

POMIARY BEZ ODDZIAŁYWANIA

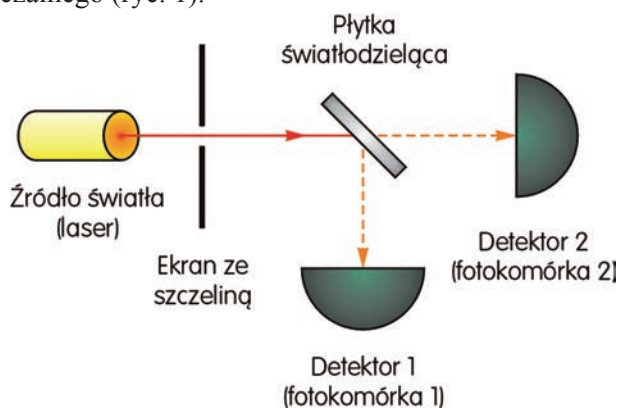
Paweł Tomasz Pęczkowski (Warszawa)

Wstęp

W doświadczeniach interferencyjnych zagadnienie można przedstawić w ten sposób, że ciemne prążki odpowiadają obszarom, w których prawdopodobieństwo padania fotonów jest małe. W doświadczeniu z dwiema szczelinami światło może dotrzeć dwiema różnymi drogami do punktu ekranu stanowiącego detektor: przejść przez jedną szczelinę albo przejść przez drugą szczelinę. Jeżeli nie podejmiemy żadnej próby określenia, którą drogę światło wybierze, to zgodnie z zasadami mechaniki kwantowej nastąpi interferencja. Gdybyśmy mogli w jakikolwiek sposób ustalić, przez którą szczelinę przeszedł foton, to interferencja nie wystąpi. Tu nasuwa się pytanie, czy wiązkę światła można rozdzielić na dwie części i czy poszczególne fotony ulegają przy tym rozszczepieniu.

Czy można rozszczepić foton?

W celu odpowiedzi na postawione pytanie przeanalizujemy doświadczenie związane z rozdzieleniem wiązki światła za pomocą zwierciadła półprzepuszczalnego (ryc. 1).



Ryc. 1. Rozszczepienie wiązki światła za pomocą zwierciadła półprzepuszczalnego (schemat układu doświadczenia). Opracowanie graficzne: Michał Pawlik

Układ doświadczenia zawiera źródło światła, ekran ze szczeliną, zwierciadło półprzepuszczalne oraz dwie fotokomórki wykrywające wiązki fal. Można tak ustawić zwierciadło rozszczepiające wiązkę, aby natężenie każdej z dwóch otrzymanych wiązek po przejściu przez nie było równe połowie natężenia wiązki wychodzącej ze szczeliny. Można tak przygotować doświadczenie, żeby fotokomórka reagowała na wiązki o energii większej niż pewna energia

progowa, a następnie tak dobrać natężenie wiązki wychodzącej ze źródła, żeby natężenie wiązki przepuszczonej do fotokomórki 2 było poniżej wartości progowej. Zgodnie z modelem klasycznym wiązka rozszczepi się na dwie części o energii równej połowie energii początkowej. Wobec tego fotokomórka 2 nie powinna zadziałać.

Wynik przewidywany przez teorię klasyczną stoi jednak w sprzeczności z faktami doświadczalnymi. Okazało się, że do fotokomórki dociera w dalszym ciągu światło, tylko szybkość zliczeń jest dwa razy mniejsza niż w przypadku usunięcia zwierciadła światłodzielnego. Wniosek z doświadczenia jest taki, że światło dochodzi do fotokomórki porcjami energii – fotonami, a fotony nie mogą się rozszczepiać.

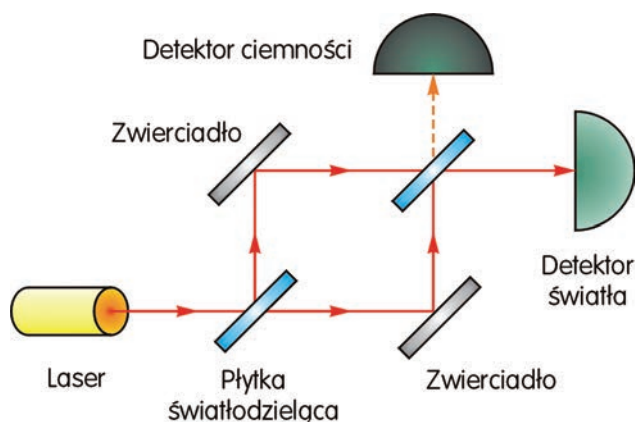
Przeprowadzono też doświadczenia badające zależność zjawiska fotoelektrycznego od odległości od źródła światła (r). Natężenie światła zgodnie z teorią falową jest proporcjonalne do $\frac{1}{r^2}$, więc zgodnie z teorią klasyczną energia wiązki docierającej do fotokomórki powinna być proporcjonalna do $\frac{1}{r^2}$. Stwierdzono doświadczalnie, że fotokomórka rejestruje światło przy wszystkich odległościach r , natomiast szybkość zliczeń maleje proporcjonalnie do $\frac{1}{r^2}$.

Wniosek, który wynika z tych doświadczeń jest taki, że fotony (kwanty światła) zachowują się inaczej niż klasyczne wiązki falowe.

Czy da się przeprowadzić pomiar „bez udziału” fotonów?

Dwaj fizycy z Uniwersytetu w Tel Awiwie, Avshalom C. Elitzur i Lev Vaidman zaprojektowali doświadczenie myślowe, które dowodzi istnienia pomiarów bez oddziaływania, czyli wykrywania obiektów „bez użycia” światła na niepadającego. Hipotetyczny układ pomiarowy stanowi interferometr Macha-Zehndera składający się z dwóch zwierciadeł i dwóch silikonowych płytek światłodzielných, ustawionych tak, jak to pokazano na ryc. 2.

Pierwsza płytka światłodzielną kieruje światło wchodzące do interferometru wzdłuż dwóch możliwych dróg – albo do jednego zwierciadła albo do drugiego. Druga płytka światłodzielną jest ustawiona w ten sposób, że obie drogi optyczne łączą się na niej i fotony zawsze trafiają do jednego z dwóch



Ryc. 2. Schemat interferometru Macha-Zehndera. Opracowanie graficzne: Michał Pawlik

detektorów. Jeden z detektorów został nazwany detektorem światła. Odpowiada on interferencji konstruktywnej w doświadczeniu z dwiema szczelinami. Drugi detektor, zwany detektorem ciemności, do którego fotony nie docierają odpowiada interferencji destruktywnej w doświadczeniu z dwiema szczelinami.

Przeanalizujemy, w jaki sposób pole elektromagnetyczne dociera do detektora. W tym celu opiszemy formalnie przejście fotonu przez interferometr. Oznaczmy stan fotonu przechodzącego przez pierwszą płytkę światłodzielną w kierunku na prawo (na ryc. 1) przez $|1\rangle$, a stan fotonu odbitego od płytki i przechodzącego w kierunku do góry przez $|2\rangle$. Każdemu odbiciu od zwierciadła, czy od płytki światłodzielną towarzyszy zmiana fazy fotonu o $\frac{\pi}{2}$. Wniosek ten wynika z optyki falowej i oznacza, że zespolone pole elektryczne $\mathbf{E}$ należy pomnożyć przy każdym odbiciu przez $(-i)$.

A więc operację przejścia fotonu przez płytkę można zapisać, jako

$$|1\rangle \rightarrow \frac{1}{\sqrt{2}}[|1\rangle + i|2\rangle] \quad (1)$$

$$|2\rangle \rightarrow \frac{1}{\sqrt{2}}[|2\rangle + i|1\rangle] \quad (2)$$

a operacja przejścia przez zwierciadło całkowicie odbijające można zapisać, jako

$$|1\rangle \rightarrow i|2\rangle \quad (3)$$

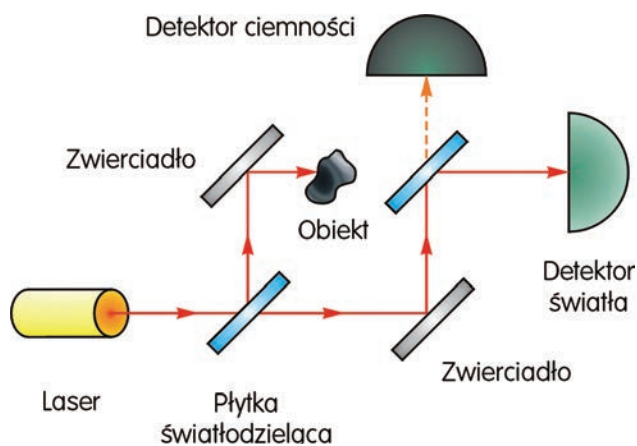
$$|2\rangle \rightarrow i|1\rangle \quad (4)$$

Jeżeli przeszkoda jest nieobecna, operację przejścia przez zwierciadło i obie płytki światłodzielną można zapisać, jako złożone operacje przejścia

$$|1\rangle \rightarrow \frac{1}{\sqrt{2}}[|1\rangle + i|2\rangle] \rightarrow \frac{1}{\sqrt{2}}[i|2\rangle + i^2|1\rangle] = \frac{1}{\sqrt{2}}[i|2\rangle - |1\rangle] \rightarrow \frac{1}{2}[i|2\rangle - |1\rangle] - \frac{1}{2}[|1\rangle + i|2\rangle] = -|1\rangle \quad (5)$$

Oznacza to, że foton jest tylko rejestrowany przez detektor światła, a nigdy nie jest rejestrowany przez detektor ciemności.

Zastanówmy się teraz, co będzie, gdy na jednej z dróg fotonu w interferometrze umieścimy obiekt-przeszkodę (ryc. 3).



Ryc. 3. Umieszczenie przeszkody (obiektu) w interferometrze Macha-Zehndera. Opracowanie graficzne: Michał Pawlik

Jeżeli na jednej z dróg optycznych umieścimy przeszkodę i wysyłamy pojedynczy foton, to możliwe są trzy przypadki:

- a) detektor światła wykryje foton;
- b) detektor ciemności wykryje foton;
- c) żaden detektor nie wykryje fotonu.

Jeżeli na jednej z dróg (np. górnej) umieścimy przeszkodę, to foton wybierając górną drogę (co zachodzi z prawdopodobieństwem 50%) nigdy nie dotrze do drugiej płytki światłodzielną. Natomiast, jeżeli foton wybierze dolną drogę to na drugiej płytce światłodzielną nie wystąpi interferencja, ponieważ foton mógł dotrzeć do tej płytki tylko jedną drogą. A zatem z prawdopodobieństwem 50% foton wybierze jedną z dwóch dalszych dróg – albo do detektora światła albo do detektora ciemności. Jeżeli foton dotrze do detektora światła (przypadek a), to nie wiemy, która sytuacja (został umieszczony obiekt, czy nie) nastąpiła. Natomiast, jeżeli foton dotrze do detektora ciemności (przypadek b), to wiemy, że w układzie był umieszczony obiekt. W tym przypadku foton nie miał styczności z obiektem, ponieważ musiał przebiegać inną drogą. Udało nam się ustalić obecność obiektu (przeprowadzić pomiar) mimo braku oddziaływania foton-obiekt. Widzimy, że sama obecność obiektu wyklucza możliwość wystąpienia interferencji, chociaż nie ma żadnego oddziaływania fotonu z obiektem. Zauważmy, że obecność przeszkody umiemy wykrywać tylko w połowie przypadków (50%), kiedy ona występuje. Zastanówmy się dlaczego? Odpowiada nam na to pytanie przypadek c. W tym przypadku

foton zostaje pochłonięty lub rozproszony przez obiekt i nigdy nie dotrze do żadnego z detektorów. Prawdopodobieństwo takiego wyniku wynosi 50%.

Opiszmy tę sytuację formalnie. Jeżeli w interferometrze znajduje się obiekt, operację przejścia można opisać jako

$$|1\rangle \rightarrow \frac{1}{\sqrt{2}}[|1\rangle + i|2\rangle] \rightarrow \frac{1}{\sqrt{2}}[i|2\rangle + i|s\rangle] = \frac{1}{2}[i|2\rangle - |1\rangle] + \frac{i}{\sqrt{2}}|s\rangle, \quad (6)$$

gdzie $|s\rangle$ oznacza stan fotonu rozproszonego przez obiekt.

Ostatnie przejście w operacji (6) ma postać

$$\frac{1}{2}[i|2\rangle - |1\rangle] + \frac{i}{\sqrt{2}}|s\rangle \rightarrow \begin{cases} |1\rangle, \\ |2\rangle, \\ |s\rangle, \end{cases} \quad (7)$$

gdzie $|1\rangle$ oznacza dojście do detektora światła (prawdopodobieństwo 0,25),

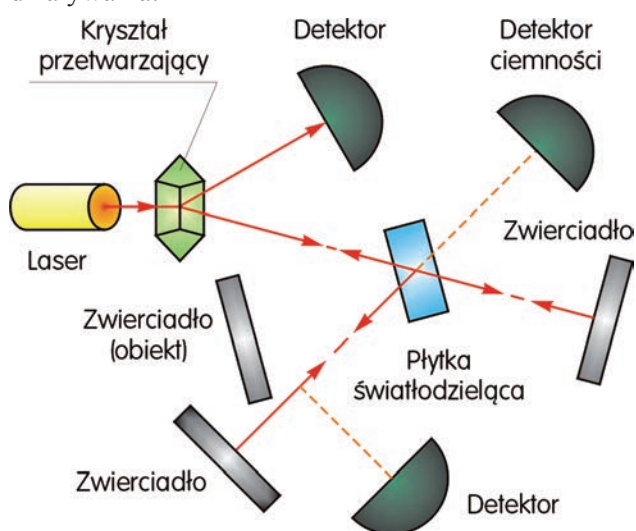
$|2\rangle$ oznacza dojście do detektora ciemności (prawdopodobieństwo 0,25),

$|s\rangle$ oznacza pochłonięcie przez obiekt (prawdopodobieństwo 0,5).

Realizacja doświadczenia myślowego

A. C. Elitzura i L. Vaidmana

Zespół P. Kwiat, H. Weinfurter, A. Zeilinger (Innsbruck) z T. Herzogiem (Genewa) przeprowadził w 1995 roku doświadczenie, które było realizacją przedstawionego eksperymentu myślowego. W ten sposób autorzy doświadczenia wykazali, że można zbudować urządzenie służące do pomiarów bez oddziaływania.



Ryc. 4. Schemat układu A. C. Elitzura i L. Vaidmana. Opracowanie graficzne: Michał Pawlik

Źródłem fotonów był specjalny kryształ nieliniowy. Pod wpływem promieni ultrafioletowych pojedyncze fotony wysyłane przez laser, kierowane na kryształ powodowały czasem wytworzenie dwóch bliźniaczych

fotonów rozbiegających się pod kątem około 30° o energii dwukrotnie mniejszej niż fotony źródłowe. W ten sposób zostaje wytworzona para fotonów. Wykrywając jeden z nich mamy absolutną pewność, że istnieje też drugi foton bliźniaczy, który został wprowadzony do układu doświadczalnego. Na ryc. 4 został przedstawiony schemat układu tego doświadczenia.

Układ różni się nieznacznie od układu z doświadczenia myślowego Elitzura-Vaidmana, ale idea doświadczenia jest taka sama. Zwierciadła i płytka światłodzielnąca były tak ustawione, że prawie wszystkie fotony wychodziły po tej samej drodze, po której weszły do interferometru. Kiedy nie było przeszkody na drodze od płytki światłodzielnącej do jednego ze zwierciadeł, to szansa, aby foton trafił do detektora ciemności była znikoma. Potem na drodze fotonu ustawiono małe zwierciadło (obiekt pełniący rolę przeszkody), które kierowało fotony do innego detektora, nazwanego detektorem przeszkody. Autorzy doświadczenia zauważyli, że mniej więcej w połowie przypadków ten detektor rejestrował foton. Zadziałał również detektor ciemności, który rejestrował foton mniej więcej co czwarty raz. W pozostałych przypadkach foton opuszczał interferometr po tej samej drodze, po której wszedł i nie dawał żadnej informacji, na temat ustawienia przeszkody. Widać, że każde zarejestrowanie fotonu przez detektor ciemności dawało informację o istnieniu obiektu pełniącego rolę przeszkody, mimo że foton nie oddziaływał z tym obiektem.

Później autorzy modyfikowali doświadczenie w ten sposób, że zmieniali stopniowo zdolność odbijającą płytki światłodzielnącej zmniejszając szansę odbicia fotonu w kierunku ścieżki, na której było umieszczone zwierciadło kierujące fotony do detektora przeszkody. Wówczas, zgodnie z przewidywaniami teoretycznymi, prawdopodobieństwa trafienia fotonów do detektora przeszkody i detektora ciemności wyrównywały się coraz bardziej. Oznacza to, że stosując bardzo słabą odbijającą płytkę światłodzielnąca można przeprowadzić około połowy pomiarów bez oddziaływania z obiektem mierzonym. Powstało pytanie, czy tej proporcji, wynoszącej 50%, nie można poprawić.

Przebywający w styczniu w 1994 roku z miesięczną wizytą w Innsbrucku M. Kasevich ze Stanford University zaprojektował eksperyment, który pozwoliłby wykrywać obiekty bez oddziaływania z nimi prawie w 100%. Doświadczenie Kasevicha związane jest efektem kwantowym zwanym **efektem lub paradoksem Zenona**, który w niniejszym artykule omówimy. Wiadomo, że układ kwantowy znajdujący się w danym stanie początkowym, pozostawiony sam sobie

może ewoluować do innego stanu. Jednak z powodu wpływu, jakie pomiary wywołują na układy kwantowe, może nastąpić „uwięzienie” układu w swoim stanie początkowym.

Kwantowy paradoks Zenona

Nazwa efektu pochodzi od słynnego paradoksu sformułowanego przez starożytnego filozofa Zenona z Elei. Zenon analizował ruch strzały wystrzelonej w kierunku uciekającego jelenia i zauważył, że strzała w każdym momencie znajduje się w określonym miejscu przestrzeni, a zatem jest ona nieruchoma („uwięziona” w tym miejscu), a więc nie może nigdy dotrzeć do jelenia. Występujący w doświadczeniu, które zaraz opiszemy, paradoks nosi nazwę **kwantowego efektu Zenona** albo **efektu pilnowanego czajnika**. To drugie określenie odnosi się do aforyzmu o gotującej się wodzie w czajniku: „obserwowany czajnik nigdy nie zagotuje wody”. Jest oczywiste, że pilnowanie czajnika nie ma wpływu na szybkość gotowania się wody. W fizyce klasycznej samo obserwowanie lub dokonywanie pomiaru nie wpływa na wynik. Okazuje się, że w mechanice kwantowej jest inaczej. Dokonywanie pomiaru niszczy stan układu. Na ogół po wykonaniu pomiaru układ może znajdować się w trudnym do przewidzenia stanie. Jest jednak klasa pomiarów, w których układ po wykonaniu pomiaru znajduje się w stanie odpowiadającym jego wynikowi. W dalszym ciągu pracy, zakładamy, że wykonujemy takie pomiary.

Effekt Zenona a efekt anty-Zenona

Zastanówmy się, dlaczego kwantowy paradoks Zenona jest nieoczekiwanym matematycznym wynikiem, który jest rozważany do dzisiejszych czasów. Zakładając, że mamy niestabilny stan kwantowy, intuicja podpowiada nam, że układ przejdzie w sposób nieodwracalny do innego stanu w ciągu pewnego czasu zwanego czasem Zenona. Jednak, jeżeli będziemy dokonywać pomiaru układu w czasie krótszym niż czas Zenona, to funkcja falowa układu ulega zmianie, zanim układ przejdzie do innego stanu. W efekcie częste pomiary układu zapobiegają przejściu do innego stanu. Co ciekawsze, jeżeli odstęp czasu między dwoma pomiarami jest dłuższy niż czas Zenona, to odsetek przejść do innego stanu zwiększa się, prowadząc do efektu nazwanego **efektem anty-Zenona**.

Wygaszenie (zanik) niestabilnego układu jest klasycznie przedstawione w postaci funkcji wykładniczej. Funkcja wykładnicza jest powszechnie stosowana do modelowania procesów, w których

znane jest prawdopodobieństwo zajścia pewnego zdarzenia (np. przejścia układu z jednego stanu do innego) w ciągu określonego czasu i to prawdopodobieństwo nie zależy od przeszłości tzn. nie zależy od tego, jak długo układ znajduje się w określonym stanie. Jednak dla bardzo krótkich i bardzo długich okresów czasu wykładnicze wygaszenie nie jest odpowiednie do zastosowania. Dla krótkich okresów wygaszenie ma charakter gaussowski.

Wychodząc z rozkładu wykładniczego do obliczenia prawdopodobieństwa przeżycia układu w ciągu czasu t dosyć nieoczekiwanie zauważamy, że kiedy stale dokonujemy pomiaru niestabilnego układu kwantowego, to zawsze pozostanie on w swoim stanie początkowym.

W praktyce nie można w sposób ciągły dokonywać pomiarów. Nieciągłe pomiary prowadzą do wyników jeszcze dziwniejszych niż kwantowy efekt Zenona.

Rozważmy klasyczny model wygaszenia układu. Niech M będzie ogólną liczbą niestabilnych układów, a γ_0 stałym prawdopodobieństwem na jednostkę czasu, że układ wygaśnie. γ_0 jest stałe i nie zależy od innych czynników takich jak zachowanie się układu w przeszłości ani od otoczenia. Wyrażmy liczbę niestabilnych układów jako funkcję czasu, $M(t)$. Wówczas wygaszenie na jednostkę czasu możemy zapisać jako

$$\frac{dM}{dt} = -M\gamma_0. \quad (8)$$

Po scałkowaniu równania (8) i oznaczeniu $M_0 = M(0)$ liczby niestabilnych układów w czasie początkowym $t = 0$ mamy

$$\frac{dM}{M} = -\gamma_0 dt, \quad (9)$$

$$\int_{M_0}^{M(t)} \frac{dM}{M} = -\gamma_0 \int_0^t dt, \quad (10)$$

stąd

$$M(t) = M_0 \exp(-\gamma_0 t). \quad (11)$$

Teraz zależne od czasu $p(t)$ prawdopodobieństwo, że układ nie wygaśnie wynosi

$$p(t) = \frac{M(t)}{M_0} = \exp(-\gamma_0 t). \quad (12)$$

To jest dobrze znany klasyczny model wygaszenia układu. Zastanówmy się teraz nad modelem kwantowym. Niech $\psi(t)$ będzie stanem funkcji falowej układu kwantowego i niech $\psi_0 = \psi(0)$ będzie

stanem funkcji falowej w czasie początkowym $t = 0$. Ewolucja czasowa układu kwantowego zadana jest przez tzw. operator ewolucji czasowej $\hat{U}(t)$, utworzony przy użyciu hamiltonianu $\hat{H}$

$$\hat{U}(t) = \exp\left(-\frac{i}{\hbar} \hat{H} t\right). \quad (13)$$

Prawdopodobieństwo, że układ nie wygaśnie jest kwadratem amplitudy przejścia, czyli

$$p(t) = |\langle \psi_0 | \hat{U}(t) | \psi_0 \rangle|^2. \quad (14)$$

To prawdopodobieństwo może być rozwinięte jako

$$p(t) \approx 1 - \frac{t^2}{\hbar^2} (\langle \psi_0 | \hat{H}^2 | \psi_0 \rangle - \langle \psi_0 | \hat{H} | \psi_0 \rangle^2) + \dots, \quad (15)$$

Oznaczając we wzorze (15)

$$\Delta \hat{H} = \sqrt{\langle \psi_0 | \hat{H}^2 | \psi_0 \rangle - \langle \psi_0 | \hat{H} | \psi_0 \rangle^2}, \quad (16)$$

otrzymujemy

$$p(t) \approx 1 - \frac{t^2}{\hbar^2} (\Delta \hat{H})^2 + \dots \quad (17)$$

Zdefiniujmy czas Zenona jako $\tau_z = \frac{\hbar}{(\Delta \hat{H})^2}$, wówczas wzór (17) możemy przepisać w postaci

$$p(t) \approx 1 - \frac{t^2}{\tau_z^2} + \dots \quad (18)$$

Jak widać ze wzoru (18) kwantowe wygaśnięcie nie jest funkcją wykładniczą. Okazuje się, że prawdopodobieństwo przeżycia układu ma rozkład Gaussa. Możemy zapisać prawdopodobieństwo $p(t)$ w postaci

$$p(t) \approx \exp\left(-\frac{t^2}{\tau_z^2}\right). \quad (19)$$

Wygaszenie niestabilnego układu kwantowego zmienia się zależnie od szybkości pomiaru zjawiska kwantowego, od charakteru kwadratowego dla krótkiego czasu dokonywania pomiaru układu poprzez wykładniczy dla średniego czasu i do potęgowego dla długiego czasu. Przejście między krótkim i średnim czasem ma charakter nieintuicyjny.

Oznaczmy przez N liczbę pomiarów dokonanych w równych odstępach czasu τ w ciągu odcinka czasowego $[0, T]$. Zatem $T = N\tau$. Prawdopodobieństwo przejścia po N pomiarach wynosi

$$p^N(T) = [p(\tau)]^N = \exp(N \ln[p(\tau)]), \quad (20)$$

Przy oznaczeniu

$$\gamma_{\text{eff}}(\tau) = \frac{1}{\tau} \ln[p(\tau)] \geq 0, \quad (21)$$

gdzie $N \rightarrow \infty$, to mamy

$$\lim_{N \rightarrow \infty} p^N(T) = \lim_{N \rightarrow \infty} [p(\tau)]^N = \lim_{N \rightarrow \infty} \exp\left[\left[-\gamma_{\text{eff}}(\tau)\right] N\tau\right] = 1 \quad (22)$$

Zatem gdy $N \rightarrow \infty$, to prawdopodobieństwo przejścia dąży do 1. Oznacza to, że ciągłe pomiary układu nigdy nie spowodują jego zanikania i układ pozostanie w stanie początkowym.

Wartość $\gamma_{\text{eff}}(\tau)$ zmienia się zgodnie z wyrażeniem

$$\gamma_{\text{eff}}(\tau) = \begin{cases} \frac{\tau}{\tau_z^2}, & \text{gdzie } \tau \rightarrow 0 \\ \gamma_0, & \text{gdzie } \tau \rightarrow \infty \end{cases}, \quad (23)$$

gdzie τ_z - czas Zenona, γ_0 - pewna stała charakteryzująca układ, prawdopodobieństwo na jednostkę czasu, że układ wygaśnie.

Załóżmy, że istnieje czas przejścia τ^* między zanikiem gaussowskim i wykładniczym taki, że

$$\gamma_{\text{eff}}(\tau) = \tau_0. \quad (24)$$

Wprowadzając pewną stałą τ_z rozważmy następujące przypadki:

1. $\tau = \tau^*$

Wówczas $\frac{\tau}{\tau_z^2} = \frac{\tau^*}{\tau_z^2}$, a więc $\gamma_{\text{eff}}(\tau) = \gamma_0$. Mamy klasyczny przypadek wygaszania wykładniczego.

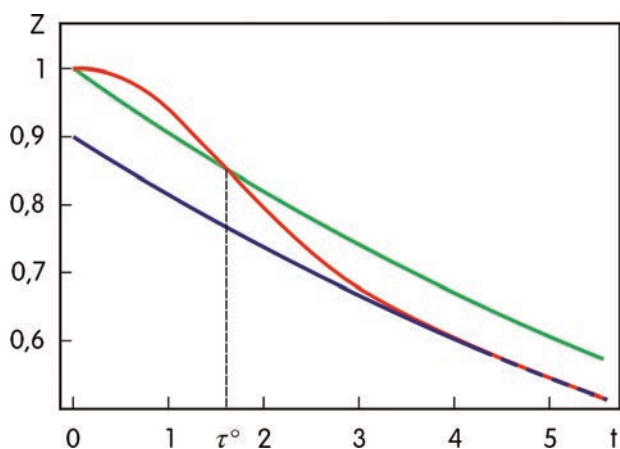
2. $\tau < \tau^*$

Wówczas $\frac{\tau}{\tau_z^2} < \frac{\tau^*}{\tau_z^2}$, a więc $\gamma_{\text{eff}}(\tau) < \gamma_0$. To oznacza, że współczynnik wygaszenia jest mniejszy w porównaniu do wygaszenia wykładniczego. Częste pomiary powodują wolniejsze przejście układu. To jest właśnie kwantowy efekt Zenona.

3. $\tau > \tau^*$

Wówczas $\frac{\tau}{\tau_z^2} > \frac{\tau^*}{\tau_z^2}$, a więc $\gamma_{\text{eff}}(\tau) > \gamma_0$. W porównaniu do kwantowego efektu Zenona wynik pokazuje, że wygaszanie układu rośnie w porównaniu do wygaszania wykładniczego. Częste pomiary powodują, że wygaszanie układu ulega przyspieszeniu. To jest efekt przeciwny do efektu Zenona, czyli efekt anty-Zenona, zwany też efektem Heraklita w nawiązaniu do słynnego powiedzenia greckiego filozofa Heraklita (*panta rhei* – wszystko płynie), który głosił, że wszystko w przyrodzie się zmienia.

Na ryc. 5. zostały przedstawione trzy funkcje ukazujące: prawdopodobieństwo przeżycia układu $p(t)$ – linia ciągła, wygaszenie wykładnicze $\exp[-\gamma_{\text{eff}}(t) \cdot t]$ – linia przerywana i kwantowo mechaniczne wygaszanie wykładnicze $Z \exp[-\gamma_{\text{eff}}(t) \cdot t]$, gdzie $Z < 1$ – linia kropkowana. Z jest parametrem służącym renormalizacji funkcji falowej w teorii kwantowej pola. $Z < 1$ dla stanów stabilnych, dla stanów niestabilnych Z nie jest ograniczone. τ^* jest punktem przecięcia się dwóch funkcji zaznaczonych na rysunku linią ciągłą i linią przerywaną. Jeżeli $Z < 1$, to istnieje τ^* , które jest warunkiem dostatecznym dla efektu Zenona i anty-Zenona.



Ryc. 5. Ilustracja funkcji pokazujących: prawdopodobieństwo przeżycia układu $p(t)$ – linia ciągła, wygaszenie wykładnicze $\exp[-\gamma_{\text{eff}}(t) \cdot t]$ – linia przerywana i kwantowo mechaniczne wygaszanie wykładnicze $Z \exp[-\gamma_{\text{eff}}(t) \cdot t]$, gdzie $Z < 1$ – linia kropkowana (na podstawie publ. Facchi, Nakazato, Pascasio, 2001 – czasopismo Phys. Rev. Lett.). Opracowanie graficzne: Michał Pawlik

Chociaż kwantowy efekt Zenona i anty-Zenona są dobrze udokumentowane z teoretycznego punktu widzenia, potrzebne są wyniki doświadczalne w celu sprawdzenia hipotetycznego modelu teoretycznego. Jedną z trudności jest odcinek czasu, w którym te efekty mogą być obserwowane. Obliczenia teoretyczne dokonane przez F. Facchi, H. Nakazato i S. Pascasio wskazują, że czas Zenona dla przejścia 2P-1S atomu wodoru wynosi około $3,59 \cdot 10^{-15}$ s. Biorąc pod uwagę czas tego przejścia wynoszący $1,595 \cdot 10^{-9}$ s, obserwacje efektu Zenona i anty-Zenona są trudnym zadaniem. Należało więc znaleźć układ, dla którego czas Zenona jest znacznie dłuższy niż $3,59 \cdot 10^{-15}$ s, aby można było zaobserwować ten efekt. Po raz pierwszy udało się znaleźć taki układ i zaobserwować kwantowy efekt Zenona w 1989 roku.

Obserwacja kwantowego efektu Zenona

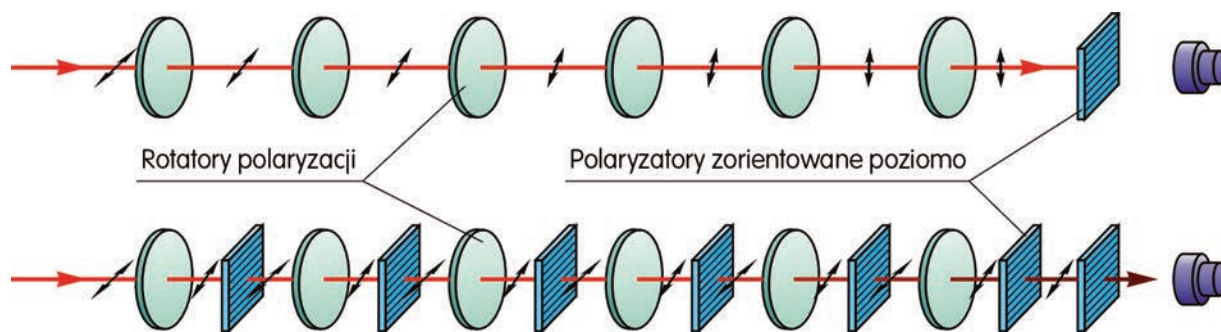
Doświadczenie to przeprowadzili W. M. Itano, D. J. Heinzen, J. J. Bollinger, D. J. Wineland w National Institute of Standards and Technology (NIST)

w Boulder w Kolorado. Eksperyment ten dotyczył pomiaru przejścia między dwoma stanami energii 5000 zamrożonych laserowo jonów berylu $^9\text{Be}^+$ zamkniętych w pułapce (tzw. pułapce Penninga). Przyjęto, że ponieważ jest bardzo mało jonów (ciśnienie w pułapce Penninga wynosiło około 10^{-8} Pa) można zaniedbać oddziaływania między jonami. Na początku eksperymentu za pomocą promieniowania laserowego o długości fali 313 nm wysyłanego przez około 5 s sprowadzono prawie wszystkie jony berylu do stanu kwantowego $|1\rangle$. Po tym czasie laser został wyłączony. Przepuszczając przez komorę z jonami wiązkę fal o częstotliwości 320,7 mHz przez około $t = 256$ ms zdołano przenieść wszystkie jony do nowego stanu wzbudzonego, oznaczonego $|2\rangle$. Oznacza to, że każdy jon przechodził ze stanu $|1\rangle$ do stanu $|2\rangle$ z prawdopodobieństwem bliskim 100%. Moment przejścia z jednego stanu kwantowego do drugiego nie jest znany i jeżeli nie dokonamy pomiaru, to nie wiemy, w którym stanie znajduje się w danej chwili jon. Mówimy wtedy, że jon znajduje się w superpozycji stanów $|1\rangle$ i $|2\rangle$. Jeżeli jednak dokonujemy pomiarów to w zależności od jego wyniku przeprowadzamy jon do stanu $|1\rangle$ albo do stanu $|2\rangle$. Zespół z NIST opracował metodę pozwalającą obserwować jon używając krótkich impulsów laserowych o długości czasu trwania 2,4 ms. Przyjmowano różną liczbę wysyłanych impulsów $n = 1, 2, 4, 8, 16, 32$ i 64 . Impulsy były dostatecznie długie, żeby zredukować funkcję falową każdego jonu do określonego stanu kwantowego bez powodowania pompowania optycznego. Na podstawie przeprowadzonych doświadczeń zauważono, że po 256 ms niemal 100% jonów znajdowało się w stanie $|2\rangle$. Teoria kwantowa mówi, że gdybyśmy popatrzyli na jony w momencie, gdy upłynie połowa z 256 ms tzn. po 128 ms, to jony zostałyby zmuszone do wyboru jednego z dwóch stanów z prawdopodobieństwem $\frac{1}{2}$ i wobec tego średnio połowa jonów znajdowałaby się w stanie $|1\rangle$ a połowa w stanie $|2\rangle$. Rzeczywiście taką sytuację zaobserwowano. Po upływie 64 ms, jony przechodzą z jednego stanu do drugiego z prawdopodobieństwem $\frac{1}{4}$, czyli pozostają w tym samym stanie z prawdopodobieństwem $\frac{3}{4}$. Jednak, gdy eksperymentatorzy obserwowali jony w odstępach co 64 ms (4 razy w ciągu 256 ms), to po zakończeniu doświadczenia $\frac{2}{3}$ (teoretycznie $1 - (\frac{3}{4})^4 = 1 - \frac{81}{256} \approx \frac{2}{3}$) jonów nadal pozostawało w stanie $|1\rangle$. Gdy obserwowano jony co 4 ms (64 razy), to po zakończeniu doświadczenia prawie wszystkie jony znajdowały się w stanie $|1\rangle$. Zatem częste obserwacje spowodowały spowolnienia przechodzenia jonów ze stanu $|1\rangle$ do stanu $|2\rangle$.

Doświadczalna realizacja efektu Zenona - doświadczenie M. A. Kasevicha

Oprócz doświadczenia przeprowadzonego przez Itano wraz z zespołem, wykonano też inne równie godne uwagi doświadczenie dotyczące efektu Zenona, które zostało przeprowadzone przez M. A. Kasevicha. Idea doświadczenia zaproponowanego przez Kasevicha wywodzi się z pomysłu, który jako

i nie ma urządzeń zmieniających kierunek polaryzacji, to do polaryzatora dochodzi światło spolaryzowane poziomo, które jest przepuszczone i zarejestrowane przez detektor. Po umieszczeniu między źródłem światła i detektorem układu sześciu rotatorów, do polaryzatora dochodzi światło spolaryzowane pionowo, które jest całkowicie absorbowane i detektor nie wykrywa żadnych fotonów.



Ryc. 6. Układ rotatorów i polaryzatorów – ilustracja kwantowego efektu Zenon. Opracowanie graficzne: Michał Pawlik

pierwszy przedstawił Asher Peres z Izraelskiego Instytutu Technologicznego Technon. W doświadczeniu wykorzystano własność polaryzacji światła. Zgodnie z zasadą rozchodzenia się fali świetlnej kierunki drgań pola magnetycznego i elektrycznego są prostopadłe do siebie i prostopadłe do kierunku rozchodzenia się fali, mogą natomiast różnić się ułożeniem tych kierunków. Światło pochodzące z typowych źródeł np. światło słoneczne lub światło żarówki może oscylować w różnych kierunkach, czyli ma różną polaryzację. Mówimy wtedy, że takie światło jest niespolaryzowane. Można uzyskać światło spolaryzowane w określonym kierunku, a przyjmując odpowiednią orientację można mieć polaryzację pionową, poziomą lub pod określonym kątem, np. 15° .

Rozważmy foton, który przechodzi przez urządzenie zmieniające kierunek polaryzacji światła o pewien kąt. Takie urządzenie nazywamy rotatorem. Jeżeli zbudujemy układ składający się z połączonych szeregowo sześciu rotatorów, z których każdy zmienia kierunek polaryzacji o 15° , to po przejściu przez ten układ kierunek polaryzacji światła zmieni się o 90° . Jeżeli promień świetlny jest na początku spolaryzowany poziomo, to po przejściu przez cały układ będzie spolaryzowany pionowo. Ustawmy na końcu układu polaryzator, czyli urządzenie przepuszczające fotony o określonej polaryzacji i zatrzymujące fotony o polaryzacji prostopadłej do niej. Załóżmy, że polaryzator przepuszcza fotony spolaryzowane poziomo, a zatrzymuje fotony spolaryzowane pionowo. Za polaryzatorem umieszczamy detektor wykrywający docierające do niego fotony. Jeżeli wiązka światła wchodząca do układu jest spolaryzowana poziomo

Ustawmy teraz po każdym rotatorze polaryzacji (skręcającym kierunek o 15°) odpowiedni polaryzator przepuszczający tylko światło o rotacji poziomej (ryc. 6). Oznacza to, że prawdopodobieństwo przejścia fotonu przez pierwszy polaryzator jest równe

$$P_1 = \cos^2(15^\circ) \approx 0,933, \text{ czyli } 93,3\%. \quad (25)$$

a prawdopodobieństwo zatrzymania fotonu wynosi

$$P_1' = \sin^2(15^\circ) \approx 0,067, \text{ czyli } 6,7\%. \quad (26)$$

Jeżeli foton przejdzie przez pierwszy polaryzator, to oznacza, że znajduje się w stanie polaryzacji o kierunku poziomym. Następnie, po przejściu przez drugi rotator polaryzacji i drugi polaryzator sytuacja jest taka sama. Prawdopodobieństwo zatrzymania fotonu przez polaryzator wynosi znów 6,7%, a przepuszczone światło będzie znów spolaryzowane poziomo. Ponieważ układ ma sześć rotatorów polaryzacji i sześć polaryzatorów, więc można obliczyć, że prawdopodobieństwo P_6 przejścia fotonu przez cały układ i dotarcie do detektora jest iloczynem prawdopodobieństw przejścia przez każdy polaryzator i wynosi

$$P_6 = \prod_{i=1}^6 P_i = [\cos^2(15^\circ)]^6 \approx \frac{2}{3}. \quad (27)$$

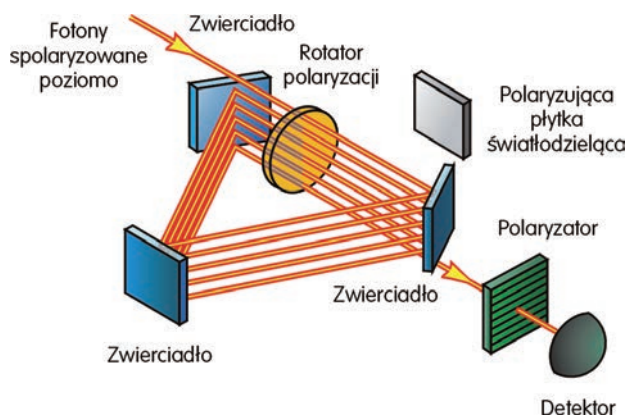
Gdy wzrasta liczba rotatorów i polaryzatorów, to prawdopodobieństwo P_Σ przejścia fotonu przez cały układ wzrasta. Gdy elementów tych jest bardzo dużo (np. kilka tysięcy), to prawdopodobieństwo P'_Σ

zatrzymania fotonu przez układ i nie wykrycia jego przez detektor jest bliskie zero (w rzeczywistym eksperymencie licząc prawdopodobieństwo zatrzymania fotonu przez układ P'_Σ należy jeszcze uwzględnić zjawisko odbicia i absorpcji) i wynosi

$$P'_\Sigma = 1 - \prod_{i=1}^n P_i = 1 - \left[\cos^2 \left(\frac{\pi}{2n} \right) \right]^n \quad (28)$$

gdzie n oznacza liczbę elementów układu. Na przykład, gdy $n = 2500$, to $P'_\Sigma \approx 0,001$.

W doświadczalnej realizacji kwantowego efektu Zenona, która pokazała, jak zatrzymać obrót kierunku polaryzacji światła użyto tego samego kryształu nieliniowego, który posłużył do otrzymania pojedynczego fotonu, ale zamiast układu rotatorów i polaryzatorów użyto tylko po jednym z tych elementów. Rotator polaryzacji jest to urządzenie, które obraca kierunek polaryzacji o określony kąt, np. o 15° . Zbudowano natomiast układ doświadczalny w ten sposób, że foton przechodził przez niego sześciokrotnie po spirali (ryc. 7). Umożliwiły to trzy zwierciadła ustawione do siebie pod kątem 60° . Jest to równoważne użyciu sześciu rotatorów i sześciu polaryzatorów.

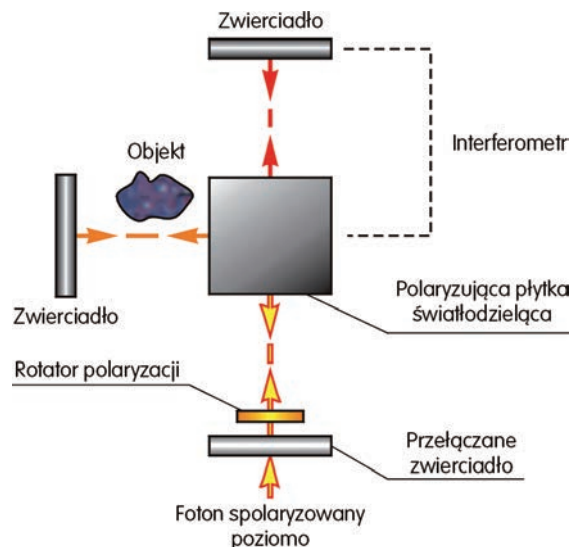


Ryc. 7. Schemat doświadczenia realizującego kwantowy efekt Zenona. Opracowanie graficzne: Michał Pawlik

Gdy w układzie nie ma polaryzatora, foton wchodzący do układu z polaryzacją poziomą ma po wyjściu z układu zmienioną polaryzację na pionową. Gdy wstawiono do układu polaryzator, który po każdym przejściu fotonów (jeden z 6 obrotów) przepuszczał tylko fotony spolaryzowane poziomo, to około $\frac{2}{3}$ fotonów przechodziło przez cały układ i docierało do detektora fotonów. Jest to wynik zgodny z obliczeniami teoretycznymi, przedstawionymi powyżej (wzór 27).

Po wykonaniu opisanego doświadczenia Kwiat, Weinfurter i Zeilinger zaprojektowali układ stanowiący połączenie zestawu do realizacji kwantowego

efektu Zenona z oryginalnym układem Elitzura-Vaidmana. Układ ten posłużył do wykonania efektywnych pomiarów bez oddziaływania (ryc. 8). W doświadczeniu użyto zwierciadeł, które można bardzo szybko włączyć lub wyłączyć np. po określonej liczbie cykli przejścia fotonu przez układ doświadczalny.



Ryc. 8. Układ do wykonania efektywnych pomiarów bez oddziaływania. Układ stanowi połączenie zestawu doświadczalnego do realizacji kwantowego efektu Zenona z oryginalnym układem Elitzura-Vaidmana. Opracowanie graficzne: Michał Pawlik

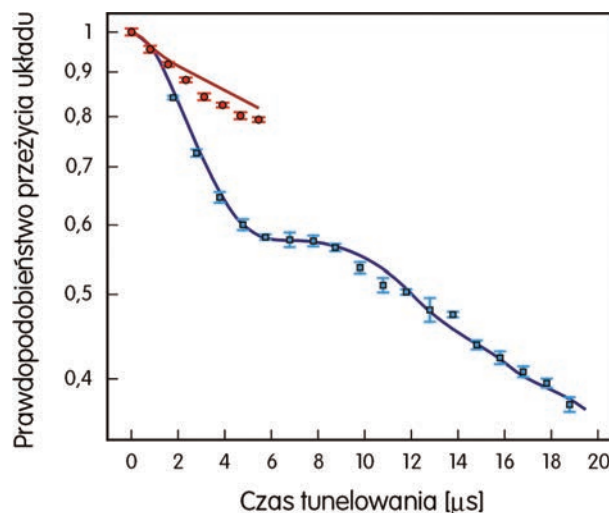
Wyobraźmy sobie, że do układu wchodzi foton spolaryzowany poziomo. Z jednej strony układu znajduje się urządzenie, które w każdym przebiegu fotonu obraca kierunek polaryzacji o 15° . Z drugiej strony został umieszczony interferometr polaryzacyjny, składający się z polaryzującej płytki światłodzielną i dwóch ramion równej długości ze zwierciadłami na końcach. Płytkę polaryzującą przepuszcza całkowicie światło spolaryzowane poziomo i całkowicie odbija światło spolaryzowane pionowo. Jeżeli w układzie nie ma innych elementów (przeszkoda), to światło jest dzielone przez płytkę polaryzującą w zależności od stanu polaryzacji, po czym odbija się od zwierciadeł umieszczonych w ramionach interferometru i wraca ponownie łącząc się w płytce światłodzielną. To tworzy się jeden cykl przebiegu fotonu. W rezultacie foton po każdym cyklu znajduje się dokładnie w tym samym stanie, w jakim był przed wejściem do interferometru, tzn. ma obrócony kierunek polaryzacji o 15° w stronę pionu. Po szóstym cyklu kierunek jego polaryzacji zmienia się z poziomego na pionowy. Sytuacja zmienia się po ustawieniu obiektu w tym ramieniu interferometru, w którym rozchodzi się światło spolaryzowane pionowo. Jest to analogiczna sytuacja jak przy wstawieniu sześciu polaryzatorów poziomych w doświadczeniu z kwantowym efektem Zenona. Po pierwszym cyklu prawdopodobieństwo, że foton mający kierunek polaryzacji

obrócony o 15° w stosunku do poziomej, wejdzie na drogę dozwoloną dla polaryzacji pionowej i zostanie zaabsorbowany wynosi 6,7%. Jeżeli absorpcja nie zachodzi, to znaczy, że foton wszedł na drogę przeznaczoną dla polaryzacji poziomej i wtedy jego polaryzacja zostanie ustawiona w kierunku poziomym. Ostatecznie po sześciu cyklach dolne zwierciadło nie zostanie wyłączone i foton opuści układ doświadczalny. Mierzając jego polaryzację okazuje się, że jest ona nadal pozioma, co oznacza, że w układzie znajduje się przeszkoda. W przeciwnym razie foton miałby polaryzację pionową.

Obserwacja kwantowego efektu anty-Zenona

Po raz pierwszy zaobserwowano jednocześnie efekt Zenona i efekt anty-Zenona w niestabilnym układzie kwantowym w 2001 roku. Dokonał tego zespół pracujący na Uniwersytecie Teksaskim w Austin (M. C. Fischer, B. Gutiérrez-Medina, M. C. Raizen). Zimne atomy sodu zostały schwytane w pułapkę magnetyczno-optyczną utworzoną przez przeciwbieżne liniowo spolaryzowane wiązki laserowe, które tworzyły falę stojącą. Dzięki oddziaływaniu atomów z tymi wiązkami zwiększało się prawdopodobieństwo tunelowania atomów z pułapki. Na początku doświadczenia liczba schwytanych atomów miała rozkład opisany niewykładniczą funkcją wygaszania, a później funkcja ta przyjmowała postać wykładniczą. Autorzy doświadczenia powtarzali pomiary liczby atomów pozostających w pułapce podczas początkowego okresu wygaszania niewykładniczego. W zależności od częstości pomiarów zaobserwowano wygaszanie mniejsze lub większe w porównaniu z układem niezakłóconym pomiarami. Celem doświadczenia było badanie wpływu pomiarów na szybkość wygaszania układu. Wielkością mierzoną była liczba atomów pozostających w pułapce. Ten pomiar mógł być dokonywany przez nagłe przerwanie tunelowania trwające przez okres 50 μ s. Po początkowym okresie powolnego wygaszania krzywa wygaszania stopniowo zanikała w sposób oscylujący. Ten zanik po pewnym czasie przechodził w wygaszanie wykładnicze. Jeżeli tunelowanie zostało przerwane to zaraz po okresie stopniowego spadku krzywej wygaszanie przebiegało szybciej niż w przypadku braku przerw. Tunelowanie zostało przerywane przez wyłączenie oddziaływania atomów z wiązką laserową. Jest to właśnie efekt przeciwny do efektu Zenona i stanowi istotę efektu anty-Zenona. Tak jak w przypadku efektu Zenona, przerwanie tunelowania powodują przejście do początkowego niewykładniczego wygaszania układu po każdym

pomiarze. Tutaj jednak okresy tunelowania między pomiarami są dłuższe niż w poprzednim przypadku. Na ryc. 9 przedstawiliśmy prawdopodobieństwo przeżycia układu (tzn. pozostania w początkowym stanie), jako funkcja czasu trwania tunelowania.



Ryc. 9. Prawdopodobieństwo przeżycia układu (tzn. pozostania w początkowym stanie) jako funkcja czasu trwania tunelowania (na podstawie publ. Fischer, Gutiérrez-Medina, Raizen, 2001 – czasopismo Phys. Rev. Lett.). Opracowanie graficzne: Michał Pawlik

Kwadraciki na ryc. 9 oznaczają punkty krzywej zanikania, gdy nie zakłócamy układu pomiarami. Kółka dotyczą sytuacji, gdy po każdym okresie tunelowania w ciągu 1 μ s następuje okres przerw trwający 50 μ s. Linie ciągłe są kwantowo-mechanicznymi symulacjami eksperymentów uzyskanymi przez numeryczne rozwiązanie równania Schrödingera. Słupki błędów oznaczają odchylenie standardowe estymatora wartości średnich. Prawdopodobieństwo przeżycia jasno pokazuje powolniejsze wygaszanie niż byłoby to w przypadku układu mierzonego bez przerywania tunelowania. Z ryc. 9 widzimy, że prawdopodobieństwo przeżycia układu maleje ze wzrostem czasu, z zachowaniem przedziału poziomego przy wartości około 7 μ s.

Przedstawiliśmy tutaj paradoksalne z punktu widzenia mechaniki kwantowej zjawiska Zenona i anty-Zenona. Są one przykładem sytuacji, gdy obserwacja układu kwantowego lub ingerencja obserwatora w ten układ wpływa na zachowanie się układu. Efekty te mogą znaleźć zastosowanie przy budowie komputerów kwantowych.

Wnioski

Przeprowadzone doświadczenia wykazały nieprawdziwość poglądu, że nie da się dokonać żadnej obserwacji bez udziału co najmniej jednego fotonu padającego na obserwowany przedmiot. Taki pogląd

wyraził w 1962 roku Dennis Gabor, odkrywca holografii (laureat Nagrody Nobla w dziedzinie fizyki w 1971 roku). Przeprowadzone doświadczenia mają wartość poznawczą i służą do weryfikacji teorii fizycznych, jednak przewiduje się, że są one także punktem wyjścia do rozwiązania wielu problemów praktycznych. Metody pomiaru bez oddziaływania mogą znaleźć zastosowanie jako sposób fotografowania, dzięki któremu uzyskuje się obraz obiektu bez wystawienia go na działanie światła. Metody te da się zastosować również w przypadku obiektu półprzezroczystego. Zastosowanie takiego sposobu fotografowania może być pożyteczne np. w medycynie przy uzyskiwaniu obrazów żywych komórek. Innym

możliwym zastosowaniem jest możliwość fotografowania chmury ultrazimnych atomów przechodzących w stan materii zwany kondensatem Bosego-Einsteina (za odkrycie kondensatu Bosego-Einsteina O. Cornell, C.E. Wieman, W. Ketterle otrzymali w 2001 roku Nagrodę Nobla w dziedzinie fizyki). W tym stanie atomy poruszają się tak powoli, że nawet pojedynczy foton może wybić je na zewnątrz niszcząc w ten sposób chmurę kondensatu. Pomiary bez oddziaływania mogą okazać się jednym sposobem sfotografowania takiego stanu atomów. Również wiąże się nadzieję z przeprowadzonymi doświadczeniami przy budowie komputerów kwantowych.

■ Dr Paweł Tomasz Pęczkowski jest pracownikiem Uniwersytetu Warszawskiego.

ZAPIS DYNAMIKI PRZEPŁYWU WODY I TRANSPORTU RUMOWISKA W CECHACH TEKSTURALNYCH ŻWIROWYCH OSADÓW KORYTOWYCH

Bartłomiej Wyżga (Kraków)

Wstęp

Cechy teksturalne osadów korytowych rzek żwirodennych mogą stanowić cenne źródło informacji o dynamice przepływu wody i transportu rumowiska. Właściwe rozpoznanie tej dynamiki jest podstawą paleogeograficznych rekonstrukcji systemów rzecznych i może umożliwiać odtwarzanie warunków środowiskowych w zlewniach. W przypadku rzek, których dno utworzone jest z materiału piaszczystego, zmiany dostawy rumowiska oraz natężenia jego transportu w rzece, wywołane zmianami środowiskowymi w zlewni, nie będą powodować istotnych zmian cech teksturalnych osadów korytowych. Z odmienną sytuacją mamy natomiast do czynienia w rzekach żwirodennych, gdzie zmiany udziału poszczególnych frakcji osadów korytowych, wynikające ze zmian dostawy i natężenia transportu rumowiska w rzece, będą znajdować wyraźne odzwierciedlenie w wysortowaniu i upakowaniu żwirów.

Cechy teksturalne żwirów w ciekach z różnych stref morfoklimatycznych

Można wymienić trzy źródła informacji wskazujących na możliwość wykorzystania cech teksturalnych żwirów rzecznych do oceny dynamiki przepływu wody i transportu rumowiska w dawnych systemach rzecznych oraz do interpretacji paleogeograficznych uwarunkowań zmian tej dynamiki. Pierwszym z nich są obserwacje wskazujące na różnice tych cech pomiędzy

osadami korytowymi wyścielającymi dno cieków w różnych strefach klimatycznych. Dobitnie ukazuje to porównanie osadów korytowych formowanych w stałych ciekach ze strefy umiarkowanej wilgotnej oraz w epizodycznych ciekach ze strefy półsuchej.

Większość stałych cieków żwirodennych charakteryzuje się obecnością powierzchniowej, gruboziarnistej warstwy bruku korytowego przykrywającej drobniejszy materiał denny (ryc. 1). Miąższość warstwy bruku korytowego odpowiada zazwyczaj 1–2 średnicom otoczków tworzących szkielet ziarnowy żwirów w podpowierzchniowej warstwie materiału dennego. Bruki takie powstają w wyniku wymywania drobniejszych ziaren z przypowierzchniowej warstwy osadu oraz mniejszej częstotliwości uruchamiania i wolniejszego transportu ziaren grubszych. Ponadto, formowane w tych ciekach żwiry cechuje zazwyczaj ciasne upakowanie i dachówkowate ułożenie (imbrykacja) otoczków tworzących szkielet ziarnowy. Widoczna na histogramach uziarnienia wyraźna przewaga frakcji żwirowej nad piaszczystą, a niekiedy wręcz unimodalny charakter osadu odzwierciedla stosunkowo dobre wysortowanie tych żwirów. Cieki okresowe i epizodyczne strefy półsuchej i suchej cechuje natomiast brak lub słaby rozwój bruków korytowych (ryc. 2). Formowane tu żwiry są zazwyczaj luźno upakowane, a otoczki tworzące szkielet ziarnowy tych osadów są rozproszone w piaszczystej masie wypełniającej. Żwiry takie wykazują bimodalny rozkład uziarnienia i bardzo złe