linii sił pola magnetycznego!). Siła Lorentza jest przyczyną przemieszczenia się ładunków w przewodniku, które dopiero po przemieszczeniu się wytwarzają to, co nazywamy napięciem a spowodowane jest niejednorodnym rozkładem ładunku elektrycznego w przewodniku.

- Tak więc, w pani rozumowaniu napięcie elektryczne jest konsekwencją przemieszczonych ładunków elektrycznych pod wpływem siły Lorentza. Natomiast w moim rozumowaniu przemieszczenie się ładunków jest konsekwencją pojawienia się napięcia elektrycznego indukowanego prawem Faradaya.

- Właśnie! W moim rozumowaniu pierw mamy ruch ładunku, a później pojawia się napięcie (związane z polem elektrycznym) w wyniku ich przemieszczenia. W pana rozumowaniu jest odwrotnie. A co będzie się działo, jak w obszarze zmieniającego się pola magnetycznego nie będzie żadnych ładunków?

- Według pani rozumowania nie będzie też indukowanego napięcia (związany z tym jest też brak pola elektrycznego), gdyż nie będzie ładunków, które mogłyby zostać przemieszczone i przez to mogłyby one wytworzyć napięcie bądź pole elektryczne.

Profesor zadumał się nad konsekwencjami tak postawionego pytania i po dłuższej chwili krzyknął.

- Ależ to absurd! Przecież to nie może być prawdą. Gdyby pani rozumowanie było prawdziwe i wchodziło w rzeczywistość przyczynę obserwowanych zjawisk, to nie mogłoby istnieć coś takiego jak fala elektromagnetyczna, która przecież przemieszcza się w próżni. We współczesnej wersji teorii Maxwella, zmienne pole magnetyczne indukuje pole elektryczne. Natomiast, zmienne pole elektryczne indukuje pole magnetyczne itd. I to wszystko dzieje się również w próżni! W ten sposób tłumaczy się powstanie i propagację fali elektromagnetycznej, która rozchodzi się z niezmiernie wielką prędkością. Z prędkością światła, dla której próżnia nie jest żadną przeszkodą.

- Rozumiem, że ta fala jest tym samym co fale radiowe, radarowe, mikrofały itd. oraz światło, o czym wcześniej już rozmawialiśmy.

- Tak, również i światło jest obecnie tłumaczone tak samo, jako fala elektromagnetyczna. Jest to klasyczne podejście do zrozumienia natury światła. Dopiero doświadczenia z przełomu wieków XIX i XX zmusiły fizyków do przyjęcia jej obecnej dualistycznej natury.

- Pamiętam, jak wspominał pan o tym wcześniej. Tyle, że doświadczenia przeprowadzone w laboratorium wykazały, że światło nie posiada własności falowych, a tym bardziej dualistycznych. A przynajmniej nie można obserwowanych zjawisk tłumaczyć odwołując się do zjawiska interferencji. Skoro tak zwane fale elektromagnetyczne chcemy traktować na równi ze światłem, to muszą być jednak czymś innym. Jeśli okazałyby się jednak czymś innym, to pańska wątpliwość co do poprawności zastosowania siły Lorentza w omawianym zadaniu byłaby nieuzasadniona.

- Owszem, nie byłyby to wówczas najlepszy kontrprzykład na wnioski wynikające z pani rozwiązania. Niemniej jednak, przy obecnych faktach doświadczalnych, zupełnym absurdem byłoby twierdzenie, że tak zwane: fale radiowe, mikrofały, itd. nie istnieją. Czy są one faktycznie falami elektromagnetycznymi, tak jak chciał tego Maxwell? {Teoria Maxwella odnosiła się do naprężeń i ruchu eteru, więc niebyły to dokładnie takie same fale jak rozumiemy to obecnie!} No cóż, na razie nie potrafimy lepiej wytłumaczyć tych zjawisk, niż właśnie poprzez odwołanie się do istnienia fal elektromagnetycznych.

- Czy zatem moje rozumowanie odwołujące się do siły Lorentza przeprowadzone w ramach tego zadania jest błędne?

- Pani rozwiązanie wygląda na bardzo dobrze przemyślane i wchodzi bardzo głęboko w istotę problemu. Znacznie bardziej niż ogólne stwierdzenie, że zmiana wartości strumienia pola magnetycznego indukuje w bliżej nieokreślonym miejscu obwodu elektrycznego pole elektryczne (siłę elektromotoryczną).

- To mi przypomina zdanie mojego byłego nauczyciela, który wspominał, że; „Fizyka składa się z wielu praw przyrody. Niektóre są bardzo ogólne i pozornie zagadkowe, ale dają się często wyprowadzić i zrozumieć z praw bardziej podstawowych i elementarnych. I to one są prawdziwą przyczyną obserwowanych zjawisk.”.

- Zgadza się pani Kamilo. Weźmy na przykład prawo Archimedesusa. Co ono opisuje?

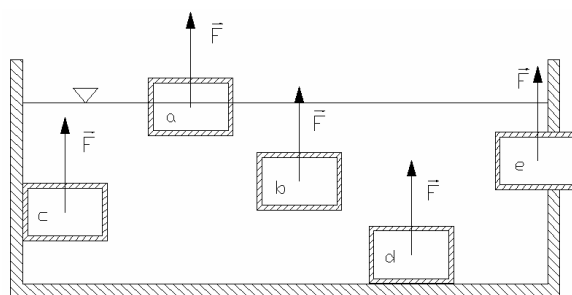
- To zna każdy panie profesorze. Prawo Archimedesusa określa siłę jaka działa na zanurzone w wodzie ciało. Jest ona równa ciężarowi wypartej przez to ciało cieczy.

- Dobrze, czy to prawo nie brzmi przypadkiem zbyt zagadkowo?

- Hm. Właściwie jeśli się nad tym troszkę głębiej zastanowić, to dla kogoś kto nie zna prawa Pascala i nie potrafi obliczyć sił parcia od ciśnienia, prawo Archimedesusa faktycznie może wydawać się zagadkowe.

- Proszę zauważyć, że w tym przypadku możemy mieć podobne dylematy dla niektórych przypadków, jak ten co miał miejsce przy bezpośrednim zastosowaniu prawa Faradaya.

Profesor odwrócił leżącą na stole kartkę na czystą stronę i narysował rysunek (Rys. 44).



Rys. 44 Przedstawienie siły Archimedesusa działającej na zanurzone w cieczy obiekty.

- Prawo Archimedesusa jak już pani wspomniała określa siłę jaka działa na ciało zanurzone w cieczy i równą co do wartości ciężarowi cieczy wypartej przez to ciało. W przypadku „a” objętość wypartej cieczy będzie częścią objętości skrzyni, którą łatwo określić. Natomiast w przypadkach „b”, „c” i „d” sprawa objętości wypartej cieczy jest banalnie prosta i wynosi ona objętość całej skrzyni. Tymczasem, dla przypadku „e”

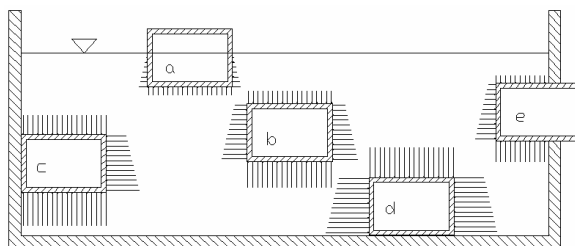
objętość wypartej cieczy jest dyskusyjna zwłaszcza, gdyby ściana przez którą przechodzi skrzynia nie byłaby pionowa.

- Podobnie jak dyskusyjną sprawą jest określenie powierzchni S przy obliczeniu wartości strumienia pola magnetycznego w przypadku obwodu otwartego, bądź nie leżącego na jednej płaszczyźnie.

- Dokładnie. Tymczasem, z prawa Archimidesa nie wynika, rzeczywiste zachowanie się skrzyni w przypadku „c”, „d” oraz „e” (Rys. 44).

- Jak to?

Profesor narysował obok podobny rozkład skrzyń, ale tym razem zamiast siły Archimidesa zaznaczył siłę parcia wynikającą z prawa Pascala i zmieniającego się wraz z głębokością ciśnienia hydrostatycznego (Rys. 45).



Rys. 45 Rozkład sił parcia działającego na zanurzone w cieczy skrzynie.

- Teraz rozumiem co miał pan na myśli. W przypadku „a” „b” czy „c” łatwo policzyć siłę wyporu, która działa w górę, poprzez dodanie wszystkich sił działających na powierzchni skrzyni od góry i dołu. Ponieważ od dołu skrzyni ciśnienie jest większe, to również siła parcia działająca na powierzchnię dolną jest większa niż na powierzchnię górną. Różnica sił powoduje pojawienie się siły wypadkowej, którą inaczej nazywamy siłą Archimidesa. Tyle, że w przypadku „c” oraz „e” siła parcia działająca na powierzchnię boczną nie jest skompensowana siłą parcia z drugiej strony. Stąd, skrzynia będzie dociskana do ściany naczynia, czego prawo Archimidesa nie przewiduje!

- Dokładnie, ale to nie wszystko. Bardzo ciekawy jest przypadek „d”. Wyobraźmy sobie, że mamy skrzynię, która posiada średnią gęstość mniejszą od wody. Jeśli będziemy chcieli ją zanurzyć w akwarium wypełnionym wodą, to będziemy musieli przeciwdziałać sile wyporu, która w tym przypadku będzie nam utrudniać przemieszczanie skrzyni w dół i będzie wypychać ją w górę. Siła wyporu będzie działać cały czas dopóki nie dojdziemy do samego dna. W tym momencie dzieje się rzecz z pozoru magiczna. Jeśli skrzynia jest gładka tak jak szklane akwarium, to można tak ją położyć na dnie, by pod nią nie było już cieczy. Wówczas okaże się, że skrzynia nie tylko przestanie nam uciekać do góry, ale jeszcze dodatkowo będzie dociskana w dół!

- Ciekawe. Nie zwróciłam na to wcześniej uwagi, ale tego faktycznie nie da się przewidzieć z prawa Archimidesa. Ten przypadek pokazuje, jak bardzo zwodnicze mogą być niektóre uogólnione prawa, które nie odwołują się bezpośrednio do podstawowych sił działających w przyrodzie. Dopiero wejście głębiej w istotę zjawiska, czyli sił które faktycznie działają, w tym przypadku na zanurzoną skrzynię, daje poprawny opis wraz z wytłumaczeniem wszelkich niuansów, które wcześniej wydawałyby się tajemnicze i niezrozumiałe w myśl prawa uogólnionego jakim jest prawo Archimidesa. Tyle, że to samo moglibyśmy odnieść do prawa Faradaya. Jego prawo jest tak samo ogólne jak prawo Archimidesa i tak samo można je wytłumaczyć odwołując się do pierwotnej przyczyny obserwowanych zjawisk, czyli sił Lorentza działających bezpośrednio na ładunki. Tak jak prawo Archimidesa można rozważać jako konsekwencję sił parcia działających bezpośrednio na powierzchnię zanurzonego w cieczy ciała. Tak prawo Faradaya nie jest niczym innym jak konsekwencją działania sił Lorentza na ładunki elektryczne. Prawami uogólnionymi są również prawa (równania) Maxwella.

- Tyle, że przy takim rozumowaniu musiałbym nie tylko stwierdzić wyższość pani rozwiązania w zadaniu z wirującym w polu magnetycznym przewodzącym prądzie, tak jak zrobiłem to wcześniej. Ale również

musiałbym przyznać, że jest to w pełni słuszne rozwiązanie, które odwołuje się do rzeczywistej przyczyny obserwowanej w takim doświadczeniu różnicy potencjału. Aczkolwiek, jest to sprzeczne z ideą fali elektromagnetycznej!

- W takim razie zastanówmy się, czy prawo Faradaya posiada jeszcze jakieś inne nieścisłości, jak te przedstawione przez pana w przypadku prawa Archimedesesa? W ogólnym rozumieniu prawa Faradaya, zmienne pole magnetyczne powoduje pojawienie się pola elektrycznego (niekoniecznie zmiennego). Zaś powstałe pole elektryczne powiązane jest z napięciem jakie pojawia się w przewodniku, które jest tym większe im szybciej zmienia się wartość strumienia tego pola. Następnie, powstałe napięcie powoduje przepływ prądu. Z prawa Ohma mamy:

$$I = \frac{U}{R} \quad (22)$$

Stąd, prąd który popłynie w przewodniku, umieszczonym w zmiennym polu magnetycznym, będzie zależał również od wartości oporu elektrycznego. Czy w takim razie w nadprzewodnikach, których opór elektryczny jest zerowy, będzie się indukować za każdym razem prąd krytyczny, niezależnie od wielkości zmian pola magnetycznego? (Teoretycznie nieskończenie duży prąd, ale takich prądów w nadprzewodnikach nie ma.)

- Nie, to byłby jakiś absurd! Gdyby tak było, to niezależnie od szybkości zbliżania magnesu do nadprzewodnika, wytwarzałby on zawsze tak samo duże pole magnetyczne. Pole to zależałoby jedynie od geometrii nadprzewodnika i gęstości prądu krytycznego, a nie od odległości i szybkości zbliżania się magnesu. Może pani o tym nie wiedzieć, ale nadprzewodniki mają niewielki opór elektryczny w przypadku zmiennych napięć. Jest on tym większy im większa jest częstotliwość zmiennego pola elektrycznego. Dlatego sądzę, że niekoniecznie od razu musi popłynąć w nadprzewodniku prąd krytyczny.

- A jednak może nie mieć pan racji profesorze. Na zajęciach rachunkowych omawialiśmy zadanie w których przewodzący pierścień był umieszczony w polu magnetycznym, które zmieniało się liniowo w czasie. Oznacza to, że strumień magnetyczny też zmieniał się liniowo w czasie (wzór 9). Zgodnie z teorią Faradaya indukowane napięcie będzie stałe. Oznacza to tyle, że można tak zbliżać magnes do nadprzewodnika, aby zgodnie z teorią indukowało się stałe napięcie, a jego wartość będzie zależała np. od siły magnesu. Wówczas, nadprzewodnik dalej miałby zerowy opór elektryczny i wcześniejsza uwaga odnośnie indukowania się prądów krytycznych jest nadal aktualna.

- Faktycznie, moja uwaga nie była w pełni przemyślana. W takim razie zastanawiam się, czy zmienne pole magnetyczne działa w przewodnikach jako źródło napięciowe, tak jak liczyła to pani na rachunkach zgodnie z teorią Faradaya, czy może jednak jako źródło prądowe, co mogłoby chyba lepiej opisywać zachowanie prądu w nadprzewodnikach? Ponadto, gdyby było to źródło prądowe, to idea że za ruch ładunków odpowiada siła Lorentza wydaje się być poprawniejszym wytłumaczeniem tego zjawiska, niż założenie, że ruch ładunków odbywa się pod wpływem siły elektrycznej pochodzącej od indukowanego pola elektrycznego. Tymczasem, nadprzewodniki mają taką własność, że często indukuje się w nich taki prąd, aby pole magnetyczne wewnątrz nadprzewodnika było zerowe, bądź wchodziło do niego w postaci tzw. worteksów (nadprzewodniki II rodzaju). Więc pani przykład z indukowaniem się napięcia w nadprzewodniku może nie być najlepszym pomysłem, gdyż sam nadprzewodnik jest niezwykle materiałem i zachowuje się nietypowo w reakcji na zewnętrzne pole magnetyczne.

Tymczasem, w przypadku zwykłych przewodników zakłada się, że elektrony przewodnictwa poruszają się w różne przypadkowe strony od zderzenia do zderzenia. Przemieszczają się one w tym czasie z dużą prędkością na odległości znacznie przekraczające odległości między atomami przewodnika. Natomiast średnia prędkość poruszania się elektronów jest zerowa. W przypadku pojawienia się dodatkowego pola elektrycznego wewnątrz przewodnika, na nośniki prądu działa dodatkowa zewnętrzna siła elektryczna. Siła ta nieco zmienia średnie wartości prędkości elektronów swobodnych nadając im niewielką prędkość wypadkową w kierunku poruszającego się w przewodniku ładunku, co my obserwujemy jako przepływ prądu. Podobny efekt uzyskamy w przypadku pojawienia się dodatkowej siły magnetycznej działającej na nośniki prądu. Siła Lorentza zastępuje w tym przypadku siłę elektryczną. Poza tym, cała reszta teorii pozostaje bez zmian. Zachowując prawo Ohma nadal w mocy.

- Ostatecznie i tak dochodzimy do wniosku panie profesorze, że indukcja Faradaya to tak naprawdę to samo co siła Lorentza i gdybyśmy chcieli w jakimś zadaniu rozważać jednocześnie prawo Faradaya (wzór 6) i siłę Lorentza działającą na ładunki swobodne to byłby to błąd.

- Oczywiście! Taki sam błąd, jakbyśmy chcieli jednocześnie uwzględnić (dodać) siłę Archimedesesa i siłę parcia od ciśnienia działającą na zanurzony obiekt. Dwa razy obliczylibyśmy to samo ale inaczej. Warto jeszcze zauważyć, że siła wypadkowa od parcia dokładnie określa nam w którym miejscu, jakie siły i w jakim kierunku działają na zanurzony w cieczy obiekt. W przeciwieństwie do wzoru na siłę Archimedesesa, gdzie uzyskujemy wynik od razu całościowy i nie zawsze poprawny (niuanse z rysunku 45). Podobnie jak siła Lorentza w pani rozumowaniu pozwala dokładnie określić, jakie napięcie elektryczne pojawi się w każdym miejscu analizowanego obwodu elektrycznego, a niekoniecznie od razu całości jak to uzyskujemy w prawie Faradaya. Tyle, że przyjęcie takiego rozumowania jako w pełni prawdziwego i faktycznie słusznego zmusza nas bezpośrednio do głębokiego przemyślenia prawdziwej natury fal radiowych, mikrofal itd. Zwłaszcza, iż przyjmujemy, że są to zjawiska tej samej natury co światło.

- Szczególnie, że doświadczenia przeprowadzone przez zemnie w laboratorium podważyły interpretację światła jako fali, a przynajmniej zasadności stosowania interpretacji odwołującej się do interferencji fal. {dodała Kamila}

- Faktycznie, pani doświadczenia są problemem dla obowiązującego obrazu światła jako fali. Niemniej jednak przypominam, że takie aspekty jak: brak masy i ściśle określona duża prędkość światła w próżni stają się jeszcze większą zagadką, gdy odrzuci się falową naturę tych zjawisk. Właśnie dlatego tak trudno, jeśli w ogóle, będzie można odrzucić ideę falowej natury światła. Nie jestem pewny, czy ludzie nauki są jeszcze gotowi na całkowite odrzucenie falowej interpretacji? Sam jeszcze nie wiem, czym mógłbym obowiązującą teorię zastąpić? Na razie jedyne co przychodzi mi do głowy, to szukanie dalszych faktów doświadczalnych i innych nieścisłości dotyczących światła i obowiązujących teorii, które dałyby nam jakieś dodatkowe wskazówki w poszukiwaniu innych racjonalnych przyczyn obserwowanych zjawisk.

- Ostatecznie, nie ma więc pan przeciwwskazań co do mojego rozwiązania w tym zadaniu?

- Nie, choć jest ono faktycznie sprzeczne z ideą fali elektromagnetycznej (A dokładniej, to indukowaniu się pola elektrycznego w próżni tylko pod wpływem zmiennego pola magnetycznego bez obecności przemieszczonych ładunków.). Nie wiem jeszcze jak to wyjaśnić, ale nie potrafię zarzucić trafności pani rozumowania.

- Rozumiem, w takim razie zobaczę, co jeszcze powie na takie rozwiązanie prowadzący ćwiczenia. ;-)
Myślę, że już czas na mnie, robi się już późno.

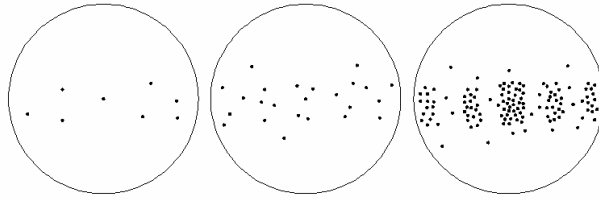
- Tak więc, do zobaczenia pani Kamilo.

- Do widzenia profesorze.

Profesor wstał i chodząc po pokoju rozmyślał nad zakończoną przed chwilą rozmową. Dobrze wiedział, że wnioski z przeprowadzonych przez Kamilę doświadczeń są sprzeczne z obecnym wyobrażeniem zjawisk mikroświata. Rozumiał, że hipotezy: dualizmu światła, fali materii, a tym bardziej fal prawdopodobieństwa są wysoce absurdalne oraz wynikają pośrednio z falowo-korpuskularnej interpretacji natury światła. Przypomniawszy sobie wielokrotnie powtarzane przez autorytety twierdzenie, że tych zjawisk nie da się wyjaśnić na bazie fizyki klasycznej (mechanistycznej).

- Czy aby na pewno? {pomyślał}

Wiedział, że w doświadczeniu Younga z pojedynczymi fotonami, bądź elektronami uzyskuje się pojedyncze z pozoru chaotyczne punkty na ekranie, które dopiero po dłuższym czasie ekspozycji kształtują się w dobrze znane prążkowe obrazy (Rys. 46).



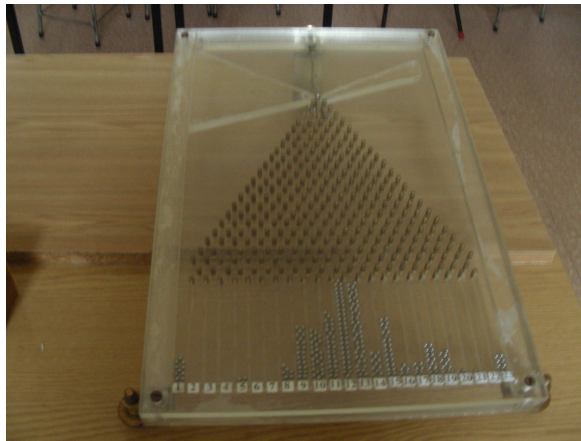
Rys. 46 Wynik naświetlania kliszy pojedynczymi fotonami, bądź elektronami zgodnie z doświadczeniem Younga [4, 5].

Taki wynik w klasycznym rozumieniu interferencji nie daje się wytłumaczyć. Pojedyncze fotony padając na ekran nie mogą interferować z innymi fotonami, które dolatują w inne miejsce ekranu i w innym momencie czasu! Stąd, obecnie obowiązująca falowa interpretacja tego doświadczenia mogłaby dawno zostać odrzucona, gdyby tylko fizycy mieli odwagę przyznać się do błędu. Zamiast tego, wyciągnięto kompletnie absurdalną hipotezę, że każdy foton bądź np. elektron interferuje sam ze sobą na zasadzie interferencji „fal prawdopodobieństwa” [5] - zlepek tych słów jest tak naprawdę pojęciowym bełkotem, który niczego nie tłumaczy! Tymczasem, doświadczenia z pojedynczymi fotonami bądź elektronami również udowadniają nam, że interferencji w tym doświadczeniu nie ma!

Niemniej jednak, profesor przypominał sobie o doświadczeniu Younga z pojedynczymi cząstkami, gdyż pokazuje ono, że ostateczne wyjaśnienie tego zjawiska musi posiadać opis probabilistyczny. Aczkolwiek, może ono posiadać charakter deterministyczny. {Ilekoć przeprowadzalibyśmy doświadczenie Younga, zawsze uzyskujemy ten sam wynik. Podobnie jest w teorii ciepła, która opisana jest w sposób probabilistyczny mimo, że jest to zjawisko w pełni deterministyczne w rozumieniu klasycznym.}

- Ale jak wyjaśnić w klasyczny sposób powstawanie prążków na ekranie w tym doświadczeniu? {zapytał siebie samego}

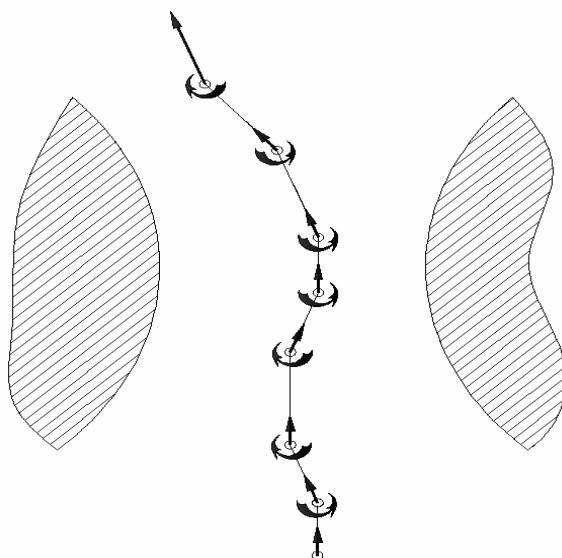
Profesor dalej rozmyślając nad tym pytaniem chodził w kółko po całym swoim gabinecie. Wtem, zauważył w rogu swojego pokoju trójkąt do wizualizacji powstawania krzywej Gausa (Rys. 47).



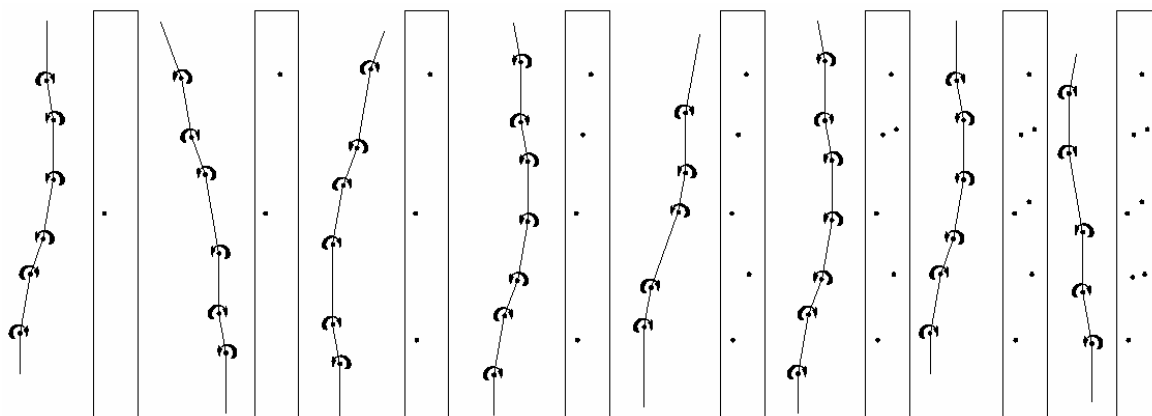
Rys. 47 Probabilistyczny rozkład kulek na trójkącie.

Szybko pojął, że ilość kulek na dole trójkąta może prezentować natężenie światła dla danego prążka z doświadczenia Younga. Na każdym etapie zmiana kierunku ruchu kulki w prawo bądź w lewo jest z naszego punktu widzenia przypadkowa. Choć w rzeczywistości jest to proces deterministyczny. Podobnie może być ze światłem. Z wcześniej przeprowadzonych doświadczeń można wyciągnąć wniosek, że światło przechodząc obok każdej krawędzi oddziałuje z materią i rozdziela się na kilka wiązek, które później obserwujemy jako prążki na ekranie. W tym momencie profesor przyjął, że na każdej powierzchni ciała

stałego istnieje pole elektryczne. Byłoby to bardzo lokalne pole pochodzące od uśrednienia oddziaływania zewnętrznych elektronów z dodatnimi rdzeniami atomów i być może posiadające znaczną wartość natężenia pola elektrycznego bądź gradientu tego pola. Wspomniane pole elektryczne posiada pewien średni rozkład wzdłuż powierzchni, który zapewne podlega przypadkowym fluktuacjom. Następnie przyjął, że kiedy światło przelatuje przez taką fluktuację lokalnego pola elektrycznego, to w zależności od własnej fazy, może zostać nieco odchylone od kierunku swojego ruchu prostopadle do krawędzi (Rys. 48). Przyjmując, że takie odchylenie jest niemal stałe dla każdego zaburzenia ruchu fotonu, powinniśmy uzyskać prążki na ekranie pochodzące z pojedynczych fotonów (Rys. 49).



Rys. 48 Schemat zachowania się fotonu w okolicy krawędzi (**tylko hipoteza!**).



Rys. 49 Probabilistyczny mechanizm zmiany kierunku ruchu fotonów na krawędzi prowadzący do powstania prążkowych obrazów.

Natężenie tych prążków miałoby rozkład podobny do rozkładu Gausa, co znacznie lepiej tłumaczyłoby zmianę jasności poszczególnych prążków niż twierdzenie, że za jasność odpowiada zmiana odległości od szczelin [2], co jest ewidentnie sprzeczne z doświadczeniem!

- Czy wymyślony prze zemnie obraz tego zjawiska jest prawdziwy? Nie wiem, ale czas pokaże. {Zastanawiał się dalej profesor} Tymczasem, model ten lepiej sobie radzi z przewidywaniem jasności poszczególnych prążków niż teoria interferencji. Wynika z niego również potencjalna zależność położenia prążków od rodzaju użytego materiału z jakiej została zrobiona krawędź na której światło ulega rozdzieleniu. Można przyjąć, że dla różnych materiałów fluktuacje pola elektrycznego na krawędzi mogą posiadać nieco inną wielkość i intensywność występowania. Również, łatwo sobie wyobrazić, że na ilość jak i siłę zaburzenia pola elektrycznego przy powierzchni miałyby wpływ krzywizna powierzchni. Możliwe, że obraz widziany na ekranie zależałby również od temperatury krawędzi, bądź szczeliny! Niemniej jednak, przy tej interpretacji profesor widział również istotny problem. Choć łatwo wyobrazić sobie istnienie bardzo lokalnego napięcia elektrycznego na powierzchni każdego ciała stałego, to jednocześnie nie znane było mu doświadczenie w którym światło oddziaływałoby bezpośrednio z polem elektrycznym. Jeśli coś takiego miałoby miejsce, to światło musiałoby oddziaływać z polami elektrycznymi znacznie większymi niż dotychczas były badane. Możliwe, że istotny jest duży gradient tego pola, a nie sama jego wartość.

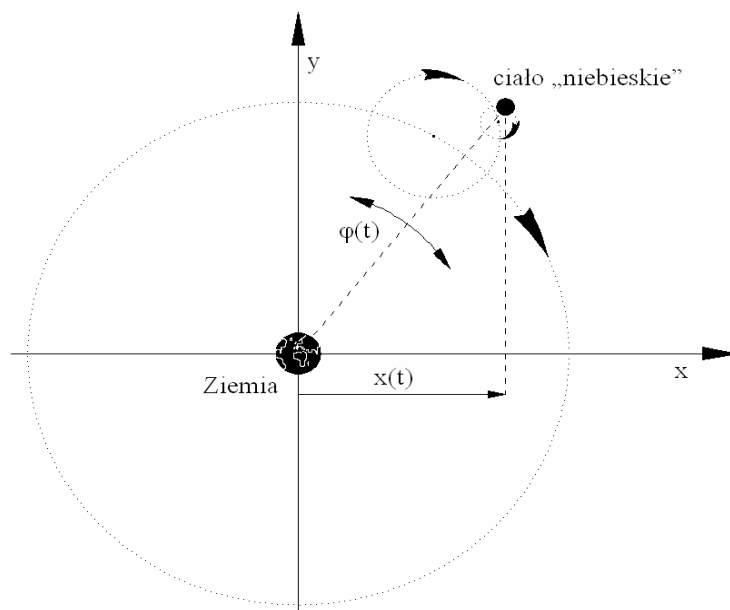
Tymczasem, w fizyce nadal przyjmuje się wysoce absurdalną hipotezę istnienia fali prawdopodobieństwa. Tak jak my niekiedy dziwimy się, jak w przeszłości można było wierzyć w istnienie fluidu ciepła bądź elektryczności, czy wielu innych pozornie dość naiwnych hipotez. Tak, następne pokolenia mogą drwić z naszych obecnych wyobrażeń mikroświata. Zwłaszcza, że tak wiele „dobrze ugruntowanych” faktów doświadczalnych dotyczących światła okazuje się być błędnych, co pokazały doświadczenia. Przyjęta hipoteza istnienia fal materii wraz z kilkoma innymi założeniami pozwoliła zbudować teorię, która niekiedy jest zadziwiająco zgodna z doświadczeniem. (Tak przynajmniej możemy przeczytać w książkach, że przewidywania mechaniki kwantowej są w doskonałej zgodności z doświadczeniem!??).

- Jak to możliwe, mimo tak absurdalnych założeń, na jakich opiera się ta teoria!

Krzyknął mimowolnie profesor i krążąc dalej po pokoju rozmyślał nad historią fizyki. Wtem przypomniał sobie, że w pierwotnej postaci teoria Kopernika gorzej przewidywała położenia ciał niebieskich niż teoria Ptolemeusza i jej rozszerzenia [6]. Mimo, że ta druga z naszego punktu widzenia jest zupełnie absurdalna i nieprawdziwa. Owszem, teoria Ptolemeusza nie radziła sobie dobrze z przewidywaniem zmian wielkości tarczy np. księżyca. Ale ówczesni astronomowie interesowali się głównie przewidywaniem położen obserwowanych przez siebie obiektów (gdzie je widzą?), jednocześnie ignorując to; jak je widzą? W modelu Ptolemeusza ciała „niebieskie” poruszały się po okręgach ruchem jednostajnym, gdyż według mniemania ówczesnych, tylko takie ruchy przystoją obiektom boskim. Tego samego przeświadczenia był Kopernik, stąd jego model nie przewidywał poprawnie położen planet i dlatego był mniej dokładny od modelu Ptolemeusza. Tak więc, model Ptolemeusza nie został odrzucony dlatego, że gorzej odtwarzał dane pomiarowe. Wręcz przeciwnie! To, że w tamtych czasach lepiej potrafił przewidzieć położenie obiektów kosmicznych niż model Kopernika, było dla wielu argumentem za jego poprawnością. Natomiast gorsze przewidywania położen planet w modelu Kopernika było argumentem przeciw jego teorii. Teoria Ptolemeusza została odrzucona głównie z powodu odkryć Galileusza i Keplera i to z bardzo dużym oporem. Podstawowym problemem był autorytet Arystotelesa i jego wizja funkcjonowania świata w której Ziemia stanowiła naturalne centrum wszystkiego i do którego wszystko co materialne w naturalny sposób podążało. Zaś obiekty kosmiczne jako boskie krążyły wiecznie ruchem jednostajnym niewrażliwe na wpływy Ziemi. Nowe teorie były bolesne dla ówczesnych elit, gdyż kazały odrzucić cały dotychczasowy światopogląd, również z kilkoma „nieomylnymi” doktrynami kościoła. Powodowało to zrozumiały opór.

- A jak byłoby dzisiaj? {Zadumał się profesor i ciągnął dalej swoje rozmyślenia}

Skoro teoria Ptolemeusza była błędna i nie mamy teraz co do tego żadnych wątpliwości, to w jaki sposób radziła sobie tak dobrze z danymi pomiarowymi? Wtem uświadomił sobie, że układ Ptolemeusza jest układem płaskim. Podobnie zresztą jak Kopernika. Dla takiego układu obserwator będący w środku na Ziemi nie ma możliwości obserwowania ruchu ciał niebieskich „z góry”. Stąd, widzi je jako rzut w kierunku Ziemi (Rys. 50).



Rys. 50 Ruch ciał niebieskich w modelu Ptolemeusza.

A przecież, rzut ruchu jednostajnego po okręgu na prostą leżącą w płaszczyźnie tego ruchu jest funkcją sinusoidalną. Zaś złożenie takich rzutów ruchów kołowych daje w efekcie sumę takich funkcji. Ostatecznie tworzy to szereg Fouriera, który posiada matematyczną zdolność dopasowania się do każdej funkcji okresowej. Niewątpliwie widziany na ziemi ruch ciał niebieskich jest z dobrym przybliżeniem ruchem okresowym i mógłby być za pomocą szeregu Fouriera całkiem dobrze opisany. Tymczasem, obserwator na Ziemi nie widzi dosłownie rzutu na prostą, ale rzut do punktu reprezentowanego przez Ziemię. Co daje ostatecznie opis położenia w postaci zmian kątowych np. $\varphi(t)$, a nie $x(t)$. Dlatego też otrzymana funkcja nie odpowiada dokładnie szeregowi Fouriera, ale posiada podobną zdolność dopasowania się do wyniku doświadczalnego poprzez dobór takich parametrów jak: promienie i prędkości kątowe dla poszczególnych ruchów kołowych. Choć w tym przypadku nie korzysta się świadomie z szeregu Fouriera, to im więcej ruchów kołowych się uwzględni do opisu danego ciała niebieskiego w modelu Ptolemeusza, tym dokładniej można wpasować się w dane pomiarowe. Profesor zrozumiał więc, że dla jakiegokolwiek rozpatrywanej teorii, zgodność z doświadczeniem nie jest wystarczającym dowodem na uznanie jej za prawdziwą. Jeśli korzysta się z takich mechanizmów matematycznych jak choćby szeregi Fouriera. Wówczas, „doskonała” zgodność z doświadczeniem wynikałaby nie z poprawności przyjętych hipotez ale z właściwości zastosowanego aparatu matematycznego, który pozwala się wpasować w dane pomiarowe.

- A jak rzecz ma się obecnie? Czy bardzo złożony aparat matematyczny mechaniki kwantowej może doskonale dopasować się do uzyskanych w doświadczeniu wyników? Czy na pewno deklarowana w książkach jej zgodność z doświadczeniem wynika z poprawności przyjętych hipotez, mimo że są one nielogiczne i absurdalne? Jak bardzo wrażliwy jest model matematyczny mechaniki kwantowej na odrzucenie hipotezy o istnieniu fal materii, skoro przeprowadzone ze światłem doświadczenia każą tą hipotezę odrzucić! Czy poprzez potraktowanie np. elektronów jako fali materii (funkcje sinusoidalne), a następnie „dodawanie” tych fal w celu obliczenia jednego wspólnego stanu kwantowego, nie generuje się przypadkiem tworu przypominającego szereg Fouriera? Być może jest to główna przyczyna pozwalająca uzyskiwać „doskonałą zgodność” z doświadczeniem! Ale w żaden sposób nie świadczy to wówczas o prawdziwości zastosowanej teorii.

- A co z falami elektromagnetycznymi? {dalej kontynuował profesor} Czy aby na pewno istnieją? W końcu, teoretycznie są traktowane jako takie same twory co światło. Ale czy światło może być tylko cząstką? Bo jeśli tak, to co z jego masą? Przecież, byłoby to sprzeczne z teorią względności, dla której żaden materialny

obiekt prędkości światła osiągnąć nie może. Czy zatem teoria Einsteina jest aby na pewno prawdziwa? Wydaje się, że jest ona dobrze ugruntowana doświadczalnie. Skoro jednak istnieją tak rażące niezgodności doświadczenia z teorią dotyczącą światła, a jest to zarazem fundament dla innych rozważań, to gdzie jeszcze moglibyśmy oczekiwać niespodzianek?

Profesor przytłoczony tymi wszystkimi wątpliwościami i pytaniami usiadł i wielce się zmartwił. Gdyż nie wiedział ostatecznie, co jest pewne w teorii, a co nie. Przeprowadzonych w laboratorium doświadczeń odrzucić nie mógł, ale nie wiedział jak je inaczej w pełni wyjaśnić. Rozumiejąc zaś znaczenie tych odkryć, zmartwił się jeszcze bardziej.

[1] Andrzej Kajetan Wróblewski, „Prawda i mity w fizyce”

[2] S. Kryszewski, „Mechanika Kwantowa”

[3] R.P. Feynman, R.B. Leighton, M. Sands, „Feynmana wykłady z fizyki”

[4] Michał Gryziński, „Sprawa Atomu”

[5] D. Halliday, R. Resnick, J. Walker, „Podstawy Fizyki”

[6] Andrzej Kajetan Wróblewski, „Historia Fizyki”

Przedstawiona praca siłą rzeczy nie daje odpowiedzi na wszystkie wątpliwości dotyczące mechaniki kwantowej i postulatów na jakich się ta teoria opiera. Wręcz przeciwnie, sporo wątpliwości zapewne narosło podczas jej czytania. Wątpliwości dotyczą zresztą większego spektrum zagadnień niż tylko mechaniki kwantowej.

Jak niemal każda intelektualna działalność człowieka, tak i ta obarczona jest ryzykiem błędzenia. Dlatego przedstawiona praca była sprawdzana w różnym czasie i na różnych etapach jej realizacji przez znajomych fizyków. Zarówno tych mocno i mniej tytułowanych jak i zdolnych studentów. Odbyły się również dwa skromne seminaria poświęcone temu zagadnieniu na Politechnice Gdańskiej by poddać wyciągnięte wnioski próbie naukowej krytyki. Na seminariach pojawiały się różne wątpliwości i zarzuty, które znajdują swoje odpowiedzi w ostatecznym kształcie prezentowanej pracy znacząco ją ubogacając. Ostatecznie, do dnia dzisiejszego praca nie została merytorycznie obalona i spotkała się u niektórych z pozytywną opinią ale jednocześnie z biernym przyjęciem.

Wszelkie rozpowszechnianie powyższej pracy jest dozwolone z zastrzeżeniem, że nie wolno bez zgody autora dokonywać w niej zmian.

**Z poważaniem
Paweł Fiertek**