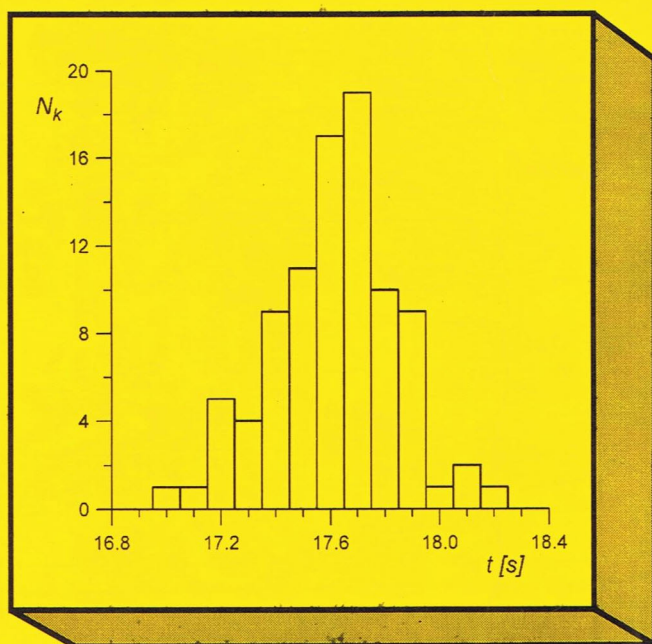


Andrzej Bielski Roman Ciuryło

PODSTAWY METOD OPRACOWANIA POMIARÓW



TORUŃ 1998

Andrzej Bielski Roman Ciuryło

Podstawy metod opracowania pomiarów

Wykład dla początkujących

Toruń 1998

Recenzenci
Bogusław Kamys, Henryk Szydłowski

ISBN 83-231-0910-9

Printed in Poland
© Copyright by Wydawnictwo Uniwersytetu Mikołaja Kopernika
Toruń 1998



M. 2718

Redaktor
Mirosława Szprenglewska

Uniwersytet Mikołaja Kopernika w Toruniu
87-100 Toruń, ul. Gagarina 11, fax (0-56) 654-29-48
Dział Wydawnictw tel. 62-14-295, kolportaż tel. 62-14-238
Wydanie I. Nakład 1000, objętość 9.0 ark. wyd.
Skład i łamanie: Jacek Jurkowski
Druk: Zakład Poligrafii UMK

Spis treści

Przedmowa	6
1. Wprowadzenie	9
1.1. Pojęcia podstawowe. Cel i zadania teorii błędów	9
1.1.1. Pojęcie pomiaru, pomiary bezpośrednie i pomiary pośrednie	9
1.1.2. Skończona dokładność pomiarów, błąd pomiaru . . .	10
1.1.3. Błąd bezwzględny i błąd względny	12
1.1.4. Przedstawianie błędów pomiarowych i zaokrąglanie wyników	13
1.1.5. Podział błędów	16
1.1.6. Zgodność wyników pomiarów	21
1.2. Znaczenie poznania niepewności pomiarowych	22
2. Ocena błędu maksymalnego	24
2.1. Oszacowanie niepewności przy odczycie skali, błąd maksymalny	24
2.2. Oszacowanie błędu maksymalnego w pomiarach pośrednich, metoda różniczki zupełnej	27
3. Wielkości charakteryzujące serię pomiarów obarczonych błędami przypadkowymi	35
3.1. Wartość średnia serii pomiarów bezpośrednich	35
3.2. Odchylenie standardowe serii pomiarów bezpośrednich . . .	40
3.3. Wartość średnia serii pomiarów pośrednich	43
3.3.1. Pomiary niezależne	43
3.3.2. Pomiary zależne	45
3.4. Odchylenie standardowe serii pomiarów pośrednich	46
3.4.1. Pomiary niezależne	46
3.4.2. Pomiary zależne	49
3.5. Metoda najmniejszych kwadratów	54
3.5.1. Dopasowanie funkcji liniowych	58
3.5.2. Dopasowanie innych krzywych	65
3.5.3. Aproksymacja metodą najmniejszych kwadratów . .	67
4. Histogramy i rozkłady, zmienna losowa	68
4.1. Histogram i rozkład zmiennej losowej skokowej	69

4.2. Histogram i rozkład zmiennej losowej ciągłej	75
5. Rozkład Gaussa i jego zastosowania	88
5.1. Rozkład Gaussa	88
5.1.1. Prawo błędów Gaussa	89
5.1.2. Rozkład Gaussa dla wyników pomiarów	91
5.1.3. Średnia arytmetyczna i wariancja serii jako najlepsze przybliżenia wartości rzeczywistej i wariancji rozkładu	95
5.1.4. Funkcja rozkładu średniej arytmetycznej	98
5.1.5. Pomiary o niejednakowej precyzji, łączenie niezależnych pomiarów, średnia ważona	100
5.1.6. Funkcja rozkładu wielkości złożonej	103
5.2. Rozkład normalny standaryzowany i jego zastosowanie do oceny błędu pomiaru	108
5.2.1. Standaryzacja rozkładu normalnego	108
5.2.2. Dystrybuenta rozkładu normalnego standaryzowanego i jej zastosowania	111
5.2.3. Ocena poziomu i przedziału ufności w przypadku pomiarów pośrednich	116
5.3. Centralne twierdzenie graniczne	120
5.4. Metoda najmniejszych kwadratów a rozkład Gaussa	123
5.4.1. Metoda najmniejszych kwadratów w przypadku pomiarów o niejednakowej precyzji	126
6. Rozkład <i>t-Studenta</i> i jego zastosowanie	128
6.1. Rozkład <i>t-Studenta</i>	128
6.2. Dystrybuenta rozkładu <i>t-Studenta</i> i jej zastosowanie	129
6.2.1. Ocena poziomu i przedziału ufności w przypadku pomiarów pośrednich	131
7. Rozkład dwumianowy (Bernoulliego) i rozkład Poissona	135
7.1. Rozkład dwumianowy	135
7.2. Rozkład Poissona i jego zastosowanie	139
7.2.1. Rozkład Poissona — zliczanie zdarzeń	139
7.2.2. Dystrybuenta rozkładu Poissona	143
8. Przedstawianie danych i graficzne oszacowanie błędu	145
8.1. Zasady zapisu wielkości fizycznych i ich jednostek	145
8.2. Tabelaryczne przedstawianie danych	146
8.2.1. Zasady sporządzania tabel	147
8.2.2. Rodzaje tabel	148
8.3. Graficzne przedstawianie danych	148

8.3.1. Najczęściej używane papiery funkcyjne	151
8.3.2. Zasady sporządzania wykresów	152
8.4. Graficzne oszacowanie błędu	156
8.4.1. Oszacowanie błędu wykreślonej krzywej	156
8.4.2. Oszacowanie błędu wartości odczytanej z wykresu	158
9. Ocena błędu w przypadku, gdy błędy przypadkowe i systematyczne są porównywalne	161
9.1. Ocena błędu maksymalnego	161
9.1.1. Pomiary bezpośrednie	161
9.1.2. Pomiary pośrednie	162
9.2. Statystyczna ocena błędu	163
9.2.1. Pomiary bezpośrednie	163
9.2.2. Pomiary pośrednie	168
Dodatki	170
A. Prawdopodobieństwo	170
B. Klasy mierników elektrycznych	177
C. Kowariancja wielkości złożonych	178
D. Wyprowadzenie wzoru (3.5.19)	180
E. Model błędów przypadkowych Laplace'a i uzasadnienie rozkładu normalnego	185
F. Obliczenie niektórych całek	189
G. Funkcja rozkładu średniej arytmetycznej	192
H. Tablice statystyczne	197
Literatura	205
Skorowidz	207

Przedmowa

Opracowanie to powstało na podstawie wykładów dla studentów I roku Wydziału Fizyki i Astronomii Uniwersytetu Mikołaja Kopernika w Toruniu. Stanowi ono, zdaniem autorów, minimalny zbiór wiadomości z teorii błędów pozwalający na poprawne opracowanie wyników pomiarów na poziomie pracowni obowiązujących na pierwszych trzech latach studiów.

Biorąc pod uwagę bardzo zróżnicowany poziom przygotowania z fizyki i matematyki studentów I roku (spowodowany różnorodnością wymiaru godzin i programów nauczania), wykład nasz staraliśmy się przygotować tak, aby stanowił on wprowadzony do teorii błędów na poziomie elementarnym. Z tego powodu założyliśmy również, że czytelnik może być słabo przygotowany z rachunku prawdopodobieństwa i statystyki. W związku z tym przyjęliśmy rzadko spotykany układ materiału. Tak więc po przedstawieniu w rozdziałach 1 i 2.1 wstępnych wiadomości z teorii błędów, do zrozumienia których nie jest potrzebna znajomość wyższej matematyki, omówiliśmy ocenę błędu maksymalnego pomiaru pośredniego metodą różniczki zupełnej (rozdział 2.2). Następnie, zamiast omawiać funkcje rozkładu prawdopodobieństwa, jak to ma miejsce w większości opracowań, przedstawiliśmy (w rozdziale 3) wielkości statystyczne charakteryzujące serię pomiarową i metodę najmniejszych kwadratów. Naszym zdaniem ten nietypowy układ materiału ułatwi studentom zapoznanie się z wielkościami charakteryzującymi pomiary obarczone błędami przypadkowymi obliczanymi bezpośrednio na podstawie wyników pomiarów, bez wgłębiania się od samego początku wykładu w rozważania statystyczne. Pomimo, że w rozdziałach 1–3 została podana większość pojęć i metod używanych przy ocenie błędu, jednak zrozumienie ich znaczenia oraz umiejętność pełnego zastosowania wymaga zapoznania się z treścią dalszych rozdziałów. Dzięki takiemu układowi materiału omówienie statystycznej oceny błędu pomiaru rozpoczyna się, gdy studenci poznają podstawy rachunku całkowego. Niezbędne do jej zrozumienia elementy rachunku prawdopodobieństwa podano w Dodatku A. Z wyżej wspomnianych względów funkcja rozkładu prawdopodobieństwa i jej dystrybuenta zostały wprowadzone dopiero w rozdziale 4. Ponieważ funkcje rozkładu prawdopodobieństwa mają wiele zastosowań w różnych działach fizyki, ich ogólnym własnościom poświęcono cały rozdział 4. W rozdziałach 5 i 6 omówiona została ocena poziomu i przedziału ufności przy użyciu, kolejno, rozkładu Gaussa i rozkładu *t-Studenta*. W rozdziale 7 zostały przedstawione dwa spośród wielu rozkładów zmiennej losowej skokowej, a mianowicie rozkład dwumianowy (Bernoulliego) i rozkład Poissona, ponieważ są one szczególnie ważne przy analizie zliczania zdarzeń przy-

padkowych. Zasady tabelarycznego i graficznego przedstawiania wyników oraz graficzna ocena błędu zostały omówione w rozdziale 8. Na zakończenie (rozdział 9) przedstawiliśmy sposoby oceny błędu maksymalnego oraz poziomu i przedziału ufności, gdy błędy przypadkowe są porównywalne z dokładnością przyrządu pomiarowego.

Ponieważ większość pomiarów wykonywanych w laboratoriach fizycznych to pomiary pośrednie, dużo uwagi poświęciliśmy ocenie błędu maksymalnego oraz poziomu i przedziału ufności wyznaczonej wielkości złożonej. Rozdziały 1, 2.1 oraz 8 zostały napisane tak, że powinny być zrozumiałe dla uczniów szkół średnich. Przygotowując to opracowanie staraliśmy się ograniczyć do minimum liczbę wprowadzanych pojęć. Wyprowadzenia wzorów zostały podane możliwie szczegółowo. Część wyprowadzeń wzorów oraz niektóre twierdzenia, których umieszczenie w tekście powodowałoby rozbicie jednolitości wykładu, zostały zamieszczone w Dodatkach. W Dodatku H zostały podane tablice statystyczne konieczne do oceny poziomu i przedziału ufności. Przykłady pomiarów na ogół nie wykraczają poza program średniej szkoły ogólnokształcącej.

Ponieważ wykład ten poświęcony jest opracowaniu wyników pomiarów staraliśmy się unikać zagadnień wykraczających poza zakres teorii błędów i przedstawiania wyników pomiarów, o niektórych z nich jedynie wzmiankowaliśmy. Z tego powodu zostały całkowicie pominięte zagadnienia związane z kontrolą jakości i tolerancją parametrów wyrobów (np. oporników). Ze względu na podstawowy charakter wykładu, pominęliśmy również weryfikację hipotez oraz testy statystyczne (w tym test χ^2).

Używaliśmy zasadniczo terminu błąd pomiaru, a nie uchyb pomiaru czy niepewność pomiaru, ponieważ termin ten jest znacznie bardziej rozpowszechniony w literaturze i to nie tylko polskiej, oraz dlatego, że jest on historycznie najstarszy, pochodzi bowiem z XVIII wieku. Ze względu na będące w powszechnym użyciu kalkulatory i komputery miejsca dziesiętne oddzielano kropką, a nie przecinkiem.

Przygotowując nasz wykład szeroko korzystaliśmy z następujących opracowań:

- H. Hänsel, *Podstawy rachunku błędów*, WN-T, Warszawa 1968;
- W. I. Romanowski, *Podstawowe zagadnienia teorii błędów*, PWN, Warszawa 1955;
- G. L. Squires, *Praktyczna fizyka*, Wydawnictwo Naukowe PWN, Warszawa 1992;
- A. Strzałkowski, A. Śliżyński, *Matematyczne metody opracowania wyników pomiarów*, PWN, Warszawa 1978;
- J. R. Taylor, *Wstęp do analizy błędu pomiarowego*, Wydawnictwo Naukowe PWN, Warszawa 1995;

- *Teoria pomiarów*, praca zb. pod red. H. Szydlowskiego, PWN, Warszawa 1981.

Zaczerpnęliśmy z nich szereg idei i wyprowadzeń wzorów. Na końcu został zamieszczony wykaz zawierający cytowane w tekście opracowania. Opracowanie nasze zostało tak przygotowane, aby czytelnik nie musiał do nich sięgać.

Dodatkowym celem naszego opracowania jest przygotowanie czytelnika do studiowania i zastosowania w praktyce pomiarowej metod statystyki matematycznej. Sądzymy więc, że nasz wykład ułatwi studiowanie bardziej zaawansowanych opracowań, takich jak:

- S. Brandt, *Metody statystyczne i obliczeniowe analizy danych*, PWN, Warszawa 1976;
- W. T. Eadie, D. Drijard, F. E. James, M. Roos, B. Sadoulet, *Metody statystyczne w fizyce doświadczalnej*, PWN, Warszawa 1989;
- W. I. Romanowski, *Podstawowe zagadnienia teorii błędów*, PWN, Warszawa 1955;
- A. Strzałkowski, A. Śliżyński, *Matematyczne metody opracowania wyników pomiarów*, PWN, Warszawa 1973;
- *Teoria pomiarów*, praca zb. pod red. H. Szydlowskiego, PWN, Warszawa 1981;
- E. B. Wilson jr, *Wstęp do badań naukowych*, PWN, Warszawa 1968.

Dziękujemy wszystkim tym, którzy przyczynili się do powstania tego opracowania. Przede wszystkim składamy podziękowanie panom profesorom Franciszkowi Bylickiemu, Tadeuszowi Marszałkowi, Józefowi Szudemu, Jarosławowi Zarembie oraz dr Ryszardowi S. Trawińskiemu i mgr Jolancie Domysławskiej za przeczytanie maszynopisu i wiele cennych uwag.

Andrzej Bielski

Roman Ciuryło

1. Wprowadzenie

1.1. Pojęcia podstawowe. Cel i zadania teorii błędów

1.1.1. Pojęcie pomiaru, pomiary bezpośrednie i pomiary pośrednie

Zasadniczym celem nauk przyrodniczych jest opisanie, analizowanie i zrozumienie zjawisk zachodzących w otaczającym nas świecie, znajdowanie warunków, w których one zachodzą oraz wykrywanie związków między poszczególnymi zjawiskami. Poznanie cech tych zjawisk pozwala na ich usystematyzowanie, ustalenie praw rządzących nimi i praktyczne ich zastosowanie.

Mierzalną cechę zjawiska lub ciała nazywamy **wielkością fizyczną**.

Wielkości fizyczne dzielimy na wielkości podstawowe i pochodne. Wielkości przyjęte za podstawowe definiuje się za pomocą określonych wzorców np. w Międzynarodowym Układzie Jednostek (SI) za jednostki podstawowe przyjmuje się: 1 m jako jednostkę długości, 1 kg — masy, 1 s — czasu, 1 A — natężenia prądu, 1 K — temperatury termodynamicznej, 1 mol — ilości materii, 1 cd — światłości [1, 2]. Natomiast wielkości pochodne definiuje się za pomocą wielkości podstawowych. Ilościowo każdą wielkość fizyczną a wyrażamy jej miarą, która składa się z dwóch członów:

$$a = \{a\} [a], \quad (1.1.1)$$

tj. liczby $\{a\}$ zwanej wielkością miary i jednostki miary $[a]$, np. prędkość $v = 36 \text{ m/s}$, w naszym przykładzie $\{v\} = 36$, $[v] = \text{m/s}$. Przejście od opisu jakościowego do opisu ilościowego jakiegokolwiek zjawiska wymaga wykonania tzw. pomiarów.

Pomiarem nazywamy zespół czynności wykonanych w celu ustalenia miary określonej wielkości fizycznej.

Pomiar polega więc na doświadczalnym porównaniu mierzonej wielkości fizycznej z wzorcem tej wielkości, przyjętym umownie za jednostkę miary.

We wszystkich naukach przyrodniczych, technicznych, rolniczych, ekonomicznych, a także produkcji przemysłowej i rolniczej spotykamy się z

pomiarami. Po wykonaniu pomiaru najważniejszym zadaniem jest ocena i wykorzystanie otrzymanych wyników. Ocenę otrzymanych wyników umożliwia teoria błędów.

Pomiary dzielimy na dwie grupy, a mianowicie na:

- 1) **pomiary bezpośrednie**, tj. takie gdy miarę wielkości fizycznej otrzymujemy jako wynik bezpośredniego porównania mierzonej wielkości z wzorcem. Bezpośrednio możemy mierzyć: przedział czasu, w którym zachodzi zjawisko fizyczne, długość, natężenie prądu, kąt itp.
- 2) **pomiary pośrednie**, czyli pomiary **wielkości złożonej**; większość wielkości fizycznych jest funkcją wielkości mierzalnych bezpośrednio. Jeżeli dokonujemy pomiaru wielkości $z = f(x_1, x_2, \dots, x_r)$, gdzie wielkości $x_1, x_2, \dots, x_r$ są mierzone bezpośrednio to mówimy, że z jest wielkością złożoną, a jej pomiar jest pomiarem pośrednim. O wielkości z mówimy, że została *wyznaczona* na podstawie zmierzonych wielkości $x_1, x_2, \dots, x_r$.

Wielkością wyznaczoną będziemy nazywali wielkość określoną za pomocą obliczeń [3] wykonanych przy użyciu wielkości zmierzonych.

Wyjaśnimy to na następującym przykładzie. Wyznaczamy przyspieszenie ziemskie g za pomocą wahadła matematycznego. Jego okres dla małych wychyleń wyraża się wzorem:

$$T = 2\pi\sqrt{\frac{l}{g}}, \quad (1.1.2)$$

a więc:

$$g = 4\pi^2 \frac{l}{T^2}. \quad (1.1.3)$$

Bezpośrednio mierzymy długość wahadła l i jego okres T . Po zmierzeniu l i T obliczamy g . Wartość g nie jest więc zmierzona bezpośrednio, jest obliczona na podstawie *zmierzonych* uprzednio wartości l i T . Wielkość g została *wyznaczona* a nie zmierzona. Jest to pomiar *pośredni*, mówimy więc, że $g = f(l, T)$ jest *wielkością złożoną*.

1.1.2. Skończona dokładność pomiarów, błąd pomiaru

Jak wynika z codziennego doświadczenia, w wyniku pomiaru, niezależnie od tego jak starannie jest on wykonany, nigdy nie otrzymujemy dokładnej wartości mierzonej wielkości fizycznej. Inaczej mówiąc każdy pomiar

jest obarczony jakąś niepewnością, czyli błędem. Błędu tego nie możemy uniknąć.

Aby zrozumieć zjawisko nieuniknionego występowania błędów pomiarowych przyjrzyjmy się jakiemukolwiek z prostych pomiarów wykonywanych w życiu codziennym, np. niech to będzie pomiar długości stołu. Najprostszym sposobem jest oszacowanie „na oko”. Przyjmijmy, że wynosi ona np. 170 cm, nie możemy jednak zaprzeczyć, jeśli ktoś inny oszacuje ją na 180 cm. Niepewność oszacowania będzie bardzo duża. Jeśli do pomiaru zastosujemy przymiar liniowy lub taśmę mierniczą z podziałką co 0.5 cm, to wynik pomiaru będzie wynosił np. 178.5 cm. Jeśli pomiaru dokonamy przymiarem liniowym z podziałką milimetrową, to otrzymamy np. 178.7 cm. Pomiar ten jest dokładniejszy od poprzednich, ale nie możemy powiedzieć, że długość stołu wynosi dokładnie 178.700 cm, a nie 178.708 cm lub 178.692 cm. Jeżeli użyjemy dokładniejszej metody pomiaru, to otrzymany przez nas wynik będzie wynosił np. 178.68 cm. Ulepszając metodę pomiaru będziemy uzyskiwali coraz dokładniejsze wartości, ale zawsze dokładność uzyskanych wyników będzie ograniczona. Jeżeli do stołu przykładamy przymiar liniowy, to powstaje pytanie, czy krawędź stołu pokrywa się z zaznaczoną podziałką, jeśli nie, to trzeba oszacować, jaka część odległości między zaznaczonymi kreskami przypada na krawędź. Jeśli nawet pokryła się, to przecież kreska podziałki ma skończoną grubość, jak określić, w którym miejscu kreski podziałki znajduje się krawędź stołu? Powstaje także pytanie czy patrzymy na krawędź stołu i przymiar dokładnie prostopadłe, jeśli nie to odczytana wartość będzie obciążona pewnym błędem zależnym od kąta, pod którym patrzymy. Kąt ten będzie się zmieniał od pomiaru do pomiaru w sposób przypadkowy. Podobnych pytań jest więcej.

Na powyższym przykładzie widać wyraźnie, że błędów pomiaru nie możemy uniknąć, tzn. że nie możemy wykonać pomiaru z absolutną dokładnością. Możemy jedynie odpowiednio postępując zmniejszać błąd pomiaru.

Z tego co powiedziano wyżej i co potwierdza doświadczenie wynika, że jeżeli dany pomiar będziemy powtarzali wielokrotnie wykonując go w ten sam sposób, w tych samych warunkach i tym samym przyrządem pomiarowym, to otrzymamy zbiór różnych wartości mierzonej wielkości fizycznej, który nazwiemy **serią pomiarową**. Każdy z tych wyników będzie obciążony jakąś **niepewnością** czyli **błędem pomiaru**. W teorii błędów termin *błąd pomiaru* nie oznacza pomyłki (jak np. w ortografii) czy też gafy. Oznacza niemożliwą do uniknięcia *niepewność* pomiaru, wynikającą z samego zjawiska pomiaru. Terminy *błąd pomiaru* i *niepewność pomiaru* będą w dalszym ciągu wykładu używane wymiennie. Skoro nie można uniknąć błędu pomiaru powstaje pytanie, jak *oszacować* wartość rzeczywistą (nazywaną również prawdziwą lub dokładną) mierzonej wielkości fizycznej x .

To oszacowanie jest celem teorii błędów, tzn. jest nim podanie nierówności:

$$\boxed{\tilde{x} - \Delta \leq x_0 \leq \tilde{x} + \Delta,} \quad (1.1.4)$$

oraz prawdopodobieństwa, że nierówność (1.1.4) jest spełniona, czyli że **wartość rzeczywista** x_0 mieści się w pewnym przedziale. W nierówności (1.1.4) $\tilde{x}$ i Δ są wielkościami wynikającymi z pomiarów; sposoby obliczania ich są celem niniejszego wykładu.

W celu znalezienia nierówności (1.1.4) i oszacowania prawdopodobieństwa, że ona zachodzi, stosuje się statystyczne metody analizy danych. W przypadkach gdy to jest niemożliwe dokonuje się oceny tzw. błędu maksymalnego i w ten sposób znajduje się przedział, w którym jest zawarta wartość rzeczywista mierzonej wielkości fizycznej.

1.1.3. Błąd bezwzględny i błąd względny

Różnicę między **wartością zmierzona** x danej wielkości fizycznej a jej wartością rzeczywistą x_0 nazywamy **błędem bezwzględnym**, czyli:

$$\boxed{\delta = x - x_0,} \quad (1.1.5)$$

(w szeregu opracowań błąd bezwzględny jest definiowany jako: $\delta = x_0 - x$). Ze wzoru (1.1.5) wynika, że **błąd bezwzględny musi być zawsze wyrażony w tych samych jednostkach co wynik pomiaru**, czyli ma zawsze ten sam wymiar co wielkość mierzona. Jeżeli wykonujemy serię pomiarów, to wówczas błąd bezwzględny i -tego pomiaru jest równy $\delta_i = x_i - x_0$, a więc otrzymujemy nie tylko zbiór wartości x_i , ale również zbiór wartości błędów bezwzględnych δ_i . *Błędu bezwzględnego wyniku danego pomiaru nigdy nie znamy*, ponieważ nie znamy wartości rzeczywistej x_0 . Stosując metody teorii błędów możemy go jednak *oszacować*.

Błędem względnym zmierzonej wielkości nazywamy wartość bezwzględną stosunku błędu bezwzględnego do wartości rzeczywistej, czyli:

$$\boxed{\text{błąd względny} = \left| \frac{\delta}{x_0} \right|.} \quad (1.1.6)$$

Jest on wielkością *bezwymiarową*. Błąd względny wyrażony w procentach bywa nazywany **błędem procentowym**

$$\boxed{\text{błąd procentowy} = \left| \frac{\delta}{x_0} \right| 100\%.}$$

1.1.4. Przedstawianie błędów pomiarowych i zaokrąglanie wyników

Jeżeli otrzymana z pomiarów wartość jakiejś wielkości fizycznej jest równa $\tilde{x}$ i błąd pomiaru wynosi Δ , to wynik pomiaru przedstawiamy w postaci:

$$x = \tilde{x} \pm \Delta. \quad (1.1.7)$$

We wzorze (1.1.7) nie precyzujemy jak wyznaczamy $\tilde{x}$ i Δ , ponieważ będzie to przedstawione w następnych rozdziałach. Należy zaznaczyć, że wzór (1.1.7) jest równoważny, przy założeniach omówionych w następnych rozdziałach, nierówności (1.1.4).

Podanie samego błędu pomiaru mija się z celem. Zawsze bowiem musimy podać wartość wielkości fizycznej wynikającą z pomiarów ($\tilde{x}$) i oszacowany błąd (Δ). Czasami podaje się wartość wynikającą z pomiarów $\tilde{x}$ z komentarzem ile wynosi błąd względny lub procentowy, np. długość stołu wynosi 181.5 cm z dokładnością 0.1%. Błąd pomiaru Δ , niezależnie od tego w jaki sposób został oszacowany, podajemy zawsze jako liczbę *dodatnią*.

Otrzymaną z pomiarów wartość jakiejś wielkości fizycznej $\tilde{x}$, jak i jej błąd Δ podaje się najczęściej w postaci liczb dziesiętnych.

Ponieważ błąd pomiaru Δ jest wielkością oszacowaną, *nie ma sensu* podawać wszystkich cyfr, które otrzymujemy z obliczeń. Podobnie *nie ma sensu* podawanie wszystkich cyfr, jakie otrzymujemy z obliczeń prowadzących do otrzymania wartości $\tilde{x}$. Dlatego obliczone wartości $\tilde{x}$ i Δ podajemy zaokrąglone. Oznacza to, że przybliżamy wartości otrzymane z obliczeń. Przy ocenie błędu pomiaru oraz w obliczeniach przybliżonych stosuje się dwa terminy pozwalające określić do ilu cyfr należy zaokrąglić $\tilde{x}$ i Δ . Są to cyfra znacząca i cyfra pewna.

Cyframi znaczącymi danej liczby różnej od zera nazywamy wszystkie jej cyfry z wyjątkiem występujących na początku zer.

Do cyfr znaczących zalicza się również zera końcowe, jeśli są one wynikiem rachunków a nie zaokrągleń. Z definicji tej wynika, że pierwsza cyfra znacząca zawsze musi być różna od zera, natomiast druga, trzecia i dalsze mogą być zerami. Na przykład liczba 0.00102 ma trzy cyfry znaczące: 1, 0 i 2. W przypadku liczby 100/3 mającej w zapisie dziesiętnym postać 33.33(3), wszystkie jej cyfry są znaczące i są trójkami.

Cyfrы pewne określamy następująco. Jeżeli błąd spowodowany przybliżeniem liczby dziesiętnej jest mniejszy od jedności na ostatnim miejscu dziesiętnym, to mówimy, że wszystkie jej cyfry są pewne.

Z określenia tego wynika, że np. jeśli w przypadku liczby sześciocyfrowej błąd trzeciej cyfry jest mniejszy od jedności, to mamy trzy cyfry pewne.

Przybliżenie dziesiętne podaje się wtedy z zachowaniem *tylko* cyfr pewnych, np. liczbę 125000 z błędem 100 zapiszemy: 125×10^3 lub 1.25×10^5 .

Przy **zaokrąglaniu wyniku pomiaru** stosujemy powszechnie przyjęte zasady zaokrągleń, tj. liczbę kończącą się cyframi:

- 0 – 4 zaokrąglamy w dół, a 5 – 9 w górę;
lub
- 0 – 4 zaokrąglamy w dół, 6 – 9 w górę, a cyfrę 5 w dół jeśli poprzedza ją liczba parzysta, zaś w górę, gdy poprzedza ją liczba nieparzysta.

Można stosować dowolną z tych zasad, ale w jednym opracowaniu wyników pomiarów należy konsekwentnie stosować tylko jedną z nich. W dalszym ciągu naszego wykładu będziemy stosowali pierwszą z nich.

Oszacowane błędy zaokrąglamy zawsze w górę. Zaokrąglenie w górę jest podyktowane tym, że w żadnym przypadku nie wolno zmniejszać błędów. Zawsze bowiem lepiej jest podać zawyżoną wartość błędu niż go niedoszacować.

Obliczenia $\tilde{x}$ jak i Δ wykonujemy zawsze z większą liczbą cyfr niż chcemy podać wynik; **zaokrąglenie dokonujemy dopiero po zakończeniu obliczeń.** Obecnie dzięki szerokiemu zastosowaniu kalkulatorów obliczenia wykonujemy wykorzystując pełną dostępną liczbę cyfr. Na przykład wyznaczyliśmy przyspieszenie ziemskie g za pomocą wahadła matematycznego i otrzymaliśmy $\tilde{g} = 981.3456 \text{ cm/s}^2$ oraz $\Delta = 3.0579102 \text{ cm/s}^2$, czyli $g = (981.3456 \pm 3.0579102) \text{ cm/s}^2$. Zarówno podanie $\tilde{g}$ jak i Δ z tak dużą liczbą cyfr nie ma sensu, ponieważ wszystkie cyfry znaczące $\tilde{g}$, poczynając od trzeciej, leżą w granicach błędu Δ . Powstaje więc pytanie jak zaokrąglić $\tilde{g}$ oraz Δ .

Przyjęto regułę, że **błędy pomiarów zaokrąglane są do pierwszej cyfry znaczącej** oraz że ostatnia cyfra znacząca w każdym wyniku pomiaru powinna stać na tym samym miejscu dziesiętnym co błąd pomiaru.

Odstępstwo od powyższej reguły stosujemy gdy zaokrąglenie niepewności powoduje jej wzrost więcej niż o 10%: wtedy błąd pomiaru zaokrąglamy do dwóch cyfr, np. błąd $\Delta = 2.025$ zaokrąglimy nie do 3 ale do 2.1 (w niektórych opracowaniach przyjęto zasadę, że gdy pierwsza cyfra niepewności jest 1 lub 2 to w zapisie błędu podajemy dwie cyfry, por. np. [4]). Tak więc wyznaczoną wartość przyspieszenia ziemskiego g zapiszemy $g = (981.3 \pm 3.1) \text{ cm/s}^2$. Gdyby błąd Δ wynosił np. 3.8542 cm/s^2 to wyznaczoną wartość g zapisalibyśmy $g = (981 \pm 4) \text{ cm/s}^2$. Otrzymana z takim błędem wartość g ma, jak łatwo zauważyć, tylko dwie cyfry pewne.

Innym często stosowanym sposobem przedstawiania niepewności pomiarowych jest podanie ich w nawiasach bezpośrednio po wyniku, np. wysokość $h = (1260 \pm 30) \text{ cm}$ zapiszemy jako $h = 1260(30) \text{ cm}$.

Tabela 1.1: Obowiązujące przedrostki dla jednostek wielokrotnych i podwielokrotnych

Przedrostek	Oznaczenie	Wielokrotność i podwielokrotność	
jotta	Y	10^{24}	= 1 000 000 000 000 000 000 000 000
zetta	Z	10^{21}	= 1 000 000 000 000 000 000 000 000
eksa	E	10^{18}	= 1 000 000 000 000 000 000 000
peta	P	10^{15}	= 1 000 000 000 000 000 000
tera	T	10^{12}	= 1 000 000 000 000 000
giga	G	10^9	= 1 000 000 000
mega	M	10^6	= 1 000 000
kilo	k	10^3	= 1 000
hekto	h	10^2	= 100
deka	da	10^1	= 10
–	–	10^0	= 1
decy	d	10^{-1}	= 0.1
centy	c	10^{-2}	= 0.01
mili	m	10^{-3}	= 0.001
mikro	μ	10^{-6}	= 0.000 001
nano	n	10^{-9}	= 0.000 000 001
piko	p	10^{-12}	= 0.000 000 000 001
femto	f	10^{-15}	= 0.000 000 000 000 001
atto	a	10^{-18}	= 0.000 000 000 000 000 001
zepto	z	10^{-21}	= 0.000 000 000 000 000 000 001
jokto	y	10^{-24}	= 0.000 000 000 000 000 000 000 001

W niektórych przypadkach, zwłaszcza przy wyznaczaniu wielkości o podstawowym znaczeniu (np. stałych uniwersalnych), podajemy *dwie pierwsze cyfry niepewności*, z tym, że ostatnia cyfra znacząca wyniku powinna stać na tym samym miejscu dziesiętnym co druga cyfra niepewności, np. stała Plancka $h = 6.6260755(40) \times 10^{-34}$ J·s [5].

Należy zaznaczyć, że jeżeli wielkość mierzona lub wyznaczona nie jest bezwymiarowa to oszacowany błąd musi mieć ten sam wymiar co mierzona wielkość i musi być wyrażony w tych samych jednostkach. Poprawny zapis ma więc postać:

$$x = \{\tilde{x} \pm \Delta\}[x], \quad (1.1.8)$$

tzn. w nawiasie okrągłym podajemy wartość liczbową wyniku pomiaru $\pm$ oszacowany błąd, za nawiasem podajemy jednostkę miary, w jakiej są wyrażone obie te wielkości, np. $v = (36 \pm 1)$ m/s.

Często zmierzoną lub wyznaczoną wartość $\tilde{x}$ wyrażamy jako liczbę a mnożoną przez 10^k , wówczas błąd Δ należy przedstawić w ten sam sposób,

tj. $\Delta = d \times 10^k$. Prawidłowy zapis będzie miał postać:

$$x = \{a \pm d\} \times 10^k [x], \quad (1.1.9)$$

np. opór $R = (2.25 \pm 0.02) \text{ M}\Omega$ lub $R = (2.25 \pm 0.02) \times 10^6 \Omega$, *niewłaściwy zaś będzie np. zapis* $R = (2.25 \times 10^6 \pm 2 \times 10^4) \Omega$. Jednostki miary, w jakich wyrażamy wynik pomiaru, jak i błąd nie muszą być jednostkami podstawowymi. Przedstawiając wyniki i ich błędy można korzystać z jednostek wielokrotnych, jak i podwielokrotnych. Obowiązujące przedrostki dla jednostek wielokrotnych i podwielokrotnych zamieszczono w tabeli 1.1.

Należy zaznaczyć, że przedrostki jednostek wielokrotnych kilo, mega, giga, tera używane w informatyce mają trochę inne znaczenie, np. mówimy że wielkość pamięci wynosi 1 kilobajt (1kb), co oznacza, że wynosi ona 2^{10}b , czyli 1024b. W tabeli 1.2 podano dokładne znaczenia tych przedrostków.

Tabela 1.2: Przedrostki dla jednostek wielokrotnych będących miarą informacji

Przedrostek	Oznaczenie	Wielokrotność
tera	T	$2^{40} = 1\,099\,511\,627\,776$
giga	G	$2^{30} = 1\,073\,741\,824$
mega	M	$2^{20} = 1\,048\,576$
kilo	k	$2^{10} = 1\,024$

Na zakończenie chcemy zaznaczyć, że tablice matematyczno-fizyczne zawierają wielkości liczbowe, których przedostatnia cyfra jest cyfrą pewną.

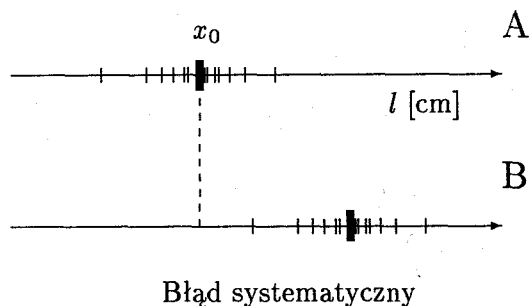
1.1.5. Podział błędów

Błędy czyli niepewności pomiarowe, ze względu na sposób w jaki wpływają na wyniki pomiarów, dzielimy na:

- błędy przypadkowe,
- błędy systematyczne,
- błędy grube.

1) Błędy przypadkowe

Jak pokazano w rozdz. 1.1.2, każdy pomiar ma skończoną dokładność. Jeżeli dany pomiar będziemy powtarzali wielokrotnie, to otrzymamy, na ogół, pewien zbiór wartości wielkości mierzonej. Inaczej mówiąc, w przypadku gdy błędy systematyczne (omówione poniżej) można pominąć, otrzymamy pewien rozrzut wyników pomiaru wokół wartości rzeczywistej x_0 . Na rysunku 1.1A przedstawiono (przy założeniu, że błędy systematyczne można pominąć) przykład takiego rozrzutu wyników pomiaru długości



Rys. 1.1: Przykład wyników pomiaru długości pręta. A) w przypadku gdy błąd systematyczny nie występuje, B) w przypadku występowania błędu systematycznego, x_0 — wartość rzeczywista

pręta. Na rysunku tym wartość rzeczywistą x_0 , wokół której rozkładają się wyniki pomiarów zaznaczono pogrubioną kreską. Powstaje więc pytanie: jakie są przyczyny rozrzutu wyników pomiaru i czy można tego uniknąć? Zanim odpowiemy na to pytanie przeanalizujemy kilka prostych przykładów.

Jeżeli wykonujemy za pomocą stopera serię pomiarów np. okresu drgań wahadła, wówczas kolejne momenty włączenia stopera, podobnie jak kolejne momenty jego wyłączenia, będą różne. Raz włączymy stoper trochę za wcześnie, a innym razem nieco za późno. Podobnie będzie z wyłączeniem stopera. Z tego powodu *zmierzony* okres drgań wahadła będzie się różnił od pomiaru do pomiaru. Występujące różnice będą jednak *przypadkowe*, czyli że błąd pomiaru (tj. różnica między wielkością rzeczywistą i zmierzoną) będzie miał również charakter *przypadkowy*. Przedstawiliśmy przykład błędu związanego z obserwatorem, tj. osobą wykonującą pomiary. Błędy związane z obserwatorem są spowodowane niedoskonałościami naszych zmysłów.

Zmieniające się w czasie pomiaru w sposób przypadkowy warunki otoczenia (np. drgania budynku, ruchy powietrza w pomieszczeniu, w którym wykonywany jest pomiar) stanowią dodatkowe źródło błędów. Mogą one zmieniać wskazania przyrządów w sposób *przypadkowy*.

W trakcie wykonywania pomiarów mogą występować przypadkowe zmiany (wahania, fluktuacje) wskazań aparatury pomiarowej. Również i one spowodują występowanie *przypadkowego* rozrzutu wyników.

Przedstawione przykłady ilustrują nam rolę czynników, które w sposób *przypadkowy* wpływają na wynik pomiaru i powodują również *przypadkowy* rozrzut wyników pomiarów. Nie są to oczywiście wszystkie możliwe czynniki i nie zawsze muszą występować równocześnie.

Wykonując pomiary nie możemy uniknąć ich wpływu i rozrzut wyników będzie występował zawsze, gdy pomiar wykonujemy wielokrotnie. Ten rozrzut wyników jest niezależny od woli obserwatora. Powstające w ten

sposób błędy noszą nazwę **błędów** lub **niepewności przypadkowych** i są nierozzerwalnie związane z samym zjawiskiem pomiaru.

Inaczej mówiąc: **błędami przypadkowymi** będziemy nazywali niepewności pomiarowe, które, jeżeli pomiar wykonujemy wielokrotnie, ujawniają się w postaci rozrzutu wyników.

Jeśli występują jedynie błędy przypadkowe, zaś błędy systematyczne można pominąć, ocenę błędu pomiaru wykonujemy metodami statystycznymi.

2) Błędy systematyczne

Każdy przyrząd pomiarowy mierzy ze skończoną dokładnością. Dokładność pomiaru, wykonanego danym przyrządem, nie może być większa niż dokładność samego przyrządu pomiarowego (o dokładności przyrządu patrz rozdział 2.1). Na przykład jeżeli śruba mikrometryczna ma dokładność 0.01 mm, to wynik *każdego* pomiaru wykonanego przy użyciu tej śruby będzie miał dokładność 0.01 mm. Jeżeli więc tą śrubą zmierzmy średnicę drutu d i otrzymamy $d = 1.25$ mm, to dokładność pomiaru oczywiście będzie 0.01 mm, co zapiszemy (1.25 ± 0.01) mm. Błąd ten będzie występował systematycznie we wszystkich pomiarach wykonanych tą śrubą. W tym przypadku mamy do czynienia z błędem spowodowanym **dokładnością przyrządu**.

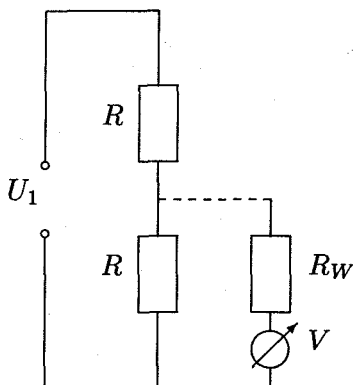
Inny rodzaj błędu, który może wystąpić we wszystkich pomiarach, jest związany z *wybozem metody pomiaru*. Na przykład wyznaczając ciepło właściwe metodą ostygania musimy uwzględnić fakt, że metoda ta daje poprawne wyniki, gdy badane ciało jest dobrym przewodnikiem ciepła. Jeżeli warunek ten nie jest spełniony, otrzymany wynik będzie obciążony błędem. Dla złych przewodników ciepła musimy stosować inne metody pomiaru.

Kolejnym źródłem błędu występującego we wszystkich pomiarach, może być *stosowanie niewłaściwego dla danego pomiaru przyrządu pomiarowego*. Aby to wyjaśnić przeanalizujmy następujący przykład:

Mamy układ dwóch jednakowych oporników, każdy o oporze R , połączonych szeregowo i zasilanych ze źródła o napięciu U_1 i o bardzo małym oporze wewnętrznym (rysunek 1.2).

Woltomierzem o oporze wewnętrznym R_W chcemy zmierzyć spadek napięcia U na oporze R . Wykonanie prostych obliczeń pokazuje, że $U = 0.5U_1$. Jeżeli jednak połączymy z końcami opornika woltomierz, to zamiast oporu R otrzymamy układ dwóch oporników połączonych równolegle o oporze wypadkowym R_Z :

$$R_Z = \frac{R_W R}{R_W + R} \quad (1.1.10)$$



Rys. 1.2: Schemat obwodu elektrycznego do pomiaru napięcia na oporze R . R_W — opór wewnętrzny woltomierza, V — woltomierz. Linia przerywaną (---) oznaczono podłączenie woltomierza

Spadek napięcia na oporze R_Z , jak łatwo pokazać korzystając z praw Kirchhoffa, jest równy:

$$U = U_1 \frac{R_Z}{R_Z + R}. \quad (1.1.11)$$

Podstawiając do (1.1.11) R_Z dane wzorem (1.1.10) otrzymujemy:

$$U = U_1 \frac{R_W}{R + 2R_W}. \quad (1.1.12)$$

Jeżeli $R_W \gg R$, tzn. natężenie prądu płynącego przez woltomierz jest bardzo małe w porównaniu z natężeniem prądu płynącego przez opór R , we wzorze (1.1.12) możemy pominąć R i otrzymamy wynik poprawny $U = 0.5U_1$.

Jeżeli np. $R_W = 5R$ to, jak wynika ze wzoru (1.1.12), $U = \frac{5}{11}U_1$. Otrzymujemy więc zaniżoną wartość napięcia. Pomiar będzie obciążony błędem, który jest możliwy do uniknięcia.

Źródłem błędu mogą być również różne *stałe bądź wielkości wyznaczone uprzednio*. Typowym przykładem jest zaokrąglenie liczby π . Jest ona liczbą niewymierną, jej wartość wynosi $\pi = 3.1415927 \dots$. Zazwyczaj w obliczeniach przyjmujemy $\pi = 3.14$, co powoduje powstanie błędu względnego około 0.0005 spowodowanego zaokrągleniem (w przypadku gdy liczba π występuje w potęgę q wynosi on $0.0005 \times q$). Błąd ten wystąpi w pomiarach pośrednich wszędzie tam, gdzie we wzorach występuje liczba π . Na przykład, gdy wyznaczamy przyspieszenie ziemskie g za pomocą wahadła matematycznego, to zaokrąglenie liczby π do 3.14 powoduje powstanie błędu względnego około 0.001, czyli błąd bezwzględny spowodowany tym

zaokrągleniem będzie wynosił około 1 cm/s^2 . *Szczególną uwagę na możliwość powstawania błędów tego typu należy zwrócić przy przeliczaniu jednostek*, gdy przechodzimy z jednego układu jednostek do drugiego, ponieważ współczynniki przeliczeniowe są przeważnie obarczone jakimś błędem. Innym przykładem może być wyznaczona przez Millikana, w 1913 roku, wartość ładunku elementarnego e . W doświadczeniu tym konieczna była znajomość współczynnika lepkości powietrza η . Millikan przyjął znaną ówczynie wartość $\eta = 1822.6 \times 10^{-8} \text{ kg/(m} \cdot \text{s)}$ co dało $e = (1.591 \pm 0.002) \times 10^{-19} \text{ C}$. Na podstawie szeregu innych pomiarów okazało się, że podana przez Millikana wartość ładunku elementarnego jest zaniżona. W latach 1936 – 1940 ponownie wykonano pomiary współczynnika lepkości powietrza i otrzymano $\eta = 1832 \times 10^{-8} \text{ kg/(m} \cdot \text{s)}$. Uwzględnienie w obliczeniach poprawionej wartości η daje $e = (1.603 \pm 0.002) \times 10^{-19} \text{ C}$. Dziś przyjmujemy $e = (1.60217733 \pm 0.00000049) \times 10^{-19} \text{ C}$ [5].

Jeszcze inną grupą błędów, występujących we wszystkich pomiarach pośrednich, jest korzystanie z niewłaściwej teorii lub przybliżonych wzorów. Na przykład gdy przyspieszenie ziemskie g wyznaczamy korzystając ze wzoru $T = 2\pi\sqrt{l/g}$, to pomiary musimy wykonać dla małych amplitud drgań wahadła. Jeżeli pomiary wykonamy przy dużych wychyleniach z położenia równowagi, to obliczona wartość g będzie obarczona błędem, bowiem dla dużych amplitud okres drgań wahadła matematycznego w przybliżeniu wynosi $T = 2\pi[1 + 0.25 \sin^2(\alpha_0/2)]\sqrt{l/g}$, gdzie α_0 jest kątem maksymalnego wychylenia (por. np. [6]).

Tego typu błędy noszą nazwę **błędów systematycznych**. Powodują one zaniżenie lub zawyżenie wartości zmierzonej wielkości fizycznej. Jeżeli będziemy powtarzali pomiary, to otrzymamy rozrzut wyników pomiarów taki sam jak dla błędów przypadkowych, ale wszystkie wartości będą zaniżone lub zawyżone. Inaczej mówiąc błędu systematycznego nie można wykryć, gdy będziemy powtarzali pomiary wielokrotnie. Osobny przypadek, omówiony w rozdziale 2.1, zachodzi, gdy błąd systematyczny jest spowodowany dokładnością przyrządu pomiarowego, a rozrzut wyników jest mniejszy niż dokładność przyrządu. Na rysunku 1.1 przedstawiono przykładowy rozrzut wyników pomiarów długości pręta w przypadku, gdy błędy systematyczne można pominąć (rysunek 1.1A) i w przypadku, gdy błąd systematyczny powoduje zawyżenie wyników (rysunek 1.1B). Pogrubiona kreska na rysunku 1.1A przedstawia wartość x_0 , zaś na rysunku 1.1B pogrubioną kreską zaznaczono wartość, wokół której grupują się wyniki pomiarów obarczonych błędem systematycznym; jak widać jest ona różna od x_0 . Przedstawione na rysunku zawyżenie wyników może być spowodowane przesunięciem zera przymiaru przez wykonującego pomiar.

Błędów systematycznych, w przeciwieństwie do błędów przypadkowych, możemy uniknąć lub ich wpływ zminimalizować stosując:

- odpowiednie przyrządy pomiarowe,
- właściwe metody pomiaru,
- przy pomiarach pośrednich używając odpowiedniej teorii.

Omawiając błędy pomiarowe wygodnie jest korzystać z następujących pojęć: dokładność i precyzja.

Mówimy, że pomiar jest **dokładny**, gdy błąd systematyczny można pominąć, a **precyzyjny**, gdy błąd przypadkowy jest mały (por. [7] s. 20).

Należy zaznaczyć, że w literaturze polskiej rozróżnienie tych dwóch pojęć nie jest przestrzegane. W dalszym ciągu wykładu będziemy ściśle przestrzegali tego rozróżnienia. W przypadku, gdy błędy systematyczne są znacznie większe niż przypadkowe, nie możemy stosować metod statystycznych oceny błędu pomiaru. Możemy w tym przypadku stosować ocenę błędu maksymalnego omówioną w rozdziale 2. Ocena wpływu błędów systematycznych na wyniki pomiarów jest zależna od eksperymentatora, od jego wiedzy, umiejętności, sumienności i uczciwości.

3) Błędy grube, czyli omyłki

Powstają zazwyczaj w wyniku niedbale wykonanych pomiarów lub omyłek w odczycie wskazań przyrządów. Występują także w wyniku sporadycznych zakłóceń, np. elektromagnetycznych (wyładowania atmosferyczne). Mogą one powstać również w wyniku uszkodzeń aparatury pomiarowej, a także jako wynik niewłaściwej kolejności czynności podczas wykonywania pomiarów. Błędy te można wykryć powtarzając pomiary.

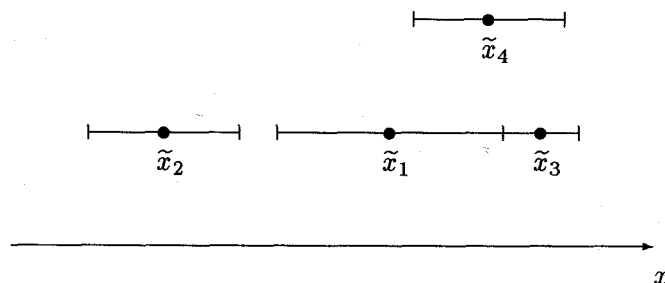
1.1.6. Zgodność wyników pomiarów

W praktyce bardzo często zdarza się, że mamy kilka wyników pomiarów jakiejś wielkości fizycznej x wykonanych niezależnie od siebie, np. w różnych laboratoriach. Mogą to być pomiary bezpośrednie, jak i pośrednie. Otrzymane z tych pomiarów wartości oznaczmy $\tilde{x}_1, \tilde{x}_2, \dots, \tilde{x}_N$, a ich błędy oznaczmy, odpowiednio, przez $\Delta_1, \Delta_2, \dots, \Delta_N$. W takiej sytuacji zawsze powstaje pytanie czy wyniki tych pomiarów są między sobą **zgodne**, czy też **sprzeczne**. Załóżmy, że mamy dwa wyniki pomiarów $\tilde{x}_1$ oraz $\tilde{x}_2$. Gdy znamy jedynie ich błędy Δ_1 i Δ_2 , mówimy, że pomiary te są zgodne, gdy przedziały $[\tilde{x}_1 - \Delta_1, \tilde{x}_1 + \Delta_1]$ i $[\tilde{x}_2 - \Delta_2, \tilde{x}_2 + \Delta_2]$ przynajmniej częściowo się pokrywają lub stykają się. **Warunek zgodności** dwóch wyników pomiarów wielkości fizycznej x wykonanych niezależnie od siebie można zapisać:

$$|\tilde{x}_1 - \tilde{x}_2| \leq \Delta_1 + \Delta_2.$$

(1.1.13)

Stosowanie warunku (1.1.13) wyjaśnimy na następującym przykładzie. Załóżmy, że wykonano w czterech laboratoriach pomiary tej samej wielkości fizycznej x . W wyniku tych pomiarów otrzymano wartości $\tilde{x}_1, \tilde{x}_2, \tilde{x}_3, \tilde{x}_4$ różniące się między sobą oraz oszacowano, że błędy ich pomiarów wynoszą odpowiednio $\Delta_1, \Delta_2, \Delta_3$ i Δ_4 . Graficznie otrzymane wyniki przedstawiono na rysunku 1.3.



Rys. 1.3: Graficzne przedstawienie zgodności wyników. $\tilde{x}_1, \tilde{x}_2, \tilde{x}_3, \tilde{x}_4$ — wartości otrzymane z pomiarów; kreskami poziomymi $\text{---}\bullet\text{---}$ oznaczono przedziały $[\tilde{x}_i - \Delta_i, \tilde{x}_i + \Delta_i]$, $i = 1, 2, 3, 4$

Jak widać z rysunku 1.3 wartości $\tilde{x}_1, \tilde{x}_3$ oraz $\tilde{x}_4$ są, na podstawie warunku (1.1.13), zgodne ze sobą, ponieważ przedziały $[\tilde{x}_1 - \Delta_1, \tilde{x}_1 + \Delta_1]$ i $[\tilde{x}_3 - \Delta_3, \tilde{x}_3 + \Delta_3]$ stykają się ze sobą, a przedział $[\tilde{x}_4 - \Delta_4, \tilde{x}_4 + \Delta_4]$ częściowo pokrywa się z każdym z nich. Wartość $\tilde{x}_2$ jest sprzeczna z pozostałymi trzema, ponieważ przedział $[\tilde{x}_2 - \Delta_2, \tilde{x}_2 + \Delta_2]$ nie styka się ani częściowo nie pokrywa żadnego z pozostałych.

1.2. Znaczenie poznania niepewności pomiarowych

W rozdziale 1.1 wykazaliśmy, że *każdy pomiar* jest obarczony niepewnością, czyli błędem pomiaru. Z tego powodu *podanie samego wyniku pomiaru* bez podania oszacowania błędu jest nieprawidłowe. Powstaje więc pytanie: czy podanie niepewności pomiarowej oprócz wyniku pomiaru jest istotne z punktu widzenia zastosowań? Zanim odpowiemy na to pytanie przedstawimy parę przykładów ilustrujących znaczenie podawania niepewności pomiarowych.

Konstruktor mostu wykonanego ze stali musi znać jej własności sprężyste i rozszerzalność cieplną. Te wielkości trzeba wyznaczyć, a więc są one obciążone niepewnościami pomiarowymi. Znajomość ich jest istotna dla projektanta, który musi uwzględnić je w obliczeniach.

Kładąc napowietrzne linie przesyłowe musimy znać wartość i błąd pomiaru współczynnika rozszerzalności liniowej drutu, który ma być użyty na położenie linii. Znajomość tych wielkości umożliwia obliczenie takiej długości drutu zawieszonego między słupami w określonej temperaturze otoczenia (np. latem), aby nie nastąpiło jego zerwanie, gdy temperatura otoczenia obniży się (np. zimą).

Przystępując do projektowania układów elektronicznych musimy znać parametry elementów, z których ma być wykonany układ. Te parametry trzeba wyznaczyć, a więc będą obciążone jakąś niepewnością pomiarową. Znając parametry i ich błędy można dokonać wyboru elementów, z których wykonamy nasz układ.

Rola znajomości niepewności pomiarowych w badaniach naukowych jest jeszcze większa. Znajomość błędów pomiarowych jest szczególnie ważna, gdy wyniki badań doświadczalnych mają służyć za podstawę do weryfikacji wyników badań teoretycznych.

Nawet w takich zagadnieniach jak produkcja np. guzików ocena błędu odgrywa istotną rolę w procesie kontroli jakości.

Z powyższych przykładów wyraźnie widać, że zawsze niezbędne jest podanie nie tylko wyniku pomiaru, ale również błędu pomiaru. W bardzo wielu przypadkach podanie jedynie samego wyniku pomiaru mija się z celem. Dlatego też nie można rozdzielać wyniku pomiaru i błędu, jakim on jest obciążony. Przykładów roli, jaką odgrywa znajomość błędu pomiaru, można podać bardzo dużo. Bardziej zainteresowanemu czytelnikowi proponujemy przeczytanie rozdziałów 1.3 i 1.4 w książce J. R. Taylora *Wstęp do analizy błędu pomiarowego* [4].

2. Ocena błędu maksymalnego

2.1. Oszacowanie niepewności przy odczycie skali, błąd maksymalny

W rozdziale 1.1 wykazaliśmy nieuchronność występowania niepewności pomiarowych, analizowaliśmy przyczyny ich występowania. Wyjaśniliśmy również, że:

- dokładność pomiaru jest zawsze skończona i nie możemy wykonać pomiaru z dokładnością większą niż dokładność przyrządu pomiarowego;
- z każdym pomiarem bezpośrednim związany jest błąd systematyczny, którego minimalna wartość jest równa dokładności przyrządu pomiarowego.

W tym rozdziale zajmujemy się oszacowaniem błędów pomiarów bezpośrednich, które sprowadzają się do odczytu skal przyrządów pomiarowych. Aby oszacować ten błąd należy zapoznać się z zasadami konstrukcji przyrządów i nanoszenia skal. Przyrządy pomiarowe są konstruowane tak, że jeżeli przyrząd jest wykonany prawidłowo i działa prawidłowo, tzn. jest sprawny, to w przypadku prawidłowo wykonanych pomiarów wyniki nie różnią się od wartości rzeczywistej x_0 więcej niż o jedną działkę, a w przypadku mierników elektrycznych o część działki określoną **klasą przyrządu** (klasy mierników elektrycznych omówiono w Dodatku B). Jeżeli klasa miernika nie jest podana, to jego dokładność jest równa wartości jednej działki i w przypadku nieliniowej skali może zmieniać się w zależności od wychylenia wskazówki. Tak więc wartość najmniejszej działki lub w przypadku mierników elektrycznych jej część określoną klasą przyrządu będziemy nazywali **dokładnością odczytu**. Inaczej mówiąc, dokładność odczytu będzie równa niepewności pomiarowej odczytu skali przyrządu. Jeśli nie będą występowały inne błędy odczytu skali przyrządu, będzie to również **błąd maksymalny**.

Błędem maksymalnym wielkości fizycznej mierzonej bezpośrednio będziemy nazywali najwyższą wartość niepewności pomiarowej, jaką jest obciążony wynik pomiaru.

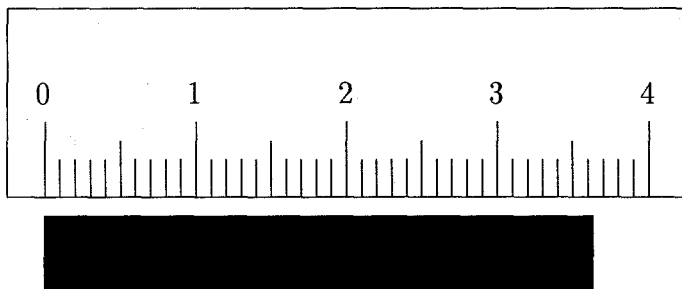
Inaczej mówiąc, *błąd maksymalny jest ograniczeniem z góry niepewności pomiaru*. Jeżeli przyrząd pomiarowy jest *sprawny* i posługujemy się nim

prawidłowo, wtedy zawsze błąd maksymalny jest równy dokładności przyrządu.

Wartość bezwzględną stosunku błędu maksymalnego do wartości rzeczywistej nazywa się **błędem maksymalnym względnym**.

Aby zilustrować powyższe rozważania rozpatrzmy następujące przykłady:

- 1) Chcemy zmierzyć długość pręta za pomocą przymiaru liniowego.



Rys. 2.1: Pomiar długości pręta przy użyciu przymiaru liniowego

W tym celu, jak pokazano na rys. 2.1, przykładamy jeden koniec pręta do początku podziałki przymiaru, tj. kreski *przyjętej za zero* podziałki i patrzymy, w którym miejscu podziałki przymiaru wypadnie drugi koniec pręta. Jeżeli podziałka jest gęsto naniesiona, to bez trudu określimy miejsce, w którym wypada koniec pręta. W przypadku przedstawionym na rys. 2.1 jest to 36 mm i część odległości między kreskami podziałki. Jeżeli koniec pręta wypadnie pomiędzy dwiema sąsiednimi kreskami, wtedy musimy oszacować *wzrokowo*, w jakiej części odległości między kreskami znajduje się koniec pręta, np. może to być w $1/3$. Powiemy wtedy, że długość pręta wynosi n -pełnych działek, np. 36, plus $1/3$ czyli 36.3 mm. Podobnie będzie z odczytem śruby mikrometrycznej czy też wskazówkowych mierników elektrycznych. Ogólnie mówiąc, postępując w ten sposób, możemy oszacować mierzoną wielkość do około $1/3$ odległości między sąsiednimi kreskami podziałki skali dowolnego przyrządu. Tak więc zawsze odczytaną (zmierzoną) wartość x możemy przedstawić jako sumę pełnej liczby działek n plus oszacowana wzrokowo część odstępu między działkami k , czyli $x = n + k$. Gdy określimy, w jakiej części odległości między sąsiednimi kreskami znajduje się koniec pręta, wtedy dokonujemy zaokrąglenia do całkowitej liczby działek. W naszym przykładzie wynik 36.3 mm zaokrąglimy do 36 mm i zapiszemy go: $x = (36 \pm 1) \text{ mm}$. Pozostaje nam jeszcze do omówienia błąd związany z określeniem położenia początku pręta względem kreski podziałki przyjętej za zero. Zwróćmy uwagę na fakt, że mierząc

długość pręta liczbę kresk podziałki liczymy *zawsze* poczynając od kreski *przyjętej za zero*. Z tego powodu jedyną niepewnością jest ustalenie czy kreska przyjęta za zero pokrywa się z początkiem pręta. Jeżeli patrzemy prostopadle, a więc możemy pominąć błąd paralaksy, to niepewność pokrycia się kreski przyjętej za zero z początkiem pręta nie będzie większa niż grubość kreski. Zazwyczaj grubość kreski nie przekracza $2/10$ odstepu między kreskami, a więc błąd spowodowany grubością kreski jest znacznie mniejszy niż błąd maksymalny i dlatego go pomijamy. Należy podkreślić, że wnioski wynikające z tego przykładu są prawdziwe wtedy, gdy pomiar wykonujemy prawidłowo.

2) *Za pomocą śruby mikrometrycznej mierzymy średnicę drutu.*

Jeżeli śruba jest wykonana prawidłowo, tzn. skok gwintu jest stały (zmiana skoku gwintu powoduje, że jeden obrót śruby będzie się różnił od drugiego) i kreski podziałki na ruchomym bębnie są naniesione w jednakowych odstępach, oraz nie występuje przesunięcie zera, to wtedy, gdy skok śruby wynosi 0.5 mm i bęben podzielono na 50 części, dokładność śruby i zarazem błąd maksymalny wynosi 0.01 mm. Jeżeli w tej samej śrubie bęben podzielono na 25 części, dokładność wynosi 0.02 mm. Gdy zmierzona wartość średnicy drutu wynosiła 2.54 mm, wynik odczytu w pierwszym przypadku zapiszemy $d = (2.54 \pm 0.01)$ mm, a w drugim $d = (2.54 \pm 0.02)$ mm.

Na zakończenie należy podkreślić, że jeżeli przyrządy są sprawne i wykonane prawidłowo, a pomiary wykonujemy poprawnie, to błąd przy odczycie skali, równy dokładności przyrządu, jest **błędem systematycznym** dla pomiarów wykonanych tym przyrządem. Jest to jednocześnie **błąd maksymalny**. Wykonując pomiary bezpośrednie często stajemy przed pytaniem: czy błąd systematyczny spowodowany dokładnością przyrządu pomiarowego jest mniejszy, czy też większy od błędów przypadkowych? Odpowiedź możemy uzyskać powtarzając pomiary. Jeżeli wykonamy serię pomiarów i rozrzut otrzymanych wyników jest w granicach dokładności przyrządu pomiarowego, to będzie to oznaczało, że błąd systematyczny jest większy niż błędy przypadkowe. Na przykład mierzymy natężenie prądu, dokładność amperomierza wynosi 0.2 A i otrzymujemy wartości 3.2 A, 3.1 A, 3.2 A, 3.3 A — rozrzut wyników jest więc zawarty w granicach dokładności amperomierza.

W praktyce (nie tylko laboratoryjnej) pomiary wielkości nieelektrycznych zastępuje się pomiarami wielkości elektrycznych stosując odpowiednie przetworniki, a mierniki elektryczne wskazówkowe są coraz częściej zastępowane miernikami cyfrowymi. Wydawałoby się, że skoro wynik pomiaru otrzymujemy w postaci cyfrowej, to jest on bezbłędny. Tak w rzeczywistości nie jest. Dokładność mierników cyfrowych jest większa niż w innych typach mierników, tym niemniej pomiary wykonane przy ich użyciu są także

obarczone błędami. Nie wnikając w zasady działania mierników cyfrowych możemy przyjąć, że dokładność ich jest równa sumie dokładności wskazania i dokładności zakresu pomiarowego. Dokładności te zależą od rodzaju wielkości mierzonej i są z reguły podawane w instrukcji miernika. W przypadku braku tych danych *musimy przyjąć, że błąd jest równy co najmniej jedności na ostatniej cyfrze wyniku odczytanego z miernika.*

2.2. Oszacowanie błędu maksymalnego w pomiarach pośrednich, metoda różniczki zupełnej

Metoda oszacowania błędu **wielkości złożonej** przedstawiona poniżej dotyczy przypadków, gdy *błędy systematyczne pomiaru wielkości mierzonych bezpośrednio są znacznie większe niż błędy przypadkowe, lub gdy wielkości te zostały zmierzone jeden raz* (jak to jest przy pomiarach kalorymetrycznych), w tych przypadkach bowiem nie można stosować metod statystycznych. Zgodnie z tym co powiedziano w rozdziale 2.1 musimy przyjąć, że pomiary wykonane jeden raz są obarczone błędem równym dokładności przyrządu pomiarowego.

Błąd pomiaru pośredniego jest zależny od błędów wielkości mierzonych bezpośrednio. W związku z tym w przypadku oceny błędu w pomiarach pośrednich sytuacja jest bardziej złożona. W celu uzyskania większej przejrzystości obliczeń przyjmijmy, że wielkość złożona z jest zależna od dwóch wielkości x i y mierzonych bezpośrednio, tzn. z jest funkcją wielkości x , y :

$$z = f(x, y). \quad (2.2.1)$$

Odpowiada to np. wyznaczaniu przyspieszenia ziemskiego g za pomocą wahadła matematycznego na podstawie wzoru (1.1.3), gdzie $g = f(l, T)$. Oznaczmy przez x_0 i y_0 wartości rzeczywiste wielkości x oraz y , a przez x_m i y_m wartości otrzymane z pomiaru wielkości fizycznych x i y . Odpowiednio niech Δx i Δy oznaczają błędy maksymalne zmierzonych wartości x_m i y_m . Wartość rzeczywistą wielkości złożonej z oznaczmy przez z_0 , zaś **wartość wyznaczoną** na podstawie pomiarów wielkości x oraz y przez z_m . Ponieważ $z = f(x, y)$, to $z_0 = f(x_0, y_0)$ oraz zakładamy, że $z_m = f(x_m, y_m)$. Wobec tego wartość bezwzględna różnicy $z_m - z_0$:

$$|z_m - z_0| = |f(x_m, y_m) - f(x_0, y_0)|, \quad (2.2.2)$$

jest błędem bezwzględnym wyznaczonej wartości wielkości złożonej z , jakim jest obarczony dany pomiar. Oszacujmy ten błąd. W tym celu rozwinijmy funkcję $f(x, y)$ na szereg Taylora funkcji dwu zmiennych w otoczeniu punktu (x_m, y_m) i obliczmy jej wartość w punkcie (x_0, y_0) (o rozwijaniu

funkcji na szereg Taylora patrz np. [8, 9]). Otrzymamy wówczas:

$$\begin{aligned} f(x_0, y_0) = & f(x_m, y_m) + \left(\frac{\partial f}{\partial x}\right)_{x_m, y_m} (x_0 - x_m) + \left(\frac{\partial f}{\partial y}\right)_{x_m, y_m} (y_0 - y_m) + \\ & + \frac{1}{2!} \left[\left(\frac{\partial^2 f}{\partial x^2}\right)_{x_m, y_m} (x_0 - x_m)^2 + 2 \left(\frac{\partial^2 f}{\partial x \partial y}\right)_{x_m, y_m} (x_0 - x_m)(y_0 - y_m) + \right. \\ & \left. + \left(\frac{\partial^2 f}{\partial y^2}\right)_{x_m, y_m} (y_0 - y_m)^2 \right] + \dots \end{aligned} \quad (2.2.3)$$

Jeżeli błędy maksymalne Δx i Δy wartości x_m i y_m zmierzonych bezpośrednio są na tyle małe, że w rozwinięciu funkcji $f(x, y)$ na szereg Taylora możemy ograniczyć się do wyrazów liniowych w zmiennych x oraz y , to wtedy:

$$f(x_0, y_0) \cong f(x_m, y_m) + \left(\frac{\partial f}{\partial x}\right)_{x_m, y_m} (x_0 - x_m) + \left(\frac{\partial f}{\partial y}\right)_{x_m, y_m} (y_0 - y_m). \quad (2.2.4)$$

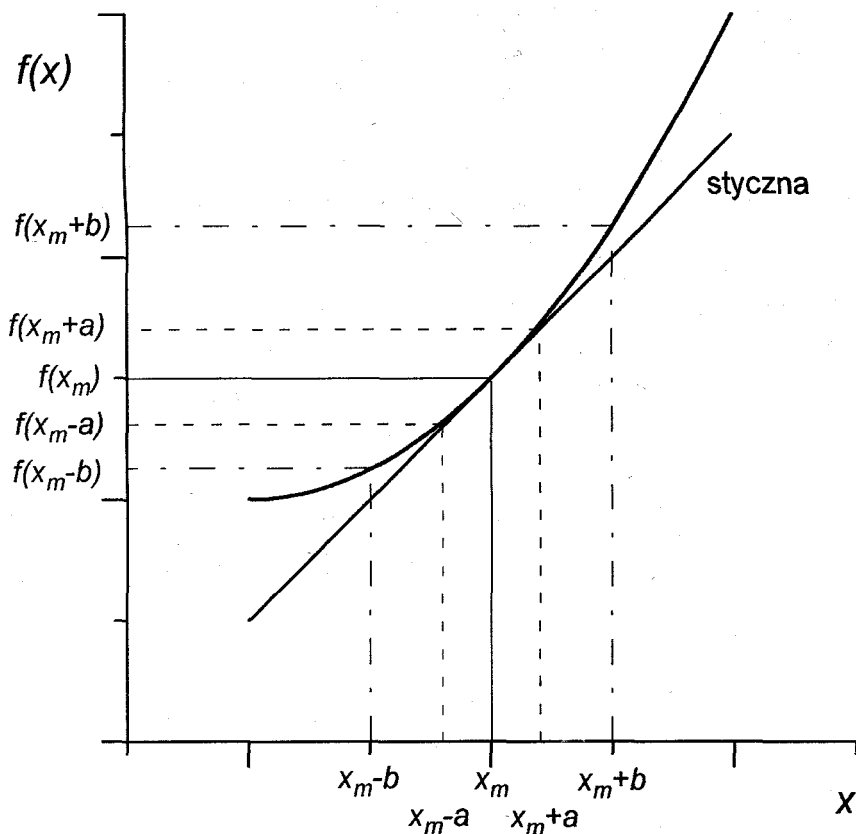
Aby zrozumieć, w jakim znaczeniu jest użyte słowo „małe”, a tym samym, aby zrozumieć znaczenie zastosowanego przybliżenia, przeanalizujmy rysunek 2.2. Na rysunku tym przedstawiono przykładowy wykres funkcji $z = f(x)$. Dla uproszczenia jest to funkcja jednej zmiennej, jednak przedstawione tu rozważania w pełni stosują się do funkcji wielu zmiennych. W punkcie x_m , w otoczeniu którego rozwijamy funkcję $f(x)$ na szereg Taylora, wykreślono styczną do krzywej $z = f(x)$. Styczna ta, jak przedstawiono na rysunku, w zakresie zmian x od $x_m - a$ do $x_m + a$ prawie pokrywa się z krzywą $z = f(x)$, zaś dla $x < x_m - a$ i $x > x_m + a$ różnice między wartościami funkcji i rzędnymi stycznej zaczynają wzrastać. Na przykład dla $x = x_m + b$ oraz dla $x = x_m - b$ różnice te są wyraźne. Wynika stąd, że dla $x_m - a \leq x \leq x_m + a$ wartości funkcji obliczone ze wzoru (2.2.4), który dla jednej zmiennej przyjmuje postać

$$f(x) \cong f(x_m) + \left(\frac{df}{dx}\right)_{x_m} (x - x_m),$$

są dobrymi przybliżeniami prawdziwej wartości funkcji. O różnicy $x - x_m$ możemy wtedy powiedzieć, że jest na tyle mała, iż w rozwinięciu funkcji na szereg Taylora możemy ograniczyć się do wyrazu liniowego. Powyższe uwagi w pełni dotyczą funkcji wielu zmiennych. W tym kontekście słowo „mały” będzie używane w dalszym ciągu wykładu.

Podstawiając (2.2.4) do (2.2.2) otrzymujemy:

$$|z_m - z_0| \cong \left| \left(\frac{\partial f}{\partial x}\right)_{x_m, y_m} (x_0 - x_m) + \left(\frac{\partial f}{\partial y}\right)_{x_m, y_m} (y_0 - y_m) \right|. \quad (2.2.5)$$



Rys. 2.2: Wyjaśnienie stosowalności rozwinięcia (2.2.4)

Ponieważ błąd maksymalny Δz jest ograniczeniem z góry wartości bezwzględnej błędu bezwzględnego danego pomiaru, zachodzi relacja:

$$|z_m - z_0| \leq \Delta z. \quad (2.2.6)$$

Aby oszacować Δz skorzystajmy z nierówności Schwartz'a: $|a+b| \leq |a|+|b|$. Z równania (2.2.5) otrzymamy wówczas:

$$|z_m - z_0| \leq \left| \left(\frac{\partial f}{\partial x} \right)_{x_m, y_m} (x_0 - x_m) \right| + \left| \left(\frac{\partial f}{\partial y} \right)_{x_m, y_m} (y_0 - y_m) \right|. \quad (2.2.7)$$

Ponieważ, zgodnie z tym co powiedziano wyżej, błędy bezwzględne wielkości x i y mierzonych bezpośrednio są nie większe od ich błędów maksymalnych, zachodzą więc nierówności:

$$\begin{aligned} |x_m - x_0| &\leq \Delta x, \\ |y_m - y_0| &\leq \Delta y. \end{aligned} \quad (2.2.8)$$

Postawiając (2.2.8) do (2.2.7) otrzymujemy:

$$|z_m - z_0| \leq \Delta z = \left| \left(\frac{\partial f}{\partial x} \right)_{x_m, y_m} \Delta x \right| + \left| \left(\frac{\partial f}{\partial y} \right)_{x_m, y_m} \Delta y \right|. \quad (2.2.9)$$

Wzór (2.2.9) można bardzo łatwo uogólnić na przypadek, gdy wielkość złożona z jest funkcją r wielkości mierzonych bezpośrednio, tzn. $z = f(x_1, x_2, \dots, x_r)$. Oznaczmy przez $x_{m,1}, x_{m,2}, \dots, x_{m,r}$ wartości zmierzone wielkości $x_1, x_2, \dots, x_r$, a przez $\Delta x_1, \Delta x_2, \dots, \Delta x_r$ ich błędy maksymalne. Wykonując rachunki identyczne jak powyżej otrzymujemy:

$$\Delta z = \sum_{i=1}^r \left| \left(\frac{\partial f}{\partial x_i} \right)_{x_{m,1}, x_{m,2}, \dots, x_{m,r}} \Delta x_i \right|. \quad (2.2.10)$$

Wyrażenie (2.2.10) pozwalające oszacować **błąd maksymalny** wielkości złożonej stanowi ilościowe ujęcie tzw. **reguły przenoszenia błędu** dla przypadku błędu maksymalnego. Przedstawiona metoda obliczania błędu maksymalnego często nosi nazwę **metody różniczki zupełnej**.

Ze wzoru (2.2.6) wynika, że:

$$z_m - \Delta z \leq z_0 \leq z_m + \Delta z. \quad (2.2.11)$$

Otrzymujemy więc związek analogiczny do nierówności (1.1.4). Możemy go również przedstawić w formie wzoru (1.1.7), tj.:

$$z_0 = z_m \pm \Delta z. \quad (2.2.12)$$

We wzorach (2.2.11) i (2.2.12) $z_m = f(x_m, y_m)$ odpowiada $\tilde{x}$ we wzorach (1.1.4) i (1.1.7), zaś błąd maksymalny Δz odpowiada Δ we wzorach (1.1.4) i (1.1.7). W ten sposób, w przypadku gdy błędy systematyczne są znacznie większe od przypadkowych, lub gdy jakakolwiek wielkość fizyczna mierzona bezpośrednio została zmierzona *jeden raz*, otrzymaliśmy odpowiedź na pytanie postawione w rozdziale 1.1.2 o cel teorii błędów.

Należy zaznaczyć, że jeżeli błędy maksymalne wielkości mierzonych bezpośrednio Δx_i zostały określone prawidłowo i nie ma innych źródeł błędów systematycznych, to wielkość Δz obliczona ze wzoru (2.2.10) określa granice (górną i dolną) przedziału, w którym zawarta jest wartość rzeczywista. Z tego co powiedziano wyżej wynika, że **błąd maksymalny względny** jest równy:

$$\left| \frac{\Delta z}{z_m} \right|.$$

Na zakończenie tego rozdziału przedstawimy przypadek, kiedy zastosowanie metody różniczki zupełnej umożliwia bardzo proste obliczenie błędu

maksymalnego względnego. Niech wielkość złożona z będzie funkcją wielkości mierzonych bezpośrednio x_i ($i = 1, 2, \dots, r$) postaci:

$$z = A \prod_{i=1}^r x_i^{a_i}, \quad (2.2.13)$$

gdzie A oraz a_i są stałymi. Logarytmując równanie (2.2.13) otrzymujemy:

$$\ln z = \ln A + \sum_{i=1}^r a_i \ln x_i. \quad (2.2.14)$$

Korzystając ze znanego wzoru $d(\ln f(x)) = d(f(x))/f(x)$ i założeń zrobionych przy wyprowadzeniu wzoru (2.2.10) otrzymujemy dla błędu względnego wyrażenie:

$$\left| \frac{\Delta z}{z_m} \right| = \sum_{i=1}^r \left| a_i \frac{\Delta x_i}{x_{m,i}} \right|. \quad (2.2.15)$$

Ten sposób obliczania **błędu maksymalnego względnego** bywa nazywany **metodą pochodnej logarytmicznej**.

Ze wzorów (2.2.10) i (2.2.15) wynikają podstawowe reguły przenoszenia błędów maksymalnych, a mianowicie jeżeli:

1. $z = x_1 + x_2$ to $\Delta z = |\Delta x_1| + |\Delta x_2|$.
2. $z = x_1 - x_2$ to $\Delta z = |\Delta x_1| + |\Delta x_2|$.
3. $z = x_1 x_2$ to $\Delta z/z = |\Delta x_1/x_1| + |\Delta x_2/x_2|$.
4. $z = x_1/x_2$ to $\Delta z/z = |\Delta x_1/x_1| + |\Delta x_2/x_2|$.

Z powyższego wynika, że: jeżeli wielkość złożona z jest sumą lub różnicą wielkości mierzonych bezpośrednio, to dodają się błędy maksymalne, a gdy iloczynem lub ilorazem to dodają się błędy maksymalne względne.

Przykłady

1. Wyznaczanie obwodu czworokąta.

Mierząc boki czworokąta otrzymujemy wartości l_1, l_2, l_3, l_4 . Obwód czworokąta $L = l_1 + l_2 + l_3 + l_4$. Niech błędy maksymalne pomiarów boków tego czworokąta wynoszą $\Delta l_1, \Delta l_2, \Delta l_3$ i Δl_4 . Wówczas błąd maksymalny wyznaczonego obwodu będzie $\Delta L = \Delta l_1 + \Delta l_2 + \Delta l_3 + \Delta l_4$. W przypadku gdy $\Delta l_1 = \Delta l_2 = \Delta l_3 = \Delta l_4 = \Delta l$ to $\Delta L = 4\Delta l$; czyli błąd wyznaczonej długości obwodu czworokąta będzie czterokrotnie większy niż błąd pomiaru każdego z boków.

2. Wyznaczanie ciepła topnienia lodu.

Ciepło topnienia lodu wyznaczamy przy użyciu kalorymetru wodnego. Oznaczmy odpowiednio przez c_k, m_k i c_w, m_w ciepło właściwe i masę kalorymetru z mieszadłkiem oraz ciepło właściwe i masę

wody. Przyjmijmy w tym przykładzie, że kalorymetr i mieszadełko są wykonane z tego samego materiału. Ćwiczenie wykonujemy następująco. Ważymy kalorymetr z mieszadełkiem — ich masa jest równa m_k . Nalewamy wodę do kalorymetru i ważymy ponownie — znajdujemy masę $m_1 = m_k + m_w$, stąd $m_w = m_1 - m_k$. Odczytujemy temperaturę początkową kalorymetru z wodą t_p . Wrzucamy do kalorymetru kawałek lodu o temperaturze 0°C . Temperatura spadnie i osiągnie wartość t_k . Ważymy ponownie kalorymetr z wodą i wodą powstałą z lodu, ich masa będzie równa $m_2 = m_1 + m_L$ (m_L - masa lodu). Ciepło topnienia lodu oznaczmy przez r . Pomijamy wymianę ciepła z otoczeniem. Spowoduje to pewien błąd systematyczny, który dla dobrze izolowanych kalorymetrów można pominąć. Dokładny opis wykonania ćwiczenia podano np. w [10, 11]. Układamy bilans cieplny

$$rm_L + c_w m_L (t_k - 0^\circ\text{C}) = c_k m_k (t_p - t_k) + c_w m_w (t_p - t_k).$$

Stąd otrzymujemy

$$r = \frac{c_k m_k (t_p - t_k) + c_w (m_1 - m_k) (t_p - t_k) - c_w (m_2 - m_1) t_k}{m_2 - m_1}. \quad (2.2.16)$$

W tym przypadku możemy jedynie ocenić błąd maksymalny, ponieważ rozrzut wyników pomiaru temperatury t_p jest zazwyczaj mniejszy od dokładności termometru, oraz dlatego, że z powodu wymiany ciepła z otoczeniem temperaturę końcową t_k możemy odczytać tylko jeden raz. Oceny błędu maksymalnego Δr możemy dokonać korzystając z wzoru (2.2.10). Jeżeli masy wyznaczamy korzystając z tej samej wagi i z tą samą dokładnością, to $\Delta m_1 = \Delta m_k = \Delta m_2 = \Delta m$, oraz, jeśli korzystamy z tego samego termometru, to $\Delta t_p = \Delta t_k = \Delta t$. Ciepło właściwe kalorymetru c_k i wody c_w odczytujemy z tablic. W celu znalezienia Δr musimy obliczyć pochodne cząstkowe:

$$\begin{aligned} \frac{\partial r}{\partial m_k} &= \frac{(c_k - c_w)(t_p - t_k)}{m_2 - m_1} = -\frac{(c_w - c_k)(t_p - t_k)}{m_2 - m_1}, \\ \frac{\partial r}{\partial m_1} &= \frac{c_w(m_2 - m_k)(t_p - t_k) + c_k m_k(t_p - t_k)}{(m_2 - m_1)^2}, \\ \frac{\partial r}{\partial m_2} &= -\frac{c_k m_k(t_p - t_k) + c_w(m_1 - m_k)(t_p - t_k)}{(m_2 - m_1)^2}, \\ \frac{\partial r}{\partial t_p} &= \frac{c_k m_k + c_w(m_1 - m_k)}{m_2 - m_1}, \\ \frac{\partial r}{\partial t_k} &= -\frac{c_k m_k + c_w(m_2 - m_k)}{m_2 - m_1}. \end{aligned}$$

Podstawiając obliczone pochodne do wzoru (2.2.10) otrzymujemy:

$$\Delta r = \left| \frac{\partial r}{\partial m_k} \Delta m \right| + \left| \frac{\partial r}{\partial m_1} \Delta m \right| + \left| \frac{\partial r}{\partial m_2} \Delta m \right| + \left| \frac{\partial r}{\partial t_p} \Delta t \right| + \left| \frac{\partial r}{\partial t_k} \Delta t \right|. \quad (2.2.17)$$

Wyznaczając ciepło topnienia lodu korzystaliśmy z kalorymetru i mieszadła wykonanego z aluminium. Jego ciepło właściwe $c_k = 0.896 \text{ kJ}/(\text{kg} \cdot \text{K})$. Wyznaczona masa kalorymetru wraz z mieszadłem $m_k = 124.2 \text{ g}$. Wyznaczone wartości pozostałych mas wynoszą: $m_1 = 207.1 \text{ g}$, $m_2 = 220.8 \text{ g}$. Pomiaru temperatury dały następujące wartości: $t_p = 21^\circ \text{C}$ i $t_k = 10^\circ \text{C}$. Masy m_k , m_1 , m_2 wyznaczono z dokładnością 0.1 g , czyli $\Delta m = 0.1 \text{ g}$, a temperatury t_p i t_k odczytano z dokładnością 0.2°C ($\Delta t = 0.2^\circ \text{C}$). Ciepło właściwe wody wynosi $c_w = 4.186 \text{ kJ}/(\text{kg} \cdot \text{K})$. Podstawiając do wzorów (2.2.16) i (2.2.17) zmierzone wartości m_k , m_1 , m_2 , t_p , t_k i odczytane z tablic wartości c_k i c_w w wyniku obliczeń otrzymujemy: $r = 326.12019 \text{ kJ/kg}$ i $\Delta r = 20.577593 \text{ kJ/kg}$. Zgodnie z tym co powiedziano w rozdziale 1.1.4 wynik pomiaru zapisujemy: $r = (326 \pm 21) \text{ kJ/kg}$, możemy podać również błąd względny:

$$\frac{\Delta r}{r} = 0.06,$$

czyli, że ciepło topnienia lodu wynosi 326 kJ/kg z dokładnością 6% .

3. Wyznaczanie gęstości cieczy nie mieszających się za pomocą naczyń połączonych.

Zakładamy, że opis naczyń połączonych i warunki równowagi cieczy w naczyniach połączonych są dobrze znane czytelnikom, dlatego pomijamy te zagadnienia. Opis wykonania ćwiczenia jest podany np. w [12].

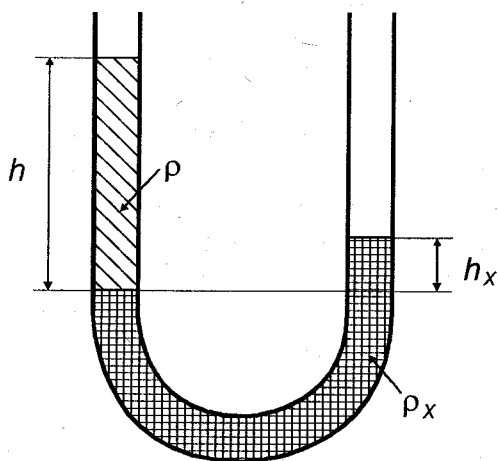
Na rysunku 2.3 przedstawiono naczynia połączone w kształcie U-rurki, w której znajdują się dwie nie mieszające się ciecze o gęstościach ρ i ρ_x . Wysokości słupów cieczy h i h_x (jak pokazano na rys. 2.3) mierzymy od poziomu ich zetknięcia się, za pomocą np. katetometru. Opis katetometru jest podany np. w [10]. W warunkach równowagi $h\rho = h_x\rho_x$, czyli

$$\rho_x = \rho \frac{h}{h_x}.$$

W tym przypadku zamiast korzystać ze wzoru (2.2.10) wygodnie jest obliczyć błąd maksymalny względny, korzystając z wzoru (2.2.15). Otrzymamy wówczas:

$$\frac{\Delta \rho_x}{\rho_x} = \left| \frac{\Delta \rho}{\rho} \right| + \left| \frac{\Delta h}{h} \right| + \left| \frac{\Delta h_x}{h_x} \right|.$$

Niech cieczą badaną o gęstości ρ_x będzie rtęć, a cieczą o znanej gęstości woda ($\rho = 1 \text{ g/cm}^3$). Pomiaru dały $h = 271.5 \text{ mm}$, $h_x = 20.4 \text{ mm}$. Dokładność wyznaczenia wysokości słupa cieczy za pomocą katetometru wynosi



Rys. 2.3: Naczynia połączone z cieciami nie mieszącymi się

0.2 mm, wobec tego $\Delta h = \Delta h_x = 0.2$ mm. Przyjmując, że znamy dokładną wartość gęstości wody ($\Delta \rho = 0$) otrzymujemy: $\rho_x = 13.3088$ g/cm³, $\Delta \rho_x / \rho_x = 0.01054$, $\Delta \rho_x = 0.14027$ g/cm³. Zaokrąglając błąd względny możemy powiedzieć, że gęstość rtęci została wyznaczona z dokładnością 1.1%. Zgodnie z tym, co powiedziano w rozdziale 1.1.4, wynik zapiszemy w postaci:

$$\rho_x = (13.31 \pm 0.15) \text{ g/cm}^3.$$

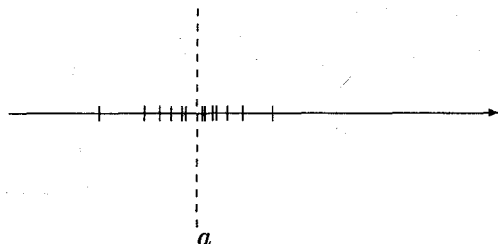
3. Wielkości charakteryzujące serię pomiarów obarczonych błędami przypadkowymi

W poprzednim rozdziale przedstawiliśmy ocenę błędu maksymalnego, która umożliwia uwzględnienie błędów systematycznych. W tym i trzech następnych rozdziałach będziemy rozpatrywać przypadki, gdy błędy systematyczne można pominąć w porównaniu z błędami przypadkowymi. Będziemy więc rozważali przypadki, gdy wykonujemy serię pomiarów i rozrzut wyników będzie znacznie większy niż dokładność przyrządu pomiarowego. Każda seria pomiarów zawiera skończoną liczbę wyników N . W praktyce liczba wyników pomiarów N wynosi od kilku do kilkudziesięciu, jest więc liczbą niezbyt dużą. W związku z tym, zanim przejdziemy do omówienia tych najprostszych zagadnień statystyki matematycznej, które mają bezpośrednie zastosowanie praktyczne przy ocenie błędu i przedstawianiu wyników pomiarów, zajmiemy się, zarówno w przypadku pomiarów bezpośrednich, jak i pośrednich, wielkościami wyznaczanymi na podstawie wyników pomiarów charakteryzującymi serię pomiarową.

3.1. Wartość średnia serii pomiarów bezpośrednich

W rozdziale 1.1.5 pokazaliśmy, że błędy przypadkowe powodują rozrzut wyników pomiaru, tzn. że jeżeli wykonamy serię N pomiarów tej samej wielkości fizycznej x , to otrzymamy ciąg jej wartości $x_1, x_2, \dots, x_N$. Wartości te będą rozłożone w pewnym przedziale, tzw. **przedziale zmienności**. Jak wykazuje praktyka pomiarowa (i co będzie dalej uzasadnione) rozrzut wyników nie jest równomierny. Przy dużej liczbie pomiarów widać, że układają się one w pobliżu pewnej wartości a , leżącej wewnątrz przedziału zmienności. Przykład takiego rozrzutu przedstawiono na rysunku 3.1. W związku z występowaniem rozrzutu wyników pomiaru w pierwszej kolejności podamy kryterium, które pozwoli znaleźć wartość a , wokół której grupują się wyniki **serii pomiarów**.

Najprostszym narzucającym się kryterium jest założenie, że suma wartości bezwzględnych różnic między poszukiwaną wartością a oraz wynikami

Rys. 3.1: Rozrzut wyników pomiarów wokół wartości a

pomiarów x_i ma być minimalna. Inaczej mówiąc funkcja

$$G_N(a) = \sum_{i=1}^N |x_i - a| . \quad (3.1.1)$$

ma osiągać minimum, czyli:

$$G_N(a) = \min . \quad (3.1.2)$$

Jednak takie kryterium nie jest poprawne, co pokażemy na następującym przykładzie.

Tabela 3.1: Wartości funkcji $G_4(a)$

a	9.0	10.0	10.5	11.0	11.5	12.0	13.0
$G_4(a)$	8.5	6.5	6.5	6.5	6.5	6.5	8.5

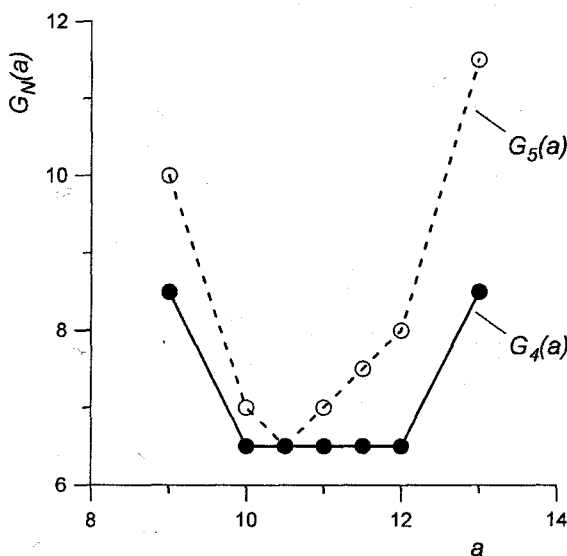
W wyniku pomiaru wielkości fizycznej x otrzymano wartości: $x_1 = 10$, $x_2 = 12$, $x_3 = 9$, $x_4 = 13.5$, czyli $N = 4$. W tabeli 3.1 przedstawiono wartości funkcji $G_4(a)$, a na rys 3.2 jej wykres (linia ciągła). Jak widać z tabeli 3.1 i rysunku 3.2 funkcja $G_4(a)$ przyjmuje wartość minimalną w przedziale od $a = 10$ do $a = 12$. Wynika z tego, że kryterium (3.1.2) z funkcją (3.1.1) nie pozwala jednoznacznie znaleźć wartości a . Jeżeli serię pomiarów powiększymy o jeszcze jeden pomiar, który np. da wartość $x_5 = 10.5$, to wówczas, jak widać z tabeli 3.2 i rysunku 3.2, funkcja $G_5(a)$ (linia przerywana) przyjmie wartość minimalną gdy $a = 10.5$.

Z powyższych przykładów wynika, że jeżeli istnieje minimum funkcji $G_N(a)$, to wystąpi ono dla a równego któremuś z wyników pomiarów.

Oba te przykłady wskazują na to, że funkcja (3.1.1) nie pozwala na znalezienie wartości a , wokół której grupują się wyniki pomiarów.

Tabela 3.2: Wartości funkcji $G_5(a)$

a	9.0	10.0	10.5	11.0	11.5	12.0	13.0
$G_5(a)$	10.0	7.0	6.5	7.0	7.5	8.0	11.5

Rys. 3.2: Wykres funkcji $G_N(a)$; — $N = 4$, - - - $N = 5$

Drugim naturalnie nasuwającym się kryterium jest założenie, że suma wszystkich różnic $x_i - a$ ma być równa zero, czyli:

$$\sum_{i=1}^N (x_i - a) = 0. \quad (3.1.3)$$

Kryterium powyższe można wyprowadzić z warunku ogólniejszego, a mianowicie:

$$A_N(a) = \sum_{i=1}^N (x_i - a)^2 = \min. \quad (3.1.4)$$

Aby znaleźć wartość a , koło której grupują się wyniki pomiarów, tzn. minimum funkcji (3.1.4), zróżniczkujemy funkcję $A_N(a)$ względem a i pochodną przyrównamy do zera, tzn.

$$\frac{dA_N(a)}{da} = -2 \sum_{i=1}^N (x_i - a) = 0, \quad (3.1.5)$$

stąd zaś mamy zależność (3.1.3) oraz

$$\left(\sum_{i=1}^N x_i \right) - Na = 0. \quad (3.1.6)$$

Ostatecznie otrzymujemy:

$$a = \frac{1}{N} \sum_{i=1}^N x_i. \quad (3.1.7)$$

Z analizy matematycznej wiadomo, że warunek zerowania się pierwszej pochodnej funkcji, której ekstremum poszukujemy, jest warunkiem koniecznym, ale niewystarczającym. Jak łatwo się przekonać, w przypadku funkcji $A_N(a)$ w punkcie a danym wzorem (3.1.7), znajduje się jej minimum. W dalszym ciągu tego opracowania wszędzie tam, gdzie będziemy poszukiwali ekstremum funkcji, ograniczymy się do znajdowania miejsc zerowych pierwszej pochodnej. We wszystkich tych przypadkach będą to ekstrema, a nie punkty przegięcia. Sprawdzenie tego pozostawiamy czytelnikowi.

Wielkość a , wokół której grupują się wyniki pomiarów obarczonych **błędami przypadkowymi**, dana równaniem (3.1.7), jest niczym innym jak **średnią arytmetyczną** $\bar{x}$ wartości x_i ($i = 1, 2, \dots, N$), czyli

$$\bar{x} = \frac{1}{N} \sum_{i=1}^N x_i. \quad (3.1.8)$$

Związek średniej arytmetycznej $\bar{x}$, wokół której grupują się wyniki serii N pomiarów obarczonych błędami przypadkowymi z wartością rzeczywistą omówimy w rozdziale 5.

W ogólności średnią arytmetyczną dowolnych wielkości nazywamy ich sumę podzieloną przez ich liczbę.

Oprócz średniej arytmetycznej zdefiniowane są inne średnie, jak np.:

– średnia kwadratowa

$$\sqrt{\frac{1}{N} \sum_{i=1}^N x_i^2},$$

– średnia harmoniczna

$$\left(\frac{1}{N} \sum_{i=1}^N \frac{1}{x_i} \right)^{-1},$$

– średnia geometryczna

$$\sqrt[N]{\prod_{i=1}^N x_i}.$$

W dalszym ciągu wykładu będziemy spotykali się z średnią arytmetyczną i średnią kwadratową.

Warunek (3.1.4) pozwala na jednoznaczne wyznaczenie wartości, wokół której grupują się wyniki pomiarów. Wyznaczone ze wzoru (3.1.7) wartości średnich arytmetycznych dla uprzednio podanych przykładowych serii pomiarowych wynoszą odpowiednio: 11.125 w przypadku 4 pomiarów i 11.000 w przypadku 5 pomiarów.

Zatem w celu znalezienia wartości, koło której grupują się wyniki pomiarów, szukamy takiej wartości a , dla której suma kwadratów różnic $x_i - a$ jest najmniejsza. Poszukiwanie wartości, dla której jakieś wyrażenie ma wartość najmniejszą nazywa się *minimalizacją*. Przedstawiony sposób minimalizacji nosi nazwę **metody najmniejszych kwadratów**. Postępowanie polegające na wykorzystaniu warunku (3.1.4) w celu znalezienia wartości a na podstawie wyników pomiarów jest najprostszym przykładem zastosowania metody najmniejszych kwadratów. Do metody tej wrócimy w rozdziale 3.5 i rozdziale 5.

Wzór (3.1.8) pozwala nam obliczyć średnią arytmetyczną serii pomiarów zawierającej N wyników. Często zdarza się, że chcemy znaleźć *wartość średnią z różnych serii pomiarowych* tej samej wielkości fizycznej x . Przypuśćmy, że wielkość fizyczna x była zmierzona w L seriach pomiarowych o jednakowej precyzji. Każda z serii pomiarowych miała inną średnią $\bar{x}_k$ ($k = 1, 2, \dots, L$) i w każdej z serii wykonano inną liczbę pomiarów N_k . Znamy tylko $\bar{x}_k$ i N_k , nie znamy natomiast wyników poszczególnych pomiarów wykonanych w każdej serii. Dla każdej z serii zgodnie z definicją średniej arytmetycznej spełniony jest związek $N_k \bar{x}_k = \sum_{j=1}^{N_k} x_{kj}$. Możemy więc znaleźć układ L równań:

$$\begin{aligned} N_1 \bar{x}_1 &= \sum_{j=1}^{N_1} x_{1,j}, \\ N_2 \bar{x}_2 &= \sum_{j=1}^{N_2} x_{2,j}, \\ &\vdots \\ N_L \bar{x}_L &= \sum_{j=1}^{N_L} x_{L,j}. \end{aligned} \tag{3.1.9}$$

Dodając do siebie stronami powyższe równania otrzymujemy:

$$\sum_{k=1}^L N_k \bar{x}_k = \sum_{j=1}^{N_1} x_{1,j} + \sum_{j=1}^{N_2} x_{2,j} + \dots + \sum_{j=1}^{N_L} x_{L,j}. \tag{3.1.10}$$

Oznaczmy

$$\sum_{k=1}^L N_k = N,$$

gdzie N jest liczbą pomiarów wykonanych we wszystkich seriach, wówczas równanie (3.1.10) przyjmie postać:

$$\sum_{k=1}^L N_k \bar{x}_k = \sum_{j=1}^N x_j. \quad (3.1.11)$$

Podzielimy obie strony równania (3.1.11) przez N :

$$\frac{1}{N} \sum_{k=1}^L N_k \bar{x}_k = \frac{1}{N} \sum_{j=1}^N x_j. \quad (3.1.12)$$

Prawa strona równania (3.1.12) jest średnią arytmetyczną pomiarów wykonanych we wszystkich seriach, a więc:

$$\bar{x} = \frac{\sum_{k=1}^L N_k \bar{x}_k}{\sum_{k=1}^L N_k} = \sum_{k=1}^L w_k \bar{x}_k. \quad (3.1.13)$$

Wzór (3.1.13) przedstawia nam wartość średnią L serii pomiarowych wielkości fizycznej x . Stosunek $\frac{N_k}{N} = w_k$ bywa, w tym przypadku, nazywany **wagą pomiaru**, a średnia dana wzorem (3.1.13) **średnią ważoną**. Do tego zagadnienia wrócimy w rozdziale 5.1.5.

Przykład

Powierznię prostokąta wyznaczono w 4 seriach pomiarowych. W pierwszej serii wykonano $N_1 = 8$, w drugiej $N_2 = 10$, trzeciej $N_3 = 4$ i czwartej $N_4 = 7$ pomiarów. Wartości średnie dla poszczególnych serii wynosiły: $\bar{x}_1 = 167.7 \text{ cm}^2$, $\bar{x}_2 = 171.6 \text{ cm}^2$, $\bar{x}_3 = 169 \text{ cm}^2$, $\bar{x}_4 = 170.3 \text{ cm}^2$. Zgodnie ze wzorem (3.1.13)

$$\bar{x} = \frac{8 \cdot 167.7 + 10 \cdot 171.6 + 4 \cdot 169 + 7 \cdot 170.3}{29} \text{ cm}^2 = 169.9 \text{ cm}^2.$$

Wagi pomiarów wynoszą dla pierwszej serii $\frac{8}{29}$, drugiej $\frac{10}{29}$, trzeciej $\frac{4}{29}$ i czwartej $\frac{7}{29}$.

3.2. Odchylenie standardowe serii pomiarów bezpośrednich

W poprzednim rozdziale stwierdziliśmy, że średnia arytmetyczna serii N pomiarów jest wartością, wokół której grupują się wyniki pomiarów. Nie

mówi ona jednak nic o rozrzucie wyników pomiarów, a więc jej wartość nie charakteryzuje precyzji pomiarów.

Różnica $x_i - \bar{x}$ jest odchyleniem pojedynczego pomiaru od wartości średniej danej serii złożonej z N pomiarów. Ważne jest więc zdefiniowanie wielkości, która będzie charakteryzowała nie pojedynczy pomiar, ale całą serię N pomiarów. Określi nam ona precyzję, z jaką została wykonana dana seria pomiarów. Obliczmy sumę różnic $x_i - \bar{x}$ ($i = 1, 2, \dots, N$). Jest ona równa:

$$\sum_{i=1}^N (x_i - \bar{x}) = \sum_{i=1}^N x_i - N\bar{x}. \quad (3.2.1)$$

Podstawiając do (3.2.1) średnią arytmetyczną $\bar{x}$ daną wzorem (3.1.8) otrzymujemy ważny związek:

$$\sum_{i=1}^N (x_i - \bar{x}) = 0. \quad (3.2.2)$$

Wyrażenie (3.2.2) jest zawsze prawdziwe, ponieważ przy jego wyprowadzeniu nie dokonywaliśmy żadnych przybliżeń. Nie może ono jednak być wykorzystane do oceny precyzji pomiaru, ponieważ dla danej serii pomiarowej jest ono spełnione tożsamościowo. Z tego powodu za miarę precyzji danej serii pomiarów przyjmujemy średnią kwadratową błędów bezwzględnych S_x zdefiniowaną następująco:

$$S_x = \sqrt{\frac{1}{N} \sum_{i=1}^N \delta_i^2} = \sqrt{\frac{1}{N} \sum_{i=1}^N (x_i - x_0)^2}, \quad (3.2.3)$$

gdzie δ_i jest błędem bezwzględnym i -tego pomiaru. Wielkość S_x będziemy nazywali w dalszym ciągu tego wykładu **odchyleniem standardowym serii** lub **odchyleniem standardowym pojedynczego pomiaru serii**. W starszych opracowaniach S_x nazywane jest **średnim błędem kwadratowym pojedynczego pomiaru**. W tym opracowaniu, aby uniknąć częstego pisania symbolu pierwiastka kwadratowego, będziemy często posługiwali się wielkością S_x^2 , którą będziemy nazywali **wariancją serii**. Przyjęcie S_x za miarę precyzji serii pomiarowej znajduje uzasadnienie, które zostanie podane w rozdziale 5.1. Wielkość S_x zdefiniowana wzorem (3.2.3) jest nieznaną, ponieważ nie jest znana wartość rzeczywista x_0 . Obliczenie S_x staje się możliwe, gdy znajdziemy związek między nieznaną sumą kwadratów błędów bezwzględnych $\sum_{i=1}^N (x_i - x_0)^2$ a sumą kwadratów odchyleń od średniej arytmetycznej $\sum_{i=1}^N (x_i - \bar{x})^2$, którą łatwo można obliczyć. Błąd bezwzględny i -tego pomiaru jest:

$$\delta_i = x_i - x_0. \quad (3.2.4)$$

Wysumujmy błędy bezwzględne i podzielmy przez liczbę pomiarów. Otrzymamy wtedy:

$$\frac{1}{N} \sum_{i=1}^N \delta_i = \frac{1}{N} \sum_{i=1}^N x_i - x_0. \quad (3.2.5)$$

Z powyższego równania wynika, że

$$x_0 = \frac{1}{N} \sum_{i=1}^N x_i - \frac{1}{N} \sum_{i=1}^N \delta_i = \bar{x} - \bar{\delta}, \quad (3.2.6)$$

gdzie

$$\bar{\delta} = \frac{1}{N} \sum_{i=1}^N \delta_i. \quad (3.2.7)$$

Ponieważ nie mamy żadnych podstaw, aby przyjąć, że $\bar{x}$ jest równe x_0 , musimy przyjąć, że w ogólnym przypadku $\bar{x} \neq x_0$ więc $\bar{\delta} \neq 0$. Podstawiając (3.2.6) do (3.2.4) otrzymujemy:

$$\delta_i = x_i - \bar{x} + \bar{\delta}.$$

Stąd:

$$x_i - \bar{x} = \delta_i - \bar{\delta}. \quad (3.2.8)$$

Podnosząc (3.2.8) obustronnie do kwadratu, sumując po i ($i = 1, 2, \dots, N$) oraz dzieląc przez liczbę pomiarów mamy:

$$\frac{1}{N} \sum_{i=1}^N (x_i - \bar{x})^2 = \frac{1}{N} \sum_{i=1}^N (\delta_i - \bar{\delta})^2 = \frac{1}{N} \sum_{i=1}^N \delta_i^2 + (\bar{\delta})^2 - \frac{2}{N} \bar{\delta} \sum_{i=1}^N \delta_i. \quad (3.2.9)$$

Wstawiając (3.2.7) do (3.2.9) i biorąc pod uwagę wzór (3.2.3) otrzymujemy:

$$\frac{1}{N} \sum_{i=1}^N (x_i - \bar{x})^2 = S_x^2 - (\bar{\delta})^2. \quad (3.2.10)$$

Ale

$$(\bar{\delta})^2 = \left(\frac{1}{N} \sum_{i=1}^N \delta_i \right)^2 = \frac{1}{N^2} \sum_{i=1}^N \delta_i^2 + \frac{1}{N^2} \sum_j \delta_j \left(\sum_{k \neq j}^N \delta_k \right). \quad (3.2.11)$$

Ponieważ iloczyny $\delta_j \delta_k$ ($j \neq k$) mają różne znaki, nie popełnimy dużego błędu, jeżeli dokonamy przybliżenia i pominiemy w równaniu (3.2.11) wyraz zawierający te iloczyny. Otrzymamy wtedy:

$$(\bar{\delta})^2 \approx \frac{1}{N^2} \sum_{i=1}^N \delta_i^2 = \frac{1}{N} S_x^2. \quad (3.2.12)$$

Podstawiając (3.2.12) do (3.2.10) znajdujemy:

$$S_x^2 = \frac{1}{N-1} \sum_{i=1}^N (x_i - \bar{x})^2. \quad (3.2.13)$$

Otrzymane wyrażenie zawiera jedynie wielkości będące wynikiem pomiaru, pozwala więc na obliczenie odchylenia standardowego serii S_x . Występująca w mianowniku wzoru (3.2.13) liczba $N-1 = k$ jest nazywana **liczbą stopni swobody**. Jest ona równa liczbie niezależnie wyznaczonych wielkości (liczbie pomiarów) minus liczba wielkości (parametrów) obliczonych przy ich użyciu. W naszym przypadku jest jedna wielkość obliczona na podstawie serii N pomiarów, jest nią średnia arytmetyczna $\bar{x}$, która jest niezbędna do obliczenia S_x . Inaczej mówiąc liczba stopni swobody $k = N - Q$, gdzie N jest liczbą niezależnie wyznaczonych wielkości, a Q liczbą równań wiążących te wielkości.

Odchylenie standardowe serii S_x charakteryzuje precyzję wykonania pomiarów w danej serii pomiarowej i możemy je uważać za średnią wartość niepewności każdego z pomiarów danej serii pomiarowej. Błędy poszczególnych pomiarów przenoszą się na obliczoną wartość średniej arytmetycznej $\bar{x}$. Wielkość charakteryzującą precyzję wyznaczonej średniej arytmetycznej $\bar{x}$ nazywamy **odchyleniem standardowym średniej arytmetycznej serii S_x** ; obliczymy je w rozdziale 3.4, w którym poznamy „regułę przenoszenia błędów” w przypadku, gdy wyniki pomiarów niezależnych obciążone są błędami przypadkowymi.

3.3. Wartość średnia serii pomiarów pośrednich

W dwóch poprzednich rozdziałach, dla mierzonej bezpośrednio wielkości fizycznej x , znaleźliśmy wartość, wokół której grupują się wyniki pomiarów, tj. średnią arytmetyczną $\bar{x}$ oraz odchylenie standardowe serii S_x . Większość wielkości fizycznych to wielkości złożone, zależne od wielkości mierzonych bezpośrednio. Ich pomiar jest **pomiarem pośrednim** i został on wyjaśniony w rozdziale 1.1.1. Dla łatwiejszego zrozumienia i większej przejrzystości rozważań przyjmijmy, że wielkość złożona z jest zależna od dwóch mierzonych bezpośrednio wielkości x i y , tzn., że

$$z = f(x, y). \quad (3.3.1)$$

3.3.1. Pomiary niezależne

Załóżmy, że obie mierzone wielkości fizyczne są od siebie całkowicie **niezależne**, to znaczy, że wynik pomiaru jednej wielkości nie zależy od wyniku

pomiaru drugiej oraz błąd pomiaru jednej wielkości nie jest związany z błędem drugiej.

Przyjmijmy, że w wyniku pomiaru wielkości x otrzymaliśmy zbiór N wartości $x_1, x_2, \dots, x_i, \dots, x_N$, a wielkości y zbiór M wartości $y_1, y_2, \dots, y_k, \dots, y_M$. Każdej parze wartości (x_i, y_k) ($i = 1, \dots, N$), ($k = 1, \dots, M$) odpowiada wartość $z_{ik} = f(x_i, y_k)$, np. parze (x_1, y_1) odpowiada z_{11} , a parze (x_4, y_1) odpowiada z_{41} . Liczba par (x_i, y_k) jest oczywiście równa NM , a więc otrzymujemy zbiór wartości wielkości złożonej z zawierający NM elementów z_{ik} . Zgodnie z tym co powiedziano o średniej arytmetycznej w rozdziale 3.1 wartość średniej arytmetycznej $\bar{z}$ serii pomiarów pośrednich wielkości $z(x, y)$ jest zdefiniowana następująco:

$$\bar{z} = \frac{1}{NM} \sum_{i=1}^N \sum_{k=1}^M z_{ik}. \quad (3.3.2)$$

Obliczenie $\bar{z}$ przy użyciu wzoru (3.3.2) prowadzi do żmudnych obliczeń numerycznych. Wyrażenie to można jednak uprościć. W tym celu rozwińmy funkcję $z = f(x, y)$ na szereg Taylora funkcji dwu zmiennych, w otoczeniu punktu $(\bar{x}, \bar{y})$, i ograniczmy się do wyrazów liniowych w zmiennych x oraz y . W tym przybliżeniu wartość funkcji $z = f(x, y)$ w punkcie (x_i, y_k) będzie równa:

$$z_{ik} = f(x_i, y_k) \cong f(\bar{x}, \bar{y}) + \left(\frac{\partial f}{\partial x}\right)_{\bar{x}, \bar{y}} (x_i - \bar{x}) + \left(\frac{\partial f}{\partial y}\right)_{\bar{x}, \bar{y}} (y_k - \bar{y}). \quad (3.3.3)$$

Przybliżenie to jest równoważne założeniu, że zmiana wartości wielkości $z = f(x, y)$ spowodowana rozrzutem wyników pomiaru w całym przedziale otrzymanych wartości x_i i y_k jest w przybliżeniu liniowa. Bardziej szczegółowe uzasadnienie stosowalności tego przybliżenia zostanie przedstawione w rozdziale 5.1.6, (porównaj także rozdział 2.2). Podstawiając (3.3.3) do (3.3.2) otrzymujemy

$$\begin{aligned} \bar{z} \cong & \frac{1}{NM} \sum_{i=1}^N \sum_{k=1}^M f(\bar{x}, \bar{y}) + \frac{1}{NM} \sum_{i=1}^N \sum_{k=1}^M \left(\frac{\partial f}{\partial x}\right)_{\bar{x}, \bar{y}} (x_i - \bar{x}) + \\ & + \frac{1}{NM} \sum_{i=1}^N \sum_{k=1}^M \left(\frac{\partial f}{\partial y}\right)_{\bar{x}, \bar{y}} (y_k - \bar{y}). \end{aligned}$$

Ponieważ $f(\bar{x}, \bar{y})$, $(\partial f / \partial x)_{\bar{x}, \bar{y}}$ i $(\partial f / \partial y)_{\bar{x}, \bar{y}}$ możemy wynieść przed znaki sum, to

$$\bar{z} \cong f(\bar{x}, \bar{y}) + \frac{1}{N} \left(\frac{\partial f}{\partial x}\right)_{\bar{x}, \bar{y}} \sum_{i=1}^N (x_i - \bar{x}) + \frac{1}{M} \left(\frac{\partial f}{\partial y}\right)_{\bar{x}, \bar{y}} \sum_{k=1}^M (y_k - \bar{y}). \quad (3.3.4)$$

Zgodnie ze wzorem (3.2.2) $\sum_{i=1}^N (x_i - \bar{x}) = 0$ i $\sum_{k=1}^M (y_k - \bar{y}) = 0$, a więc z równania (3.3.4) otrzymujemy *przybliżoną wartość*

$$\bar{z} = f(\bar{x}, \bar{y}). \quad (3.3.5)$$

Tak więc **wartość średnia wielkości złożonej** $z = f(x, y)$ jest funkcją średnich arytmetycznych niezależnych wielkości mierzonych bezpośrednio pod warunkiem, że *możemy w rozwinięciu (3.3.3) ograniczyć się do wyrazów liniowych w zmiennych x i y .*

Gdy wielkość złożona z jest funkcją r niezależnych wielkości mierzonych bezpośrednio $x_1, x_2, \dots, x_r$ to, w tym przybliżeniu, zachodzi związek:

$$\bar{z} = f(\bar{x}_1, \bar{x}_2, \dots, \bar{x}_r). \quad (3.3.6)$$

Dowód pozostawiamy czytelnikowi.

3.3.2. Pomiaru zależne

Rozpatrzmy teraz przypadek, gdy zmierzona wartość wielkości fizycznej y jest zależna od wyniku pomiaru innej wielkości fizycznej x . Pomiaru takie będziemy nazywali **pomiarami zależnymi**. W takim przypadku każdej zmierzonej wartości x_i odpowiada wartość y_i będąca wynikiem pomiaru wielkości y .

Przyjmijmy, że wykonaliśmy N odpowiadających sobie pomiarów dwóch wielkości x oraz y , w wyniku których otrzymaliśmy N par (x_i, y_i) ($i = 1, 2, \dots, N$). Na ich podstawie można wyznaczyć N wartości $z_i = f(x_i, y_i)$ ($i = 1, 2, \dots, N$). Zgodnie z tym co powiedziano o średniej arytmetycznej w rozdziale 3.1 wartość **średniej arytmetycznej serii pomiarów pośrednich** będzie w tym przypadku wyrażać się wzorem:

$$\bar{z} = \frac{1}{N} \sum_{i=1}^N z_i. \quad (3.3.7)$$

Podobnie jak w przypadku pomiarów niezależnych rozwińmy funkcję $z = f(x, y)$ na szereg Taylora funkcji dwóch zmiennych w otoczeniu punktu $(\bar{x}, \bar{y})$ i ograniczmy się do wyrazów liniowych w zmiennych x oraz y . W tym przybliżeniu wartość funkcji w punkcie (x_i, y_i) jest

$$z_i = f(x_i, y_i) \cong f(\bar{x}, \bar{y}) + \left(\frac{\partial f}{\partial x}\right)_{\bar{x}, \bar{y}} (x_i - \bar{x}) + \left(\frac{\partial f}{\partial y}\right)_{\bar{x}, \bar{y}} (y_i - \bar{y}). \quad (3.3.8)$$

Podstawiając (3.3.8) do (3.3.7) i powtarzając obliczenia wykonane dla pomiarów niezależnych otrzymujemy

$$\bar{z} = f(\bar{x}, \bar{y}). \quad (3.3.9)$$

Wzór (3.3.9) może być stosowany wówczas, gdy rozwinięcie (3.3.8) jest dobrym przybliżeniem lub związek między mierzonymi wielkościami jest liniowy. Jest on identyczny ze wzorem (3.3.5) otrzymanym dla pomiarów niezależnych. W przypadku gdy wielkość złożoną wyznaczamy na podstawie pomiarów bezpośrednich r wielkości zależnych otrzymamy również wzór (3.3.6).

3.4. Odchylenie standardowe serii pomiarów pośrednich

3.4.1. Pomiary niezależne

Odchylenie standardowe serii pomiarów pośrednich S_z wielkości $z = f(x, y)$ definiujemy podobnie jak wielkości mierzonej bezpośrednio. Zgodnie ze wzorem (3.2.3) jego kwadrat wynosi:

$$S_z^2 = \frac{1}{NM} \sum_{i=1}^N \sum_{k=1}^M (z_{ik} - z_0)^2, \quad (3.4.1)$$

gdzie $z_0 = f(x_0, y_0)$ jest wartością rzeczywistą wielkości złożonej z , pozostałe oznaczenia jak w rozdziale 3.3. Ponieważ nieznana jest wartość rzeczywista z_0 nie możemy, ze wzoru (3.4.1), obliczyć S_z^2 . Aby obliczyć *przybliżoną wartość* S_z^2 rozwińmy funkcję $z = f(x, y)$ na szereg Taylora funkcji dwu zmiennych, w *otoczeniu punktu* $(\bar{x}, \bar{y})$ i ograniczmy się do wyrazów liniowych w zmiennych x i y (przybliżenie to zostanie szczegółowo omówione w rozdziale 5.1.6, por. także rozdział 2.2), wówczas wartości funkcji $f(x, y)$ w punktach (x_i, y_k) oraz (x_0, y_0) wyrażą się następująco:

$$z_{ik} = f(x_i, y_k) \cong f(\bar{x}, \bar{y}) + \left(\frac{\partial f}{\partial x}\right)_{\bar{x}, \bar{y}} (x_i - \bar{x}) + \left(\frac{\partial f}{\partial y}\right)_{\bar{x}, \bar{y}} (y_k - \bar{y}), \quad (3.4.2)$$

$$z_0 = f(x_0, y_0) \cong f(\bar{x}, \bar{y}) + \left(\frac{\partial f}{\partial x}\right)_{\bar{x}, \bar{y}} (x_0 - \bar{x}) + \left(\frac{\partial f}{\partial y}\right)_{\bar{x}, \bar{y}} (y_0 - \bar{y}). \quad (3.4.3)$$

Odejmując stronami równania (3.4.2) i (3.4.3) znajdujemy:

$$z_{ik} - z_0 \cong \left(\frac{\partial f}{\partial x}\right)_{\bar{x}, \bar{y}} (x_i - x_0) + \left(\frac{\partial f}{\partial y}\right)_{\bar{x}, \bar{y}} (y_k - y_0). \quad (3.4.4)$$

Podstawiając (3.4.4) do (3.4.1) i wynosząc przed znaki sum wielkości niezależne od wskaźników sumowania otrzymujemy:

$$\begin{aligned} S_z^2 \cong & \frac{1}{N} \left(\frac{\partial f}{\partial x}\right)_{\bar{x}, \bar{y}}^2 \sum_{i=1}^N (x_i - x_0)^2 + \frac{1}{M} \left(\frac{\partial f}{\partial y}\right)_{\bar{x}, \bar{y}}^2 \sum_{k=1}^M (y_k - y_0)^2 + \\ & + \frac{2}{NM} \left(\frac{\partial f}{\partial x}\right)_{\bar{x}, \bar{y}} \left(\frac{\partial f}{\partial y}\right)_{\bar{x}, \bar{y}} \left(\sum_{i=1}^N (x_i - x_0)\right) \left(\sum_{k=1}^M (y_k - y_0)\right). \end{aligned} \quad (3.4.5)$$

Ostatnią sumę w wyrażeniu (3.4.5) możemy zapisać jako:

$$\left(\sum_{i=1}^N (x_i - x_0) \right) \left(\sum_{k=1}^M (y_k - y_0) \right) = \sum_{i=1}^N \sum_{k=1}^M (x_i - x_0)(y_k - y_0). \quad (3.4.6)$$

Ponieważ w podwójnej sumie (3.4.6) występują iloczyny małych i **niezależnych** wielkości, których znaki mogą być różne, to zgodnie z założeniem o błędach przypadkowych możemy przyjąć, że jest ona równa zero. Wobec tego ostatni składnik we wzorze (3.4.5) może zostać pominięty w porównaniu z poprzednimi wyrazami. Przybliżenie to jest tym lepsze im więcej wykonamy pomiarów. Jeżeli skorzystamy z definicji odchylenia standardowego (3.2.3), to wzór (3.4.5) przyjmie postać:

$$S_z^2 \cong \left(\frac{\partial f}{\partial x} \right)_{\bar{x}, \bar{y}}^2 S_x^2 + \left(\frac{\partial f}{\partial y} \right)_{\bar{x}, \bar{y}}^2 S_y^2. \quad (3.4.7)$$

W przypadku gdy wielkość złożona z jest funkcją r niezależnych wielkości mierzonych bezpośrednio $x_1, x_2, \dots, x_r$ to w tym *przybliżeniu*:

$$S_z^2 \cong \sum_{j=1}^r \left(\frac{\partial f}{\partial x_j} \right)_{\bar{x}_1, \dots, \bar{x}_r}^2 S_{x_j}^2. \quad (3.4.8)$$

Wyrażenie (3.4.8) stanowi ilościowe ujęcie **reguły przenoszenia błędów** przypadkowych dla pomiarów niezależnych. Wynikają z niego następujące podstawowe reguły, a mianowicie jeżeli:

1. $z = x_1 + x_2$, to $S_z^2 = S_{x_1}^2 + S_{x_2}^2$,
2. $z = x_1 - x_2$, to $S_z^2 = S_{x_1}^2 + S_{x_2}^2$,
3. $z = x_1 x_2$, to $(S_z/z)^2 = (S_{x_1}/x_1)^2 + (S_{x_2}/x_2)^2$,
4. $z = x_1/x_2$, to $(S_z/z)^2 = (S_{x_1}/x_1)^2 + (S_{x_2}/x_2)^2$.

Mówimy więc o kwadratowym przenoszeniu błędów przypadkowych, podczas gdy dla błędów maksymalnych mamy przenoszenie liniowe (por. rozdział 2.2). Jeżeli wielkość złożona z będzie funkcją r niezależnych wielkości mierzonych bezpośrednio postaci

$$z = A \prod_{j=1}^r x_j^{a_j},$$

gdzie A oraz a_j są stałymi, to łatwo pokazać, że:

$$\frac{S_z^2}{z^2} = \sum_{j=1}^r \left(a_j \frac{S_{x_j}}{x_j} \right)^2.$$

Korzystając z „reguły przenoszenia błędów” (3.4.8) znajdziemy teraz odchylenie standardowe średniej arytmetycznej serii N pomiarów bezpośrednich. Zgodnie ze wzorem (3.1.8) średnia arytmetyczna jest funkcją wyników wszystkich N pomiarów, możemy to zapisać:

$$\bar{x} = f(x_1, x_2, \dots, x_N) = \frac{1}{N} \sum_{i=1}^N x_i. \quad (3.4.9)$$

Zgodnie ze wzorem (3.4.8) i (3.4.9) $S_{\bar{x}}^2$ jest

$$S_{\bar{x}}^2 = \sum_{i=1}^N \left(\frac{\partial \bar{x}}{\partial x_i} \right)_{x_i}^2 S_{x_i}^2. \quad (3.4.10)$$

Wszystkie wartości S_{x_i} są sobie równe, ponieważ jest to odchylenie standardowe serii i jest stałe dla danej serii. Aby obliczyć $S_{\bar{x}}$ musimy obliczyć pochodne $\partial \bar{x} / \partial x_i$. Różniczkując (3.4.9) względem x_i otrzymujemy $\partial \bar{x} / \partial x_i = 1/N$, czyli że wszystkie pochodne są sobie równe. Wobec tego wzór (3.4.10) przyjmie postać:

$$S_{\bar{x}}^2 = \frac{1}{N^2} N S_x^2$$

A więc **odchylenie standardowe średniej arytmetycznej serii pomiarów** wielkości mierzonej bezpośrednio jest dane wzorem:

$$S_{\bar{x}} = \frac{S_x}{\sqrt{N}}. \quad (3.4.11)$$

Wzór (3.4.8) wyraża wartość odchylenia standardowego serii pomiarów pośrednich. Błędy przypadkowe wielkości mierzonych bezpośrednio przenoszą się, jak przedstawiono powyżej, na ich wartości średnie, zaś błędy średnich arytmetycznych tych wielkości przenoszą się na wartość $\bar{z} = f(\bar{x}_1, \bar{x}_2, \dots, \bar{x}_r)$. Tak więc, podobnie jak w przypadku pomiarów bezpośrednich, miarą precyzji wyznaczenia wartości $\bar{z}$ będzie $S_{\bar{z}}$. Obliczymy teraz odchylenie standardowe wartości średniej serii pomiarów pośrednich $S_{\bar{z}}$. W tym celu skorzystamy z przybliżonej zależności $\bar{z} = f(\bar{x}_1, \bar{x}_2, \dots, \bar{x}_r)$. W związku z tym znaleziona wartość $S_{\bar{z}}$ będzie dotyczyła wartości $\bar{z}$ danej przybliżonym wzorem (3.3.6), a nie wzorem (3.3.2). Dla uproszczenia rachunków przyjmijmy, że $z = f(x, y)$. Wtedy, zgodnie z wzorem (3.3.5) $\bar{z} = f(\bar{x}, \bar{y})$, oraz pamiętając, że $\bar{x}$ i $\bar{y}$ wyrażają się wzorami:

$$\bar{x} = \frac{1}{N}(x_1 + x_2 + \dots + x_N) \quad \text{ i } \quad \bar{y} = \frac{1}{M}(y_1 + y_2 + \dots + y_M), \quad (3.4.12)$$

mamy:

$$\bar{z} = f(\bar{x}(x_1, x_2, \dots, x_N), \bar{y}(y_1, y_2, \dots, y_M)). \quad (3.4.13)$$

Aby obliczyć pochodne $\partial \bar{z} / \partial x_i$ i $\partial \bar{z} / \partial y_k$ musimy we wzorze (3.4.8) uwzględnić to, że funkcja $\bar{z}$ jest funkcją złożoną. Otrzymamy wówczas:

$$S_{\bar{z}}^2 = \sum_{i=1}^N \left(\frac{\partial f}{\partial \bar{x}} \frac{\partial \bar{x}}{\partial x_i} S_{x_i} \right)^2 + \sum_{k=1}^M \left(\frac{\partial f}{\partial \bar{y}} \frac{\partial \bar{y}}{\partial y_k} S_{y_k} \right)^2. \quad (3.4.14)$$

Ponieważ $\partial \bar{x} / \partial x_i = 1/N$ i $\partial \bar{y} / \partial y_k = 1/M$, zaś $S_{x_1} = S_{x_2} = \dots = S_{x_N} = S_x$ oraz $S_{y_1} = S_{y_2} = \dots = S_{y_M} = S_y$ a więc wzór (3.4.14) możemy zapisać:

$$S_{\bar{z}}^2 = \left(\frac{\partial f}{\partial \bar{x}} \right)_{\bar{x}, \bar{y}}^2 \frac{S_x^2}{N} + \left(\frac{\partial f}{\partial \bar{y}} \right)_{\bar{x}, \bar{y}}^2 \frac{S_y^2}{M}. \quad (3.4.15)$$

Stąd po skorzystaniu z wzoru (3.4.11) otrzymujemy:

$$S_{\bar{z}}^2 = \left(\frac{\partial f}{\partial \bar{x}} \right)_{\bar{x}, \bar{y}}^2 S_{\bar{x}}^2 + \left(\frac{\partial f}{\partial \bar{y}} \right)_{\bar{x}, \bar{y}}^2 S_{\bar{y}}^2. \quad (3.4.16)$$

W przypadku gdy wielkość złożona z jest funkcją r niezależnych wielkości mierzonych bezpośrednio $x_1, x_2, \dots, x_r$, wzór (3.4.16) przyjmie postać:

$$S_z^2 = \sum_{j=1}^r \left(\frac{\partial f}{\partial \bar{x}_j} \right)_{\bar{x}_1, \bar{x}_2, \dots, \bar{x}_r}^2 S_{\bar{x}_j}^2. \quad (3.4.17)$$

3.4.2. Pomiary zależne

Przejdziemy teraz do obliczenia odchylenia standardowego serii pomiarów pośrednich, gdy wielkości mierzone bezpośrednio są zależne. Podobnie jak w przypadku wielkości niezależnych przyjmijmy, że wielkość złożona z jest funkcją dwóch zależnych wielkości mierzonych bezpośrednio, tzn. $z = f(x, y)$. Przyjmijmy również, że wykonaliśmy N pomiarów wielkości x oraz y i otrzymaliśmy N par (x_i, y_i) ($i = 1, 2, \dots, N$) wyników pomiarów i N odpowiadających im wartości $z_i = f(x_i, y_i)$ ($i = 1, 2, \dots, N$). W tym przypadku **wariancja serii pomiarów pośrednich** S_z^2 , zgodnie z definicją (3.2.3), wyraża się wzorem:

$$S_z^2 = \frac{1}{N} \sum_{i=1}^N (z_i - z_0)^2. \quad (3.4.18)$$

Różnicę $z_i - z_0$ obliczamy, podobnie jak w przypadku pomiarów niezależnych (patrz rozdział 3.4.1), korzystając z rozwinięcia funkcji $z = f(x, y)$

na szereg Taylora w otoczeniu punktu $(\bar{x}, \bar{y})$ i ograniczając rozwinięcie do wyrazów liniowych w zmiennych x i y . Otrzymamy wyrażenie analogiczne do wzoru (3.4.4), a mianowicie

$$z_{ik} - z_0 = \left(\frac{\partial f}{\partial x}\right)_{\bar{x}, \bar{y}} (x_i - x_0) + \left(\frac{\partial f}{\partial y}\right)_{\bar{x}, \bar{y}} (y_k - y_0). \quad (3.4.19)$$

Podstawiając (3.4.19) do (3.4.18) otrzymujemy:

$$\begin{aligned} S_z^2 = & \frac{1}{N} \left(\frac{\partial f}{\partial x}\right)_{\bar{x}, \bar{y}}^2 \sum_{i=1}^N (x_i - x_0)^2 + \frac{1}{N} \left(\frac{\partial f}{\partial y}\right)_{\bar{x}, \bar{y}}^2 \sum_{i=1}^N (y_i - y_0)^2 \\ & + \frac{2}{N} \left(\frac{\partial f}{\partial x}\right)_{\bar{x}, \bar{y}} \left(\frac{\partial f}{\partial y}\right)_{\bar{x}, \bar{y}} \sum_{i=1}^N (x_i - x_0)(y_i - y_0). \end{aligned} \quad (3.4.20)$$

Przerwijmy teraz dalsze obliczanie S_z^2 i zanim będziemy je kontynuować zajmijmy się ostatnim składnikiem w wyrażeniu (3.4.20). Składnik ten, w rozpatrywanym przypadku **pomiarów zależnych**, nie może być pominięty, jak to wystąpiło, gdy pomiary wielkości mierzonych bezpośrednio były niezależne. Na początku założmy, że błędy przypadkowe mierzonych wielkości x i y można pominąć. Wtedy poszczególne składniki sumy

$$\sum_{i=1}^N (x_i - x_0)(y_i - y_0)$$

będą miały ten sam znak, ponieważ zmiana wielkości x powoduje zmianę wielkości y (pomiary są zależne!) i odwrotnie. Wszystkie składniki będą miały znak *plus*, gdy wzrostowi (spadkowi) wartości x odpowiada wzrost (spadek) wartości y . W przypadku gdy wzrostowi x odpowiada spadek y lub na odwrót, to wszystkie składniki sumy będą miały znak *minus*. Jeżeli zakres zmian wielkości zależnych mierzonych bezpośrednio jest taki, że rozwinięcie (3.4.20) jest dobrym przybliżeniem lub zależność między nimi jest liniowa, to suma ta ma wartość maksymalną i wnosi istotny wkład do wartości S_z^2 . W przypadku gdy błędów przypadkowych nie możemy pominąć, składniki tej sumy mogą mieć różne znaki i wartość jej będzie mniejsza. Jednak pominąć jej nie będzie można.

Po tych uwagach przejdźmy do obliczenia S_z^2 . Korzystając z definicji (3.2.3) odchylenia standardowego serii pomiarów bezpośrednich znajdujemy

$$S_z^2 = \left(\frac{\partial f}{\partial x}\right)_{\bar{x}, \bar{y}}^2 S_x^2 + \left(\frac{\partial f}{\partial y}\right)_{\bar{x}, \bar{y}}^2 S_y^2 + 2 \left(\frac{\partial f}{\partial x}\right)_{\bar{x}, \bar{y}} \left(\frac{\partial f}{\partial y}\right)_{\bar{x}, \bar{y}} C_{x,y}, \quad (3.4.21)$$

gdzie

$$C_{x,y} = \frac{1}{N} \sum_{i=1}^N (x_i - x_0)(y_i - y_0), \quad (3.4.22)$$

i nazywa się **kowariancją** serii pomiarów wielkości x i y .

Kowariancji serii pomiarów wielkości x i y danej wzorem (3.4.22) nie możemy jednak obliczyć, ponieważ nie są znane wartości x_0 i y_0 . Postępując tak jak w rozdziale 3.2, można pokazać, że gdy zastąpimy z_0 przez $\bar{z}$, to wariancja serii S_z^2 wyrazi się wzorem:

$$S_z^2 = \frac{1}{N-1} \sum_{i=1}^N (z_i - \bar{z})^2. \quad (3.4.23)$$

Jeżeli teraz wartość wielkości złożonej $z_i = f(x_i, y_i)$ obliczymy z rozwinięcia funkcji $z = f(x, y)$ na szereg Taylora funkcji dwóch zmiennych w otoczeniu punktu $(\bar{x}, \bar{y})$ i ograniczymy się do wyrazów liniowych w zmiennych x oraz y , to podstawiając otrzymany wynik do (3.4.23) dostaniemy

$$C_{x,y} = \frac{1}{N-1} \sum_{i=1}^N (x_i - \bar{x})(y_i - \bar{y}). \quad (3.4.24)$$

Kowariancja serii $C_{x,y}$ jest miarą zależności zmiennych x i y . Może ona przyjmować dowolne wartości i z tego powodu jest niewygodna w użyciu. W związku z tym wprowadza się inną bezwymiarową wielkość

$$\rho = \frac{C_{x,y}}{S_x S_y}, \quad (3.4.25)$$

która nosi nazwę **współczynnika korelacji liniowej** serii pomiarów wielkości x i y .

W przypadku gdy rozwinięcie (3.4.20) jest dobrym przybliżeniem lub gdy zależność między wielkościami fizycznymi x oraz y mierzonymi bezpośrednio jest liniowa i błędy pomiaru można pominąć, to jak można pokazać $\rho = \pm 1$. Znak plus dotyczy przypadku, gdy wzrostowi wielkości x odpowiada wzrost wielkości y i na odwrót; znak minus występuje wówczas, gdy wzrostowi x odpowiada spadek y lub odwrotnie. Jeżeli wielkości x i y są niezależne to $\rho = 0$. Błędy pomiaru zmniejszają wartość współczynnika korelacji liniowej ρ . Przy dużych błędach pomiaru $|\rho|$ może być znacznie mniejsza od 1, chociaż warunki stosowalności rozwinięcia (3.4.20) są spełnione i wielkości x oraz y są zależne. W bardziej zaawansowanych opracowaniach pokazuje się, że współczynnik korelacji liniowej serii pomiarów, gdy występują błędy przypadkowe, spełnia nierówność:

$$|\rho| \leq 1. \quad (3.4.26)$$

Współzależność wielkości fizycznych *może, ale nie musi*, być znana w postaci zależności funkcyjnej. Na przykład: wielkość plonów zbóż jest zależna od wielkości wiosennych opadów deszczu. W tym przypadku występuje współzależność, ale nie znamy jej postaci funkcyjnej. Ogólnie mówiąc, współzależność między wielkościami fizycznymi niezależnie od tego, czy jej przyczyny bądź postać funkcyjną znamy, czy nie, nazywamy **korelacją**, a wielkości **skorelowanymi**. W celu stwierdzenia, czy mierzone wielkości są skorelowane, musimy obliczyć współczynnik korelacji ρ .

Gdy wielkość złożona z jest funkcją r zależnych wielkości mierzonych bezpośrednio $x_1, \dots, x_r$, to

$$S_z^2 = \sum_{j=1}^r \left(\frac{\partial f}{\partial x_j} \right)_{\bar{x}_1, \dots, \bar{x}_r}^2 S_{x_j}^2 + \sum_{k \neq j} \left(\frac{\partial f}{\partial x_j} \right)_{\bar{x}_1, \dots, \bar{x}_r} \left(\frac{\partial f}{\partial x_k} \right)_{\bar{x}_1, \dots, \bar{x}_r} C_{x_j, x_k}. \quad (3.4.27)$$

Przykłady

1. Wyznaczanie przyspieszenia ziemskiego przy użyciu wahadła matematycznego — pomiary niezależne

W celu ilustracji rozważań dotyczących pomiarów niezależnych przeanalizujemy wyznaczanie przyspieszenia ziemskiego g za pomocą wahadła matematycznego. Zgodnie ze wzorem (1.1.3)

$$g = 4\pi^2 \frac{l}{T^2},$$

gdzie l jest długością wahadła, a T jego okresem. Aby wyznaczyć g musimy wykonać pomiary długości wahadła l i jego okresu T . Mierząc długość wahadła otrzymaliśmy następujące wartości:

$$l[\text{cm}] = 100.1, 99.8, 100.2, 100.0, 100.3, 99.8, 99.9,$$

czyli że seria pomiarów zawierała $N = 7$ wartości. Pomiary okresu drgań wahadła wykonaliśmy $M = 10$ razy i otrzymaliśmy następujące wyniki:

$$T[\text{s}] = 2.02, 1.98, 2.01, 2.00, 2.03, 1.99, 2.01, 2.01, 2.00, 2.02.$$

Tak dokładne wartości możemy uzyskać mierząc czas np. 10 okresów elektronicznym miernikiem czasu o dokładności 0.01 s. Wszystkie obliczenia wykonujemy za pomocą kalkulatora. W obliczeniach przyjęto $\pi = 3.1415927$. Ponieważ w rozdziale 5 zostaną omówione związki wartości średniej i odchyłeń standardowych z wartością rzeczywistą i błędem pomiaru przedstawione poniżej wyniki obliczeń zostały podane bez zaokrągleń, tzn. tak jak zostały odczytane z kalkulatora. Korzystając ze wzoru (3.1.8) znajdujemy:

$$\bar{l} = 100.01428 \text{ cm} \quad \text{oraz} \quad \bar{T} = 2.007 \text{ s},$$

zaś na podstawie wzoru (3.2.13):

$$S_l = 0.1951802 \text{ cm} \quad \text{i} \quad S_T = 0.0149443 \text{ s},$$

oraz na podstawie wzoru (3.4.11):

$$S_{\bar{l}} = 0.0737712 \text{ cm} \quad \text{i} \quad S_{\bar{T}} = 0.0047258 \text{ s}.$$

Zgodnie ze wzorem (3.3.5) wartość średnia $\bar{g}$ jest równa: $\bar{g} = 4\pi^2 \bar{l} / \bar{T}^2$. Podstawiając obliczone powyżej wartości $\bar{l}$ i $\bar{T}$ otrzymujemy:

$$\bar{g} = 980.22503 \text{ cm/s}^2.$$

Dla porównania, $\bar{g}$ obliczone ze wzoru (3.3.2) wynosi $\bar{g} = 980.37505 \text{ cm/s}^2$, różnica jak widać jest niewielka, a więc stosowane przybliżenie przy wprowadzeniu wzoru (3.3.5) jest w tym przypadku uzasadnione. Obliczając pochodne $\partial g(\bar{l}, \bar{T}) / \partial l$ oraz $\partial g(\bar{l}, \bar{T}) / \partial T$, a następnie podstawiając ich wartości liczbowe i obliczone wartości S_l , S_T do wzoru (3.4.8), otrzymujemy:

$$S_{\bar{g}} = \frac{\partial g(\bar{l}, \bar{T})}{\partial l} S_l + \frac{\partial g(\bar{l}, \bar{T})}{\partial T} S_T \quad S_{\bar{g}} = 14.722557 \text{ cm/s}^2.$$

Odchylenie standardowe wartości $\bar{g}$ obliczamy ze wzoru (3.4.17) podstawiając obliczone uprzednio wartości pochodnych i odchyłeń standardowych $S_{\bar{l}}$ oraz $S_{\bar{T}}$. Wynosi ono:

$$S_{\bar{g}} = 4.6690808 \text{ cm/s}^2.$$

Do przykładu tego wrócimy w rozdziale 5.1.6, 5.2 i 6.3.

2. Wyznaczanie oporu elektrycznego przy użyciu prawa Ohma — pomiary zależne

W celu wyznaczenia wartości oporu omowego skorzystano z prawa Ohma $R = U/I$, gdzie U — spadek napięcia na oporniku, I — natężenie płynącego prądu. Pomiary wykonano zmieniając wartość natężenia prądu I (w takim zakresie, aby rozwinięcie (3.3.8) było dobrym przybliżeniem) i wykonując pomiary U oraz I . Przy takim sposobie wykonania pomiarów jest to pomiar wielkości zależnych. Otrzymano następujące wyniki:

I [A]	0.43	0.46	0.49	0.51	0.52	0.55
U [V]	11.0	11.5	12	12.5	12.3	13.8
R [Ω]	25.6	25.0	24.5	24.5	23.6	25.1

Obliczone ze wzoru (3.1.8) średnie arytmetyczne wynoszą $\bar{I} = 0.493 \text{ A}$, $\bar{U} = 12.18 \text{ V}$, $\bar{R} = 24.72 \Omega$. Średnia wartość $\bar{R}$ obliczona ze wzoru (3.3.7) $\bar{R} = \bar{U} / \bar{I} = 24.71 \Omega$. Odchylenia standardowe serii pomiarów bezpośrednich napięcia i natężenia prądu obliczone ze wzoru (3.2.13) wynoszą: $S_I = 0.043 \text{ A}$ oraz $S_U = 0.96 \text{ V}$. Kowariancję $C_{I,U}$ obliczamy ze wzoru

(3.4.25), wynosi ona $C_{I,U} = 0.04064 \text{ V}\cdot\text{A}$. Odchylenie standardowe serii pomiarów pośrednich otrzymamy korzystając ze wzoru (3.4.21). Wynosi ono

$$S_R = \sqrt{\left(\frac{1}{\bar{I}}\right)^2 S_U^2 + \left(\frac{\bar{U}}{\bar{I}^2}\right)^2 S_I^2 - 2 \frac{\bar{U}}{\bar{I}^3} C_{I,U}}.$$

Podstawiając obliczone wartości $\bar{I}$, $\bar{U}$, S_U , S_I oraz $C_{I,U}$ otrzymujemy

$$S_R = 0.42 \Omega.$$

Obliczmy jeszcze ze wzoru (3.4.25) współczynnik korelacji ρ , jego wartość wynosi 0.98.

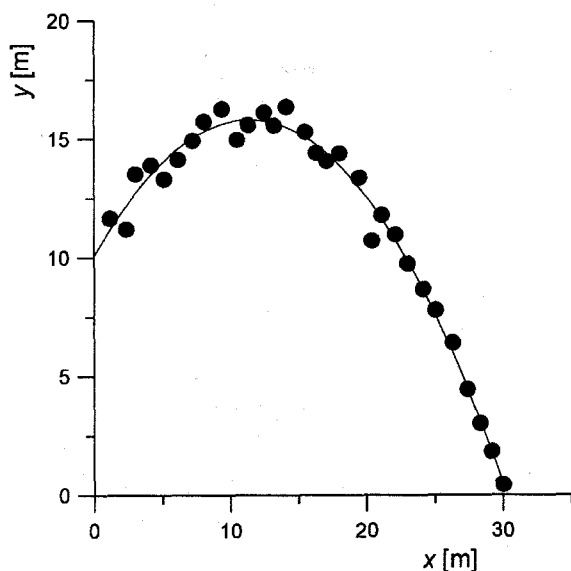
3.5. Metoda najmniejszych kwadratów

W rozdziale 3.1 korzystając z metody najmniejszych kwadratów znaleźliśmy wartość, wokół której grupują się wyniki pomiarów danej wielkości fizycznej obciążonej błędami przypadkowymi. W tym rozdziale podamy inne jej zastosowania.

W wyniku serii N pomiarów dwóch wielkości fizycznych x i y , o których sądzimy, że są zależne, otrzymujemy pary liczb (x_i, y_i) ($i = 1, 2, \dots, N$). Jeżeli te wielkości są zależne i pomiary zostały wykonane w dużym zakresie zmian tych wielkości, to na wykresie w układzie współrzędnych XY wyniki pomiarów powinny układać się wokół jakiejś krzywej $y = f(x)$. Ogólna postać funkcji $y = f(x)$ może być *znana* lub na podstawie wykresu zależności $y = f(x)$ możemy *sądzić* jaka jest postać tej zależności. W obu przypadkach do analizy zależności $y = f(x)$ możemy stosować metodę najmniejszych kwadratów. Omawiając zastosowania metody najmniejszych kwadratów, na początku założymy, że *ogólną postać zależności funkcyjnej znamy*, a w rozdziale 3.5.3 krótko omówimy przypadek, gdy rzeczywista zależność $y = f(x)$ nie jest znana i będziemy ją przybliżać funkcją, której postać zakładamy.

Zanim przejdziemy do przedstawienia przypadku, gdy znamy ogólną postać funkcji $y = f(x)$, rozpatrzmy następujący przykład.

Z wieży o wysokości H_0 został rzucony kamień pod kątem α do poziomu z prędkością początkową v_0 . Dokonujemy pomiaru zależności wysokości y kamienia nad powierzchnią ziemi od odległości x od podstawy wieży, tzn., że przy ustalonych odległościach $x_1, x_2, \dots, x_N$ wykonujemy pomiary wysokości $y_1, y_2, \dots, y_N$. Inaczej mówiąc dla danej odległości od podstawy wieży x_i mierzymy wysokość y_i , na której znajduje się kamień. Mierzmy więc dwie wielkości i jako wynik pomiaru otrzymujemy pary liczb (x_i, y_i) ($i = 1, 2, \dots, N$). Przykład wyników pomiarów przedstawiono na rysunku 3.3.



Rys. 3.3: Przykład wyników pomiarów położenia w rzucie ukośnym z wieży o wysokości $H_0 = 10$ m. • punkty zmierzone, linia ciągła przedstawia obliczony z równań ruchu tor kamienia

Jak widać z rysunku 3.3, wyniki pomiarów układają się *wokół* pewnej krzywej. Z rozwiązania równań ruchu dla rzutu ukośnego wynika, że tor kamienia (jeżeli pominiemy opór powietrza) jest parabolą o równaniu $y = c_1 + c_2x + c_3x^2$, gdzie współczynniki c_1, c_2, c_3 są związane z wielkościami H_0, v_0, α, g , a mianowicie: $c_1 = H_0, c_2 = \text{tg}\alpha, c_3 = g/(2v_0^2 \cos^2 \alpha)$. Aby wyznaczyć tę parabolę musimy znać wartości współczynników c_j ($j = 1, 2, 3$). Jeżeli znamy wartości H_0, v_0, α i g , to możemy je obliczyć korzystając z podanych wyżej zależności; w przypadku gdy któreś z tych wielkości nie są znane, współczynniki c_j możemy wyznaczyć metodą najmniejszych kwadratów na podstawie zmierzonych wartości x_i, y_i . Metoda ta wymaga bowiem, aby suma kwadratów odchyłań zmierzonych wartości y_i od nieznannej funkcji $c_1 + c_2x + c_3x^2$ była minimalna, tzn. aby

$$\sum_{i=1}^N [y_i - (c_1 + c_2x_i + c_3x_i^2)]^2 = \min.$$

Wartość tej sumy jest zależna od wartości współczynników c_j , tzn. jest ich funkcją $G(c_1, c_2, c_3)$. Wobec tego wyrażenie analogiczne do (3.1.4) przyjmie postać

$$G(c_1, c_2, c_3) = \sum_{i=1}^N [y_i - (c_1 + c_2x_i + c_3x_i^2)]^2 = \min. \quad (3.5.1)$$

Oznacza to, że musimy tak *dobrać* współczynniki c_j , aby $G(c_1, c_2, c_3) = \min$.

Zwróćmy uwagę, że z przytoczonego przykładu wynika, iż w przypadku ogólnym funkcja $f(x)$ jest zależna nie tylko od zmiennej x , ale również od Q współczynników $c_1, c_2, \dots, c_Q$, czyli $y = f(x, c_1, c_2, \dots, c_Q)$. W ogólnym więc przypadku wyrażenie (3.5.1) należy zapisać:

$$G(c_1, c_2, \dots, c_Q) = \sum_{i=1}^N [y_i - f(x_i, c_1, c_2, \dots, c_Q)]^2 = \min. \quad (3.5.2)$$

Ponieważ minimalizujemy wyrażenie (3.5.2) względem współczynników c_j ($j = 1, 2, \dots, Q$) są one **parametrami funkcji** $f(x, c_1, c_2, \dots, c_Q)$. Minimalizacji funkcji $G(c_1, c_2, \dots, c_Q)$ dokonujemy znaną z analizy matematycznej metodą, tzn. różniczkując ją względem parametrów c_j ($j = 1, 2, \dots, Q$) i przyrównując pochodne do zera. Otrzymujemy wtedy układ Q równań algebraicznych:

$$\begin{aligned} \left(\frac{\partial G}{\partial c_1} \right)_{c_1, c_2, \dots, c_Q} &= 0 \\ &\vdots \\ \left(\frac{\partial G}{\partial c_j} \right)_{c_1, c_2, \dots, c_Q} &= 0 \\ &\vdots \\ \left(\frac{\partial G}{\partial c_Q} \right)_{c_1, c_2, \dots, c_Q} &= 0 \end{aligned} \quad (3.5.3)$$

Rozwiązując go znajdujemy na podstawie zmierzonych wartości x_i , y_i ($i = 1, 2, \dots, N$) współczynniki c_j ($j = 1, 2, \dots, Q$) i tym samym znajdujemy funkcję $f(x, c_1, c_2, \dots, c_Q)$, wokół której układają się wyniki pomiarów. Inaczej mówiąc *dopasowujemy funkcję* $f(x, c_1, c_2, \dots, c_Q)$ do wyników pomiarów. Funkcja $f(x, c_1, c_2, \dots, c_Q)$ ze współczynnikami c_j obliczonymi z równań (3.5.3) została tak dopasowana do wyników doświadczalnych, aby zależność (3.5.2) była spełniona. Wyznaczone współczynniki c_j są funkcjami zmierzonych wartości x_i , y_i , czyli $c_j = c_j(x_1, x_2, \dots, x_N, y_1, y_2, \dots, y_N)$ i dlatego są obarczone błędami spowodowanymi niepewnościami pomiarowymi mierzonych wielkości. Z tego powodu w ogólności nie są one równe rzeczywistym współczynnikom. Jak się przekonamy w rozdziale 5.4, stanowią one *najlepsze przybliżenie współczynników rzeczywistych* i dlatego w dalszych rozważaniach będziemy je oznaczali przez $\hat{c}_j$.

Przedstawione powyżej rozważania są słuszne w przypadku, gdy błędy pomiaru wartości x_i są na tyle małe, że możemy je pominąć. W dalszych

rozważaniach tego rozdziału będziemy przyjmowali, że założenie to jest spełnione.

Odchylenia standardowe $S_{\hat{c}_j}$ możemy obliczyć korzystając z reguły przenoszenia błędów przypadkowych (3.4.8), ponieważ przyjęliśmy powyżej, że błędy zmierzonych wartości x_i można pominąć, zaś pomiary wielkości y są między sobą niezależne. Wobec tego na wyznaczone współczynniki $\hat{c}_j$ przenoszą się jedynie błędy pomiaru wartości y_i , a więc wzór (3.4.8) przyjmie postać:

$$S_{\hat{c}_j}^2 = \sum_{k=1}^N \left(\frac{\partial \hat{c}_j}{\partial y_k} \right)^2 S_y^2, \quad (3.5.4)$$

gdzie

$$S_y^2 = \frac{1}{N} \sum_{i=1}^N (y_i - f(x_i, c_1, c_2, \dots, c_Q))^2. \quad (3.5.5)$$

Wyznaczone z układu równań (3.5.3) wartości współczynników $\hat{c}_j$ ($j = 1, 2, \dots, Q$) są od siebie zależne ponieważ zostały obliczone na podstawie tych samych wielkości mierzonych bezpośrednio. Na podstawie tego co powiedziano w Dodatku C **kowariancja** pary współczynników $\hat{c}_i, \hat{c}_j$, w rozpatrywanym przypadku, gdy błędy pomiaru wielkości fizycznej x można pominąć, wyraża się wzorem:

$$C_{\hat{c}_i, \hat{c}_j} = \sum_{k=1}^N \left(\frac{\partial \hat{c}_i}{\partial y_k} \right) \left(\frac{\partial \hat{c}_j}{\partial y_k} \right) S_y^2, \quad (3.5.6)$$

gdzie S_y^2 dane jest wzorem (3.5.5). Różnica $y_i - f(x_i, c_1, c_2, \dots, c_Q) = \delta_i$ jest błędem bezwzględnym zmierzonej wartości y_i , bo $f(x_i, c_1, c_2, \dots, c_Q)$ jest wartością rzeczywistą wielkości fizycznej y . Ponieważ nie znamy rzeczywistych wartości współczynników c_j , musimy je zastąpić wartościami przybliżonymi $\hat{c}_j$. Zgodnie z tym, co powiedziano w rozdziale 3.2, liczba stopni swobody $k = N - Q$, ponieważ mamy Q równań, z których wyznaczamy współczynniki $\hat{c}_j$ ($j = 1, 2, \dots, Q$). Wobec tego

$$S_y^2 = \frac{1}{N - Q} \sum_{i=1}^N (y_i - f(x_i, \hat{c}_1, \hat{c}_2, \dots, \hat{c}_Q))^2. \quad (3.5.7)$$

Ze wzoru (3.5.7) wynika, że liczba stopni swobody $k = N - Q \geq 1$, aby odchylenie standardowe serii pomiarów wielkości zależnych S_y było określone. Tak więc, jeśli dobieramy Q parametrów do N danych doświadczalnych, to liczba N musi spełniać nierówność $N \geq Q + 1$. Odchylenia standardowe

$S_{\hat{c}_j}$ mówią nam, jak dobrymi przybliżeniami wartości rzeczywistych współczynników c_j są wyznaczone wartości $\hat{c}_j$. Odchylenia te określają, jak błędy przypadkowe wartości y_i przenoszą się na wartości współczynników c_j .

Wyznaczona zależność $\hat{y} = f(x, \hat{c}_1, \hat{c}_2, \dots, \hat{c}_Q)$ może być użyta do obliczenia wartości $\hat{y}(x)$ w innych punktach niż były wykonywane pomiary. Na obliczoną wartość $\hat{y}(x)$ przenoszą się błędy wyznaczenia współczynników $\hat{c}_1, \hat{c}_2, \dots, \hat{c}_Q$, spowodowane, jak wykazano wyżej, błędami pomiarów wielkości y_i ($i = 1, 2, \dots, N$) mierzonych bezpośrednio. Aby obliczyć $S_{\hat{y}}^2$, musimy skorzystać ze wzoru (3.4.27), który jest matematycznym sformulowaniem reguły przenoszenia błędów przypadkowych dla wielkości zależnych. Otrzymamy wówczas:

$$S_{\hat{y}}^2 = \sum_{j=1}^Q \left(\frac{\partial f}{\partial \hat{c}_j} \right)_{\hat{c}_1, \hat{c}_2, \dots, \hat{c}_Q}^2 S_{\hat{c}_j}^2 + \sum_{i \neq j}^Q \left(\frac{\partial f}{\partial \hat{c}_i} \right)_{\hat{c}_1, \dots, \hat{c}_Q} \left(\frac{\partial f}{\partial \hat{c}_j} \right)_{\hat{c}_1, \dots, \hat{c}_Q} C_{\hat{c}_i, \hat{c}_j}. \quad (3.5.8)$$

3.5.1. Dopasowanie funkcji liniowych

Dopasowanie funkcji liniowej do wyników doświadczalnych jest najprostszym i bardzo często spotykanym w praktyce przypadkiem zastosowania metody najmniejszych kwadratów. Jeżeli do zbioru punktów (x_i, y_i) ($i = 1, 2, \dots, N$) chcemy metodą najmniejszych kwadratów dopasować prostą o równaniu:

$$y = ax + b, \quad (3.5.9)$$

to wyrażenie (3.5.2) przyjmie postać:

$$G(a, b) = \sum_{i=1}^N (y_i - ax_i - b)^2 = \min. \quad (3.5.10)$$

Różniczkując funkcję $G(a, b)$ względem parametrów a oraz b i przyrównując pochodne do zera, otrzymamy:

$$\begin{aligned} \frac{\partial G(a, b)}{\partial a} &= 2 \sum_{i=1}^N (-x_i)(y_i - ax_i - b) = 0 \\ \frac{\partial G(a, b)}{\partial b} &= 2 \sum_{i=1}^N (-1)(y_i - ax_i - b) = 0. \end{aligned} \quad (3.5.11)$$

Przekształcając powyższy układ równań algebraicznych otrzymujemy:

$$\begin{aligned} a \sum_{i=1}^N x_i^2 + b \sum_{i=1}^N x_i - \sum_{i=1}^N x_i y_i &= 0 \\ a \sum_{i=1}^N x_i + bN - \sum_{i=1}^N y_i &= 0. \end{aligned} \quad (3.5.12)$$

Rozwiązując układ równań (3.5.12) znajdujemy przybliżone wartości współczynników, tj. $\hat{a}$ i $\hat{b}$. Są one równe:

$$\hat{a} = \frac{N \sum_{i=1}^N x_i y_i - \left(\sum_{i=1}^N x_i \right) \left(\sum_{i=1}^N y_i \right)}{N \sum_{i=1}^N x_i^2 - \left(\sum_{i=1}^N x_i \right)^2}, \quad (3.5.13)$$

$$\hat{b} = \frac{\left(\sum_{i=1}^N x_i^2 \right) \left(\sum_{i=1}^N y_i \right) - \left(\sum_{i=1}^N x_i \right) \left(\sum_{i=1}^N x_i y_i \right)}{N \sum_{i=1}^N x_i^2 - \left(\sum_{i=1}^N x_i \right)^2}. \quad (3.5.14)$$

Współczynniki $\hat{a}$ i $\hat{b}$ są funkcjami zależnymi od zmierzonych wartości x_i , y_i ($i = 1, 2, \dots, N$), tj. $\hat{a} = g_1(x, y)$ i $\hat{b} = g_2(x, y)$. Zmierzone wartości x_i, y_i ($i = 1, 2, \dots, N$) są obarczone błędami przypadkowymi, wobec tego również wyznaczone wartości współczynników $\hat{a}$ i $\hat{b}$ będą obarczone błędami. Ponieważ przyjęliśmy, że błąd pomiaru zmiennej niezależnej (w naszym przypadku wielkości fizycznej x) jest tak mały, iż możemy go pominąć, to zgodnie ze wzorem (3.5.4):

$$S_{\hat{a}}^2 = \sum_{k=1}^N \left(\frac{\partial \hat{a}}{\partial y_k} \right)^2 S_y^2 = S_y^2 \sum_{k=1}^N \left(\frac{\partial \hat{a}}{\partial y_k} \right)^2, \quad (3.5.15)$$

oraz

$$S_{\hat{b}}^2 = \sum_{k=1}^N \left(\frac{\partial \hat{b}}{\partial y_k} \right)^2 S_y^2 = S_y^2 \sum_{k=1}^N \left(\frac{\partial \hat{b}}{\partial y_k} \right)^2, \quad (3.5.16)$$

a zgodnie ze wzorem (3.5.6)

$$C_{\hat{a}, \hat{b}} = \sum_{k=1}^N \left(\frac{\partial \hat{a}}{\partial y_k} \right) \left(\frac{\partial \hat{b}}{\partial y_k} \right) S_y^2 = S_y^2 \sum_{k=1}^N \left(\frac{\partial \hat{a}}{\partial y_k} \right) \left(\frac{\partial \hat{b}}{\partial y_k} \right), \quad (3.5.17)$$

gdzie:

$$S_y^2 = \frac{1}{N} \sum_{i=1}^N (y_i - a x_i - b)^2. \quad (3.5.18)$$

Ponieważ wartości rzeczywistych współczynników a i b nie znamy, musimy więc zastąpić je przez wartości przybliżone $\hat{a}$ oraz $\hat{b}$. Zgodnie z tym, co

powiedziano wyżej i w rozdziale 3.2, w tym przypadku **liczba stopni swobody** będzie $k = N - 2$ (ponieważ na podstawie pomiarów wyznaczamy dwie wielkości, $\hat{a}$ oraz $\hat{b}$), a więc:

$$S_y^2 = \frac{1}{N-2} \sum_{i=1}^N (y_i - \hat{a}x_i - \hat{b})^2 = \frac{1}{N-2} \sum_{i=1}^N d_i^2, \quad (3.5.19)$$

gdzie $d_i = y_i - \hat{a}x_i - \hat{b}$.

Wyprowadzenie wzoru (3.5.19) podano w Dodatku D. Obliczmy teraz sumy kwadratów pochodnych cząstkowych $\sum_{k=1}^N \left(\frac{\partial \hat{a}}{\partial y_k}\right)^2$ i $\sum_{k=1}^N \left(\frac{\partial \hat{b}}{\partial y_k}\right)^2$. Różniczkując (3.5.13), otrzymujemy:

$$\frac{\partial \hat{a}}{\partial y_k} = \frac{Nx_k - \sum_{i=1}^N x_i}{N \sum_{i=1}^N x_i^2 - \left(\sum_{i=1}^N x_i\right)^2}. \quad (3.5.20)$$

A zatem:

$$\begin{aligned} \sum_{k=1}^N \left(\frac{\partial \hat{a}}{\partial y_k}\right)^2 &= \sum_{k=1}^N \left[\frac{Nx_k - \sum_{i=1}^N x_i}{N \sum_{i=1}^N x_i^2 - \left(\sum_{i=1}^N x_i\right)^2} \right]^2 \\ &= \frac{1}{\left[N \sum_{i=1}^N x_i^2 - \left(\sum_{i=1}^N x_i\right)^2\right]^2} \sum_{k=1}^N \left(N^2 x_k^2 - 2Nx_k \sum_{i=1}^N x_i - \left(\sum_{i=1}^N x_i\right)^2 \right) \\ &= \frac{N^2 \sum_{k=1}^N x_k^2 - 2N \left(\sum_{k=1}^N x_k\right) \sum_{i=1}^N x_i - N \left(\sum_{i=1}^N x_i\right)^2}{\left[N \sum_{i=1}^N x_i^2 - \left(\sum_{i=1}^N x_i\right)^2\right]^2}. \end{aligned}$$

Biorąc pod uwagę, że wartość sumy jest niezależna od wskaźnika sumowania (tzn. $\sum_{k=1}^N x_k = \sum_{i=1}^N x_i$) oraz wnosząc w liczniku N przed nawias i skracając ułamek ostatecznie dostajemy:

$$\sum_{k=1}^N \left(\frac{\partial \hat{a}}{\partial y_k}\right)^2 = \frac{N}{N \sum_{i=1}^N x_i^2 - \left(\sum_{i=1}^N x_i\right)^2}. \quad (3.5.21)$$

Wstawiając (3.5.19) i (3.5.21) do (3.5.15) znajdujemy:

$$S_a^2 = S_y^2 \frac{N}{N \sum_{i=1}^N x_i^2 - \left(\sum_{i=1}^N x_i \right)^2} = \frac{N \sum_{i=1}^N d_i^2}{(N-2) \left[N \sum_{i=1}^N x_i^2 - \left(\sum_{i=1}^N x_i \right)^2 \right]}. \quad (3.5.22)$$

Aby obliczyć S_b^2 różniczkujemy (3.5.14). Pochodna $\partial \hat{b} / \partial y_k$ jest równa:

$$\frac{\partial \hat{b}}{\partial y_k} = \frac{\sum_{i=1}^N x_i^2 - x_k \sum_{i=1}^N x_i}{N \sum_{i=1}^N x_i^2 - \left(\sum_{i=1}^N x_i \right)^2}. \quad (3.5.23)$$

Wobec tego:

$$\begin{aligned} \sum_{k=1}^N \left(\frac{\partial \hat{b}}{\partial y_k} \right)^2 &= \sum_{k=1}^N \frac{\left(\sum_{i=1}^N x_i^2 \right)^2 - 2x_k \left(\sum_{i=1}^N x_i \right) \left(\sum_{i=1}^N x_i^2 \right) + x_k^2 \left(\sum_{i=1}^N x_i \right)^2}{\left[N \sum_{i=1}^N x_i^2 - \left(\sum_{i=1}^N x_i \right)^2 \right]^2} \\ &= \frac{N \left(\sum_{i=1}^N x_i^2 \right)^2 - 2 \left(\sum_{k=1}^N x_k \right) \left(\sum_{i=1}^N x_i \right) \left(\sum_{i=1}^N x_i^2 \right) + \left(\sum_{k=1}^N x_k^2 \right) \left(\sum_{i=1}^N x_i \right)^2}{\left[N \sum_{i=1}^N x_i^2 - \left(\sum_{i=1}^N x_i \right)^2 \right]^2} \\ &= \frac{N \left(\sum_{i=1}^N x_i^2 \right)^2 - \left(\sum_{i=1}^N x_i \right)^2 \left(\sum_{i=1}^N x_i^2 \right)}{\left[N \sum_{i=1}^N x_i^2 - \left(\sum_{i=1}^N x_i \right)^2 \right]^2} \end{aligned}$$

Wynosząc w liczniku $\sum_{i=1}^N x_i^2$ przed nawias i skracając ułamek znajdujemy:

$$\sum_{k=1}^N \left(\frac{\partial \hat{b}}{\partial y_k} \right)^2 = \frac{\sum_{i=1}^N x_i^2}{N \sum_{i=1}^N x_i^2 - \left(\sum_{i=1}^N x_i \right)^2}. \quad (3.5.24)$$

Wstawiając (3.5.19) i (3.5.24) do (3.5.16), otrzymujemy:

$$S_b^2 = S_y^2 \frac{\sum_{i=1}^N x_i^2}{N \sum_{i=1}^N x_i^2 - \left(\sum_{i=1}^N x_i \right)^2} = \frac{N \left(\sum_{i=1}^N d_i^2 \right) \left(\sum_{i=1}^N x_i^2 \right)}{(N-2) \left[N \sum_{i=1}^N x_i^2 - \left(\sum_{i=1}^N x_i \right)^2 \right]}. \quad (3.5.25)$$

Kowariancję $C_{\hat{a}, \hat{b}}$ współczynników $\hat{a}$ i $\hat{b}$ obliczamy wstawiając (3.5.20) i (3.5.23) do (3.5.17). Otrzymujemy:

$$\begin{aligned} C_{\hat{a}, \hat{b}} &= S_y^2 \sum_{k=1}^N \left(\frac{\partial \hat{a}}{\partial y_k} \right) \left(\frac{\partial \hat{b}}{\partial y_k} \right) = S_y^2 \sum_{k=1}^N \frac{\left(N x_k - \sum_{i=1}^N x_i \right) \left(\sum_{i=1}^N x_i^2 - x_k \sum_{i=1}^N x_i \right)}{\left[N \sum_{i=1}^N x_i^2 - \left(\sum_{i=1}^N x_i \right)^2 \right]^2} \\ &= S_y^2 \frac{N \left(\sum_{k=1}^N x_k \right) \left(\sum_{i=1}^N x_i^2 \right) - N \left(\sum_{i=1}^N x_i \right) \left(\sum_{k=1}^N x_k^2 \right)}{\left[N \sum_{i=1}^N x_i^2 - \left(\sum_{i=1}^N x_i \right)^2 \right]^2} \\ &\quad + S_y^2 \frac{-N \left(\sum_{k=1}^N x_k^2 \right) \left(\sum_{i=1}^N x_i \right) + \left(\sum_{k=1}^N x_k \right) \left(\sum_{i=1}^N x_i^2 \right)}{\left[N \sum_{i=1}^N x_i^2 - \left(\sum_{i=1}^N x_i \right)^2 \right]^2} \\ &= S_y^2 \frac{\sum_{i=1}^N x_i \left[\left(\sum_{i=1}^N x_i \right)^2 - N \sum_{i=1}^N x_i^2 \right]}{\left[N \sum_{i=1}^N x_i^2 - \left(\sum_{i=1}^N x_i \right)^2 \right]^2}. \end{aligned}$$

Stąd

$$C_{\hat{a}, \hat{b}} = -S_y^2 \frac{\sum_{i=1}^N x_i}{N \sum_{i=1}^N x_i^2 - \left(\sum_{i=1}^N x_i \right)^2} = - \frac{\left(\sum_{i=1}^N d_i^2 \right) \left(\sum_{i=1}^N x_i \right)}{(N-2) \left[N \sum_{i=1}^N x_i^2 - \left(\sum_{i=1}^N x_i \right)^2 \right]}.$$

(3.5.26)

Przedstawiona metoda wyznaczania współczynników $\hat{a}$ i $\hat{b}$ nosi nazwę **regresji liniowej**, a otrzymana prosta — **prostej regresji zmiennych y i x** .

W przypadku, gdy błędy pomiarowe wartości y_i są do pominięcia, a błędów wartości x_i nie możemy pominąć, to dopasowujemy nie prostą $y = ax + b$, ale prostą $x = a'y + b'$. Należy zaznaczyć, że w ogólności proste $y = y(x)$ i $x = x(y)$ dopasowane do tych samych wyników pomiarów nie będą się pokrywały.

Gdy znamy prostą regresji zmiennych x i y możemy, korzystając z niej, obliczyć wartości zmiennych x i y w innych punktach niż zmierzone. Obliczone wartości tych zmiennych będą jednak obarczone pewnym błędem, ponieważ współczynniki $\hat{a}$ i $\hat{b}$ prostej regresji są jedynie przybliżeniami wartości rzeczywistych współczynników a oraz b . Poniżej znajdziemy odchylenia standardowe wartości zmiennych x oraz y obliczonych przy użyciu prostej regresji. Z zależności:

$$y = \hat{a}x + \hat{b} \quad (3.5.27)$$

możemy dla danej wartości x obliczyć $\hat{y} = \hat{a}x + \hat{b}$. Odchylenie standardowe $S_{\hat{y}}$ obliczonej wartości $\hat{y}$ znajdujemy ze wzoru (3.5.8), a mianowicie:

$$S_{\hat{y}}^2 = \left(\frac{\partial y}{\partial a}\right)_{\hat{a}, \hat{b}}^2 S_{\hat{a}}^2 + \left(\frac{\partial y}{\partial b}\right)_{\hat{a}, \hat{b}}^2 S_{\hat{b}}^2 + 2\left(\frac{\partial y}{\partial a}\right)_{\hat{a}, \hat{b}} \left(\frac{\partial y}{\partial b}\right)_{\hat{a}, \hat{b}} C_{\hat{a}, \hat{b}}. \quad (3.5.28)$$

Pochodne znajdujemy różniczkując (3.5.27); po ich podstawieniu do (3.5.28) otrzymujemy:

$$S_{\hat{y}}^2 = x^2 S_{\hat{a}}^2 + S_{\hat{b}}^2 + 2x C_{\hat{a}, \hat{b}}. \quad (3.5.29)$$

Korzystając z zależności (3.5.27) możemy obliczyć również wartość $\hat{x}$ odpowiadającą danej wartości y , tj.

$$\hat{x} = \frac{y - \hat{b}}{\hat{a}}. \quad (3.5.30)$$

Odchylenie standardowe $S_{\hat{x}}$ obliczonej wartości $\hat{x}$ znajdujemy ze wzoru (3.5.8), a mianowicie

$$S_{\hat{x}}^2 = \left(\frac{\partial x}{\partial a}\right)_{\hat{a}, \hat{b}}^2 S_{\hat{a}}^2 + \left(\frac{\partial x}{\partial b}\right)_{\hat{a}, \hat{b}}^2 S_{\hat{b}}^2 + 2\left(\frac{\partial x}{\partial a}\right)_{\hat{a}, \hat{b}} \left(\frac{\partial x}{\partial b}\right)_{\hat{a}, \hat{b}} C_{\hat{a}, \hat{b}}. \quad (3.5.31)$$

Pochodne obliczamy różniczkując (3.5.30); po podstawieniu ich do wzoru (3.5.31) znajdujemy:

$$S_{\hat{x}}^2 = \left(\frac{y - \hat{b}}{\hat{a}^2}\right)^2 S_{\hat{a}}^2 + \left(\frac{1}{\hat{a}}\right)^2 S_{\hat{b}}^2 + 2\left(\frac{y - \hat{b}}{\hat{a}^2}\right) \left(\frac{1}{\hat{a}}\right) C_{\hat{a}, \hat{b}}, \quad (3.5.32)$$

lub korzystając z (3.5.30):

$$S_{\hat{x}}^2 = \left(\frac{\hat{x}}{\hat{a}}\right)^2 S_{\hat{a}}^2 + \left(\frac{1}{\hat{a}}\right)^2 S_b^2 + 2\hat{x}\left(\frac{1}{\hat{a}}\right)^2 C_{\hat{a},b}. \quad (3.5.33)$$

Jak widać ze wzoru (3.5.29) odchylenie standardowe $S_{\hat{y}}$ obliczonej (a więc przybliżonej) wartości $\hat{y}$ jest nieliniową funkcją zmiennej x . Ze wzoru (3.5.33) wynika, że odchylenie standardowe $S_{\hat{x}}$ wartości $\hat{x}$ obliczonej ze wzoru (3.5.30) jest jej nieliniową funkcją.

Przykład

Ciało porusza się ruchem jednostajnie przyspieszonym po linii prostej. Dokonujemy pomiaru zależności prędkości ciała v od czasu t , który upłynął od momentu rozpoczęcia obserwacji ruchu. Dokonujemy więc pomiaru dwóch wielkości v oraz t i otrzymujemy *pary* wyników (v_i, t_i) ($i = 1, 2, \dots, N$). Otrzymane wyniki pomiarów zebrano w tabeli 3.3 i przedstawiono na rysunku 3.4.

Tabela 3.3: Wyniki pomiarów zależności prędkości v w ruchu jednostajnie przyspieszonym od czasu t

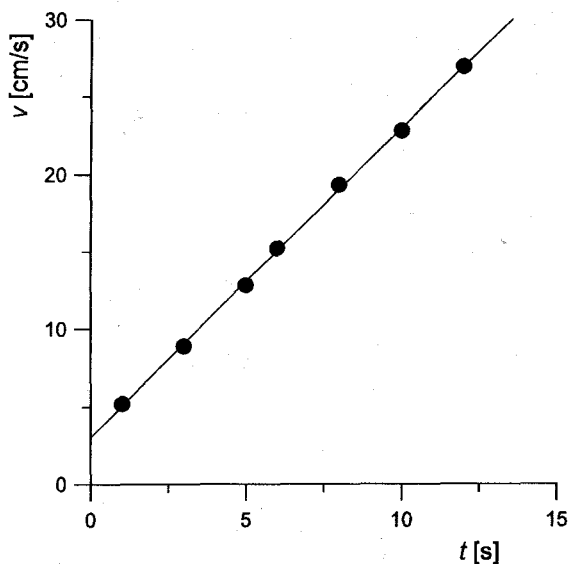
t [s]	1	3	5	6	8	10	12
v [cm/s]	5.2	8.9	12.8	15.2	19.3	22.8	26.9

Jak widać z rysunku 3.4, wyniki układają się *wokół* prostej o równaniu $v = v_0 + at$. Aby wyznaczyć tę prostą, musimy znać współczynniki v_0 i a . Ponieważ ich nie znamy, wyznaczymy je metodą najmniejszych kwadratów. W tym celu skorzystamy z wzorów (3.5.13) i (3.5.14). Współczynnik a w równaniu prostej (3.5.9) w naszym przypadku odpowiada przyspieszeniu, zaś współczynnik b prędkości początkowej v_0 , zmienna x czasowi obserwacji ruchu t , a zmienna y prędkości ciała v . Występujące we wzorach (3.5.13) i (3.5.14) sumy wynoszą:

$$\sum_{i=1}^7 t_i = 45.00 \text{ s}; \quad \sum_{i=1}^7 v_i = 11.1 \text{ cm/s}; \quad \sum_{i=1}^7 t_i^2 = 379.00 \text{ s}^2; \quad \sum_{i=1}^7 t_i v_i = 892.3 \text{ cm},$$

liczba pomiarów $N = 7$. Podstawiając obliczone sumy do wzorów (3.5.13) i (3.5.14) otrzymujemy: $\hat{v}_0 = 3.1105096 \text{ cm/s}$ oraz $\hat{a} = 1.985032 \text{ cm/s}^2$.

Obliczmy jeszcze $S_{\hat{a}}$, $S_{\hat{v}_0}$ i $C_{\hat{a},\hat{v}_0}$. W tym celu skorzystajmy ze wzorów (3.5.22) – (3.5.26). Aby z nich skorzystać musimy obliczyć S_v^2 . Podstawiając do wzoru (3.5.19) zmierzone wartości v_i, t_i oraz obliczone $\hat{v}_0$ i $\hat{a}$ otrzymujemy $S_v^2 = 0.049694 \text{ cm}^2/\text{s}^2$. Podstawiając do wzorów (3.5.22), (3.5.25) i (3.5.26) obliczone sumy oraz S_v^2 otrzymujemy: $S_{\hat{v}_0} = 0.173178 \text{ cm/s}$, $S_{\hat{a}} = 0.0235354 \text{ cm/s}^2$ oraz $C_{\hat{a},\hat{v}_0} = -0.00356089 \text{ cm}^2/\text{s}^3$.



Rys. 3.4: Wyniki pomiarów prędkości ciała i dopasowania metodą najmniejszych kwadratów prostej o równaniu $v(t) = v_0 + at$

Obliczmy teraz przy użyciu wyznaczonych wartości $\hat{v}_0$ i $\hat{a}$ prędkość, jaką ciało uzyska po upływie czasu $t = 11$ s. Wynosi ona $\hat{v}_{11} = 24.984095$ cm/s. Jej odchylenie standardowe $S_{\hat{v}_{11}}$ znajdujemy ze wzoru (3.5.29) przy użyciu wyznaczonych wartości $S_{\hat{v}_0}$, $S_{\hat{a}}$ i $C_{\hat{a}, \hat{v}_0}$. Wynosi ono $S_{\hat{v}_{11}} = 0.1366627$ cm/s. Z kolei obliczmy, przy użyciu wyznaczonych wartości $\hat{v}_0$, $\hat{a}$, $S_{\hat{v}_0}$, $S_{\hat{a}}$ i $C_{\hat{a}, \hat{v}_0}$, czas, po upływie którego prędkość ciała będzie wynosić $v = 20$ cm/s, oraz jego odchylenie standardowe. Czas ten jest $\hat{t}_{20} = 8.5084222$ s. Jego odchylenie standardowe obliczamy ze wzoru (3.5.32) lub (3.5.33). Wynosi ono $S_{\hat{t}_{20}} = 0.0490927$ s.

W przykładzie tym wszystkie obliczone wartości zostały podane bez zaokrągleń, tzn. tak jak zostały odczytane z kalkulatora. Przedstawienie powyższych wyników tak, aby były zgodne z tym, co powiedziano w rozdziale 1.1.4, stanie się jasne po przeczytaniu rozdziału 5, gdy omówimy związek obliczonych wartości $\hat{a}$ i $\hat{b}$ z wartościami rzeczywistymi tych współczynników oraz odchyłeń standardowych z błędem pomiaru.

3.5.2. Dopasowanie innych krzywych

Stosując metodę najmniejszych kwadratów możemy dopasowywać różne funkcje do wyników pomiarów. Rozwiązanie analityczne układu równań (3.5.3) jest jednak możliwe wtedy, gdy równania te są liniowe ze względu na współczynniki c_j . Równania (3.5.3) są liniowe, gdy dopasowujemy na przy-

kład wielomiany, jednak trudności rachunkowe szybko wzrastają ze wzrostem stopnia wielomianu. Dla przykładu znajdziemy układ równań (3.5.3) w przypadku, gdy dopasowujemy parabolę o równaniu $y = c_1 + c_2x + c_3x^2$. Warunek (3.5.2) przyjmie w tym przypadku postać:

$$\sum_{i=1}^N (y_i - c_1 - c_2x_i - c_3x_i^2)^2 = \min ,$$

Wówczas układ równań (3.5.3) będzie następujący:

$$\begin{aligned} Nc_1 + c_2 \sum_{i=1}^N x_i + c_3 \sum_{i=1}^N x_i^2 &= \sum_{i=1}^N y_i \\ c_1 \sum_{i=1}^N x_i + c_2 \sum_{i=1}^N x_i^2 + c_3 \sum_{i=1}^N x_i^3 &= \sum_{i=1}^N y_i x_i \\ c_1 \sum_{i=1}^N x_i^2 + c_2 \sum_{i=1}^N x_i^3 + c_3 \sum_{i=1}^N x_i^4 &= \sum_{i=1}^N y_i x_i^2 . \end{aligned}$$

Odchylenia standardowe wyznaczonych współczynników $\hat{c}_j$ obliczamy podobnie, jak w przypadku dopasowania linii prostej, pamiętając, że liczba stopni swobody $k = N - 3$. Czytelnikowi pozostawiamy ocenę stopnia trudności obliczeń współczynników $\hat{c}_1$, $\hat{c}_2$ i $\hat{c}_3$ oraz ich odchyłeń standardowych.

Znacznie trudniejszy jest przypadek, gdy układ równań (3.5.3) jest nieliniowy ze względu na nieznane współczynniki c_j (np. gdy dopasowujemy funkcję $f(x, c_1, c_2, c_3) = c_1 + c_2e^{c_3x}$). W takim przypadku, na ogół, układu równań (3.5.3) nie możemy rozwiązać analitycznie i trzeba stosować metody numeryczne. W tym opracowaniu takimi zagadnieniami nie będziemy się zajmowali.

Szczególnie łatwe jest dopasowanie funkcji, które przez odpowiednie podstawienie można sprowadzić do funkcji liniowej. Tak przekształconą funkcję dopasowuje się sposobem przedstawionym w rozdziale 3.5.1.

Przykłady

1) *Funkcja $y = ax^\alpha + b$ ($\alpha \neq 0$ jest znaną liczbą).*

Stosując podstawienie $x^\alpha = z$ otrzymujemy funkcję liniową $y = az + b$.

2) *Funkcja $y = be^{ax^\alpha}$ ($\alpha \neq 0$ jest znaną liczbą).*

Logarytmując otrzymujemy $\ln y = \ln b + ax^\alpha$. Oznaczmy: $\ln y = z$, $\ln b = c$, $x^\alpha = t$, wtedy otrzymamy $z = at + c$, a więc funkcję liniową w układzie współrzędnych zt .

Należy zaznaczyć, że błędy, z jakimi są wyznaczone parametry funkcji po sprowadzeniu do postaci liniowej, przenoszą się na parametry funkcji danej w postaci pierwotnej. Aby znaleźć odchylenia standardowe parametrów funkcji pierwotnej musimy skorzystać ze wzoru (3.4.21).

3.5.3. Aproksymacja metodą najmniejszych kwadratów

Dotychczas rozpatrywaliśmy przypadek, gdy zależność funkcyjna $y = f(x, c_1, c_2, \dots, c_Q)$ była **rzeczywistą zależnością funkcyjną** między mierzonymi bezpośrednio wielkościami x oraz y i na podstawie wyników pomiarów wyznaczaliśmy parametry $\hat{c}_j$ ($j = 1, 2, \dots, N$). Teraz zajmijmy się innym jej zastosowaniem.

Przypuśćmy, że w wyniku pomiarów wielkości fizycznych x i y otrzymaliśmy N par punktów (x_i, y_i) ($i = 1, 2, \dots, N$). Stwierdziliśmy, że wielkości te są zależne, tzn. że istnieje między nimi jakaś zależność funkcyjna. Nie znamy postaci tej zależności funkcyjnej. Chcemy jednak tę zależność przedstawić w postaci jakiejś funkcji. W takim przypadku, najlepiej na podstawie wykresu zależności $y = y(x)$, musimy przyjąć jakąś postać zależności funkcyjnej. Oznaczmy ją $\gamma(x, c_1, \dots, c_Q)$. Inaczej mówiąc *postulujemy ogólną postać zależności funkcyjnej*. Po przyjęciu ogólnej postaci zależności funkcyjnej postępujemy tak, jak to zostało opisane na początku rozdziału 3.5. Takie postępowanie nazywamy **aproksymacją** (przybliżeniem) rzeczywistej zależności między mierzonymi wielkościami x i y funkcją, której ogólną postać *przyjeliśmy*. Należy jednak zaznaczyć, że wielkość S_y^2 dana wzorem (3.5.5) będzie zawierała nie tylko odstępstwa od postulowanej funkcji $\gamma(x, c_1, \dots, c_Q)$ spowodowane rozrzutem wyników pomiarów, ale również odstępstwa od rzeczywistej zależności między mierzonymi wielkościami. Zagadnieniami tymi nie będziemy się dalej zajmowali, ponieważ wykracza to poza ramy naszego wykładu. W tym miejscu sygnalizujemy jedynie to zastosowanie metody najmniejszych kwadratów.

Na zakończenie należy zaznaczyć, że metodą najmniejszych kwadratów *nie możemy znaleźć postaci funkcji dopasowywanej* $\gamma(x, c_1, c_2, \dots, c_Q)$, możemy *jedynie wyznaczyć wartości parametrów* c_j ($j = 1, 2, \dots, Q$). Oznacza to, że musimy *znać* postać funkcji $\gamma(x, c_1, c_2, \dots, c_Q)$. W praktyce postać funkcji $\gamma(x, c_1, c_2, \dots, c_Q)$ albo znamy z teorii, albo na podstawie tego, jak układają się na wykresie punkty doświadczalne, *postulujemy postać zależności funkcyjnej* i wyznaczamy wartości parametrów c_j . O takiej zależności mówimy, że jest **zależnością empiryczną**.

4. Histogramy i rozkłady, zmienna losowa

Błędy przypadkowe są spowodowane wieloma czynnikami zmieniającymi się w sposób przypadkowy i dlatego do ich analizy musimy stosować metody rachunku prawdopodobieństwa i statystyki matematycznej. W tym rozdziale zajmiemy się najprostszymi zagadnieniami statystyki matematycznej, które tworzą podstawę jej zastosowań w praktycznej ocenie błędu pomiaru.

W statystyce matematycznej bardzo ważne są pojęcia **populacji generalnej** lub krótko **populacji** oraz **próby**. Nie będziemy podawali pełnych definicji tych pojęć, ograniczymy się jedynie do podania ich określeń w formie użytecznej dla zrozumienia tego wykładu.

Populacją generalną będziemy nazywali zbiór wszystkich możliwych wartości danej wielkości fizycznej.

Przykłady

- 1) Wielkością fizyczną jest wzrost L mężczyzn w wieku 21 lat zamieszkających w województwie toruńskim. Populacją generalną tej wielkości fizycznej jest zbiór *wszystkich* wartości wzrostu L otrzymanych w wyniku pomiarów wzrostu *wszystkich* mężczyzn w wieku 21 lat zamieszkających w województwie toruńskim. Liczba elementów tego zbioru może być duża, ale jest skończona.
- 2) Długość stołu jest wielkością fizyczną. Zbiór *wszystkich* możliwych wyników pomiarów długości stołu będzie *populacją generalną wyników pomiarów długości stołu*. W tym przypadku wszystkich możliwych wyników pomiaru jest nieskończenie wiele, a więc populacja generalna jest zbiorem zawierającym nieskończoną ilość elementów.

Populacja generalna wyników pomiarów określonej wielkości fizycznej będzie zawierała wszystkie możliwe wyniki pomiarów. Populacja generalna może być skończona (jak w przykładzie 1) lub nieskończona (jak w przykładzie 2). Najczęściej nie jest możliwe wykonanie badań na całej populacji (np. nie jest możliwe wykonanie nieskończonej ilości pomiarów długości stołu).

Dlatego z populacji generalnej wybieramy pewien podzbiór, który nazywamy **próbą**. Jeżeli podzbiór ten wybieramy losowo, to nazywamy go **próbą losową**.

Celem statystyki matematycznej jest uzyskanie informacji o populacji generalnej na podstawie próby (np. prognozy wyborcze, gdzie na podstawie ankiet zebranych od pewnej grupy ludzi wybranych losowo, np. 1000 osób, wnioskujemy o wynikach głosowania w całym kraju). Powstaje pytanie: czym jest seria pomiarowa? Na podstawie tego, co powiedziano wyżej, seria pomiarowa jest próbą i jak się przekonamy dalej jest próbą losową. Liczbę pomiarów N w serii nazywamy jej **liczebnością**. Ponieważ w praktyce laboratoryjnej na ogół nie używa się terminu próba, w dalszym ciągu naszego wykładu będziemy mówili o serii pomiarowej.

Jeżeli mamy serię N pomiarów wielkości fizycznej x , to, niezależnie od tego czy są to pomiary bezpośrednie, czy pośrednie, na ich podstawie możemy obliczyć: średnią arytmetyczną $\bar{x}$, odchylenie standardowe serii S_x , odchylenie standardowe średniej serii $S_{\bar{x}}$. Wielkości te charakteryzują serię pomiarową jako całość, są miarą precyzji z jaką zostały wykonane pomiary. Na ich podstawie jednak nie można powiedzieć, jak w **przedziale zmienności** rozkładają się wyniki pomiarów, np. jaka wartość x_i występuje najczęściej.

4.1. Histogram i rozkład zmiennej losowej skokowej

Na początku rozpatrzmy następujący przykład. W fizyce atomowej i jądrowej często zachodzi konieczność liczenia cząstek. Przypuśćmy, że zliczaliśmy cząstki α . Czas zliczania wynosił 1 min i był znacznie krótszy od średniego czasu życia jąder promieniotwórczych, pomiar powtarzaliśmy 20 razy. Dla kolejnych pomiarów otrzymaliśmy następujące wartości:

$$5, 1, 7, 2, 4, 11, 8, 5, 6, 3, 4, 7, 5, 10, 2, 7, 4, 6, 5, 2 \quad (4.1.1)$$

cząstek/min. Tak zapisany zbiór wyników nie daje zbyt wielu informacji. Gdyby było ich znacznie więcej otrzymalibyśmy zbiór liczb, z którego byłoby bardzo trudno uzyskać jakiegokolwiek dodatkowe informacje poza obliczeniem $\bar{x}$, S_x oraz $S_{\bar{x}}$. Pierwszym krokiem w celu uporządkowania takiego zbioru liczb może być ich przedstawienie w kolejności rosnącej lub malejącej. Na przykład przy uporządkowaniu rosnącym otrzymujemy:

$$1, 2, 2, 2, 3, 4, 4, 4, 5, 5, 5, 5, 6, 6, 7, 7, 7, 8, 10, 11.$$

Znacznie wygodniejszym i bardziej użytecznym sposobem uporządkowania zbioru wyników pomiarów będzie podanie wartości x_k oraz liczby jej wystąpień N_k . Tak uporządkowane wyniki przedstawiono w tabeli 4.1. Obliczmy

Tabela 4.1: Uporządkowanie występowania określonej liczby zliczeń

x_k	1	2	3	4	5	6	7	8	9	10	11
N_k	1	3	1	3	4	2	3	1	0	1	1

teraz, korzystając ze wzoru (3.1.8), **średnią arytmetyczną** $\bar{x}$ liczby zliczeń w naszym przykładzie (zbiór 4.1.1):

$$\bar{x} = \frac{1}{20} \sum_{i=1}^{20} x_i = \frac{1}{20} (5 + 1 + 7 + \dots + 5 + 2) = 5.2,$$

jest to dokładnie to samo co:

$$\bar{x} = \frac{1}{20} [1 + (2 \cdot 3) + (3 \cdot 1) + (4 \cdot 3) + \dots + (9 \cdot 0) + (10 \cdot 1) + (11 \cdot 1)].$$

Tak więc $\bar{x}$ możemy wyrazić następująco:

$$\bar{x} = \frac{1}{N} \sum_{k=1}^M N_k x_k, \quad (4.1.2)$$

gdzie M jest liczbą *różnych* wartości x_k , w tym przykładzie $M = 11$. Wielkości N_k noszą nazwę **liczebności występowania wyniku** x_k . Zwróćmy uwagę na to, że

$$\sum_{k=1}^M N_k = N. \quad (4.1.3)$$

Równanie (4.1.2) można przedstawić w postaci:

$$\bar{x} = \sum_{k=1}^M \frac{N_k}{N} x_k = \sum_{k=1}^M f_k x_k, \quad (4.1.4)$$

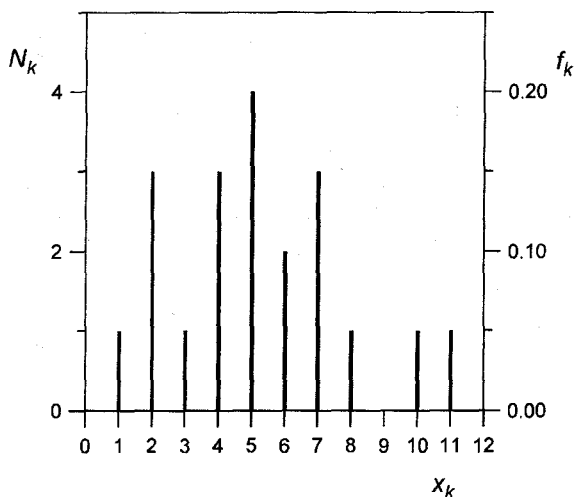
gdzie $f_k = N_k/N$ nazywamy **częstością występowania wyniku** x_k . Częstość f_k mówi nam jaka część spośród wszystkich N wyników pomiarów ma wartość x_k , czyli jak często wartość ta występuje. Częstość określa więc jak w danej serii wyniki pomiarów *rozłożyły* się pomiędzy różne możliwe wartości x_k . Inaczej mówiąc częstość f_k określa nam **rozkład wyników pomiarów serii**. Z równania (4.1.3) wynika, że

$$\sum_{k=1}^M f_k = 1.$$

Relacja ta nosi nazwę **warunku normalizacji**. Graficznie rozkład wyników pomiarów można przedstawić za pomocą tzw. **histogramu**.

Histogramem jest wykres zależności N_k lub f_k od x_k , przy czym wartości x_k są odłożone na osi odciętych (osi poziomej).

Na rysunku 4.1 przedstawiono, jako przykład, histogram zamieszczonych w tabeli 4.1 wyników pomiarów liczby cząstek α . Histogram w postaci rysunku 4.1 bywa nazywany **histogramem słupkowym**, ponieważ rozkład wyników przedstawiany jest w postaci pionowych słupków bezpośrednio nad wartościami x_k .



Rys. 4.1: Histogram wyników pomiarów liczby zliczeń cząstek α zebranych w tabeli 4.1

Przedstawianie wyników badań za pomocą histogramów jest pogładowe i dlatego dane przedstawione w ten sposób są szybko i łatwo przyswajalne.

Jeżeli liczba pomiarów $N \rightarrow \infty$, tzn. próba zmierza do populacji generalnej, to

$$\lim_{N \rightarrow \infty} f_k = p_k, \quad (4.1.5)$$

jest prawdopodobieństwem tego, że w wyniku pojedynczego pomiaru otrzymamy wartość x_k (por. Dodatek A). Prawdopodobieństwo p_k , podobnie jak częstość f_k jest funkcją x_k , czyli $p_k = p(x_k)$. Oznacza to, że wielkość mierzoną należy traktować jak *zmienną*. Zmienna ta ma charakterystyczną własność: przed pomiarem nie znamy jej wartości. Wartość jej poznajemy dopiero po wykonaniu pomiaru, czyli jak mówimy losowo. Dlatego wielkość taką nazywa się **zmienną losową**.

Zmienną losową nazywamy zmienną, która przyjmuje wartości z określonym prawdopodobieństwem. Jeżeli zmienna losowa przyjmuje tylko określone wartości, które mogą być ponumerowane za pomocą liczb naturalnych, to mówimy o niej, że jest **zmienną losową skokową** lub **dyskretną**.

Z określenia tego wynika, że wartości, które mogą być przyjmowane przez zmienną losową skokową, można przedstawić w postaci ciągu. Zbiór tych wartości jest zbiorem przeliczalnym.

Przykładami zmiennych losowych skokowych mogą być: liczba cząstek, liczba fotonów, liczba ziaren w kłosie zboża, liczba dzieci w przypadkowo wybranej rodzinie itp. W naszym przykładzie *kolejne* wyniki pomiarów różnią się o wielokrotność jedności. Należy zaznaczyć, że nie zawsze różnica między wartościami zmiennej losowej skokowej musi być wielokrotnością jedności.

Ponieważ wynik każdego pomiaru ma charakter przypadkowy, to seria składająca się z N niezależnych pomiarów ma również charakter przypadkowy. Oznacza to, że seria pomiarowa jest próbą pobraną z populacji generalnej w sposób przypadkowy. Jest więc **próbą losową**. Z tego wynika, że jeżeli będziemy powtarzali serię pomiarową, to wyniki otrzymane w każdej serii będą w ogólności różne.

Jeżeli liczba pomiarów $N \rightarrow \infty$ to $f_k \rightarrow p_k$ i histogram narysowany w układzie (f_k, x_k) przechodzi w wykres zależności $p(x_k)$. Prawdopodobieństwo p_k otrzymania w wyniku pojedynczego pomiaru wartości x_k jest więc funkcją x_k . Funkcję $p(x_k)$ nazywamy **rozkładem prawdopodobieństwa** zmiennej losowej skokowej lub krótko **rozkładem zmiennej losowej skokowej**. Tak więc w granicy, przy $N \rightarrow \infty$, histogram zmiennej losowej skokowej przechodzi w jej rozkład. Inaczej mówiąc histogram przedstawia nam rozkład wyników dla danej serii pomiarowej, a wykres funkcji rozkładu prawdopodobieństwa rozkład jaki otrzymalibyśmy, gdyby wykonać pomiary dla całej populacji generalnej.

Obliczmy teraz średnią arytmetyczną w przypadku granicznym, gdy liczba pomiarów $N \rightarrow \infty$. Ze wzoru (4.1.4) wynika, że

$$\lim_{N \rightarrow \infty} \bar{x} = \lim_{N \rightarrow \infty} \sum_{k=1}^M f_k x_k.$$

Ponieważ uzyskana w wyniku pomiaru wartość x_k jest niezależna od liczby pomiarów N , to korzystając z (4.1.5), otrzymujemy:

$$\lim_{N \rightarrow \infty} \bar{x} = \sum_{k=1}^M x_k \lim_{N \rightarrow \infty} f_k = \sum_{k=1}^M x_k p_k.$$

Wielkość

$$\bar{x}_R = \sum_{k=1}^M x_k p_k, \quad (4.1.6)$$

nazywamy **średnią rozkładu zmiennej losowej** albo krótko **średnią rozkładu** lub **wartością oczekiwaną zmiennej losowej**. Zwróćmy uwagę, że

$$\lim_{N \rightarrow \infty} \bar{x} = \bar{x}_R,$$

tzn. że gdy liczba pomiarów $N \rightarrow \infty$, to średnia arytmetyczna serii pomiarów $\bar{x}$ zmierza do średniej rozkładu $\bar{x}_R$, czyli do średniej populacji generalnej.

Rozkład prawdopodobieństwa zmiennej losowej skokowej będziemy charakteryzowali dwoma parametrami, a mianowicie **średnią rozkładu** $\bar{x}_R$ i **wariancją**

$$\sigma^2 = \sum_{k=1}^M (x_k - \bar{x}_R)^2 p_k. \quad (4.1.7)$$

Podnieśmy we wzorze (4.1.7) wyrażenie w nawiasie do kwadratu. Otrzymamy wtedy:

$$\sigma^2 = \sum_{k=1}^M x_k^2 p_k - 2\bar{x}_R \left(\sum_{k=1}^M x_k p_k \right) + \bar{x}_R^2 \left(\sum_{k=1}^M p_k \right),$$

stąd, uwzględniając to, że $\sum_{k=1}^M p_k = 1$ oraz (4.1.6), znajdujemy:

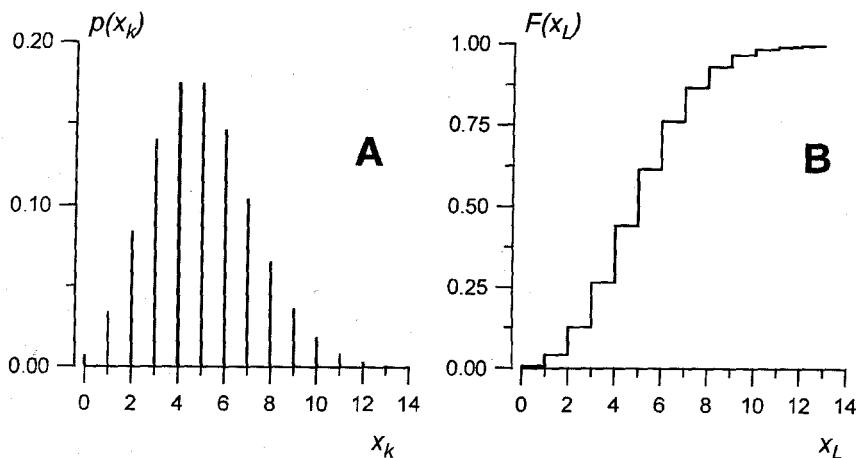
$$\sigma^2 = \sum_{k=1}^M x_k^2 p_k - (\bar{x}_R)^2.$$

Oznaczmy $\sum_{k=1}^M x_k^2 p_k = \overline{x_R^2}$; wtedy wzór powyższy przyjmie postać:

$$\sigma^2 = \overline{x_R^2} - (\bar{x}_R)^2. \quad (4.1.8)$$

Wzór (4.1.8) jest związkiem ogólnym, niezależnym od rodzaju rozkładu i charakteru zmiennej losowej.

Innymi parametrami charakteryzującymi rozkład nie będziemy się zajmowali. Na rysunku 4.2A przedstawiono przykład rozkładu zmiennej losowej skokowej. Wielkość $+\sqrt{\sigma^2}$ nazywa się **odchyleniem standardowym** i oznaczamy ją przez σ .



Rys. 4.2: Przykład rozkładu i dystrybuanty zmiennej losowej skokowej

Prawdopodobieństwo $P(x_k < x_L)$ tego, że wynik pojedynczego pomiaru x_k jest mniejszy od x_L (tzn. $x_k < x_L$) nazywa się **dystrybuantą rozkładu** zmiennej losowej x . Jest ono równe

$$P(x_k < x_L) = \sum_{j=1}^{L-1} p_j = F(x_L). \quad (4.1.9)$$

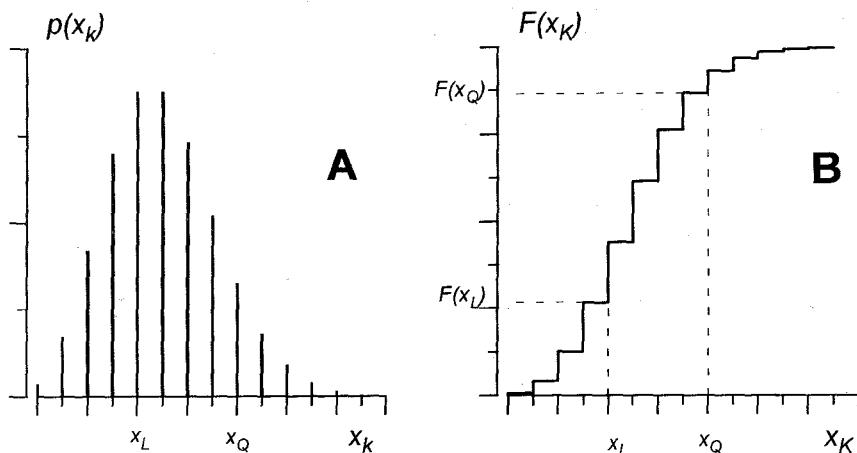
Na rysunku 4.2B przedstawiono wykres dystrybuanty dla rozkładu podanego na rysunku 4.2A. Jak widać z rysunku 4.2B dystrybuanta rozkładu zmiennej losowej skokowej jest funkcją skokową. Dystrybuanta jest funkcją ograniczoną, nieujemną, niemalejącą i dla dużych wartości x_L zmierza asymptotycznie do jedności.

Korzystając z funkcji $p(x_k)$ możemy obliczyć prawdopodobieństwo $P(x_L \leq x_s < x_Q)$ tego, że wynik pojedynczego pomiaru x_s będzie zawarty między x_L a x_Q , tj. $x_L \leq x_s < x_Q$. Na rysunku 4.3A przedstawiono przykładowy rozkład zmiennej skokowej $p(x_k)$. Z rysunku tego widać, że poszukiwane prawdopodobieństwo $P(x_L \leq x_s < x_Q) = p(x_L) + p(x_{L+1}) + \dots + p(x_{Q-1})$, czyli że

$$P(x_L \leq x_s < x_Q) = \sum_{j=L}^{Q-1} p_j. \quad (4.1.10)$$

Korzystając z definicji (4.1.9) i wzoru (4.1.10) otrzymujemy ważny w teorii błędów związek, a mianowicie:

$$P(x_L \leq x_s < x_Q) = F(x_Q) - F(x_L). \quad (4.1.11)$$



Rys. 4.3: Przykładowy rozkład (A) i dystrybuenta (B) zmiennej losowej skokowej; ilustracja do wyprowadzenia wzorów (4.1.9) i (4.1.10)

Prawdopodobieństwo to oznaczać będziemy $F(x_L, x_Q)$. Na rysunku 4.3B przedstawiono przykładowy wykres dystrybuenta z zaznaczonymi wartościami x_L i x_Q oraz odpowiadającymi im wartościami $F(x_L)$, $F(x_Q)$.

Do rozkładów zmiennej losowej skokowej wrócimy w rozdziale 7.

4.2. Histogram i rozkład zmiennej losowej ciągłej

Zmienną losową, która może przyjąć każdą wartość z pewnego przedziału nazywamy **zmienną losową ciągłą**.

Przedział ten w ogólności może być nieskończony. Przykładami zmiennych losowych ciągłych mogą być: czas trwania zjawiska, długość, masa, wzrost człowieka itp. Tak jak nigdy nie możemy zliczyć w określonym czasie np. 5.5 cząstek, tak pomiar czasu może dać nam każdą wartość z przedziału zmienności. W praktyce ciągły charakter mierzonej zmiennej losowej jest ograniczony dokładnością przyrządu pomiarowego. Wyjaśnimy to na następującym przykładzie.

Wykonujemy pomiary czasu t za pomocą miernika cyfrowego o dokładności 0.01 s. W tym przypadku zmierzone wartości t będą różniły się o wielokrotność dokładności miernika, tj. 0.01 s. Otrzymamy na przykład 18.12 s, 17.99 s, 17.88 s, ale nigdy nie otrzymamy wartości 18.118 s, 17.986 s, 17.882 s. Gdyby stosować miernik o dokładności np. 0.001 s, to zmierzone wartości t różniłyby się o wielokrotność tej dokładności. Ten pseudoskokowy

charakter jest spowodowany *dokładnością miernika*, a nie skokowym charakterem zmiennej. Ze względu na to, że zmienna losowa ciągła może przyjmować każdą wartość z pewnego przedziału, jej funkcja rozkładu prawdopodobieństwa, jak i dystrybuanta będą miały inny charakter niż dla zmiennej losowej skokowej; tworząc histogramy musimy brać pod uwagę jej ciągłość. Aby to wyjaśnić rozpatrzmy następujący przykład.

Jedna z metod wyznaczania współczynnika lepkości powietrza η polega na tym, że mając w zbiorniku (o znanej objętości V) powietrze pod ciśnieniem p_1 większym od ciśnienia atmosferycznego p_0 wykonujemy pomiary czasu t , w którym na skutek wypływu powietrza przez kapilarę (o znanej długości l i średnicy d) do atmosfery, ciśnienie obniży się do wartości $p_2 > p_0$. Współczynnik lepkości obliczamy ze wzoru:

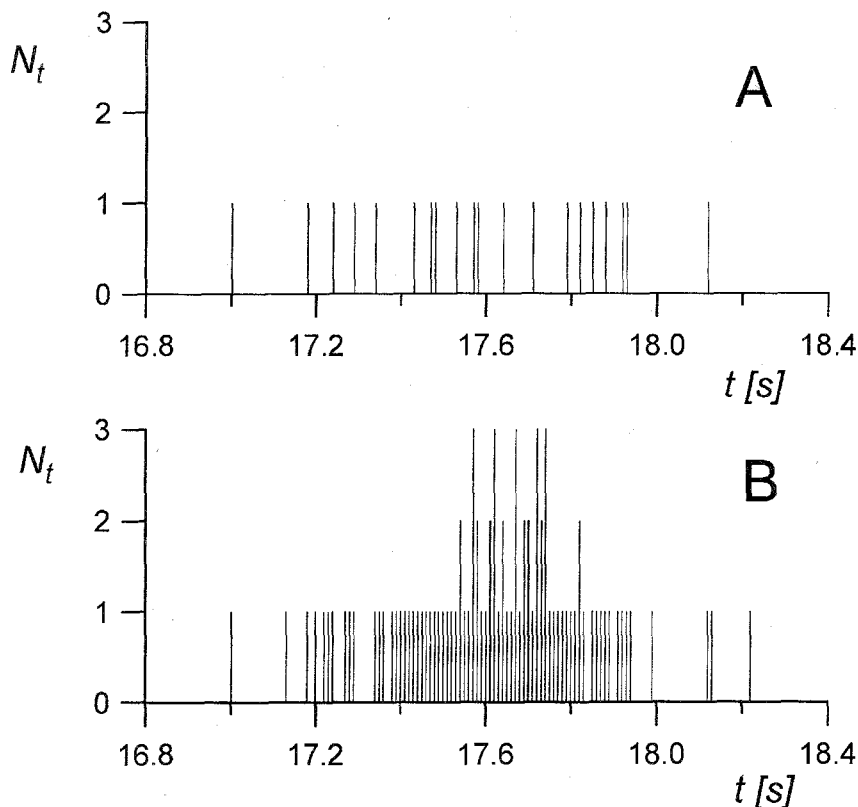
$$\eta = \frac{\pi d^4 t}{128 l V} \frac{p_0}{\ln \frac{p_1 - p_0}{p_1 + p_0} - \ln \frac{p_2 - p_0}{p_2 + p_0}}.$$

Szczegółowy opis metody podano np. w [10]. W celu dokładnego wyznaczenia η wykonano 90 pomiarów czasu t , używając miernika cyfrowego o dokładności 0.01 s. W pomiarach tych objętość zbiornika wynosiła 20 l, długość kapilary 235 mm, a jej średnica 0.9 mm. W czasie pomiarów ciśnienie $p_0 = 100797$ Pa, $p_1 = 106797$ Pa; i $p_2 = 104997$ Pa. Otrzymane wartości czasów wypływu t przedstawiono w tabeli 4.2.

Tabela 4.2: Wyniki pomiarów czasu wypływu t_i w sekundach

17.24	17.64	18.22	17.72	17.61	17.22	17.69	17.73	17.23
17.88	17.53	17.68	17.77	17.44	17.70	17.49	17.87	17.78
17.00	17.82	17.74	17.28	17.54	17.57	17.35	17.58	17.52
17.93	17.43	17.83	17.72	17.74	17.27	17.41	17.60	17.81
17.85	17.79	17.91	17.59	17.38	17.86	17.57	17.74	17.39
17.34	18.12	17.64	17.67	17.62	17.73	17.50	17.61	17.51
17.71	17.48	17.13	17.42	17.45	17.63	17.20	17.99	17.65
17.47	17.57	17.67	17.69	17.94	17.62	17.36	18.13	17.76
17.58	17.29	17.54	17.80	17.67	17.66	17.89	17.75	17.55
17.92	17.18	17.40	17.70	17.46	17.72	17.56	17.62	17.82

Do przykładu tego wrócimy jeszcze w rozdziale 5. Wyniki pomiarów przedstawione w takiej postaci jak w tabeli 4.2 nie tylko nic nie mówią o ich rozkładzie, ale są po prostu mało „czytelne”. Innymi słowy taki sposób przedstawienia wyników pomiarów jest nieprzejrzysty. Na rysunku 4.4A przedstawiono rozkład 20 pierwszych pomiarów zamieszczonych w dwóch pierwszych kolumnach tabeli 4.2; na osi rzędnych odłożono krotność N_t



Rys. 4.4: Rozkład wyników pomiarów czasu wypływu t_i zamieszczonych w tabeli 4.2.
A — 20 pierwszych pomiarów, B — wszystkie pomiary

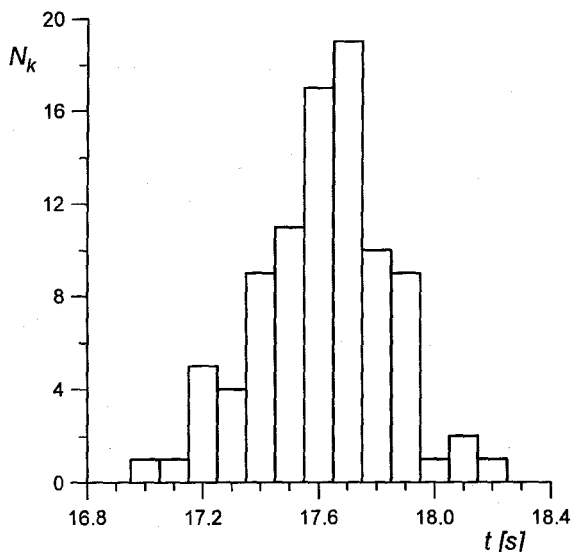
występowania czasu wypływu t_i . Na rysunku tym nie widać, aby wyniki grupowały się koło jakiejś wartości. Na rysunku 4.4B zamieszczono rozkład dla wszystkich 90 pomiarów, widać, że niektóre wartości t występują dwu- lub trzykrotnie. Zgodnie z tym, co powiedziano powyżej, jest to jednak spowodowane dokładnością miernika. Widać jednakże, że wyniki grupują (zagęszczają) się w pobliżu $t \approx 17.6$ s. Na przykładzie przedstawionym na rysunku 4.4 pokazaliśmy, że w przypadku zmiennej losowej ciągłej przedstawienie wyników w postaci histogramu słupkowego (jak dla zmiennej losowej skokowej) nie daje nic poza graficznym przedstawieniem, w jakiej części przedziału zmienności grupuje się najwięcej wyników pomiarów. Aby stwierdzić, czy występują jakieś prawidłowości w rozkładzie zmiennej losowej ciągłej, musimy postąpić inaczej niż w przypadku zmiennej losowej skokowej. Celem stwierdzenia występowania prawidłowości w rozkładzie zmiennej losowej ciągłej tworzymy tzw. **szereg rozdzielczy zmiennej losowej ciągłej**. W pierwszej kolejności określamy tak zwany

Tabela 4.3: Podział na przedziały klasowe wyników pomiarów czasu wpływu t_i zamieszczonych w tabeli 4.2

Nr przedziału klasowego k	Przedział Klasowy $[t_k, t_{k+1})$		Liczebność przedziału N_k	Częstość g_k
1	[16.95,17.05)		1	0.011
2	[17.05,17.15)		1	0.011
3	[17.15,17.25)		5	0.055
4	[17.25,17.35)		4	0.044
5	[17.35,17.45)		9	0.099
6	[17.45,17.55)		11	0.122
7	[17.55,17.65)		17	0.189
8	[17.65,17.75)		19	0.211
9	[17.75,17.85)		10	0.111
10	[17.85,17.95)		9	0.099
11	[17.95,18.05)		1	0.011
12	[18.05,18.15)		2	0.022
13	[18.15,18.25)		1	0.011

rozstęp, tzn. różnicę między wynikiem maksymalnym $x_{\max}$ a minimalnym $x_{\min}$, czyli całą długość *otrzymanego* w pomiarach przedziału zmienności. W naszym przykładzie $t_{\max} = 18.22$ s, $t_{\min} = 17.00$ s, a rozstęp wynosi 1.22 s. Następnie całą długość rozstępu wziętą z nadmiarem dzielimy na przedziały nazywane **przedziałami klasowymi** [13]. Przedziały muszą być rozłączne, jednostronnie domknięte, pokrywać cały rozstęp i powinny być jednakowej długości. Liczba przedziałów klasowych nie powinna być mniejsza niż 6. Nie ma ogólnej reguły określającej **długość przedziału klasowego**, ani ich liczby. Przy dobieraniu długości przedziału klasowego należy więc zachować pewną ostrożność. Jeżeli przedziały będą za szerokie, wówczas większość wyników znajdzie się w jednym lub paru (dwóch, trzech) przedziałach i nie uzyskamy informacji o rozkładzie. Jeżeli przedziały będą za wąskie, to możemy otrzymać rozkład podobny do przedstawionego na rysunku 4.4. Długość przedziału trzeba dobierać tak, aby w większości przedziałów mieściło się przynajmniej po kilka wyników. Im większa liczba pomiarów N , tym długość przedziału klasowego może być mniejsza, tzn. że jeżeli będziemy zwiększali liczbę pomiarów, to jednocześnie możemy zmniejszać długość przedziału klasowego. Wróćmy teraz do naszego przykładu. Podzielmy przedział zmienności na przedziały klasowe

równej długości. Jako długość przedziału klasowego przyjmijmy 0.1 s i podziel na przedziały klasowe rozpoczniemy od wartości $t = 16.95$ s. Wyniki pomiarów podzielone na przedziały klasowe przedstawia tabela 4.3.



Rys. 4.5: Histogram komórkowy czasów wypływu t ; wykonany na podstawie tabeli 4.3

Dopiero tak przygotowane wyniki pomiarów możemy przedstawić za pomocą histogramu. Tworzymy go jednak inaczej niż histogram zmiennej losowej skokowej przedstawiony na rysunku 4.1, a mianowicie na osi rzędnych odkładamy liczebność przedziału klasowego N_k , tj. liczbę wyników pomiarów w k -tym przedziale klasowym, a na osi odciętych odkładamy wartość zmiennej losowej; następnie rozkład wyników przedstawiamy za pomocą pionowego słupa, którego podstawa jest równa długości przedziału klasowego. Histogram taki bywa nazywany **histogramem komórkowym**. Na rysunku 4.5 przedstawiono histogram wyników pomiarów, zamieszczonych w tabeli 4.3.

Stosunek liczebności przedziału klasowego N_k do liczby wszystkich pomiarów N jest **częstością** g_k występowania wyniku z k -tego przedziału klasowego, czyli

$$g_k = \frac{N_k}{N}. \quad (4.2.1)$$

Wielkości tej nie należy mylić z wprowadzoną w rozdziale 4.1, dla przypadku zmiennej losowej skokowej z częstością f_k występowania wyniku pomiaru o wartości x_k . Tak zdefiniowana częstość g_k jest zależna od długości przedziału klasowego Δ_k . W analizowanym przykładzie dwukrotne zwiększenie długości przedziału spowoduje duże zmiany częstości, proponujemy

czytnikowi obliczyć g_k dla $\Delta_k = 0.2$ s. Częstość g_k spełnia warunek normalizacyjny, tj.

$$\sum_{k=1}^M g_k = 1, \quad (4.2.2)$$

gdzie M jest liczbą przedziałów klasowych. Relacja (4.2.2) jest niezależna od długości przedziału klasowego Δ_k i liczby pomiarów N (dowód pozostawiamy czytelnikowi).

Gdy liczba pomiarów $N \rightarrow \infty$, to, dla danej długości przedziału klasowego Δ_k ,

$$\lim_{N \rightarrow \infty} g_k = P(x_k \leq x < x_{k+1})$$

jest prawdopodobieństwem tego, że w wyniku pojedynczego pomiaru otrzymamy wartość należącą do k -tego przedziału klasowego, tzn. że zmierzona wartość x będzie zawarta w przedziale $[x_k, x_{k+1})$. Prawdopodobieństwo to jest zależne od długości przedziału klasowego. Zwróćmy uwagę, że $P(x_k \leq x < x_{k+1})$ jest równe różnicy prawdopodobieństw $P(x < x_{k+1})$ otrzymania $x < x_{k+1}$ i $P(x < x_k)$ otrzymania $x < x_k$, czyli

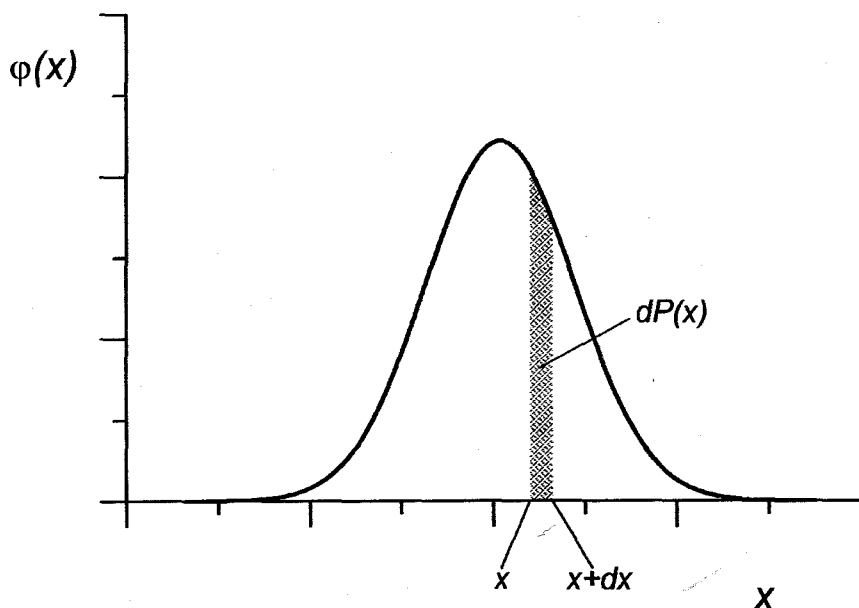
$$P(x_k \leq x < x_{k+1}) = P(x < x_{k+1}) - P(x < x_k). \quad (4.2.3)$$

Ponieważ $P(x_k \leq x < x_{k+1})$ jest zależne od długości przedziału klasowego Δ_k , wygodnie jest wprowadzić wielkość, która będzie opisywała to prawdopodobieństwo niezależnie od długości przedziału klasowego. Jak pokazuje dokładna analiza, taką wielkością jest prawdopodobieństwo przypadające na jednostkę długości przedziału klasowego. Zanim znajdziemy to prawdopodobieństwo zwróćmy uwagę na to, że w granicy, gdy liczba pomiarów $N \rightarrow \infty$, możemy zawężać przedziały klasowe tak, aby $\Delta_k \rightarrow 0$, a co za tym idzie $M \rightarrow \infty$. Wtedy wartości x_k określające granice przedziałów klasowych, które mają charakter zmiennych skokowych, będą zbliżały się do siebie na nieskończenie małe odległości. Przedziały klasowe stają się wtedy nieskończenie wąskie i zmienne x_k stają się zmiennymi ciągłymi, a histogram przechodzi w linię ciągłą. Aby znaleźć prawdopodobieństwo przypadające na jednostkę długości przedziału klasowego założmy, że liczba pomiarów $N \rightarrow \infty$ i obliczmy granicę stosunku $P(x_k \leq x < x_{k+1})/\Delta_k$, gdy $\Delta_k \rightarrow 0$. Korzystając z (4.2.3) i pamiętając, że $x_{k+1} = x_k + \Delta_k$ otrzymujemy:

$$\begin{aligned} \lim_{\Delta_k \rightarrow 0} \frac{P(x_k \leq x < x_{k+1})}{\Delta_k} &= \lim_{\Delta_k \rightarrow 0} \frac{P(x < x_{k+1}) - P(x < x_k)}{\Delta_k} \\ &= \frac{dP(x \leq x_k)}{dx_k} = \varphi(x_k). \end{aligned} \quad (4.2.4)$$

Ponieważ, jak wspomniano wyżej, zmienna x_k przy przejściach granicznych $N \rightarrow \infty$, $\Delta_k \rightarrow 0$ jest zmienną ciągłą, opuścimy w dalszej części wykładu

niepotrzebny już wskaźnik k i będziemy pisali po prostu $\varphi(x)$. Funkcję $\varphi(x)$ nazywamy **gęstością prawdopodobieństwa**; jest ona prawdopodobieństwem przypadającym na jednostkowy przedział zmiennej losowej ciągłej. Wobec tego $dP(x) = \varphi(x)dx$ jest prawdopodobieństwem, że wykonując pojedynczy pomiar otrzymamy wartość zmiennej losowej (wartość otrzymaną z pomiaru) zawartą w *nieskończenie wąskim przedziale* między x a $x + dx$. Funkcja $\varphi(x)$ jest wielkością mianowaną, jej wymiar jest odwrotnością wymiaru zmiennej losowej. Nazywana jest także **funkcją rozkładu prawdopodobieństwa** albo **funkcją rozkładu** lub krótko **rozkładem**. Trzeba zaznaczyć, że statystycy często nazywają dystrybuantę, daną wzorem (4.2.7), funkcją rozkładu. W fizyce jednak funkcją rozkładu nazywamy funkcję $\varphi(x)$. Przykładem tego może być makswelewski rozkład prędkości. Na rysunku 4.6 przedstawiono przykładową funkcję rozkładu zmiennej losowej ciągłej. Zakreskowana powierzchnia jest równa $\varphi(x)dx$ i przedstawia prawdopodobieństwo dP tego, że wynik pomiaru da nam wartość zawartą między x a $x + dx$.



Rys. 4.6: Przykład rozkładu zmiennej losowej ciągłej; objaśnienia w tekście

Z niezależności warunku normalizacyjnego (4.2.2) od liczby pomiarów i długości przedziału klasowego wynika warunek normalizacyjny prawdo-

podobieństwa $P(x_k \leq x < x_{k+1})$, a mianowicie:

$$\sum_{k=1}^M P(x_k \leq x < x_{k+1}) = 1. \quad (4.2.5)$$

Oznacza to, że otrzymanie w pojedynczym pomiarze wartości z dowolnego przedziału klasowego jest pewnością. Wyrażenie (4.2.5) łatwo jest zastosować w przypadku granicznym, gdy $\Delta \rightarrow 0$. Skoro $dP(x) = \varphi(x)dx$ jest prawdopodobieństwem tego, że wynik pojedynczego pomiaru będzie zawarty między x a $x + dx$, to jeżeli wysumujemy $dP(x)$ po wszystkich wartościach x z całego obszaru zmienności wielkości mierzonej, wtedy otrzymamy prawdopodobieństwo, że wynik pojedynczego pomiaru przyjmie dowolną wartość z całego obszaru zmienności wielkości mierzonej. Otrzymane w takim przypadku prawdopodobieństwo jest pewnością, czyli jest równe 1. Jak wiadomo z analizy matematycznej sumowanie, jakie należy wykonać w rozpatrywanym przypadku zmiennej ciągłej, sprowadza się do obliczenia odpowiedniej całki. Całkowanie to musimy wykonać po całym obszarze zmienności zmiennej losowej. Dla wygody granice całkowania rozszerzamy od $-\infty$ do $+\infty$. Postępowanie takie jest uzasadnione tym, że poza obszarem zmienności zmiennej losowej jej funkcja rozkładu $\varphi(x) \equiv 0$ i nie wnosi żadnego wkładu do wartości całki. W rezultacie otrzymujemy **warunek normalizacji** dla funkcji rozkładu:

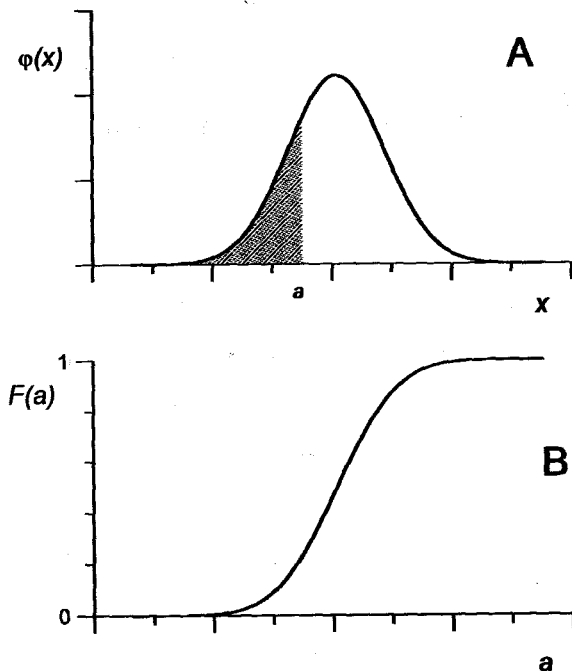
$$\int_{-\infty}^{+\infty} \varphi(x)dx = 1. \quad (4.2.6)$$

Prawdopodobieństwo $P(x < a)$ tego, że w wyniku pomiaru otrzymamy wartość $x < a$ (czyli, że $-\infty < x < a$), wyraża się analogicznym wzorem jak w przypadku zmiennej losowej skokowej (wzór 4.1.9), z tym, że sumowanie musimy zastąpić całkowaniem, a więc

$$P(x < a) = \int_{-\infty}^a \varphi(x)dx = F(a).$$

(4.2.7)

Prawdopodobieństwo to jest funkcją górnej granicy całkowania. Nazywamy je **dystrybuentą zmiennej losowej ciągłej**. Wzory definiujące dystrybuenty zmiennej losowej ciągłej i skokowej są więc podobne. Dystrybuenta zmiennej losowej ciągłej jest funkcją: ciągłą, nieujemną, monotoniczną, asymptotycznie zmierzającą do jedności. Na rysunku 4.7A przedstawiono przykładową funkcję rozkładu $\varphi(x)$; zakreskowana powierzchnia jest wartością dystrybuenty dla punktu a . Na rysunku 4.7B zamieszczono dystrybuentę $F(a)$ rozkładu $\varphi(x)$ pokazanego na rysunku 4.7A.



Rys. 4.7: Przykład rozkładu (A) i jego dystrybucyj (B) zmiennej losowej ciągłej; objaśnienia w tekście

Obliczmy teraz **prawdopodobieństwo** $P(b \leq x < c)$ tego, że w wyniku pojedynczego pomiaru otrzymamy wartość x zmiennej losowej ciągłej taką, że $b \leq x < c$. Aby obliczyć to prawdopodobieństwo zauważmy, że z prawa dodawania prawdopodobieństw wynika, iż

$$P(x < c) = P(x < b) + P(b \leq x < c),$$

stąd

$$P(b \leq x < c) = P(x < c) - P(x < b).$$

Korzystając z definicji dystrybucyj (4.2.7) znajdujemy:

$$\begin{aligned} P(b \leq x < c) &= F(c) - F(b) = \int_{-\infty}^c \varphi(x) dx - \int_{-\infty}^b \varphi(x) dx \\ &= \int_b^c \varphi(x) dx. \end{aligned} \quad (4.2.8)$$

Podobnie jak w przypadku zmiennej losowej skokowej będziemy je oznaczali $F(b, c)$.

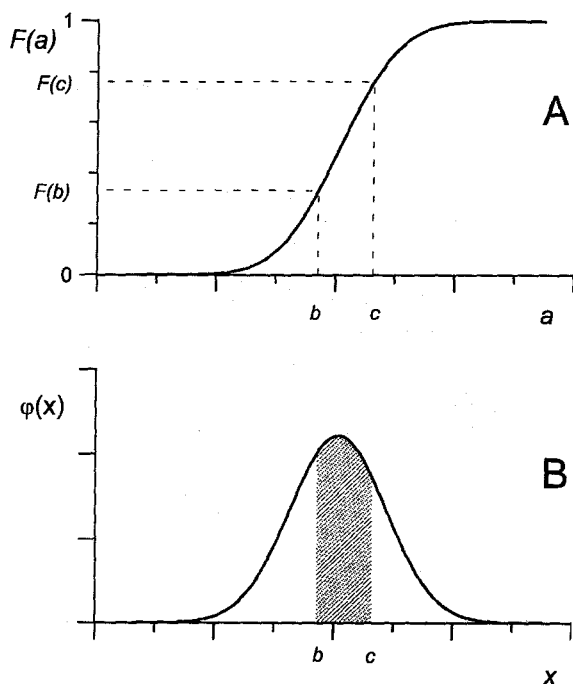
Ponieważ dla zmiennej losowej ciągłej zachodzą związki (patrz np. [14]):

$$P(b \leq x < c) = P(b \leq x \leq c) = P(b < x \leq c) = P(b < x < c), \quad (4.2.9)$$

w dalszym ciągu tego opracowania będziemy przyjmowali:

$$F(b, c) = P(b \leq x \leq c) = \int_b^c \varphi(x) dx. \quad (4.2.10)$$

Na rysunku 4.8A przedstawiono przykładową dystrybuantę $F(a)$, zaznaczono wartości b oraz c i odpowiadające im prawdopodobieństwa $F(b)$ i $F(c)$; zaś na rysunku 4.8B funkcję rozkładu $\varphi(x)$ — zakreślona powierzchnia jest równa $P(b \leq x \leq c)$.



Rys. 4.8: Ilustracja do wzorów (4.2.8) i (4.2.10); objaśnienia w tekście

Ze wzoru (4.2.10) wynika, że prawdopodobieństwo, iż zmienna losowa ciągła przyjmie określoną wartość $P(x = b)$, wynosi zero.

Rozkład zmiennej losowej ciągłej będziemy charakteryzowali, podobnie jak zmiennej losowej skokowej, dwoma parametrami, a mianowicie **wartością średnią rozkładu**, którą będziemy oznaczać $\bar{x}_R$ i **wariancją** σ^2 . Aby

otrzymać te wyrażenia wystarczy we wzorach (4.1.6) i (4.1.7) sumowanie zastąpić całkowaniem. Otrzymamy wtedy

$$\bar{x}_R = \int_{-\infty}^{+\infty} x \varphi(x) dx, \quad (4.2.11)$$

oraz

$$\sigma^2 = \int_{-\infty}^{+\infty} (x - \bar{x}_R)^2 \varphi(x) dx. \quad (4.2.12)$$

Innymi parametrami charakteryzującymi rozkład zmiennej losowej ciągłej nie będziemy się tu zajmowali.

Obliczmy teraz średnią arytmetyczną serii N pomiarów zmiennej losowej ciągłej. Zgodnie ze wzorem (3.1.8):

$$\bar{x} = \frac{1}{N} \sum_{i=1}^N x_i.$$

Gdy liczba pomiarów N będzie duża i jeżeli dokonamy podziału wyników pomiaru na przedziały klasowe, których długość Δ_k będzie mała, to średnią arytmetyczną możemy zapisać:

$$\bar{x} = \sum_{k=1}^M x_k \frac{N_k}{N} = \sum_{k=1}^M x_k g_k. \quad (4.2.13)$$

Jeżeli $N \rightarrow \infty$ to $g_k \rightarrow P(x_k \leq x < x_{k+1})$, zaś prawdopodobieństwo to, zgodnie z (4.2.3), jest równe:

$$P(x_k \leq x < x_{k+1}) = P(x < x_{k+1}) - P(x < x_k) = \Delta P_k.$$

Wobec tego wzór (4.2.13) przyjmie postać:

$$\lim_{\substack{N \rightarrow \infty \\ \Delta_k \rightarrow 0}} \bar{x} = \sum_{k=1}^M x_k \Delta P_k. \quad (4.2.14)$$

Jeżeli teraz, przy $N \rightarrow \infty$, będziemy zawężali przedziały klasowe tak, że $\Delta_k \rightarrow 0$ ($M \rightarrow \infty$), to korzystając z definicji całki oznaczonej otrzymamy:

$$\lim_{\substack{N \rightarrow \infty \\ \Delta_k \rightarrow 0}} \bar{x} = \lim_{\Delta_k \rightarrow 0} \sum_{k=1}^M x_k \Delta P_k = \int_{-\infty}^{+\infty} x dP(x).$$

Ale $dP(x) = \varphi(x)dx$, zatem:

$$\lim_{\substack{N \rightarrow \infty \\ \Delta_k \rightarrow 0}} \bar{x} = \int_{-\infty}^{+\infty} x\varphi(x)dx = \bar{x}_R. \quad (4.2.15)$$

Stąd wynika, że granicą średniej arytmetycznej serii N pomiarów, gdy $N \rightarrow \infty$ jest średnia rozkładu, dana wzorem (4.2.11).

Wariancja serii S_x^2 , zgodnie ze wzorem (3.2.3) wynosi

$$S_x^2 = \frac{1}{N} \sum_{i=1}^N (x_i - x_0)^2.$$

Założmy, że znamy wartość rzeczywistą x_0 i dokonajmy podziału błędu bezwzględnego $\delta_i = x_i - x_0$ na przedziały klasowe. Jeżeli długości przedziałów klasowych pozostaną niezmiennione, to otrzymamy histogram takiego samego kształtu jak w przypadku wyników pomiarów, tyle że wszystkie jego komórki zostaną przesunięte o x_0 . Częstość g_k nie ulegnie zmianie. Powtarzając rozumowanie przeprowadzone przy obliczaniu średniej arytmetycznej otrzymamy:

$$\lim_{N \rightarrow \infty} S_x^2 = \sum_{k=1}^M (x_k - x_0)^2 \Delta P_k. \quad (4.2.16)$$

Jeżeli liczba pomiarów będzie bardzo duża ($N \rightarrow \infty$) i będziemy zawężali przedziały klasowe, tzn. $\Delta_k \rightarrow 0$ ($M \rightarrow \infty$), to korzystając z definicji całki oznaczonej dostajemy:

$$\lim_{\substack{N \rightarrow \infty \\ \Delta_k \rightarrow 0}} S_x^2 = \lim_{\Delta_k \rightarrow 0} \sum_{k=1}^M (x_k - x_0)^2 \Delta P_k = \int_{-\infty}^{+\infty} (x - x_0)^2 \varphi(x) dx. \quad (4.2.17)$$

Z powyższego wzoru wynika, że **wariancja** serii S_x^2 zmierza do **wariancji rozkładu** σ^2 , gdy liczebność serii $N \rightarrow \infty$. Podobnie odchylenie standardowe serii S_x zmierza do odchylenia standardowego σ , tj.

$$\lim_{N \rightarrow \infty} S_x = \sigma. \quad (4.2.18)$$

Rozpatrzmy jeszcze przypadek gdy *zmienna losowa skokowa* przyjmuje *duże* wartości. Na przykład zliczamy fotony emitowane z lampy żarowej w czasie 1 s. Pomiar powtarzaliśmy 20 razy i otrzymaliśmy następujące wyniki:

58457, 58389, 58144, 58752, 58491, 58777, 58488, 58491, 58152, 58421,
58561, 58533, 58721, 58324, 53411, 58112, 58583, 58354, 58762, 58405

fotonów/s. Jeżeli będziemy chcieli wykonać histogram słupkowy, tak jak podano w rozdziale 4.1, to otrzymamy wykres podobny do przedstawionego na rysunku 4.4A. Zmienna losowa skokowa będzie wykazywać **pseudociągły charakter**. W takich przypadkach w celu otrzymania histogramu *musimy* postępować tak jak w przypadku zmiennej losowej ciągłej. A więc musimy utworzyć szereg rozdzielnicy i dokonać podziału na przedziały klasowe, w wyniku czego otrzymamy histogram komórkowy. Do zagadnienia tego wrócimy w rozdziale 7.

Należy zaznaczyć, że *wielkość złożona, która jest funkcją zmiennych losowych, jest również zmienną losową*. Funkcje rozkładu zmiennych losowych mierzonych bezpośrednio i wielkości złożonej, która jest ich funkcją, są w znakomitej większości przypadków *różne*.

Na zakończenie pragniemy zwrócić uwagę czytelnika na to,

- że powiedzenia, iż wyniki podlegają lub zmienna losowa podlega, lub zmienna losowa (wielkość fizyczna) ma rozkład $\varphi(x)$ oznaczają, że rozkład populacji generalnej jest $\varphi(x)$;
- że w języku używanym potocznie w laboratoriach odchyleniem standardowym nazywa się odchylenie standardowe serii S_x oraz odchyleniem standardowym średniej nazywa się odchylenie standardowe średniej serii $S_{\bar{x}}$, podobnie jest w przypadku wariancji.

5. Rozkład Gaussa i jego zastosowania

Informacje dotyczące zmiennych losowych, histogramów i funkcji rozkładu prawdopodobieństwa podane w rozdziale 4 są całkowicie ogólne i prawdziwe dla wszystkich funkcji rozkładu. W tym i następnym rozdziale zajmiemy się dwiema, szczególnie ważnymi przy ocenie błędu pomiaru, funkcjami rozkładu prawdopodobieństwa **zmiennej losowej ciągłej**, a mianowicie rozkładem Gaussa (nazywanym również: prawem błędów Gaussa, rozkładem normalnym, normalnym rozkładem błędów) i rozkładem *t-Studenta*.

5.1. Rozkład Gaussa

W rozdziale 4 pokazaliśmy, że wielkość mierzona jest zmienną losową oraz że funkcje zmiennych losowych są także zmiennymi losowymi. Z tego wynika, że błąd bezwzględny $\delta = x - x_0$, spowodowany błędami przypadkowymi, jest również zmienną losową.

Aby to wyjaśnić wróćmy do przykładu omawianego w rozdziale 4.2. Jeżeli od każdej wartości czasu wypływu t_i ($i = 1, 2, \dots, N$) odejmiemy jakąś stałą wartość, to wówczas wszystkie wyniki pomiarów zostaną jednakowo przesunięte. Przyjmijmy, że znamy wartość rzeczywistą czasu wypływu t_0 . Niech tą stałą wartością, którą odejmujemy będzie t_0 , tzn. zamiast wartości t_i będziemy mieli $\delta_i = t_i - t_0$. Powtarzając rozważania przedstawione w rozdziale 4.2 stwierdzamy, że:

- przedziały klasowe czasów wypływu $[t_k, t_{k+1})$ ($k = 1, 2, \dots, M$) podane w tabeli 4.2 staną się przedziałami klasowymi błędów bezwzględnych czasów wypływu $[\delta_k, \delta_{k+1})$, tzn. zostaną przesunięte o wartość t_0 ;
- histogram przedstawiony na rysunku 4.5, który jest w układzie (g_k, t_k) , będzie histogramem w układzie (g_k, δ_k) , tzn. zachowa swój kształt, ale wszystkie jego komórki zostaną przesunięte o t_0 ; oczywiście jeśli długości przedziałów klasowych będą takie same;
- jego granicą, gdy $N \rightarrow \infty$ i $\Delta_k \rightarrow 0$ ($M \rightarrow \infty$), będzie funkcja $\varphi(\delta)$.

Funkcja $\varphi(\delta) = \varphi(t - t_0)$ jest więc funkcją rozkładu prawdopodobieństwa wystąpienia błędu bezwzględnego zawartego w przedziale $[\delta, \delta + d\delta]$. Wobec tego prawdopodobieństwo wystąpienia błędu bezwzględnego zawartego między δ a $\delta + d\delta$ wynosi $\varphi(\delta)d\delta$. Z tego wynika, że znalezienie funkcji rozkładu błędów bezwzględnych jest równoważne znalezieniu funkcji rozkładu wyników pomiarów. Rozważania, które zostaną przedstawione w dalszym

ciągu tego rozdziału, będą poświęcone znalezieniu funkcji rozkładu błędów bezwzględnych i, zgodnie z tym co powiedziano w rozdziale 4.2, będą dotyczyć *bardzo dużej liczby pomiarów*, czyli sytuacji, gdy będziemy mogli przyjąć, że $N \rightarrow \infty$. Oznacza to, że będziemy zajmowali się funkcją rozkładu **populacji generalnej**. Celem tych rozważań jest, oprócz znalezienia funkcji rozkładu populacji generalnej, podanie sposobu oszacowania przy jej użyciu nierówności (1.1.4).

5.1.1. Prawo błędów Gaussa

Na podstawie tego, co zostało powiedziane o **błędach przypadkowych** w rozdziale 1.1, można podać, jakie warunki musi spełniać funkcja rozkładu błędów przypadkowych $\varphi(\delta)$, a mianowicie musi być:

- nieujemna, tj. $\varphi(\delta) \geq 0$;
- symetryczna, czyli $\varphi(-\delta) = \varphi(\delta)$; oznacza to, że prawdopodobieństwo wystąpienia błędu zawartego między $-\delta$ a $-(\delta + d\delta)$ jest takie samo jak błędu zawartego między δ a $\delta + d\delta$;
- malejącą funkcją δ ; czyli $\varphi(\delta_1) < \varphi(\delta_2)$, gdy $|\delta_1| > |\delta_2|$;
- asymptotycznie malejąca do 0, gdy $|\delta| \rightarrow \infty$, tzn. $\varphi(\pm\infty) \rightarrow 0$.

Z warunków tych wynika, że funkcja rozkładu $\varphi(\delta)$ musi posiadać maksimum w punkcie $\delta = 0$, tzn. $\varphi(0) = \max$.

Ogólna postać jednej spośród funkcji spełniających powyższe warunki jest następująca:

$$\varphi(\delta) = C e^{-h^2 \delta^2}, \quad (5.1.1)$$

zaś występujące stałe muszą być: $C > 0$ i h — liczbą rzeczywistą. Warunek, że $C > 0$ wynika z konieczności, aby funkcja (5.1.1) była nieujemna, zaś warunek, że h musi być liczbą rzeczywistą skończoną wynika z konieczności, aby funkcja $\varphi(\delta)$ była rzeczywista i asymptotycznie malejąca do zera, gdy $\delta \rightarrow \infty$. Ponadto funkcja (5.1.1) musi spełniać warunek normalizacji (4.2.6), który przyjmuje postać:

$$\int_{-\infty}^{+\infty} \varphi(\delta) d\delta = C \int_{-\infty}^{+\infty} e^{-h^2 \delta^2} d\delta = 1. \quad (5.1.2)$$

Korzystając z warunku normalizacji (5.1.2) możemy znaleźć wartość stałej C . W tym celu zastosujemy podstawienie $s = h\delta$; wtedy wzór (5.1.2) możemy zapisać:

$$\frac{C}{h} \int_{-\infty}^{+\infty} e^{-s^2} ds = 1, \quad (5.1.3)$$

całka we wzorze (5.1.3) jest równa $\sqrt{\pi}$ (obliczenie całki podano w Dodatku F). W rezultacie otrzymujemy:

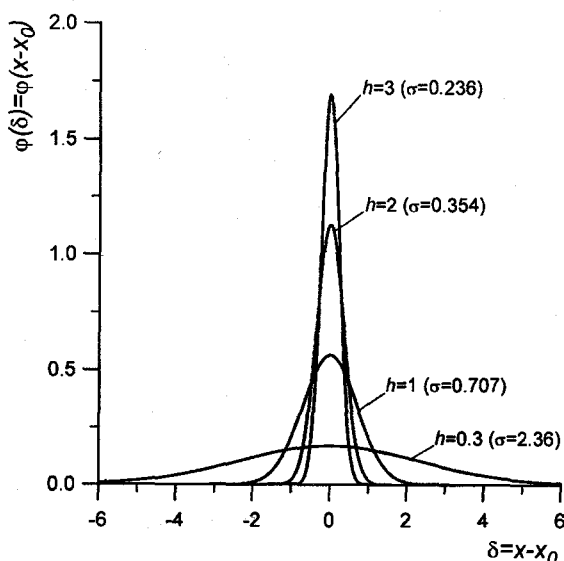
$$\frac{C}{h}\sqrt{\pi} = 1, \quad \text{stąd} \quad C = \frac{h}{\sqrt{\pi}}$$

i wzór (5.1.1) przyjmie postać:

$$\varphi(\delta) = \frac{h}{\sqrt{\pi}} e^{-h^2 \delta^2}. \quad (5.1.4)$$

Rozkład ten nazywa się **rozkładem Gaussa** lub **prawem błędów Gaussa** bądź **rozkładem normalnym** albo **normalnym rozkładem błędów**. Stałą h nazywamy miarą precyzji pomiarów.

O słuszności stosowania rozkładu Gaussa w ocenie błędu pomiarowego napisano wiele. To, co zostało napisane na ten temat doskonale charakteryzuje następująca wypowiedź: „...eksperymentatorzy stosują funkcję Gaussa, gdyż myślą, że matematyka może to udowodnić, a matematycy gdyż myślą, iż zostało to ustalone eksperymentalnie...” (por. [7], s. 40). Tak źle jednak nie jest i rozkład Gaussa można otrzymać stosując różne podejścia. Jedno z nich, wykorzystujące model błędów przypadkowych Laplace’a, przedstawimy w Dodatku E, który radzimy przeczytać po poznaniu rozkładu dwumianowego (rozdział 7.1).



Rys. 5.1: Wykresy funkcji rozkładu błędów Gaussa dla różnych wartości h , w nawiasach podano wariancję rozkładów

Na rysunku 5.1 przedstawiono przykładowe wykresy funkcji rozkładu błędów Gaussa (wzór (5.1.4)) dla czterech różnych wartości h . Jak widać z rysunku 5.1 dla dużych wartości h funkcja $\varphi(\delta)$ jest wąska z ostrym maksimum. Taki pomiar jest pomiarem o *wysokiej precyzji* i rozrzut błędów bezwzględnych jest *mały*. Im wartość h jest mniejsza, tym maksimum staje się niższe, a krzywa jest szersza — rozrzut wartości błędów bezwzględnych wzrasta. Dla małej wartości h ($h = 0.3$) wykres funkcji $\varphi(\delta)$ jest płaski, rozrzut wartości błędów bezwzględnych będzie duży — o takim pomiarze powiemy, że jest *mało precyzyjny*. Jak widać h określa szerokość funkcji rozkładu.

5.1.2. Rozkład Gaussa dla wyników pomiarów

Zgodnie z tym, co powiedziano w rozdziale 4.2 i 5.1.1, **prawdopodobieństwo** tego, że błąd bezwzględny jest zawarty między δ a $\delta + d\delta$, wynosi:

$$dP(\delta) = \varphi(\delta)d\delta, \quad (5.1.5)$$

gdzie $\varphi(\delta)$ jest funkcją rozkładu daną wzorem (5.1.4). Wiedząc, że zgodnie z wzorem (1.1.5) $\delta = x - x_0$, pytamy teraz jakie jest prawdopodobieństwo tego, że *wynik pojedynczego pomiaru* będzie zawarty między x a $x + dx$. W celu uzyskania odpowiedzi na postawione pytanie różniczkujemy wzór (1.1.5) i otrzymujemy $d\delta = dx$; ten wynik oraz δ dane wzorem (1.1.5) podstawiamy do (5.1.5). W rezultacie otrzymamy prawdopodobieństwo:

$$dP(x) = \varphi(x)dx. \quad (5.1.6)$$

Wobec tego funkcja rozkładu (5.1.4) przyjmie postać:

$$\varphi(x) = \frac{h}{\sqrt{\pi}} e^{-h^2(x-x_0)^2}. \quad (5.1.7)$$

Z porównania wzorów (5.1.4) i (5.1.7) wynika, że wykresy obu funkcji, przy tym samym h , są identyczne, z tym że wykres funkcji rozkładu $\varphi(x - x_0)$ będzie przesunięty o wartość x_0 . Oba wykresy pokrywają się dla $x_0 = 0$.

Korzystając z funkcji rozkładu (5.1.7) obliczmy **średnią rozkładu**. Podstawiając funkcję rozkładu $\varphi(x)$ do wzoru (4.2.11) otrzymujemy:

$$\bar{x}_R = \frac{h}{\sqrt{\pi}} \int_{-\infty}^{+\infty} x e^{-h^2(x-x_0)^2} dx.$$

Aby obliczyć $\bar{x}_R$ odejmijmy i dodajmy do wyrażenia pod całką $x_0 \exp[-h^2(x - x_0)^2]$; wówczas:

$$\bar{x}_R = \frac{h}{\sqrt{\pi}} \int_{-\infty}^{+\infty} (x - x_0) e^{-h^2(x-x_0)^2} dx + \frac{h}{\sqrt{\pi}} x_0 \int_{-\infty}^{+\infty} e^{-h^2(x-x_0)^2} dx.$$

Zastosujmy teraz podstawienie $y = h(x - x_0)$, wówczas $dy = hdx$ i powyższe wyrażenie przyjmie postać:

$$\bar{x}_R = \frac{1}{h\sqrt{\pi}} \int_{-\infty}^{+\infty} ye^{-y^2} dy + \frac{1}{\sqrt{\pi}} x_0 \int_{-\infty}^{+\infty} e^{-y^2} dy.$$

Pierwsza całka jest równa zero, a druga $\sqrt{\pi}$ (obliczenia tych całek podano w Dodatku F), a więc:

$$\bar{x}_R = x_0. \quad (5.1.8)$$

Równość ta oznacza, że w przypadku gdy wyniki pomiarów podlegają rozkładowi Gaussa (5.1.7), średnia rozkładu $\bar{x}_R$ jest równa wartości rzeczywistej x_0 . Jeżeli dodatkowo weźmiemy pod uwagę wyniki uzyskane w rozdziale 4.2, tzn. wzór (4.2.15), to stwierdzimy, że:

$$\lim_{N \rightarrow \infty} \bar{x} = x_0. \quad (5.1.9)$$

W ten sposób uzyskaliśmy bardzo ważny wynik, który oznacza, że jeżeli rozkład wyników pomiarów jest rozkładem Gaussa (5.1.7), to dla bardzo dużej liczby pomiarów ($N \rightarrow \infty$) ich średnia arytmetyczna dąży do wartości rzeczywistej, oczywiście, gdy błędy systematyczne można pominąć.

Obliczmy teraz **wariancję rozkładu Gaussa**, pamiętając, że zgodnie z (5.1.8) $\bar{x}_R = x_0$. Ze wzoru (4.2.12) znajdujemy:

$$\sigma^2 = \frac{h}{\sqrt{\pi}} \int_{-\infty}^{+\infty} (x - x_0)^2 e^{-h^2(x-x_0)^2} dx. \quad (5.1.10)$$

Wyrażenie to jest identyczne ze wzorem (4.2.17). Oznacza to, że gdy mierzona wielkość fizyczna ma rozkład Gaussa, to w przypadku bardzo dużej liczby pomiarów ($N \rightarrow \infty$) S_x^2 zmierza do wariancji rozkładu, czyli

$$\lim_{N \rightarrow \infty} S_x^2 = \sigma^2. \quad (5.1.11)$$

Aby obliczyć σ^2 dokonajmy podstawienia $y = h(x - x_0)$; wtedy $dy = hdx$ i wzór (5.1.10) przyjmie postać:

$$\sigma^2 = \frac{1}{h^2\sqrt{\pi}} \int_{-\infty}^{+\infty} y^2 e^{-y^2} dy.$$

Ponieważ występująca tu całka jest równa $\sqrt{\pi}/2$ (obliczenie całki podano w Dodatku F), otrzymujemy:

$$\sigma^2 = \frac{1}{2h^2}, \quad \text{stad} \quad h = \frac{1}{\sigma\sqrt{2}}.$$

Podstawiając wartość h do wzoru (5.1.7), możemy go wyrazić za pomocą wariancji. Wówczas funkcja rozkładu przyjmie postać:

$$\varphi(x) = \frac{1}{\sigma\sqrt{2\pi}} e^{-\frac{(x-x_0)^2}{2\sigma^2}}. \quad (5.1.12)$$

Jest to najczęściej używana postać **rozkładu Gaussa (normalnego)**. Funkcja rozkładu normalnego jest jednoznacznie określona dwoma parametrami, a mianowicie: **wartością rzeczywistą** x_0 i **wariancją** σ^2 . Funkcję tę oznaczamy $N(x_0, \sigma)$, np. zapis $N(20, 2)$ znaczy, że wartość rzeczywista wynosi 20 i odchylenie standardowe równe σ wynosi 2.

Na rys. 5.1 przedstawiono wykresy funkcji rozkładu $\varphi(\delta) = \varphi(x - x_0)$ dla różnych wartości miary precyzji h . Jak pokazano powyżej $\sigma = 1/h\sqrt{2}$, jest więc miarą wartości w maksimum i szerokości funkcji rozkładu. Jeżeli rozrzut wyników pomiarów jest duży, wtedy również σ jest duże, a funkcja rozkładu jest szeroka i o małej wartości w maksimum, ponieważ $\varphi(0) = 1/\sigma\sqrt{2\pi}$. Jeżeli zaś rozrzut wyników pomiarów jest mały, wtedy mamy sytuację odwrotną, tj. wartość funkcji rozkładu w maksimum jest duża, a szerokość mała.

Na zakończenie obliczmy dla przykładu prawdopodobieństwo dP , że wynik pojedynczego pomiaru da nam wartość zawartą między $x_0 + 1.5\sigma$ a $x_0 + 1.5\sigma + dx$. Oczywiście nie jesteśmy w stanie wykonać obliczeń dla nieskończenie wąskiego przedziału dx . Możemy je wykonać jedynie dla bardzo wąskiego przedziału, ale o skończonej szerokości Δx . Otrzymamy więc przybliżoną wartość prawdopodobieństwa, którą oznaczmy ΔP . Przyjmijmy, że $\Delta x = 0.001\sigma$. Prawdopodobieństwo to na podstawie wzorów (5.1.6) i (5.1.12) jest równe $\Delta P(x_0 + 1.5\sigma) = \frac{1}{\sqrt{2\pi}} e^{-1.125} \cdot 0.001 \approx 0.000129$. Czyli, że jeżeli wykonamy 10 000 pomiarów, to średnio 1 pomiar będzie zawarty w tym przedziale. Jeżeli $\Delta x = 0.01\sigma$, to $\Delta P(x_0 + 1.5\sigma) \approx 0.00129$, czyli będzie wzrastać ze wzrostem szerokości przedziału.

Zgodnie ze wzorem (4.2.7) **dystrybuanta rozkładu normalnego** $N(x_0, \sigma)$, którą będziemy oznaczać symbolem $\Phi(a)$, wyrazi się następująco:

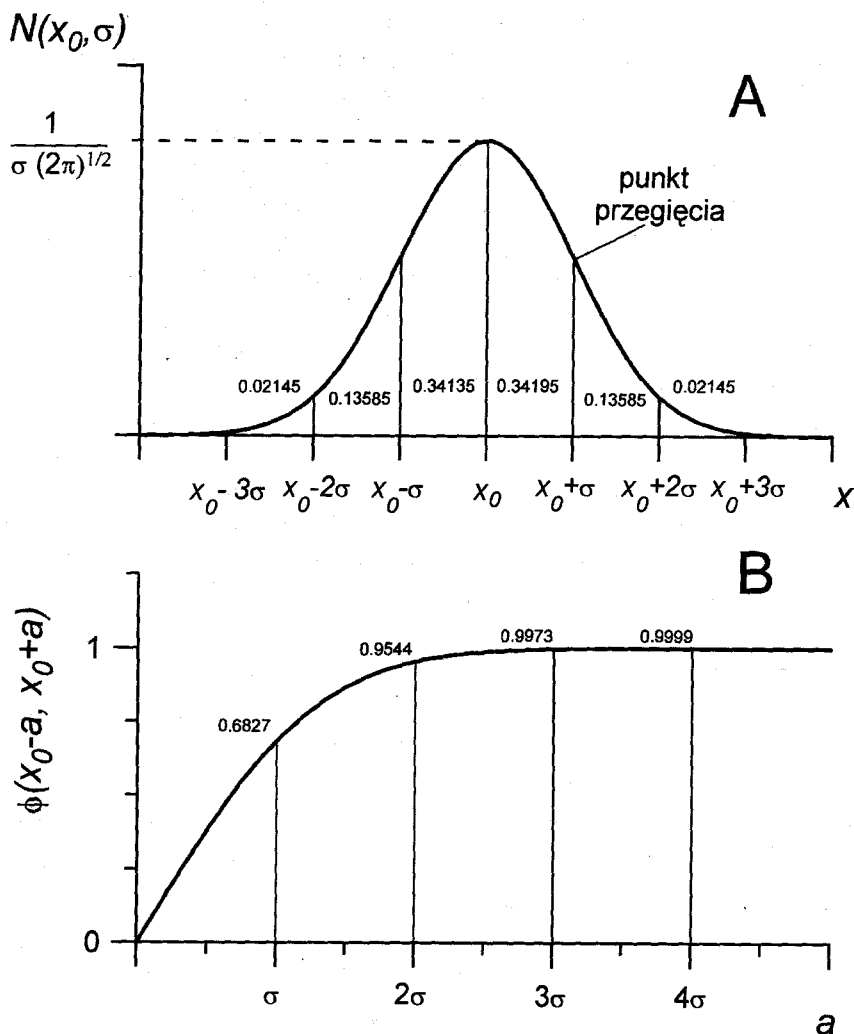
$$\Phi(a) = \frac{1}{\sigma\sqrt{2\pi}} \int_{-\infty}^a e^{-\frac{(x-x_0)^2}{2\sigma^2}} dx, \quad (5.1.13)$$

zaś prawdopodobieństwo $P(a \leq x \leq b)$ tego, że wynik pojedynczego po-

miaru będzie zawarty w przedziale $a \leq x \leq b$, jest:

$$P(a \leq x \leq b) = \Phi(b) - \Phi(a) = \Phi(a, b) = \frac{1}{\sigma\sqrt{2\pi}} \int_a^b e^{-\frac{(x-x_0)^2}{2\sigma^2}} dx. \quad (5.1.14)$$

Szczególnie często jest używana wielkość $P(x_0 - a \leq x \leq x_0 + a)$, tzn. **prawdopodobieństwo**, że wynik pomiaru jest zawarty w przedziale symetrycznym względem wartości rzeczywistej. Dla przykładu obliczmy prawdopo-



Rys. 5.2: A) Przykładowy wykres funkcji $N(x_0, \sigma)$. B) Wykres zależności prawdopodobieństwa $\Phi(x_0 - a, x_0 + a)$ od a . Objaśnienia w tekście

dobieństwo tego, że wynik pojedynczego pomiaru jest zawarty w przedziale $[x_0 - \sigma, x_0 + \sigma]$, tzn.

$$P(x_0 - \sigma \leq x \leq x_0 + \sigma) = \Phi(x_0 - \sigma, x_0 + \sigma) = \frac{1}{\sigma\sqrt{2\pi}} \int_{x_0 - \sigma}^{x_0 + \sigma} e^{-\frac{(x-x_0)^2}{2\sigma^2}} dx. \quad (5.1.15)$$

Całki tej nie możemy obliczyć analitycznie, można jednak jej wartość obliczyć numerycznie. W powyższym przypadku wynosi ona 0.6827, tzn. że jeżeli wykonamy serię N pomiarów, to średnio 68.27% wyników będzie zawarte w przedziale $[x_0 - \sigma, x_0 + \sigma]$. Pozostałe 31.73% wyników pomiarów da nam wartości mniejsze od $x_0 - \sigma$ i większe od $x_0 + \sigma$.

Na rys. 5.2A przedstawiono przykładową funkcję rozkładu $N(x_0, \sigma)$. Na rysunku tym zaznaczono również przedziały $[x_0 - 3\sigma, x_0 - 2\sigma], \dots, [x_0 + 2\sigma, x_0 + 3\sigma]$ i podano wartości całek z funkcji rozkładu $N(x_0, \sigma)$ obliczone dla poszczególnych przedziałów. Inaczej mówiąc podano **prawdopodobieństwa**, że w wyniku pojedynczego pomiaru wielkości fizycznej x otrzymamy wartość zawartą w danym **przedziale**. Na rysunku 5.2B przedstawiono wykres funkcji $\Phi(x_0 - a, x_0 + a)$ i podano wartości, jakie przyjmuje ta funkcja dla $a = k\sigma$ ($k = 1, 2, 3, 4$). Jak widać z rysunków 5.2A i 5.2B w przedziale $[x_0 + 3\sigma, x_0 - 3\sigma]$ są zawarte *prawie wszystkie wyniki pomiarów*. Może się oczywiście zdarzyć sytuacja, że jakiś pomiar znajdzie się poza tym przedziałem, ale prawdopodobieństwo takiego zdarzenia wynosi zaledwie 27/10000. W rozdziale 5.2 poznamy również inną postać funkcji $\Phi(a)$ wygodniejszą w bezpośrednich zastosowaniach oraz sposób oceny błędu przy użyciu rozkładu normalnego.

5.1.3. Średnia arytmetyczna i wariancja serii jako najlepsze przybliżenia wartości rzeczywistej i wariancji rozkładu

Wyniki przedstawione w rozdziałach 4.2 i 5.1.2 dotyczą przypadków, gdy liczba pomiarów N jest bardzo duża, tzn. taka, iż można przyjąć, że $N \rightarrow \infty$. W praktyce jednak N jest liczbą niezbyt dużą. Wykonujemy bowiem 5, 10, 50 czy 100 pomiarów. Załóżmy, że wykonaliśmy serię N pomiarów wielkości fizycznej x i otrzymaliśmy wartości $x_1, x_2, \dots, x_N$. Pytamy się, jakie będzie najlepsze przybliżenie wartości rzeczywistej x_0 i wariancji rozkładu σ^2 , jeżeli wyniki pomiarów podlegają rozkładowi normalnemu $N(x_0, \sigma)$.

Prawdopodobieństwo tego, że w wyniku pojedynczego pomiaru otrzymamy wartość zawartą między x_i a $x_i + dx_i$ jest $dP(x_i) = \varphi(x_i)dx_i$. Poprawnie wykonane pomiary są od siebie niezależne, tzn. że wynik jednego pomiaru nie wywiera wpływu na wartość otrzymaną w innym pomiarze. Z tego

powodu pomiary traktujemy jak **zdarzenia niezależne**, zaś prawdopodobieństwo $dP(x_1, x_2, \dots, x_N)$, że w serii N pomiarów otrzymamy wartości zawarte w przedziałach $[x_1, x_1 + dx_1]$, $[x_2, x_2 + dx_2]$, $\dots$, $[x_N, x_N + dx_N]$ jest równe iloczynowi $dP(x_1)dP(x_2)\dots dP(x_N)$. Gdy wyniki pomiarów podlegają rozkładowi normalnemu to

$$dP(x_1, x_2, \dots, x_N) = \prod_{i=1}^N dP(x_i) = \prod_{i=1}^N \frac{1}{\sigma\sqrt{2\pi}} e^{-\frac{(x_i - x_0)^2}{2\sigma^2}} dx_i.$$

Biorąc pod uwagę to, że przy mnożeniu wykładniki potęg dodajemy, oraz że można przyjąć, iż $dx_1 = dx_2 = \dots = dx_N = dx = \text{const}$, otrzymujemy:

$$dP(x_1, x_2, \dots, x_N) = \left(\frac{1}{\sigma\sqrt{2\pi}} \right)^N e^{-\frac{1}{2\sigma^2} \sum_{i=1}^N (x_i - x_0)^2} (dx)^N. \quad (5.1.16)$$

Wielkości $x_1, \dots, x_N$ występujące we wzorze (5.1.16) są *znane*, są to bowiem wyniki pomiarów. Prawdopodobieństwo $dP(x_1, \dots, x_N)$ jest prawdopodobieństwem uzyskania w serii N -pomiarów wielkości fizycznej x jej wartości zawartych w przedziałach $[x_1, x_1 + dx]$, $[x_2, x_2 + dx]$, $\dots$, $[x_N, x_N + dx]$; możemy je obliczyć, gdy *znamy* lub *założymy* wartości x_0 i σ^2 . Ponieważ wartości x_0 i σ^2 są *nieznane* problem nasz sprowadza się do znalezienia **najbardziej wiarygodnych wartości** $\hat{x}_0$ i $\hat{\sigma}^2$, tzn. takich, dla których prawdopodobieństwo $dP(x_1, x_2, \dots, x_N)$ osiąga *wartość maksymalną*. Wartości te będziemy oznaczali $\hat{x}_0$ i $\hat{\sigma}^2$. Poszukujemy więc takich wartości $\hat{x}_0$ i $\hat{\sigma}^2$, dla których otrzymanie w serii pomiarów wartości $x_1, x_2, \dots, x_N$ jest *najbardziej prawdopodobne*. Te wartości $\hat{x}_0$ i $\hat{\sigma}^2$ będą *najlepszymi przybliżeniami* wartości rzeczywistej i wariancji rozkładu *niezależnie od liczby pomiarów*.

Postępując w podobny sposób można otrzymać najbardziej wiarygodne przybliżenia wartości parametrów określających rozkład wyników pomiarów, także w przypadku, gdy podlegają one innemu rozkładowi niż rozkład normalny. Takie postępowanie nazywa się **metodą największej wiarygodności**.

Aby znaleźć najlepsze przybliżenie wartości rzeczywistej x_0 zwróćmy uwagę na fakt, że prawdopodobieństwo (5.1.16) osiąga *największą* wartość, gdy wykładnik potęgi, tj.

$$\frac{1}{2\sigma^2} \sum_{i=1}^N (x_i - x_0)^2, \quad (5.1.17)$$

jest *najmniejszy*. Aby znaleźć minimum wyrażenia (5.1.17) zróżniczkujmy je względem x_0 i przyrównajmy do zera pierwszą pochodną, otrzymamy wówczas:

$$\sum_{i=1}^N (x_i - x_0) = 0,$$

stąd znajdujemy wartość $\hat{x}_0$, dla której prawdopodobieństwo (5.1.16) osiąga maksimum,

$$\hat{x}_0 = \frac{1}{N} \sum_{i=1}^N x_i = \bar{x}. \quad (5.1.18)$$

Ze wzoru (5.1.18) wynika, że **średnia arytmetyczna serii N pomiarów podlegających rozkładowi Gaussa jest zawsze najlepszym przybliżeniem wartości rzeczywistej**, niezależnie od tego jak duża jest liczba pomiarów N . Wynika stąd, że gdy wyniki pomiarów podlegają rozkładowi normalnemu, to przy *dowolnej* liczbie pomiarów N ich średnia arytmetyczna $\bar{x}$, a nie jakaś inna wielkość x_{in} , obliczona na podstawie wyników pomiarów, **najlepiej przybliży wartość rzeczywistą**, czyli że poza szczególnymi przypadkami, zachodzi nierówność $|\bar{x} - x_0| < |x_{in} - x_0|$. Co więcej, ponieważ, jak wykazano w rozdziale 5.1.2 (wzór (5.1.9)), $\bar{x} \rightarrow x_0$, gdy $N \rightarrow \infty$, to średnia arytmetyczna serii N pomiarów będzie tym lepszym przybliżeniem wartości rzeczywistej im liczba pomiarów N będzie większa. Z przedstawionych powodów średnia arytmetyczna jest bardziej wiarygodnym przybliżeniem wartości rzeczywistej niż jakikolwiek pojedynczy pomiar.

Przejdźmy teraz do znalezienia najlepszego przybliżenia wariancji. W tym celu zróżniczkujemy względem σ wyrażenie (5.1.16) i tak obliczoną pochodną przyrównajmy do zera. Otrzymamy:

$$\begin{aligned} -N \frac{1}{\sigma^{N+1}} \left(\frac{1}{\sqrt{2\pi}} \right)^N \exp \left[-\frac{1}{2\sigma^2} \sum_{i=1}^N (x_i - x_0)^2 \right] + \\ \left(\frac{1}{\sigma\sqrt{2\pi}} \right)^N \frac{1}{\sigma^3} \sum_{i=1}^N (x_i - x_0)^2 \exp \left[-\frac{1}{2\sigma^2} \sum_{i=1}^N (x_i - x_0)^2 \right] = 0. \end{aligned}$$

Z powyższego równania znajdujemy:

$$\frac{1}{\sigma^2} \sum_{i=1}^N (x_i - x_0)^2 = N,$$

a stąd wartość $\hat{\sigma}^2$, dla której prawdopodobieństwo (5.1.16) osiąga wartość maksymalną, wynosi

$$\hat{\sigma}^2 = \frac{1}{N} \sum_{i=1}^N (x_i - x_0)^2 = S_x^2. \quad (5.1.19)$$

Otrzymaliśmy wyrażenie identyczne ze wzorem (3.2.3). Wynika z tego, że najlepszym przybliżeniem wariancji σ^2 jest S_x^2 , a odchylenia standardowego σ jest S_x . Ponieważ w praktyce N jest liczbą niezbyt dużą i wartość

rzeczywista x_0 jest nieznana, we wzorach (5.1.19) i (3.2.3) musimy użyć $\bar{x}$ zamiast x_0 . Wówczas wzór (5.1.19) przejdzie w (3.2.13), tj.

$$S_x^2 = \frac{1}{N-1} \sum_{i=1}^N (x_i - \bar{x})^2. \quad (5.1.20)$$

Tak więc w obliczeniach musimy posługiwać się wzorem (5.1.20). W przypadku dużej liczby pomiarów, gdy $N \gg 1$, wówczas $\bar{x} \rightarrow x_0$ i jedynekę w mianowniku wzoru (5.1.20) możemy pominąć, a wtedy (zgodnie z tym co powiedziano w rozdziale 5.1.2), $S_x^2 \rightarrow \sigma^2$. Oznacza to, że gdy wyniki podlegają rozkładowi Gaussa, to S_x^2 **jest zawsze najlepszym przybliżeniem wariancji** σ^2 , tym lepszym, im większa jest liczba pomiarów N .

Z tego, co dotychczas powiedziano w tym rozdziale, wynika, że wielkości charakteryzujące serię pomiarową (próbę losową) $\bar{x}$ i S_x^2 przybliżają, czyli estymują, wartości x_0 i σ^2 charakteryzujące rozkład normalny, któremu podlegają wyniki pomiarów. Noszą one w statystyce matematycznej nazwę **estymatorów** parametrów rozkładu normalnego.

Ponieważ $\bar{x}$ i S_x^2 są najlepszymi przybliżeniami x_0 i σ^2 , wobec tego najlepszym przybliżeniem rozkładu $N(x_0, \sigma)$ jest rozkład $N(\bar{x}, S_x)$ i zapis $N(10, 1.5)$ będzie oznaczał $\bar{x} = 10$, $S_x = 1.5$.

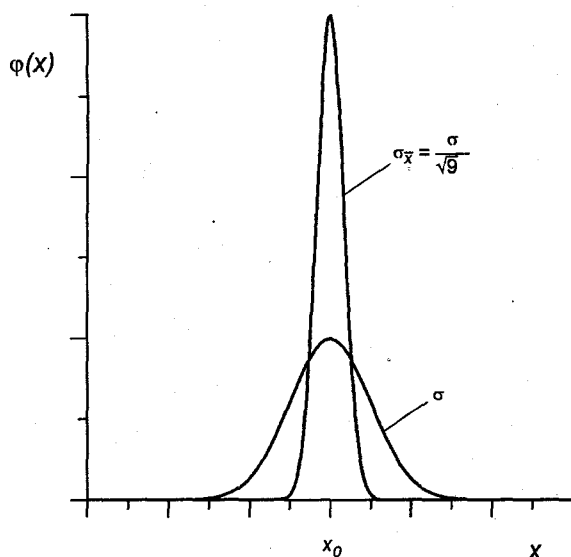
5.1.4. Funkcja rozkładu średniej arytmetycznej

W poprzednim rozdziale wykazaliśmy, że gdy pomiary podlegają rozkładowi Gaussa z wartością rzeczywistą x_0 i wariancją σ^2 , to średnia arytmetyczna

$$\bar{x} = \frac{x_1 + x_2 + \dots + x_N}{N} \quad (5.1.21)$$

jest najlepszym przybliżeniem wartości rzeczywistej x_0 . Ponieważ mierzona wielkość fizyczna x jest zmienną losową, jej wartości zmierzone w serii N pomiarów są otrzymane losowo. Wobec tego ich średnia arytmetyczna $\bar{x}$ będzie miała również wartość otrzymaną losowo, czyli też będzie zmienną losową. Oznacza to, że jeżeli będziemy powtarzali wielokrotnie, przy użyciu tych samych przyrządów pomiarowych, serię N pomiarów (czyli zrobimy ciąg doświadczeń polegających na wykonaniu N pomiarów i obliczeniu na ich podstawie $\bar{x}$), obliczone wartości $\bar{x}$ będą różniły się między sobą. Otrzymamy więc zbiór wartości $\bar{x}_j$. Wartości te będą podlegać pewnemu rozkładowi o wartości rzeczywistej x_0 .

Jak pokazano w Dodatku G rozkład ten, gdy wyniki pomiarów podlegają



Rys. 5.3: Funkcja rozkładu wyników pomiarów o wariancji σ^2 i funkcja rozkładu ich średniej arytmetycznej o wariancji $\sigma_{\bar{x}}^2 = \sigma^2/N$ dla $N = 9$

rozkładowi normalnemu $N(x_0, \sigma)$, jest rozkładem normalnym

$$\varphi(\bar{x}) = \frac{1}{\sigma_{\bar{x}}\sqrt{2\pi}} \exp \left[-\frac{(\bar{x} - x_0)^2}{2\sigma_{\bar{x}}^2} \right], \quad (5.1.22)$$

gdzie

$$\sigma_{\bar{x}}^2 = \frac{\sigma^2}{N}, \quad (5.1.23)$$

N jest liczbą pomiarów. Postępując podobnie jak w rozdziale 5.1.3 można pokazać, że najlepszym przybliżeniem $\sigma_{\bar{x}}^2$ jest $S_{\bar{x}}^2$ dane wzorem (3.4.11), tj.

$$S_{\bar{x}}^2 = \frac{1}{N} S_x^2 = \frac{1}{N(N-1)} \sum_{i=1}^N (x_i - \bar{x})^2.$$

Na rysunku 5.3 przedstawiono funkcję rozkładu normalnego $N(x_0, \sigma)$ i **funkcję rozkładu średniej arytmetycznej** $N(x_0, \sigma_{\bar{x}})$. Zgodnie z wzorami (5.1.22) i (5.1.12) obie funkcje rozkładu mają maksimum przy tej samej wartości x_0 , ale różnią się szerokością. Szerokość rozkładu $N(x_0, \sigma_{\bar{x}})$ jest znacznie mniejsza niż rozkładu wyników pomiarów $N(x_0, \sigma)$. Oznacza to, że rozrzut wartości $\bar{x}$ będzie zawsze zawarty w mniejszym przedziale niż rozrzut wyników pomiarów.

5.1.5. Pomiary o niejednakowej precyzji, łączenie niezależnych pomiarów, średnia ważona

W rezultacie wykonania serii N pomiarów wielkości fizycznej x otrzymujemy ciąg wartości $x_1, x_2, \dots, x_N$. Na podstawie tych wartości możemy obliczyć ich średnią arytmetyczną, odchylenie standardowe serii S_x i odchylenie standardowe średniej serii $S_{\bar{x}}$. Odchylenie standardowe serii S_x jest średnią wartością niepewności każdego z pomiarów danej serii. Ponieważ wszystkie pomiary danej serii mają *tę samą średnią niepewność*, S_x mówimy o nich, że są *jednakowo precyzyjne*. Wszystkie dotychczasowe rozważania przedstawione w tym rozdziale dotyczyły pomiarów o jednakowej precyzji.

Inna jest sytuacja, gdy jeden obserwator wykona pomiary w szeregu niezależnych od siebie seriach pomiarowych, czyli uzyska szereg różnych wartości $\bar{x}$, lub gdy dysponujemy wynikami (wartościami $\bar{x}$) uzyskanymi przez różnych obserwatorów (np. mogą to być wyniki otrzymane w różnych laboratoriach). Wyniki te będą na ogół różniły się od siebie wartościami $\bar{x}$, S_x i $S_{\bar{x}}$. O pomiarach tych powiemy, że są pomiarami o **niejednakowej precyzji**, ponieważ każda z serii pomiarowych ma inną wartość $S_{\bar{x}}$.

Z problemem pomiarów o niejednakowej precyzji spotykamy się często, ponieważ wiele wielkości fizycznych jest mierzonych wielokrotnie i w różnych laboratoriach. Przypuśćmy, że trzech obserwatorów uzyskało niezależnie od siebie następujące wartości $\bar{x}_j$ i $S_{\bar{x}_j}$:

obserwator 1 : $\bar{x}_1, S_{\bar{x}_1}$,

obserwator 2 : $\bar{x}_2, S_{\bar{x}_2}$,

obserwator 3 : $\bar{x}_3, S_{\bar{x}_3}$.

Powstaje pytanie, jak na podstawie powyższych wartości znaleźć najlepsze przybliżenie $\bar{x}_w$ wartości rzeczywistej x_0 ? Na pierwszy rzut oka wydaje się, że należałoby wziąć ich średnią arytmetyczną $(\bar{x}_1 + \bar{x}_2 + \bar{x}_3)/3$. Tak jednak nie jest, ponieważ $S_{\bar{x}_j}$ ($j = 1, 2, 3$) są różne, a przyjęcie średniej arytmetycznej jest równoważne założeniu, że są to pomiary jednakowo precyzyjne. Skoro wartości $S_{\bar{x}_j}$ są w ogólności różne, to słuszne jest takie postępowanie, w którym pomiary o największej precyzji (najmniejsze $S_{\bar{x}_j}$) powinny mieć największy wpływ na wartość poszukiwanego przybliżenia.

W celu znalezienia najlepszego przybliżenia wartości rzeczywistej x_0 założmy, że znamy L różnych średnich arytmetycznych jakiejś wielkości fizycznej wyznaczonych w różnych i **niezależnych seriach pomiarowych** $\bar{x}_1, \bar{x}_2, \dots, \bar{x}_L$ i odchylenia standardowe średnich serii $S_{\bar{x}_1}, S_{\bar{x}_2}, \dots, S_{\bar{x}_L}$, przyjmijmy również, że liczba pomiarów w każdej serii była na tyle duża, że $S_{\bar{x}_j}^2$ jest dobrym przybliżeniem $\sigma_{\bar{x}_j}^2$. Prawdopodobieństwo tego, że w pojedynczej serii pomiarowej otrzymamy wartość średniej arytmetycznej za-

wartą między $\bar{x}_j$ a $\bar{x}_j + d\bar{x}_j$, jest dane wzorem (por. rozdział 5.1.3 i 5.1.4):

$$dP(\bar{x}_j) = \frac{1}{\sigma_{\bar{x}_j} \sqrt{2\pi}} \exp \left[-\frac{(\bar{x}_j - x_0)^2}{2\sigma_{\bar{x}_j}^2} \right] d\bar{x}_j.$$

Wobec tego prawdopodobieństwo $dP(\bar{x}_w)$, że w L seriach pomiarowych różnych i niezależnych uzyskamy wartości $\bar{x}_1, \bar{x}_2, \dots, \bar{x}_L$, jest iloczynem prawdopodobieństw $dP(\bar{x}_1)dP(\bar{x}_2) \dots dP(\bar{x}_L)$. Powtarzając rozważania z rozdziału 5.1.3 otrzymujemy:

$$dP(\bar{x}_1, \bar{x}_2, \dots, \bar{x}_L) = \frac{1}{(\sqrt{2\pi})^L \prod_{j=1}^L \sigma_{\bar{x}_j}} \exp \left[-\sum_{j=1}^L \frac{(\bar{x}_j - x_0)^2}{2\sigma_{\bar{x}_j}^2} \right] (d\bar{x})^L.$$

Prawdopodobieństwo to osiąga maksimum wówczas, gdy wykładnik potęgi ma wartość minimalną, tzn.

$$\frac{1}{2} \sum_{j=1}^L \frac{(\bar{x}_j - x_0)^2}{\sigma_{\bar{x}_j}^2} = \min. \quad (5.1.24)$$

W celu znalezienia minimum wyrażenia (5.1.24) zróżniczkujemy je względem x_0 i pierwszą pochodną przyrównamy do zera. Otrzymamy wówczas

$$\sum_{j=1}^L \frac{(\bar{x}_j - x_0)}{\sigma_{\bar{x}_j}^2} = 0. \quad (5.1.25)$$

Rozwiązując równanie (5.1.25) znajdujemy najbardziej prawdopodobną wartość x_0 , którą oznaczmy $\bar{x}_w$ i która jest najlepszym przybliżeniem wartości rzeczywistej, a mianowicie:

$$\bar{x}_w = \frac{\sum_{j=1}^L w_j \bar{x}_j}{\sum_{j=1}^L w_j}, \quad (5.1.26)$$

gdzie

$$w_j = \frac{1}{\sigma_{\bar{x}_j}^2}. \quad (5.1.27)$$

Średnią $\bar{x}_w$ daną wzorem (5.1.26) nazywamy **średnią ważoną**, a wielkości w_j **wagami pomiarów**. W przypadku pomiarów o jednakowej precyzji $\sigma_{\bar{x}_j} = \text{const}$ i średnia ważona jest równa średniej arytmetycznej. Ze wzoru (5.1.26) wynika, że pomiary o dużym odchyleniu standardowym średniej $\sigma_{\bar{x}_j}$, znacznie większym niż innych pomiarów, mało wpływają na otrzymaną wartość średniej ważonej, tj. na wynik końcowy. Obliczmy teraz odchylenie

standardowe średniej ważonej. W tym celu skorzystajmy z wzoru (3.4.10), który w naszym przypadku przyjmie postać:

$$\sigma_w^2 = \sum_{j=1}^L \left(\frac{\partial \bar{x}_w}{\partial \bar{x}_j} \right)^2 \sigma_{\bar{x}_j}^2. \quad (5.1.28)$$

Obliczając ze wzoru (5.1.26) pochodne $\partial \bar{x}_w / \partial \bar{x}_j$ i podstawiając do (5.1.28) otrzymujemy:

$$\sigma_w^2 = \sum_{j=1}^L \frac{w_j^2 \sigma_{\bar{x}_j}^2}{(\sum_{k=1}^L w_k)^2} = \frac{\sum_{j=1}^L w_j}{(\sum_{k=1}^L w_k)^2} = \frac{1}{\sum_{k=1}^L w_k}, \quad (5.1.29)$$

lub

$$\frac{1}{\sigma_w^2} = \sum_{j=1}^L \frac{1}{\sigma_{\bar{x}_j}^2}. \quad (5.1.30)$$

Należy zaznaczyć, że przy obliczeniach średniej ważonej $\bar{x}_w$ odchylenia standardowe $\sigma_{\bar{x}_j}$ zastępujemy przez ich najlepsze przybliżenia, tj. przez odchylenia standardowe średnich serii $S_{\bar{x}_j}$.

Łączenie niezależnych pomiarów odgrywa istotną rolę przy określaniu najlepszych przybliżeń różnych stałych fizycznych, uniwersalnych i materiałowych, czyli przy tworzeniu tablic wielkości fizycznych. Łączenia pomiarów niezależnych możemy dokonać obliczając średnią ważoną. Poprawnego obliczenia średniej ważonej nie możemy jednak dokonać bez wstępnej analizy posiadanych danych, która polega na stwierdzeniu czy wartości $\bar{x}_j$ ($j = 1, 2, \dots, L$) są zgodne. Na podstawie tego, co powiedziano o zgodności wyników pomiarów w rozdziale 1.1.6, wyniki będziemy uważali za **zgodne**, jeżeli rozbieżność $|\bar{x}_{\max} - \bar{x}_{\min}|$ spełnia nierówność:

$$|\bar{x}_{\max} - \bar{x}_{\min}| \leq 3S_{\bar{x}_{\max}} + 3S_{\bar{x}_{\min}},$$

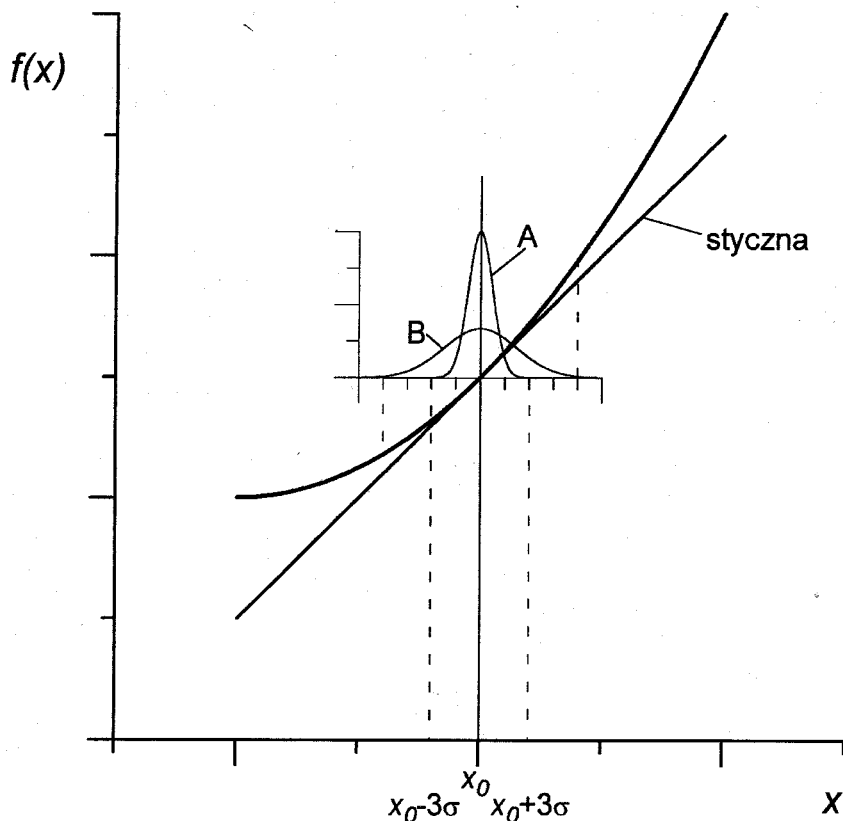
($\bar{x}_{\max}$ i $\bar{x}_{\min}$ oznaczają odpowiednio największą i najmniejszą wartość ze zbioru L wartości $\bar{x}_j$). Jeżeli rozbieżność $|\bar{x}_{\max} - \bar{x}_{\min}|$ jest *istotnie większa* niż $3S_{\bar{x}_{\max}} + 3S_{\bar{x}_{\min}}$, to wyniki uważamy za **sprzeczne**. Wtedy należy wyniki dokładnie przeanalizować, ponieważ któryś z nich może być obciążony błędem systematycznym. Wynika to z faktu, że funkcja rozkładu (5.1.12) dla $x_0 + 3S_{\bar{x}}$ bądź $x_0 - 3S_{\bar{x}}$ jest prawie równa zero i w tych granicach powinna być zawarta wartość rzeczywista oraz z tego, że wielkość $3S_{\bar{x}}$ możemy uważać za **błąd maksymalny** (wyjaśnimy to w rozdziale 5.2.2). Jeżeli wyniki są zgodne to możemy przystąpić do obliczania średniej ważonej.

5.1.6. Funkcja rozkładu wielkości złożonej

Niech z będzie wielkością złożoną. Przyjmijmy dla większej przejrzystości rachunków, że wielkość złożona z jest funkcją dwóch **niezależnych zmiennych losowych** x i y oraz że każda z tych zmiennych losowych podlega **rozkładowi Gaussa**, a ich wartości rzeczywiste wynoszą odpowiednio x_0 i y_0 , zaś wartość rzeczywista wielkości złożonej wynosi $z_0 = z(x_0, y_0)$. Ponieważ wielkość złożona $z = z(x, y)$ jest funkcją zmiennych losowych, to i ona sama jest zmienną losową. **W ogólności funkcja rozkładu prawdopodobieństwa zmiennej losowej $z(x, y)$, pomimo że zmienne losowe x i y podlegają rozkładowi Gaussa, nie będzie rozkładem Gaussa.** Jak można wykazać, na podstawie twierdzeń podanych w Dodatku G, jedynie w szczególnym przypadku, gdy wielkość złożona z jest liniową funkcją zmiennych losowych x i y , jej rozkład jest rozkładem normalnym. Jak pokazano w rozdziale 2.2 (wzór (2.2.4)) i rozdziale 3.3 (wzór (3.3.3)) wielkość złożoną $z = z(x, y)$ można przybliżyć funkcją liniową wielkości x i y w otoczeniu dowolnego punktu, np. punktu (x_0, y_0) , jeśli różnice $x - x_0$ i $y - y_0$ są na tyle małe, że rozwijając $z(x, y)$ na szereg Taylora możemy ograniczyć się do wyrazów liniowych w zmiennych x oraz y , czyli w przypadku gdy:

$$z(x, y) \cong z(x_0, y_0) + \left(\frac{\partial z}{\partial x}\right)_{x_0, y_0} (x - x_0) + \left(\frac{\partial z}{\partial y}\right)_{x_0, y_0} (y - y_0). \quad (5.1.31)$$

Ogólne warunki stosowalności tego przybliżenia zostały omówione w rozdziale 2.2. Poniżej przedstawimy **warunki stosowalności rozwinięcia** (5.1.31) w przypadku, gdy wielkości mierzone bezpośrednio mają rozkład normalny. Jeżeli wyniki pomiarów danej wielkości fizycznej x podlegają rozkładowi Gaussa, to prawie wszystkie otrzymane wartości są zawarte między $x_0 - 3\sigma_x$ a $x_0 + 3\sigma_x$. Z tego wynika, że o dokładności przybliżenia (5.1.31) decyduje wielkość $3\sigma_x$, tzn. że dla $x - x_0$ zawartych w przedziale $[-3\sigma_x, +3\sigma_x]$ rozwinięcie (5.1.31) musi być dobrym przybliżeniem. Na rysunku 5.4 przedstawiono to w przypadku, gdy wielkość złożona jest funkcją jednej zmiennej losowej x . Wnioski, które z tego wynikają, stosują się do przypadku wielu zmiennych losowych, do każdej zmiennej osobno. Jak widać z rysunku 5.4, w przypadku gdy rozkład wyników pomiarów jest dany krzywą A , stosowanie rozwinięcia (5.1.31) jest uzasadnione. W przypadku rozkładu B stosowanie rozwinięcia (5.1.31) jest nieuzasadnione, ponieważ na skrzydłach funkcji rozkładu styczna poprowadzona w punkcie x_0 zaczyna istotnie różnić się od krzywej i rozkład zmiennej losowej złożonej w ogólności nie będzie dobrze przybliżony rozkładem normalnym. W dalszych rozważaniach tego rozdziału przyjmujemy, że zachodzi przypadek odpowiadający rozkładowi A .



Rys. 5.4: Wyjaśnienie stosowalności rozwinięcia $z = f(x)$ ograniczonego do wyrazu liniowego

Aby rozwinięcie (5.1.31) było dobrym przybliżeniem w naszym przypadku wystarczy, żeby obliczona wartość $z(\bar{x} \pm 3S_x, \bar{y})$ i wartość $z_R(\bar{x} \pm 3S_x, \bar{y})$ obliczona przy użyciu rozwinięcia (5.1.31) różniły się tylko nieznacznie. Formalnie można to wyrazić następująco:

$$\left| \frac{z(\bar{x} \pm 3S_x, \bar{y}) - z_R(\bar{x} \pm 3S_x, \bar{y})}{3S_x \left(\frac{\partial z}{\partial x} \right)_{\bar{x}, \bar{y}}} \right| \ll 1. \quad (5.1.32)$$

Nierówność ta powinna być spełniona ze względu na każdą zmienną losową, od której jest zależna wielkość złożona z . Jest to warunek konieczny, lecz niewystarczający. Szczególnymi przypadkami funkcji $z(x_1, x_2, \dots, x_r)$, kiedy powyższy warunek jest niewystarczający, nie będziemy się zajmowali. W przeważającej części przypadków spotykanych w praktyce wystarczy jedynie posłużyć się powyższą nierównością. Należy zaznaczyć, że wybór granicy, poniżej której można przyjąć, że lewa strona nierówności (5.1.32) jest

dużo mniejsza od jedności, jest w pewnym sensie „dowolny”. Zależy on od charakteru zależności funkcyjnej $z(x_1, x_2, \dots, x_r)$, wymagań stawianych wynikom pomiarów, oraz możliwości uzyskania jak największej ich precyzji. W praktyce, w jednych przypadkach granica ta może być 0.01, w innych np. 0.001, a w jeszcze innych wystarczy tylko 0.1. Ta „dowolność” jest usprawiedliwiona tym, że *błąd pomiaru oszacowujemy*, a nie ściśle obliczamy.

Korzystając z warunku (5.1.32) zbadamy poprawność stosowania rozwinięcia (3.3.3) i wzorów otrzymanych przy jego użyciu w przypadku wyników pomiarów podanych w przykładzie z rozdziału 3.4. W przykładzie tym wyznaczaliśmy przyspieszenie ziemskie g przy użyciu wahadła matematycznego i dokonaliśmy obliczenia wielkości statystycznych charakteryzujących błędy przypadkowe. Przyspieszenie ziemskie $g = 4\pi l/T^2$ jest więc liniową funkcją długości wahadła l . Wobec tego rozwinięcie (3.3.3) dla zmiennej losowej l jest zawsze dokładne i należy jedynie stwierdzić, czy warunek (5.1.32) jest spełniony dla zmiennej losowej T — okresu wahadła. Musimy zbadać, czy jest on spełniony w punktach końcowych przedziału $[\bar{T} - 3S_T, \bar{T} + 3S_T]$. W rozpatrywanym przypadku nierówność (5.1.32) ma postać:

$$\left| \frac{\frac{4\pi^2 \bar{l}}{(\bar{T} \pm 3S_T)^2} - \left(\frac{4\pi^2 \bar{l}}{\bar{T}^2} \pm \frac{4\pi^2 \bar{l}}{\bar{T}^3} 6S_T \right)}{-\frac{4\pi^2 \bar{l}}{\bar{T}^3} 6S_T} \right| = \left| -\frac{\bar{T}^3}{6S_T(\bar{T} \pm 3S_T)^2} + \frac{\bar{T}}{6S_T} \pm 1 \right| \ll 1.$$

Podstawiając do lewej strony tej nierówności $\bar{T}$ i S_T obliczone w rozdziale 3.4, a mianowicie $\bar{T} = 2.007$ s i $S_T = 0.0149$ s, otrzymujemy, że jest ona równa:

- dla punktu końcowego przedziału $\bar{T} - 3S_T$ 0.034;
- dla punktu końcowego przedziału $\bar{T} + 3S_T$ 0.032.

Wykonując analogiczne obliczenia dla zmiennej losowej $\bar{l}$ otrzymujemy po lewej stronie nierówności zero. Widać stąd, że warunek (5.1.32) jest spełniony i możemy w tym przypadku stosować rozwinięcie (3.3.3).

Poniżej wykazemy, że jeżeli pomiary niezależnych wielkości fizycznych x i y podlegają rozkładowi normalnemu oraz jeżeli rozwinięcie (5.1.31) wielkości złożonej $z(x, y)$ jest dobrym przybliżeniem, to wyznaczona wielkość z podlega również rozkładowi normalnemu w otoczeniu punktu (x_0, y_0) . Niech błędy bezwzględne wielkości x , y i z będą odpowiednio równe $\delta_x = x - x_0$, $\delta_y = y - y_0$ i $\delta_z = z - z_0$. Korzystając ze wzoru (5.1.31) znajdujemy

$$\delta_z = \left(\frac{\partial z}{\partial x} \right)_{x_0, y_0} \delta_x + \left(\frac{\partial z}{\partial y} \right)_{x_0, y_0} \delta_y. \quad (5.1.33)$$

Ponieważ każda ze zmiennych losowych δ_x i δ_y podlega rozkładowi normalnemu (por. rozdział 5.1.1) o wariancji odpowiednio równej σ_x^2 i σ_y^2 oraz pochodne $(\partial z/\partial x)_{x_0, y_0}$ i $(\partial z/\partial y)_{x_0, y_0}$ są wielkościami stałymi, to z twierdzenia 4 w Dodatku G wynika, że każda ze zmiennych pomocniczych

$$\begin{aligned}\tilde{\delta}_x &= \left(\frac{\partial z}{\partial x}\right)_{x_0, y_0} \delta_x, \\ \tilde{\delta}_y &= \left(\frac{\partial z}{\partial y}\right)_{x_0, y_0} \delta_y,\end{aligned}\tag{5.1.34}$$

ma rozkład normalny o wariancji odpowiednio równej

$$\begin{aligned}\tilde{\sigma}_x^2 &= \left(\frac{\partial z}{\partial x}\right)_{x_0, y_0}^2 \sigma_x^2, \\ \tilde{\sigma}_y^2 &= \left(\frac{\partial z}{\partial y}\right)_{x_0, y_0}^2 \sigma_y^2.\end{aligned}\tag{5.1.35}$$

Ze wzorów (5.1.33) i (5.1.34) otrzymujemy:

$$\delta_z = \tilde{\delta}_x + \tilde{\delta}_y.\tag{5.1.36}$$

Wobec tego, zgodnie z twierdzeniem 1 w Dodatku G, rozkład zmiennej losowej δ_z (a tym samym i z) będzie **spłotem** rozkładów zmiennych losowych $\tilde{\delta}_x$ i $\tilde{\delta}_y$. Ponieważ rozkłady zmiennych losowych $\tilde{\delta}_x$ i $\tilde{\delta}_y$ są normalne, a więc zgodnie z twierdzeniem 2 w Dodatku G, rozkład zmiennej losowej δ_z będzie rozkładem normalnym o wariancji

$$\sigma_z^2 = \tilde{\sigma}_x^2 + \tilde{\sigma}_y^2.\tag{5.1.37}$$

Podstawiając (5.1.34) do (5.1.37) otrzymujemy

$$\sigma_z^2 = \left(\frac{\partial z}{\partial x}\right)_{x_0, y_0}^2 \sigma_x^2 + \left(\frac{\partial z}{\partial y}\right)_{x_0, y_0}^2 \sigma_y^2.$$

Zgodnie z tym, co powiedziano w rozdziale 5.1.3, gdy posługujemy się funkcją rozkładu normalnego, nieznane wartości $z_0 = z(x_0, y_0)$ i σ_z^2 musimy zastąpić przez ich najlepsze przybliżenia, tj. $\bar{z}$ i S_z^2 , które w przypadku r niezależnych zmiennych losowych jest dane wzorem (3.4.8). Oczywiście jeśli rozwinięcie (5.1.31) jest dobrym przybliżeniem, to $\bar{z} = z(\bar{x}, \bar{y})$ (por. rozdział 3.3).

Jak można pokazać na podstawie twierdzeń podanych w Dodatku G, funkcja rozkładu prawdopodobieństwa wartości $\bar{z} = z(\bar{x}, \bar{y})$ będzie rozkładem normalnym z wariancją $\sigma_{\bar{z}}^2$, którą w obliczeniach należy zastąpić jej najlepszym przybliżeniem $S_{\bar{z}}^2$, danym wzorem (3.4.17). Otrzymane wyniki

są słuszne dla r niezależnych zmiennych losowych $x_1, x_2, \dots, x_r$ pod warunkiem, że każda z nich ma rozkład normalny oraz, że rozwinięcie (5.1.31) jest dobrym przybliżeniem.

Należy jeszcze raz podkreślić, że **funkcja rozkładu prawdopodobieństwa wielkości złożonej $z(x_1, x_2, \dots, x_r)$ jest rozkładem normalnym w otoczeniu punktu $(x_{0,1}, x_{0,2}, \dots, x_{0,r})$, jeśli niezależne zmienne losowe mają rozkład normalny oraz rozwinięcie (5.1.31) jest dobrym przybliżeniem.** Jak łatwo pokazać wariancja rozkładu σ_z^2 jest dana wzorem:

$$\sigma_z^2 = \sum_{i=1}^r \left(\frac{\partial z}{\partial x_i} \right)_{x_{0,1}, x_{0,2}, \dots, x_{0,r}}^2 \sigma_{x_i}^2. \quad (5.1.38)$$

Na zakończenie zajmiemy się przypadkiem, gdy jedno lub więcej z założeń przyjętych na początku tego rozdziału nie jest spełnione, a mimo tego wielkość złożona z ma rozkład normalny w całym zakresie zmian wielkości mierzonych bezpośrednio. Oczywiście w takim przypadku nie można stosować otrzymanych powyżej wzorów. Ponieważ wielkość złożona z ma rozkład normalny, całe zagadnienie sprowadza się do obliczania wielkości $\bar{z}$, S_z^2 oraz $S_{\bar{z}}$ przybliżających parametry tego rozkładu. Przyjmijmy dla większej przejrzystości wzorów, że $z = z(x, y)$, tzn. zależy od dwóch niezależnych wielkości mierzonych bezpośrednio. Przyjmijmy również, że wielkość x zmierzylśmy N razy i otrzymaliśmy zbiór wartości $x_1, x_2, \dots, x_N$, a wielkość y zmierzylśmy M razy i otrzymaliśmy wartości $y_1, y_2, \dots, y_M$. Wówczas każdej parze x_i, y_k odpowiada wartość $z_{ik} = z(x_i, y_k)$, otrzymamy więc NM wartości z_{ik} . Ich średnia arytmetyczna, zgodnie z tym co powiedziano w rozdziale 3.3 (wzór (3.3.2)), jest:

$$\bar{z} = \frac{1}{NM} \sum_{i=1}^N \sum_{k=1}^M z_{ik}. \quad (5.1.39)$$

Postępując tak samo jak w przypadku wielkości mierzonej bezpośrednio (patrz rozdział 3.2) łatwo pokazać, że

$$S_z^2 = \frac{1}{NM-1} \sum_{i=1}^N \sum_{k=1}^M (z_{ik} - \bar{z})^2 \quad (5.1.40)$$

oraz (patrz rozdział 3.4):

$$S_{\bar{z}} = \frac{S_z}{\sqrt{NM}}. \quad (5.1.41)$$

Jak widać z wyżej przedstawionych wywodów, w przypadku gdy wielkość złożona z ma rozkład normalny $N(z_0, \sigma_z)$ w całym zakresie zmian wielkości mierzonych bezpośrednio, to aby obliczyć jego parametry należy postępować *tak samo* jak w przypadku wielkości mierzonej bezpośrednio.

5.2. Rozkład normalny standaryzowany i jego zastosowanie do oceny błędu pomiaru

5.2.1. Standaryzacja rozkładu normalnego

Dystrybuanta rozkładu normalnego (5.1.13) nie wyraża się przez funkcje elementarne, tzn. całki tej nie możemy obliczyć analitycznie. Możemy natomiast obliczyć jej wartość numerycznie. Wyrażenie podcałkowe, tj. funkcja rozkładu normalnego $N(x_0, \sigma)$, ma jednak niedogodność polegającą na tym, że gdy wartość rzeczywistą x_0 przybliżymy przez średnią arytmetyczną $\bar{x}$, a wariancję σ^2 przez S_x^2 , to parametry rozkładu będą funkcjami wartości zmiennej losowej. Dlatego też gdybyśmy chcieli stosować dystrybuantę (5.1.13) do oceny błędu pomiaru, to musielibyśmy każdorazowo obliczać tę całkę. Aby uniknąć konieczności każdorazowego obliczania całki (5.1.13) wprowadźmy nową *bezwymiarową zmienną*

$$u_i = \frac{x_i - x_0}{\sigma}. \quad (5.2.1)$$

Zmienna ta jest funkcją zmiennej losowej i z tego powodu sama jest zmienną losową. Nosi ona nazwę **zmiennej losowej standaryzowanej**, ponieważ wyraża różnicę $x_i - x_0$ w standardowych jednostkach, a mianowicie w jednostkach σ . Jednostki te są więc identyczne dla *dowolnych zmiennych losowych* podlegających rozkładowi Gaussa. Jak wynika ze wzoru (5.2.1), zmienna losowa standaryzowana przyjmuje zarówno wartości dodatnie, jak i ujemne; jest zawarta w przedziale $(-\infty, +\infty)$.

Prawdopodobieństwo $dP(x)$ tego, że wynik pomiaru jest zawarty między x a $x + dx$ jest $dP(x) = \varphi(x)dx$ (por. rozdział 5.1.2), zaś prawdopodobieństwo $dP(u)$, że zmienna losowa standaryzowana przyjmie wartość zawartą między u a $u + du$ jest $dP(u) = \varphi_u(u)du$. Funkcja $\varphi_u(u)$ jest funkcją rozkładu prawdopodobieństwa zmiennej losowej standaryzowanej. Ponieważ zmienna losowa standaryzowana u jest związana ze zmienną losową x zależnością (5.2.1), $dP(x) = dP(u)$, a więc

$$dP(u) = \varphi(x)dx. \quad (5.2.2)$$

Różnicę $x - x_0$ znajdujemy ze wzoru (5.2.1); wynosi ona

$$x - x_0 = \sigma u. \quad (5.2.3)$$

Różniczkując (5.2.3) znajdujemy

$$dx = \sigma du. \quad (5.2.4)$$

Podstawiając (5.2.3) i (5.2.4) do (5.2.2), otrzymujemy

$$dP(u) = \varphi(\sigma u) \sigma du = \varphi_u(u) du,$$

czyli

$$\varphi_u(u) = \varphi(\sigma u) \sigma.$$

Ponieważ $\varphi(x)$ dane jest wzorem (5.1.12), to funkcja rozkładu prawdopodobieństwa zmiennej losowej standaryzowanej przyjmuje postać

$$\varphi_u(u) = \frac{1}{\sqrt{2\pi}} e^{-\frac{u^2}{2}}. \quad (5.2.5)$$

Rozkład ten nazywamy **rozkładem normalnym standaryzowanym**. Obliczmy teraz średnią $\bar{u}_R$ i wariancję σ_u^2 rozkładu (5.2.5). Zgodnie ze wzorem (4.2.11)

$$\bar{u}_R = \int_{-\infty}^{+\infty} u \varphi_u(u) du = \frac{1}{\sqrt{2\pi}} \int_{-\infty}^{+\infty} u e^{-\frac{u^2}{2}} du.$$

Występująca tu całka jest równa zero (patrz Dodatek F), a więc

$$\bar{u}_R = 0. \quad (5.2.6)$$

Wariancja σ_u^2 , zgodnie ze wzorem (4.2.12), wynosi

$$\sigma_u^2 = \int_{-\infty}^{+\infty} (u - \bar{u}_R)^2 \varphi_u(u) du,$$

ale $\bar{u}_R = 0$, a więc:

$$\sigma_u^2 = \int_{-\infty}^{+\infty} u^2 \varphi_u(u) du = \frac{1}{\sqrt{2\pi}} \int_{-\infty}^{+\infty} u^2 e^{-\frac{u^2}{2}} du. \quad (5.2.7)$$

Całka we wzorze (5.2.7) wynosi $\sqrt{2\pi}$ (patrz Dodatek F), wobec tego

$$\sigma_u^2 = 1. \quad (5.2.8)$$

Ponieważ $\bar{u}_R = 0$ i $\sigma_u^2 = 1$, funkcję rozkładu zmiennej losowej standaryzowanej oznaczamy jako $N(0, 1)$.

Ze wzorów (5.2.5), (5.2.6) i (5.2.8) wynika, że niezależnie od rozrzutu wyników pomiarów:

- szerokość rozkładu jest stała,
- wartość funkcji rozkładu w maksimum $\varphi_u(0) = \frac{1}{\sqrt{2\pi}}$.

Dotychczas omawialiśmy zmienną losową standaryzowaną i jej rozkład dla pojedynczych pomiarów podlegających rozkładowi normalnemu $N(x_0, \sigma)$. Podobnie jak dla pojedynczych pomiarów można zdefiniować **zmienną losową standaryzowaną dla średniej arytmetycznej**, a mianowicie

$$u = \frac{\bar{x} - x_0}{\sigma_{\bar{x}}}. \quad (5.2.9)$$

W rozdziale 5.1.4 pokazaliśmy, że funkcja rozkładu średniej arytmetycznej, gdy pomiary podlegają rozkładowi normalnemu $N(x_0, \sigma)$, jest rozkładem normalnym $N(x_0, \sigma_{\bar{x}})$. Z tego wynika, że zmienna u dana wyrażeniem (5.2.9) ma funkcję rozkładu daną wzorem (5.2.5); dowód pozostawiamy czytelnikowi.

Pozostało nam jeszcze do omówienia zagadnienie pomiarów pośrednich. Przyjmijmy dla większej przejrzystości obliczeń, że wielkość złożona z jest funkcją dwóch niezależnych wielkości x i y mierzonych bezpośrednio, tzn. $z = z(x, y)$. Jeżeli wielkość złożona $z = z(x, y)$ *ma rozkład normalny*, lub gdy są *spełnione warunki podane w rozdziale 5.1.6*, to zmienną standaryzowaną u możemy zdefiniować analogicznie do (5.2.1) jako

$$u = \frac{z(x, y) - z(x_0, y_0)}{\sigma_z}, \quad (5.2.10)$$

a zmienną standaryzowaną dla wartości średniej $\bar{z}$ analogicznie do (5.2.9):

$$u = \frac{\bar{z} - z(x_0, y_0)}{\sigma_{\bar{z}}}. \quad (5.2.11)$$

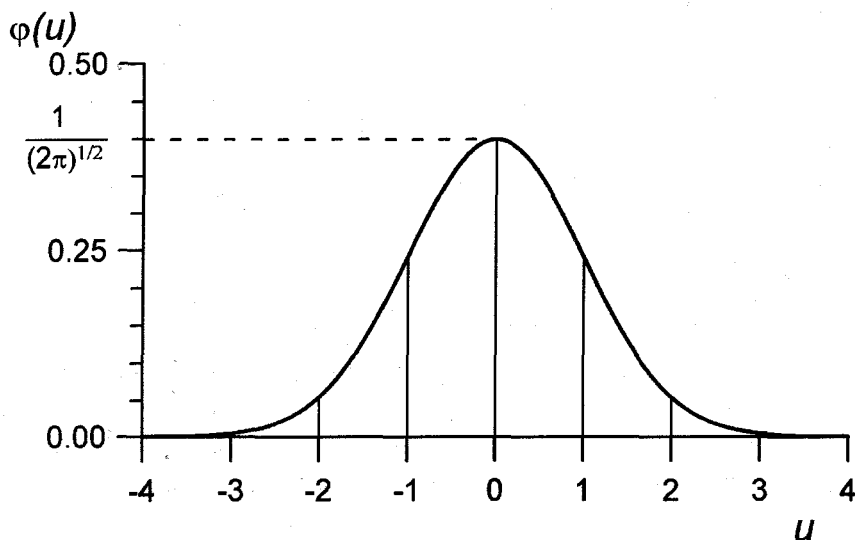
Funkcja rozkładu będzie wtedy identyczna z daną wzorem (5.2.5).

Funkcja rozkładu normalnego standaryzowanego ma charakter uniwersalny, ponieważ *nie zależy* od zmiennej losowej i jej odchylenia standardowego oraz może być stosowana do pomiarów bezpośrednich i pośrednich wtedy, gdy podlegają one rozkładowi normalnemu.

Ze względu na swoje własności funkcja rozkładu $\varphi_u(u)$ jest stabelaryzowana. Tablice tej funkcji podano w Dodatku H, zaś na rysunku 5.5 przedstawiono jej wykres.

Czasami bywa używana inna postać funkcji rozkładu normalnego standaryzowanego, a mianowicie:

$$\varphi_s(s) = \frac{1}{\sqrt{\pi}} e^{-s^2}. \quad (5.2.12)$$



Rys. 5.5: Funkcja rozkładu normalnego standaryzowanego

Dlatego korzystając z tablic funkcji rozkładu normalnego standaryzowanego musimy *zawsze sprawdzić* czy stabelaryzowana jest funkcja (5.2.5), czy (5.2.12).

5.2.2. Dystrybuanta rozkładu normalnego standaryzowanego i jej zastosowania

Zgodnie ze wzorem (4.2.7) **dystrybuanta $\Phi_u(a)$ rozkładu normalnego standaryzowanego** wyraża się następująco:

$$\Phi_u(a) = \frac{1}{\sqrt{2\pi}} \int_{-\infty}^a e^{-\frac{u^2}{2}} du. \quad (5.2.13)$$

Całka (5.2.13) nie wyraża się przez funkcje elementarne, wartości jej możemy jedynie obliczyć numerycznie i z tego powodu została ona stabelaryzowana. Tablice wartości funkcji $\Phi_u(a)$ podano w Dodatku H.

Prawdopodobieństwo tego, że zmienna losowa standaryzowana u będzie zawarta w przedziale $[a, b]$ wynosi

$$P_u(a \leq u \leq b) = \Phi_u(b) - \Phi_u(a) = \Phi_u(a, b). \quad (5.2.14)$$

Korzystając z własności całek oznaczonych i wzoru (5.2.13) otrzymujemy:

$$P_u(a \leq u \leq b) = \Phi_u(a, b) = \frac{1}{\sqrt{2\pi}} \int_a^b e^{-\frac{u^2}{2}} du. \quad (5.2.15)$$

Tak więc prawdopodobieństwo, że wartość zmiennej standaryzowanej u spełnia nierówność

$$a \leq u \leq b, \quad (5.2.16)$$

wynosi $P_u(a \leq u \leq b)$. Wzory (5.2.13) – (5.2.16) są całkowicie ogólne i mają różne zastosowania, zależne od tego, czy zmienna standaryzowana u jest zdefiniowana wzorem (5.2.1) lub (5.2.10), czy też wzorem (5.2.9) lub (5.2.11). Rozważania ograniczymy do przypadku dla nas najważniejszego, tj. do zagadnień związanych z oszacowaniem błędu pomiaru. Funkcja rozkładu prawdopodobieństwa $\varphi_u(u)$ jak i $\varphi(x)$ oraz $\varphi(\delta)$ są funkcjami symetrycznymi. Oznacza to, że $\varphi_u(-u) = \varphi_u(u)$ czy też $\varphi(-\delta) = \varphi(\delta)$, inaczej mówiąc oznacza to, że jednakowo prawdopodobne jest wystąpienie błędu $+\delta$ i $-\delta$. Z tego powodu, przy ocenie błędu pomiaru *najczęściej* poszukujemy prawdopodobieństwa, że $|u|$ różni się od zera co najwyżej o a , czyli że spełniona jest nierówność:

$$-a \leq u \leq a. \quad (5.2.17)$$

Wobec tego całka (5.2.15) przyjmie postać:

$$P_u(-a \leq u \leq a) = \frac{1}{\sqrt{2\pi}} \int_{-a}^a e^{-\frac{u^2}{2}} du. \quad (5.2.18)$$

Ponieważ funkcja podcałkowa jest symetryczna, to

$$\Phi_u(-a, a) = P_u(-a \leq u \leq a) = \sqrt{\frac{2}{\pi}} \int_0^a e^{-\frac{u^2}{2}} du.$$

(5.2.19)

Tabela tej całki została podana w Dodatku H (tablica III). Bardzo często jest używana inna postać funkcji (5.2.19), którą nazywamy **funkcją błędu** bądź **normalną całką błędu** i oznaczamy $\text{erf}(t)$. Jest ona równa:

$$\text{erf}(t) = \frac{2}{\sqrt{\pi}} \int_0^t e^{-s^2} ds. \quad (5.2.20)$$

Funkcję $\text{erf}(t)$ otrzymujemy z funkcji (5.2.19) przez podstawienie $s = u/\sqrt{2}$, zaś $t = a/\sqrt{2}$. Z tego powodu *zawsze*, gdy korzystamy z tablic funkcji błędu musimy sprawdzić, czy stabilizowana jest funkcja (5.2.19), czy (5.2.20).

Podstawiając do (5.2.17) u dane wzorem (5.2.9) lub (5.2.11) znajdujemy, że

$$-a \leq \frac{\bar{x} - x_0}{\sigma_{\bar{x}}} \leq a. \quad (5.2.21)$$

Rozwiązując powyższą nierówność otrzymujemy:

$$\bar{x} - a\sigma_{\bar{x}} \leq x_0 \leq \bar{x} + a\sigma_{\bar{x}}. \quad (5.2.22)$$

Zgodnie z tym co powiedziano w rozdziale 5.1.3, w nierówności (5.2.22) nieznaną wartość $\sigma_{\bar{x}}$ musimy przybliżyć przez $S_{\bar{x}}$; wówczas nierówność ta przyjmie postać:

$$\boxed{\bar{x} - aS_{\bar{x}} \leq x_0 \leq \bar{x} + aS_{\bar{x}}.} \quad (5.2.23)$$

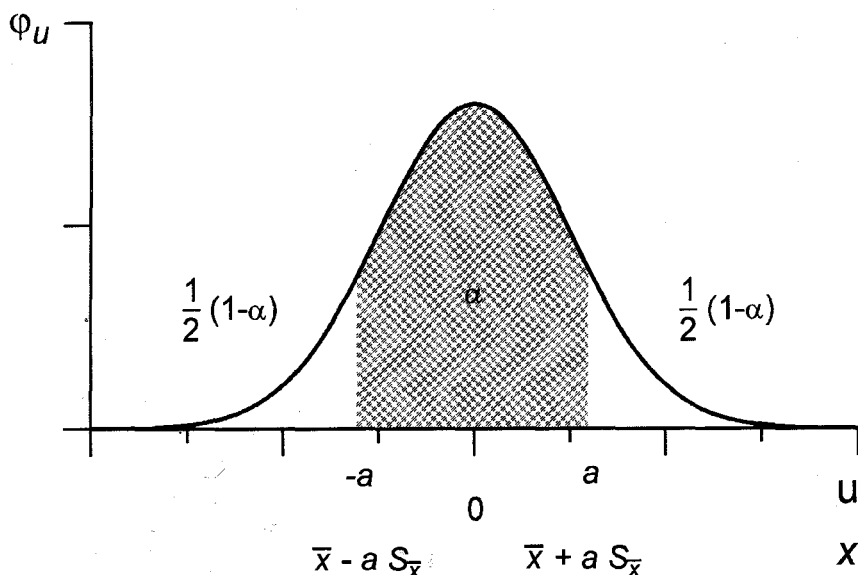
Z nierówności tej, i z tego co powiedziano wyżej, wynika, że wartość rzeczywista x_0 mierzonej wielkości fizycznej jest zawarta w przedziale $[\bar{x} - aS_{\bar{x}}, \bar{x} + aS_{\bar{x}}]$ z prawdopodobieństwem $\Phi_u(-a, a)$. Jeżeli $a = 3$, to, jak wynika z tablicy III w Dodatku H, $\Phi_u(-3, 3) = 0.9973$, co oznacza, że wartość rzeczywista x_0 jest zawarta w przedziale $[\bar{x} - 3S_{\bar{x}}, \bar{x} + 3S_{\bar{x}}]$ z prawdopodobieństwem 0.9973. Z tego powodu można przyjąć, że wielkość $3S_{\bar{x}}$ jest **błędem maksymalnym średniej arytmetycznej**, spowodowanym błędami przypadkowymi. *Nie wyklucza* to jednak możliwości, iż może się zdarzyć, że błąd będzie większy niż $3S_{\bar{x}}$, ale prawdopodobieństwo otrzymania takiego wyniku wynosi mniej niż 0.003.

W ten sposób, w przypadku gdy wyniki pomiarów bezpośrednich podlegają rozkładowi normalnemu i liczba pomiarów jest duża, otrzymaliśmy odpowiedź na pytanie postawione w rozdziale 1.1 o cel teorii błędów. We wzorze (5.2.22) $\bar{x}$ odpowiada $\tilde{x}$ we wzorze (1.1.4), zaś $aS_{\bar{x}}$ odpowiada Δ we wzorze (1.1.4).

Przedział określony nierównością (5.2.23) nazywamy **przedziałem ufności**, a prawdopodobieństwo $P_u(-a \leq u \leq a) = \alpha$, że spełniona jest nierówność (5.2.23) **poziomem** lub **współczynnikiem ufności**.

Należy zaznaczyć, że z symetrii rozkładu Gaussa wynika, iż wybór symetrycznego przedziału ufności *zapewnia przy danym poziomie ufności najmniejszą jego szerokość*. Mogą być również sytuacje, kiedy zachodzi konieczność posługiwania się niesymetrycznym przedziałem ufności (określonym nierównością (5.2.16)), nie będziemy ich jednak analizowali.

Wielkość $1 - \alpha$ nazywamy **poziomem** lub **współczynnikiem istotności**. Współczynnik istotności jest prawdopodobieństwem tego, że wartość rzeczywista znajduje się poza przedziałem ufności. W niektórych podręcznikach (patrz np. [13]) poziom ufności jest oznaczany przez $1 - \alpha$, a współczynnik istotności przez α . Dlatego korzystając z tablic funkcji $\Phi_u(-a, a)$ lub $\text{erf}(t)$ należy *zawsze zwracać uwagę* na to, jak oznaczony jest poziom



Rys. 5.6: Funkcja rozkładu $\varphi_u(u)$. Zakreskowana powierzchnia jest równa poziomowi ufności α dla przedziału $\bar{x} \pm a S_{\bar{x}}$

ufności, tzn. czy przez α , czy przez $1 - \alpha$, oraz która z tych wielkości jest tablicowana.

Na rysunku 5.6 przedstawiono wykres funkcji $\varphi_u(u)$. Zakreskowana część powierzchni pod krzywą, między $u = -a$ i $u = +a$, jest równa $\Phi_u(-a, a) = \alpha$. Każda z części niezakreskowanych jest równa $\frac{1}{2}(1 - \alpha)$, tj. połowie współczynnika istotności.

W celu znalezienia przedziału i poziomu ufności możemy postępować dwojako:

- *ustalamy granice przedziału ufności, tzn. współczynnik a i następnie z tablic dystrybuanty $\Phi_u(a)$ lub funkcji (5.2.18) czy (5.2.19) odczytujemy poziom ufności α ,*
- *ustalamy poziom ufności α i z tablic funkcji (5.2.18) lub (5.2.19) znajdujemy współczynnik a , który następnie wstawiamy do nierówności (5.2.23).*

Przykłady

1) Pomiary czasów wypływu przedstawione w tabeli 4.2 dają $\bar{t} = 17.6153$ s, $S_t = 0.2280$ s, $S_{\bar{t}} = 0.02403$ s. Przyjmijmy, że $a = 1.6$, czyli że $a S_{\bar{t}} = 0.0384$ s. Będziemy więc poszukiwać prawdopodobieństwa, że wartość rzeczywista spełnia nierówność $(17.6153 - 0.0384) \text{ s} \leq t_0 \leq (17.6153 + 0.0384) \text{ s}$

lub $17.5769 \text{ s} \leq t_0 \leq 17.6537 \text{ s}$. Zgodnie z tym co powiedziano w rozdziale 1.1.4 wynik ten zapiszemy: $t = (17.62 \pm 0.04) \text{ s}$ lub $t = 17.62(4) \text{ s}$.

Obliczymy teraz poziom ufności. Korzystając z tablicy II w Dodatku H otrzymujemy

$$\begin{aligned}\alpha &= P_u(-1.6 \leq u \leq 1.6) = \Phi_u(1.6) - \Phi_u(-1.6) \\ &= 0.94520 - 0.05480 = 0.8904,\end{aligned}$$

lub korzystając z tablicy III w Dodatku H otrzymamy

$$\alpha = P_u(-1.6 \leq u \leq 1.6) = \Phi_u(-1.6, 1.6) = 0.8904,$$

czyli że $t = (17.62 \pm 0.04) \text{ s}$ z prawdopodobieństwem 0.89.

Z dotychczasowych rozważań wynika, że publikując wyniki wystarczy podać $\bar{x}$ i $S_{\bar{x}}$ stosując zasady przedstawiania wyników pomiarów omówione w rozdziale 1.1.4, jednak obowiązkiem autora jest wyraźne zaznaczenie, że podano $S_{\bar{x}}$. Na przykład wynik pomiaru czasu wypływu możemy zapisać: $t = 17.615(24) \text{ s}$ z uwagą: w nawiasie podano odchylenie standardowe średniej. Takie postępowanie jest uzasadnione tym, że ktoś, kto chce użyć ten wynik, może znaleźć przedział i poziom ufności w miarę swoich potrzeb, zaś $S_{\bar{x}}$ jako granica przedziału ufności odpowiada poziomowi ufności 0.68.

2) Przykład zastosowania rozkładu normalnego standaryzowanego do innych zagadnień niż ocena błędu pomiaru [13]. Wzrost osobników jednej płci pewnej grupy ludzi ma rozkład normalny. Chcemy znaleźć prawdopodobieństwo $P_u(a \leq u \leq b)$ tego, że przypadkowo wybrany osobnik z tej grupy będzie miał wzrost między 190 i 200 cm, gdy średni wzrost całej grupy $\bar{x} = 172 \text{ cm}$, a odchylenie standardowe $S_x = 6 \text{ cm}$. W tym celu korzystamy ze wzoru (5.2.14), tj. $P_u(a \leq u \leq b) = \Phi_u(b) - \Phi_u(a)$. Granice całkowania a i b , które są wartościami zmiennej standaryzowanej u dla $x = 190$ i 200 cm , obliczamy je ze wzoru (5.2.1) zastępując w nim x_0 i σ ich najlepszymi przybliżeniami, tj. $\bar{x}$ i S_x (patrz rozdział 5.1.3). Otrzymamy:

$$a = \frac{190 - 172}{6} = 3.00, \quad b = \frac{200 - 172}{6} = 4.67.$$

Z tabeli dystrybucyj rozkładu normalnego standaryzowanego, zamieszczonej w Dodatku H, znajdujemy $\Phi_u(4.67) = 0.99997$ i $\Phi_u(3.00) = 0.99865$, a więc $P_u(3.00 \leq u \leq 4.67) = 0.00132$. Oznacza to, że średnio wśród 10000 osobników z tej grupy 13 będzie miało wzrost w granicach 190–200 cm. W ten sam sposób można obliczyć prawdopodobieństwo tego, że przypadkowo wybrany osobnik z tej grupy będzie miał wzrost między 180 i 190 cm. Wynosi ono $P_u(a \leq u \leq b) = 0.99865 - 0.9071 = 0.0916$, co oznacza, że średnio wśród 10000 osobników z tej grupy 916 będzie miało wzrost w granicach 180 – 190 cm. Sprawdzenie poprawności obliczeń pozostawiamy czytelnikom. Należy zwrócić uwagę na praktyczne znaczenie takich badań dla przemysłu, np. odzieżowego.

5.2.3. Ocena poziomu i przedziału ufności w przypadku pomiarów pośrednich

Rozważania przedstawione w rozdziale 5.2.2 w pełni stosują się do oceny poziomu i przedziału ufności w przypadku pomiarów bezpośrednich, gdy wielkości mierzone mają rozkład normalny i liczba pomiarów jest duża. Dystrybucję rozkładu normalnego standaryzowanego można również, w pewnych przypadkach, stosować do oceny poziomu i przedziału ufności wielkości złożonej. Zagadnienie to jest jednak w przypadku pomiarów pośrednich bardziej skomplikowane, ponieważ rozkład wielkości złożonej na ogół nie jest rozkładem normalnym w całym zakresie zmian wielkości mierzonych bezpośrednio, nawet wtedy, gdy te ostatnie mają rozkład normalny. Dlatego będziemy je omawiali wraz z uwagami dotyczącymi wyznaczania wielkości złożonej związanymi ze sposobami przeprowadzenia samego pomiaru i opracowania wyników. Zajmiemy się więc, w pierwszej kolejności, wyznaczaniem wielkości złożonej, gdy wielkości mierzone bezpośrednio są *niezależne*. Następnie przejdziemy do przypadku, gdy wielkości mierzone bezpośrednio są *zależne*.

Zanim przejdziemy do omówienia przypadków, gdy oceny poziomu i przedziału ufności wielkości złożonej możemy dokonać przy użyciu dystrybucji rozkładu normalnego standaryzowanego, przypomnimy, kiedy wielkości mierzone bezpośrednio są niezależne, a kiedy te same wielkości należy uważać za zależne. Zrobimy to na przykładzie wyznaczania przyspieszenia ziemskiego g przy użyciu wahadła matematycznego. Okres drgań wahadła matematycznego $T = 2\pi\sqrt{l/g}$ (l — długość wahadła), czyli $g = 4\pi^2 l/T^2$. Jest to przykład, gdy wielkość złożona jest zależna od dwóch wielkości mierzonych bezpośrednio, co możemy zapisać $g = g(l, T)$. Pomiar przyspieszenia ziemskiego g przy użyciu wahadła matematycznego możemy wykonać dwojako.

- i) Ustalamy długość wahadła l . Wykonujemy N pomiarów długości l i M pomiarów okresu T . W tak przeprowadzonych pomiarach wynik pomiaru długości wahadła l nie ma wpływu na wynik pomiaru okresu drgań wahadła T . Oznacza to, że błąd pomiaru długości l nie wpływa na błąd pomiaru okresu T . Takie pomiary są *niezależne* od siebie.
- ii) Dla określonej długości wahadła l wykonujemy jej pomiar (raz lub więcej) i pomiar okresu drgań T (raz lub więcej). Następnie zmieniamy długość wahadła l i ponownie mierzymy jego długość l i okres T . Powtarzamy to N razy. W wyniku takiego postępowania otrzymujemy N par (l_i, T_i) ($i = 1, 2, \dots, N$). Znaleźliśmy więc zależność $T(l)$, na podstawie której możemy wyznaczyć przyspieszenie ziemskie g . W tym przypadku nie możemy mówić, że pomiary wielkości l i T

są niezależne. W tak przeprowadzonym doświadczeniu są one bowiem **zależne**.

Jak widać z tego przykładu zagadnienie czy pomiary wielkości mierzonych bezpośrednio są zależne, czy też nie, w pewnych przypadkach może zależeć od sposobu przeprowadzenia pomiaru, musi więc być uwzględnione przy opracowaniu wyników oraz ocenie poziomu i przedziału ufności wielkości złożonej.

Przejdziemy teraz do omówienia oceny **poziomu i przedziału ufności wielkości złożonej**. Dla przejrzystości przyjmijmy, że wielkość złożona $z = z(x, y)$, tzn. jest funkcją dwóch wielkości mierzonych bezpośrednio. W pierwszej kolejności omówimy ocenę poziomu i przedziału ufności, gdy wielkości x i y są niezależne, rozpatrzymy tu trzy przypadki. Następnie zajmiemy się oceną poziomu i przedziału ufności, gdy wielkości x i y są zależne.

1. *Pomiary niezależne*

Przyjmijmy, że wielkość fizyczną x zmierzylśmy N razy i otrzymaliśmy zbiór wartości $x_1, x_2, \dots, x_N$, a w wyniku M -krotnego pomiaru wielkości y otrzymaliśmy zbiór wartości $y_1, y_2, \dots, y_M$.

A) *Każda z wielkości x oraz y mierzonych bezpośrednio ma rozkład normalny i można stosować rozwinięcie (5.1.31).*

Wtedy, zgodnie z tym co powiedziano w rozdziale 3.4, średnia arytmetyczna wielkości złożonej $\bar{z} = z(\bar{x}, \bar{y})$, zaś S_z^2 , dane wzorem (3.4.8), można zapisać w postaci:

$$S_z^2 = \left(\frac{\partial z}{\partial x} \right)_{\bar{x}, \bar{y}}^2 S_x^2 + \left(\frac{\partial z}{\partial y} \right)_{\bar{x}, \bar{y}}^2 S_y^2. \quad (5.2.24)$$

Analogicznie $S_{\bar{z}}^2$ dane wzorem (3.4.17) możemy wyrazić jako:

$$S_{\bar{z}}^2 = \left(\frac{\partial z}{\partial x} \right)_{\bar{x}, \bar{y}}^2 S_{\bar{x}}^2 + \left(\frac{\partial z}{\partial y} \right)_{\bar{x}, \bar{y}}^2 S_{\bar{y}}^2. \quad (5.2.25)$$

Jak wynika z rozdziału 5.1.6, w rozpatrywanym przypadku w *otoczeniu punktu* (x_0, y_0) , rozkład wielkości złożonej $z = z(x, y)$ jest normalny. Wobec tego zmienna standaryzowana u dana wzorem (5.2.10) przyjmie postać:

$$u = \frac{z(x_i, y_k) - z(x_0, y_0)}{\sigma_z}, \quad (5.2.26)$$

a zmienna standaryzowana dla wartości średniej, dana wzorem (5.2.11), będzie:

$$u = \frac{z(\bar{x}, \bar{y}) - z(x_0, y_0)}{\sigma_{\bar{z}}}. \quad (5.2.27)$$

Mając dane $z(\bar{x}, \bar{y})$ oraz $\sigma_{\bar{z}}$ możemy przystąpić do oceny poziomu i przedziału ufności. Ocenę tę przeprowadzamy dalej tak samo jak w przypadku wielkości mierzonej bezpośrednio (por. rozdział 5.2.2). Oczywiście w obliczeniach wartości σ_z i $\sigma_{\bar{z}}$ musimy zastąpić przez ich najlepsze przybliżenia, tj. S_z oraz $S_{\bar{z}}$.

B) *Nie możemy stosować rozwinięcia (5.1.31), ale wielkość złożona z ma rozkład normalny.*

W tym przypadku dla każdej pary zmierzonych wartości (x_i, y_k) ($i = 1, 2, \dots, N$, $k = 1, 2, \dots, M$) wyznaczamy wartości $z_{ik} = z(x_i, y_k)$. Wartość średniej arytmetycznej $\bar{z}$ i odchylenia standardowego serii S_z obliczamy ze wzorów (5.1.39) i (5.1.40), natomiast na mocy (5.1.41) odchylenie standardowe średniej serii $S_{\bar{z}}$ wyraża się następująco:

$$S_{\bar{z}} = \sqrt{\frac{1}{NM(NM-1)} \sum_{i=1}^N \sum_{k=1}^M (z_{ik} - \bar{z})^2}. \quad (5.2.28)$$

Ponieważ wielkość złożona $z(x, y)$ *ma rozkład normalny*, to zmienna standaryzowana u dla wielkości z będzie dana wzorem (5.2.26), a zmienna standaryzowana dla średniej arytmetycznej przyjmie postać:

$$u = \frac{\bar{z} - z_0}{\sigma_{\bar{z}}}, \quad (5.2.29)$$

gdzie $\bar{z}$ dane jest wzorem (5.1.39), a z_0 jest wartością rzeczywistą wielkości złożonej z . W obliczeniach $\sigma_{\bar{z}}$ zastępujemy przez najlepsze przybliżenie, tj. $S_{\bar{z}}$ (wzór (5.2.28)). Mając zmienną standaryzowaną u (wzór (5.2.29)), oceny poziomu i przedziału ufności dokonujemy tak, jak dla wielkości mierzonej bezpośrednio. Należy zaznaczyć, że średnie arytmetyczne wielkości złożonej, dane wzorami (3.3.2) i (3.3.5), są liczone w różny sposób. W konsekwencji odchylenia standardowe serii S_z i średniej serii $S_{\bar{z}}$ obliczone ze wzorów (3.4.16) i (5.1.41) nie są sobie równe. Podkreśmy, że jeżeli nie można stosować rozwinięcia (5.1.31), to musimy stwierdzić, czy wielkość złożona z ma rozkład normalny. O tym, czy dana wielkość fizyczna ma rozkład normalny, można się przekonać stosując odpowiednie testy, nie omówione w tym opracowaniu (patrz np. [4, 13]).

C) *Nie możemy stosować rozwinięcia (5.1.31) i rozkład wielkości złożonej jest inny niż rozkład normalny.*

W takim przypadku dla wszystkich par wyników pomiarów (x_i, y_k) ($i = 1, 2, \dots, N$, $k = 1, 2, \dots, M$) wyznaczamy wartości $z_{ik} = z(x_i, y_k)$. Następnie ze wzoru (5.1.39) obliczamy *średnią arytmetyczną*. Dystrybuantę rozkładu normalnego standaryzowanego możemy w tym przypadku

stosować do oceny poziomu i przedziału ufności średniej arytmetycznej serii pomiarów pośrednich, jeśli są spełnione warunki, przy których można stosować centralne twierdzenie graniczne podane w rozdziale 5.3. Ocenę przeprowadzamy tak jak dla wielkości mierzonej bezpośrednio.

2. Pomiary zależne

Przyjmijmy, że dla N różnych zmierzonych wartości x_i ($i = 1, 2, \dots, N$) wykonaliśmy pomiary odpowiadających im wartości y_i . W rezultacie otrzymaliśmy N par wyników (x_i, y_i) .

A) *Rozkład każdej z wielkości mierzonych bezpośrednio jest normalny.*

Jeżeli wyrażenie $z = z(x, y)$ możemy przekształcić tak, że znajdziemy zależność funkcyjną $y = f(x, z)$ (lub $x = g(y, z)$), to do otrzymanych wyników pomiarów wielkości x oraz y dopasowujemy tę zależność metodą najmniejszych kwadratów przedstawioną w rozdziale 3.5 (patrz także rozdział 5.4). Z wyznaczonych w ten sposób współczynników i ich odchyień standardowych znajdujemy najlepsze przybliżenie $\hat{z}$ (wartość $\bar{z}$) oraz S_z^2 . Jeśli pozwalają na to własności dopasowywanej funkcji, to do oceny poziomu i przedziału ufności wyznaczonej wielkości z można użyć dystrybucy anty rozkładu normalnego standaryzowanego. W tym przypadku zmienna standaryzowana u jest dana wzorem:

$$u = \frac{\hat{z} - z_0}{\sigma_z}. \quad (5.2.30)$$

Oceny poziomu oraz przedziału ufności dokonujemy jak dla wielkości mierzonej bezpośrednio, zastępując σ_z przez najlepsze przybliżenie, tj. S_z .

B) *Nie stosujemy metody najmniejszych kwadratów, rozkład wielkości złożonej jest normalny.*

Jeżeli z jakichkolwiek powodów nie stosujemy metody najmniejszych kwadratów, wówczas dla każdej pary (x_i, y_i) ($i = 1, 2, \dots, N$) wyznaczamy $z_i = z(x_i, y_i)$. Następnie znajdujemy średnią arytmetyczną

$$\bar{z} = \frac{1}{N} \sum_{i=1}^N z_i \quad (5.2.31)$$

i odchylenie standardowe średniej serii $S_{\bar{z}}$, które, jak łatwo pokazać, wynosi

$$S_{\bar{z}} = \sqrt{\frac{1}{N(N-1)} \sum_{i=1}^N (z_i - \bar{z})^2}. \quad (5.2.32)$$

Jeżeli wielkość złożona z ma rozkład normalny, wtedy dalej postępujemy tak, jak podano w punkcie 1B.

C) *Nie stosujemy metody najmniejszych kwadratów i rozkład wielkości złożonej nie jest rozkładem normalnym.*

W tym przypadku, podobnie jak w poprzednim, dla każdej pary (x_i, y_i) ($i = 1, 2, \dots, N$) wyznaczamy $z_i = z(x_i, y_i)$ i znajdujemy średnią arytmetyczną (5.2.31). Jeżeli są spełnione warunki, przy których można stosować centralne twierdzenie graniczne, podane w rozdziale 5.3, to do oceny poziomu i przedziału ufności możemy zastosować dystrybuantę rozkładu normalnego standaryzowanego. Ocenę przeprowadzamy tak, jak dla wielkości mierzonej bezpośrednio.

Przykłady

1) Wróćmy teraz do przedstawionego w rozdziale 3.4 przykładu z wyznaczaniem przyspieszenia ziemskiego przy użyciu wahadła matematycznego. W przykładzie tym długość wahadła l była stała. Wykonywaliśmy jedynie jej pomiary i pomiary okresu drgań wahadła T . W tym przypadku długość wahadła l i jego okres T występują jako wielkości niezależne. Ocenę poziomu i przedziału ufności należy więc przeprowadzić tak jak podano w punkcie 1A. Otrzymaliśmy tam $\bar{g} = 980.22503 \text{ cm/s}^2$ i $S_{\bar{g}} = 4.6690808 \text{ cm/s}^2$. Wiedząc teraz jak $S_{\bar{g}}$ jest związane z błędem pomiaru i stosując zasady podane w rozdziale 1.1.4 wynik ten możemy zapisać: $g = 980(5) \text{ cm/s}^2$ lub $g = 980.2(4.7) \text{ cm/s}^2$. W nawiasie podano odchylenie standardowe średniej. Znajdziemy jeszcze poziom ufności dla przedziału ufności np. $\bar{g} \pm 1.8S_{\bar{g}}$. Z tablicy III w Dodatku H znajdujemy $\alpha = 0.92814$, czyli $971.82 \text{ cm/s}^2 \leq g \leq 988.63 \text{ cm/s}^2$ z prawdopodobieństwem $\alpha = 0.92814$.

2) W celu wyznaczenia przyspieszenia ziemskiego, za pomocą wahadła matematycznego, wykonano pomiary zależności okresu drgań T od jego długości l . Otrzymano następujące wyniki

$l \text{ [cm]}$	50	99	151	202	248	301
$T \text{ [s]}$	1.41	2.02	2.44	2.84	3.18	3.46

Okres drgań wahadła matematycznego $T = 2\pi\sqrt{l/g}$. Wyrażenie to można przekształcić do postaci $4\pi^2l = gT^2$. Oznaczmy $4\pi^2l = y$ oraz $T^2 = x$, wtedy otrzymamy $y = gx$, czyli równanie prostej przechodzącej przez początek układu współrzędnych. Stosując regresję liniową przedstawioną w rozdziale 3.5 otrzymujemy $\hat{g} = 984.3661 \text{ cm/s}^2$ i $S_{\hat{g}} = 6.3895 \text{ cm/s}^2$, co zapisujemy: $g = 984.4(6.4) \text{ cm/s}^2$. W nawiasie podano odchylenie standardowe dopasowania metodą najmniejszych kwadratów.

5.3. Centralne twierdzenie graniczne

Wykonując wielokrotnie pomiary jakiejś wielkości fizycznej x , otrzymujemy ciąg jej wartości $x_1, x_2, \dots, x_n, \dots$. W praktyce często się zdarza, że o ich

rozkładzie nie mamy żadnych informacji. Nawet jeśli wykonamy dużą liczbę pomiarów np. $n = 100$, to często na ich podstawie nie możemy znaleźć funkcji rozkładu. Nie znając funkcji rozkładu nie jesteśmy w stanie dokonać oceny poziomu i przedziału ufności średniej rozkładu $\bar{x}_R$. Okazuje się jednak, że w wielu przypadkach można dokonać tej oceny przy użyciu rozkładu normalnego standaryzowanego $N(0, 1)$. O tym kiedy takie postępowanie jest możliwe mówią tzw. **centralne twierdzenia rachunku prawdopodobieństwa** nazywane również **centralnymi twierdzeniami granicznymi**. Twierdzeń tych nie będziemy tu przytaczali, ponieważ wykraczałoby to poza ramy tego opracowania. Zainteresowany czytelnik znajdzie je podane w przystępny sposób wraz z dowodami np. w [14, 15]. W opracowaniu tym ograniczymy się jedynie do podania — w uproszczonej postaci — twierdzenia najważniejszego dla oceny poziomu i przedziału ufności, a mianowicie centralnego twierdzenia granicznego Lindeberga-Levy'ego.

Zanim je podamy wprowadzimy wielkości, którymi będziemy się posługiwać. Załóżmy, że w ustalonych warunkach wykonujemy serię pomiarów jakiejś wielkości fizycznej x . Oznaczmy przez $B(x)$ *nieznany* rozkład prawdopodobieństwa mierzonej wielkości fizycznej x . Rozkład ten jest oczywiście spowodowany występowaniem błędów przypadkowych. Niech

$$y_n = \frac{\sqrt{n}(\bar{x}_n - \bar{x}_R)}{\sigma} \quad (5.3.1)$$

będzie standaryzowanym odchyleniem $\bar{x}_n = (x_1 + x_2 + \dots + x_n)/n$ średniej arytmetycznej z n kolejnych pomiarów od $\bar{x}_R$ średniej rozkładu $B(x)$, gdzie σ jest odchyleniem standardowym rozkładu $B(x)$. Wariancja σ^2 rozkładu $B(x)$, zgodnie z (4.2.12), wyraża się wzorem:

$$\sigma^2 = \int_{-\infty}^{+\infty} (x - \bar{x}_R)^2 B(x) dx. \quad (5.3.2)$$

Oznaczmy funkcję rozkładu wielkości y_n przez $A_n(y_n)$. Funkcja ta jest jednak nieznana, ponieważ nieznana jest funkcja rozkładu $B(x)$. Dystrybuantę $\Psi_n(a)$ rozkładu $A_n(y_n)$, zgodnie ze wzorem (4.2.7), można zapisać w postaci:

$$\Psi_n(a) = \int_{-\infty}^a A_n(y_n) dy_n. \quad (5.3.3)$$

Przejdźmy teraz do podania (bez dowodu) centralnego twierdzenia granicznego Lindeberga-Levy'ego w uproszczonej postaci, przydatnej do oceny poziomu i przedziału ufności średniej rozkładu $\bar{x}_R$.

Niech $x_1, x_2, \dots, x_n, \dots$ stanowią ciąg zmierzonych wartości wielkości fizycznej x , podlegających rozkładowi $B(x)$.

Twierdzenie Lindeberga-Levy'ego stanowi, że ciąg dystrybuant $\Psi_n(a)$ zmierza do dystrybuanty rozkładu normalnego standaryzowanego $\Phi_u(a)$, tzn.

$$\lim_{n \rightarrow \infty} \Psi_n(a) = \lim_{n \rightarrow \infty} \int_{-\infty}^a A_n(y_n) dy_n = \frac{1}{\sqrt{2\pi}} \int_{-\infty}^a e^{-\frac{1}{2}u^2} du = \Phi_u(a) \quad (5.3.4)$$

pod warunkiem, że rozkład $B(x)$ ma skończoną średnią rozkładu $\bar{x}_R$ i skończoną różną od zera wariancję σ^2 .

Inaczej mówiąc: dystrybuanta rozkładu wielkości y_n zmierza do dystrybuanty rozkładu normalnego standaryzowanego, gdy liczba zmierzonych wartości wielkości fizycznej x , użytych do obliczenia średniej arytmetycznej, zmierza do nieskończoności. To oznacza również, że funkcja rozkładu wielkości y_n w granicy, gdy $n \rightarrow \infty$ zmierza do rozkładu normalnego standaryzowanego $N(0, 1)$, czyli

$$\lim_{n \rightarrow \infty} A_n(y_n) = \frac{1}{\sqrt{2\pi}} e^{-\frac{1}{2}y_n^2}.$$

Stąd wynika, że dla dostatecznie dużych n funkcję rozkładu średniej arytmetycznej $\bar{x}_n$ możemy przybliżyć rozkładem normalnym o średniej $\bar{x}_R$ i wariancji $\sigma_{\bar{x}}^2 = \sigma^2/n$ (por. [15] s. 146):

$$\varphi(\bar{x}_n) = \frac{1}{\sigma_{\bar{x}} \sqrt{2\pi}} \exp \left[-\frac{(\bar{x}_n - \bar{x}_R)^2}{2\sigma_{\bar{x}}^2} \right]. \quad (5.3.5)$$

Tak więc, jeśli spełnione są wyżej przedstawione warunki, to, pomimo że o rozkładzie wielkości fizycznej x nie możemy nic powiedzieć, **rozkład średniej arytmetycznej** $\bar{x}$ zmierzonych jej wartości możemy przybliżyć rozkładem normalnym. Stąd wynika, że oceny poziomu i przedziału ufności średniej rozkładu $\bar{x}_R$ w przypadku bardzo dużej liczby pomiarów ($n \rightarrow \infty$) możemy dokonać przy użyciu dystrybuanty $\Phi_u(a)$, tak jak to przedstawiono w rozdziale 5.2.2. Musimy jednak pamiętać, że możemy tego dokonać tylko wtedy, gdy średnia $\bar{x}_R$ rozkładu $B(x)$ jest skończona, jak również skończona jest wariancja $\sigma^2 \neq 0$ rozkładu $B(x)$.

Wariancja $\sigma_{\bar{x}}^2$ we wzorze (5.3.5) jest równa

$$\sigma_{\bar{x}}^2 = \frac{\sigma^2}{n}, \quad (5.3.6)$$

zaś σ^2 jest dane wzorem (5.3.2). W praktyce σ^2 we wzorze (5.3.6) zastępujemy przez $S_{\bar{x}}^2$, ponieważ nie znamy rozkładu $B(x)$.

Centralne twierdzenie graniczne Lindeberga-Levy'ego stosuje się zarówno do pomiarów bezpośrednich jak i pomiarów pośrednich.

W większości wykonywanych pomiarów (obojętnie, czy to będą pomiary bezpośrednie, czy pośrednie) warunki stosowalności centralnego twierdzenia granicznego są spełnione. Pozwala to na stosowanie do oceny poziomu i przedziału ufności **średniej artmetycznej** dużej serii pomiarowej rozkładu normalnego. Powstaje więc pytanie, kiedy liczba pomiarów jest na tyle duża, że możemy korzystać z centralnego twierdzenia granicznego i stosować rozkład normalny? Odpowiedź nie jest jednoznaczna. Oczywiście, im więcej pomiarów, tym lepiej spełnione są warunki, o których była mowa powyżej. W praktyce, w przypadku zmiennej losowej ciągłej, gdy liczba pomiarów użytych do obliczenia średniej arytmetycznej $n > 30$, do oszacowania przedziału i poziomu ufności możemy stosować rozkład normalny [13].

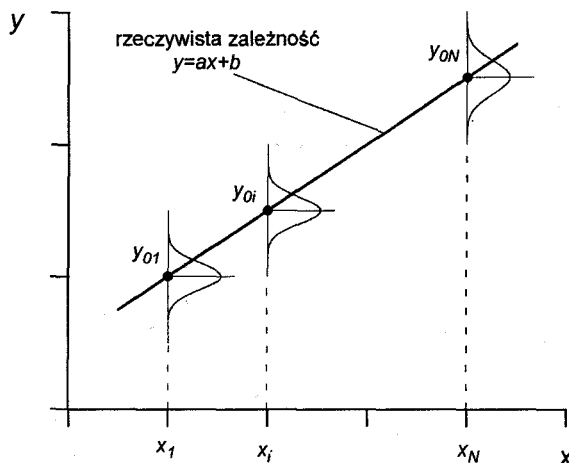
5.4. Metoda najmniejszych kwadratów a rozkład Gaussa

W rozdziale 3.5 przedstawiliśmy zastosowanie metody najmniejszych kwadratów do dopasowywania krzywych do punktów doświadczalnych. Ponieważ metoda najmniejszych kwadratów należy do najczęściej stosowanych metod statystycznych, w tym rozdziale podamy statystyczną interpretację wyznaczonych tą metodą parametrów dopasowywanej funkcji i jej uogólnienie na przypadek pomiarów o niejednakowej precyzji. Do tego celu wykorzystamy metodę największej wiarygodności stosowaną w rozdziale 5.1.3.

Przyjmijmy, że wykonujemy serię N pomiarów wartości dwóch wielkości fizycznych x i y oraz dla każdej wartości x_i otrzymujemy odpowiadającą jej wartość y_i ($i = 1, 2, \dots, N$). W układzie współrzędnych XY otrzymujemy więc punkt (x_i, y_i) . O wielkościach x i y *wiemy*, że są związane zależnością funkcyjną $y = f(x)$. Funkcja $f(x)$, jak pokazano w rozdziale 3.5, w ogólności jest zależna od Q współczynników c_j ($j = 1, 2, \dots, Q$), czyli $y = f(x, c_1, c_2, \dots, c_Q)$. Przyjmijmy, podobnie jak w rozdziale 3.5, że błędy pomiaru wartości x_i ($i = 1, 2, \dots, N$) są tak *małe*, iż *możemy je pominąć*. Wówczas wartość rzeczywista $y_{0,i}$ wielkości fizycznej y będzie dana wzorem:

$$y_{0,i} = f(x_i, c_1, c_2, \dots, c_Q). \quad (5.4.1)$$

Ponieważ na początku przyjęliśmy, iż wyniki pomiarów wielkości y podlegają rozkładowi normalnemu, co oznacza, że jeżeli dla ustalonej wartości x_i wykonamy bardzo dużo pomiarów wielkości y , to otrzymamy rozkład wyników pomiarów będzie rozkładem normalnym o wartości rzeczywistej $y_{0,i}$ i wariancji σ_y^2 . Na rysunku 5.7 przedstawiono to dla przypadku, gdy zmienna y jest liniową funkcją zmiennej x , tzn. $y = ax + b$. Gdyby liczba



Rys. 5.7: Bardzo duża liczba pomiarów wielkości y przy ustalonych x_i . Dla każdej wartości x_i wyniki pomiarów wielkości y dają pewien rozkład wokół wartości rzeczywistej $y_{0,i} = ax_i + b$, [7]

pomiarów wielkości y dla ustalonej wartości x_i była bardzo duża, wówczas byłoby można wyznaczyć z bardzo dobrym przybliżeniem wartości $y_{0,i}$. Gdyby również liczba ustalonych wartości wielkości x była bardzo duża, to z bardzo dobrym przybliżeniem byłoby można wyznaczyć zależność funkcyjną $y_{0,i} = f(x_i)$. W praktyce najczęściej dla każdej wartości x_i wykonujemy jeden pomiar wielkości y . Wobec tego otrzymana z pomiaru wartość y_i będzie odchyłać się od wartości rzeczywistej $y_{0,i}$ z prawdopodobieństwem określonym jej rozkładem. Ponieważ, jak założyliśmy powyżej, wyniki pomiarów y_i podlegają rozkładowi normalnemu o wartości rzeczywistej $y_{0,i}$ i wariancji σ_y^2 , tzn. zakładamy, że wszystkie pomiary wykonujemy z jednakową precyzją, wtedy prawdopodobieństwo otrzymania w wyniku pomiaru wielkości y wartości zawartej między y_i a $y_i + dy_i$ jest:

$$dP(y_i) = \frac{1}{\sigma_y \sqrt{2\pi}} \exp \left[-\frac{(y_i - y_{0,i})^2}{2\sigma_y^2} \right] dy_i.$$

Podstawiając $y_{0,i}$ ze wzoru (5.4.1) otrzymujemy:

$$dP(y_i) = \frac{1}{\sigma_y \sqrt{2\pi}} \exp \left[-\frac{[y_i - f(x_i, c_1, c_2, \dots, c_Q)]^2}{2\sigma_y^2} \right] dy_i. \quad (5.4.2)$$

Ponieważ wynik pomiaru wielkości y i jego błąd dla danej wartości wielkości x nie są w żaden sposób związane z wynikiem pomiaru wielkości y i jego błędem dla innej wartości x , to zgodnie z tym, co powiedziano w rozdziale

5.1.3, prawdopodobieństwo tego, że w wyniku N takich pomiarów otrzymamy zbiór N wartości (x_i, y_i) ($i = 1, 2, \dots, N$), wyraża się następująco:

$$\begin{aligned} dP(y_1, y_2, \dots, y_N) &= \\ &= \left(\frac{1}{\sigma_y \sqrt{2\pi}} \right)^N \exp \left[-\frac{1}{2\sigma_y^2} \sum_{i=1}^N [y_i - f(x_i, c_1, c_2, \dots, c_Q)]^2 \right] (dy)^N. \end{aligned} \quad (5.4.3)$$

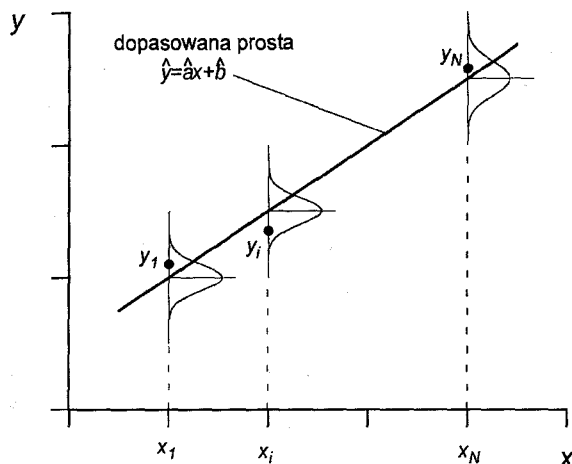
Prawdopodobieństwo to osiąga maksimum, gdy wykładnik potęgi jest minimalny, tzn. gdy

$$\frac{1}{2\sigma_y^2} \sum_{i=1}^N [y_i - f(x_i, c_1, c_2, \dots, c_Q)]^2 = G(c_1, c_2, \dots, c_Q) = \min. \quad (5.4.4)$$

Musimy więc tak dobrać współczynniki c_j , aby suma kwadratów odchyłeń $y_i - f(x_i, c_1, c_2, \dots, c_Q)$ była najmniejsza. Wyrażenie (5.4.4) minimalizujemy ze względu na **nieznane współczynniki** c_j ($j = 1, 2, \dots, Q$) wykorzystując warunki na minimum funkcji wielu zmiennych. Różniczkując (5.4.4) i przyrównując pochodne do zera otrzymujemy układ równań algebraicznych (3.5.3), tj.

$$\begin{aligned} \left(\frac{\partial G}{\partial c_1} \right)_{c_1, c_2, \dots, c_Q} &= 0, \\ &\vdots \\ \left(\frac{\partial G}{\partial c_j} \right)_{c_1, c_2, \dots, c_Q} &= 0, \\ &\vdots \\ \left(\frac{\partial G}{\partial c_Q} \right)_{c_1, c_2, \dots, c_Q} &= 0. \end{aligned} \quad (5.4.5)$$

Rozwiązując ten układ równań algebraicznych znajdujemy wartości współczynników (które będziemy oznaczać $\hat{c}_j$ ($j = 1, 2, \dots, Q$)), dla których prawdopodobieństwo (5.4.2) ma wartość największą. Wyznaczone w ten sposób wartości $\hat{c}_j$ współczynników są więc **wartościami najbardziej wiarygodnymi**, czyli stanowią *najlepsze przybliżenie* wartości rzeczywistych współczynników c_j i dlatego oznaczyliśmy je przez $\hat{c}_j$. Funkcja $\hat{y} = f(x, \hat{c}_1, \hat{c}_2, \dots, \hat{c}_Q)$ jest więc **najbardziej prawdopodobną funkcją** wiążącą ze sobą zmierzone wartości x_i, y_i ($i = 1, 2, \dots, N$), która spełnia warunek (5.4.4). Jest więc ona *najlepszym przybliżeniem* rzeczywistej zależności $y = f(x, c_1, c_2, \dots, c_Q)$, a wartości $\hat{y}_i = f(x_i, \hat{c}_1, \hat{c}_2, \dots, \hat{c}_Q)$ są **najlepszymi przybliżeniami wartości rzeczywistych** $y_{0,i}$. Na rysunku 5.8 przedstawiono jako ilustrację dopasowania metodą najmniejszych kwadratów dopasowanie funkcji liniowej $y = ax + b$ w przypadku, gdy dla każdej



Rys. 5.8: Przykład dopasowania, gdy dla każdej wartości x_i wykonujemy jeden pomiar wielkości y

wartości x_i wykonano jeden pomiar wielkości y i otrzymano jej wartość y_i .

Reasumując należy stwierdzić, że gdy wyniki pomiarów wielkości fizycznej y mają rozkład Gaussa i dopasowywana funkcja $y = f(x, c_1, \dots, c_Q)$ jest ogólną postacią rzeczywistej zależności funkcyjnej między mierzonymi wielkościami, to metoda najmniejszych kwadratów jest równoważna metodzie największej wiarygodności. Wyznaczone zaś współczynniki $\hat{c}_1, \hat{c}_2, \dots, \hat{c}_Q$ są najlepszymi przybliżeniami rzeczywistych współczynników $c_1, c_2, \dots, c_Q$. Wyznaczone współczynniki $\hat{c}_j$ są obarczone błędami, które są spowodowane niepewnościami zmierzonych wartości y_i . Sposób obliczenia odchyleń standardowych $S_{\hat{c}_j}$ podano w rozdziale 3.5.

5.4.1. Metoda najmniejszych kwadratów w przypadku pomiarów o niejednakowej precyzji

W przypadku, gdy każdy z pomiarów y_i podlega rozkładowi normalnemu o wartości rzeczywistej $y_{0,i}$ i wariancji $\sigma_{y_i}^2$, to, jak wynika z rozważań przedstawionych w rozdziale 5.1.5, prawdopodobieństwo (5.4.4) wyrazi się wzorem:

$$dP(y_1, y_2, \dots, y_N) = \quad (5.4.6)$$

$$= \frac{1}{\prod_{i=1}^N \sigma_{y_i}} \left(\frac{1}{\sqrt{2\pi}} \right)^N \exp \left[-\frac{1}{2} \sum_{i=1}^N \frac{1}{\sigma_{y_i}^2} [y_i - f(x_i, c_1, c_2, \dots, c_Q)]^2 \right] (dy)^N.$$

Warunek (5.4.4) przyjmie postać:

$$\frac{1}{2} \sum_{i=1}^N w_i [y_i - f(x_i, c_1, c_2, \dots, c_Q)]^2 = \min, \quad (5.4.7)$$

gdzie **waga pomiaru** $w_i = 1/\sigma_{y_i}^2$. W przypadku pomiarów o jednakowej precyzji $w_i = \text{const}$ możemy wynieść przed znak sumy i otrzymamy wzór (5.4.4). Minimalizacji wyrażenia (5.4.7) dokonujemy tak samo, jak w przypadku pomiarów o jednakowej precyzji.

6. Rozkład *t-Studenta* i jego zastosowanie

6.1. Rozkład *t-Studenta*

Oszacowanie przedziału i poziomu ufności podane w rozdziale 5.2 jest słuszne, gdy liczba pomiarów N jest na tyle duża, że można przyjąć, iż $S_{\bar{x}}^2 \approx \sigma_{\bar{x}}^2$, lub gdy znamy wariancję rozkładu $\sigma_{\bar{x}}^2$. W praktyce, z różnych powodów takich jak czasochłonność pomiarów, koszty przeprowadzenia doświadczeń itp., wykonujemy przeważnie kilka do kilkunastu pomiarów. Wówczas $S_{\bar{x}}^2$ może znacznie się różnić od wariancji $\sigma_{\bar{x}}^2$ rozkładu średnich arytmetycznych. W takim przypadku stosowanie dystrybucyj rozkładu $N(0, 1)$ prowadzi do błędnego oszacowania przedziału i poziomu ufności. Gosset, używający pseudonimu *Student* wykazał, że jeżeli wartości zmiennej losowej $x_1, x_2, \dots, x_N$ mają **rozkład normalny**, to dla $N \geq 2$ wielkość

$$t = \frac{\bar{x} - x_0}{S_{\bar{x}}}, \quad (6.1.1)$$

ma funkcję rozkładu prawdopodobieństwa

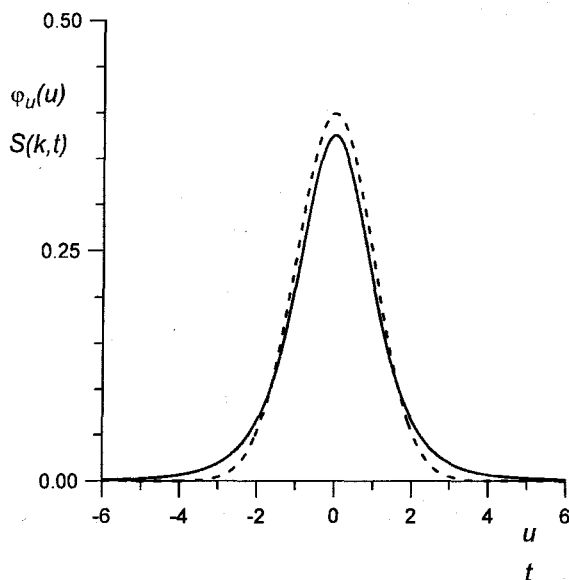
$$S(t, k) = \frac{\Gamma(\frac{k+1}{2})}{\sqrt{k\pi}\Gamma(\frac{k}{2})} \left(1 + \frac{t^2}{k}\right)^{-\frac{k+1}{2}}, \quad (6.1.2)$$

gdzie $k = N - 1$ jest **liczbą stopni swobody**, $S_{\bar{x}}$ dane jest wzorem (3.4.11), a $\Gamma(k)$ jest funkcją Γ -Eulera i wyraża się wzorem:

$$\Gamma(k) = \int_0^{\infty} e^{-s} s^{k-1} ds.$$

Należy podkreślić, że wielkość t jest zdefiniowana dla **średniej arytmetycznej** danej serii pomiarowej. Funkcja rozkładu $S(t, k)$ dana wzorem (6.1.2) nosi w literaturze nazwę rozkładu *t-Studenta*. Wyprowadzenie wzoru (6.1.2) podano np. w [13, 16]. Funkcja rozkładu $S(t, k)$ jest symetryczna względem zmiennej t , a mianowicie $S(-t, k) = S(t, k)$. Liczba stopni swobody k jest parametrem funkcji rozkładu.

Na rysunku 6.1 przedstawiono porównanie krzywej rozkładu normalnego i rozkładu *t-Studenta*. Funkcja rozkładu *t-Studenta* jest bardziej spłaszczona



Rys. 6.1: Porównanie rozkładu normalnego standaryzowanego (linia przerywana) i rozkładu *t-Studenta* dla $k = 4$ (linia ciągła)

w maksimum i wolniej zmierza do zera, gdy $|t| \rightarrow \infty$ niż funkcja rozkładu normalnego standaryzowanego $N(0, 1)$. Wraz ze wzrostem liczby stopni swobody k rozkład $S(t, k)$ szybko zmierza do rozkładu normalnego standaryzowanego $N(0, 1)$. W praktyce, gdy liczba pomiarów $N > 30$, można przyjąć, że rozkład *t-Studenta* i rozkład normalny dają te same oszacowania przedziału i poziomu ufności [13].

6.2. Dystrybuanta rozkładu *t-Studenta* i jej zastosowanie

Zgodnie ze wzorem (4.2.7), dystrybuanta rozkładu *t-Studenta* wyrazi się następująco (patrz także rozdział 5.2.2):

$$F(a) = \int_{-\infty}^a S(t, k) dt. \quad (6.2.1)$$

Prawdopodobieństwo α , że wartość zmiennej t będzie zawarta między ustalonymi wartościami $-t_\alpha$ i t_α ($t_\alpha > 0$), czyli

$$-t_\alpha \leq t \leq t_\alpha, \quad (6.2.2)$$

jest równe (por. rozdział 4.2 i 5.2.2):

$$\alpha = P(-t_\alpha \leq t \leq t_\alpha) = \int_{-t_\alpha}^{t_\alpha} S(t, k) dt = 2 \int_0^{t_\alpha} S(t, k) dt \quad (6.2.3)$$

i zależy od liczby stopni swobody k , czyli $\alpha = \alpha(k)$. Z nierówności (6.2.2) otrzymujemy **przedział ufności**:

$$\bar{x} - t_\alpha S_{\bar{x}} \leq x_0 \leq \bar{x} + t_\alpha S_{\bar{x}}, \quad (6.2.4)$$

przy poziomie ufności α . Ponieważ poziom ufności α , dany wzorem (6.2.3), zależy od liczby stopni swobody k , to również współczynnik t_α , wyznaczony dla ustalonego α , będzie zależny od k . Otrzymaliśmy więc oszacowanie nierówności (1.1.4) w przypadku małej liczby pomiarów, gdy mierzona wielkość fizyczna podlega rozkładowi normalnemu.

Ponieważ zarówno prawdopodobieństwo α , jak i współczynnik t_α są zależne od parametru rozkładu k , zazwyczaj nie podaje się tablic całki (6.2.3), ale tablice współczynnika t_α dla różnych wartości k i α . Tablice te noszą często nazwę **tablic Studenta-Fishera** i zostały zamieszczone w Dodatku H (tablica IV).

W przypadku małej liczby pomiarów, gdy do oszacowania przedziału i poziomu ufności stosujemy rozkład *t-Studenta*, nie możemy podać niezależnej od liczby pomiarów granicy przedziału ufności, którą można by przyjąć za błąd maksymalny (w przypadku rozkładu normalnego było to $3S_{\bar{x}}$). W tym przypadku należy z góry określić poziom ufności α . W naszym opracowaniu, przez analogię do rozkładu normalnego, przyjmijmy $\alpha = 0.997$, ponieważ jest to poziom ufności odpowiadający granicy przedziału ufności $3S_{\bar{x}}$. Wtedy za **błąd maksymalny** przyjmijmy $t_{0.997} S_{\bar{x}}$; współczynnik $t_{0.997}$ dla każdej liczby stopni swobody k musimy jednak odczytać z tablicy IV w Dodatku H.

Przykład

Wróćmy do przykładu z tabeli 4.2 i oszacujmy przedział ufności czasu wpływu przy użyciu pierwszych sześciu pomiarów tam zamieszczonych. Ich średnia arytmetyczna $\bar{t} = 17.540$ s, $S_t = 0.396$ s, a $S_{\bar{t}} = 0.162$ s. Przyjmijmy, że szukamy przedziału ufności dla poziomu ufności $\alpha = 0.9$. Z tablicy IV w Dodatku H znajdujemy $t_\alpha = 2.015$ dla $k = 5$. Wobec tego $t_\alpha S_{\bar{t}} = 0.326$ s ≈ 0.33 s, a więc $(17.54 - 0.33)$ s $\leq t_0 \leq (17.54 + 0.33)$ s, czyli 17.21 s $\leq t_0 \leq 17.87$ s z prawdopodobieństwem 0.9.

Obliczmy teraz to samo przy użyciu rozkładu normalnego standaryzowanego $N(0, 1)$. Z tablicy III w Dodatku H znajdujemy, że poziomowi ufności

$\alpha = 0.9$ odpowiada $a = 1.604$, a więc $aS_{\bar{t}} = 0.266 \text{ s} \approx 0.27 \text{ s}$ i przedział ufności jest $17.27 \text{ s} \leq t_0 \leq 17.81 \text{ s}$.

Z przytoczonego przykładu widać, że oszacowanie błędu pomiaru (przedziału ufności) będzie znacznie bardziej optymistyczne przy użyciu rozkładu normalnego niż rozkładu *t-Studenta*. Należy zaznaczyć, że tak będzie we wszystkich podobnych przypadkach.

6.2.1. Ocena poziomu i przedziału ufności w przypadku pomiarów pośrednich

Rozkład *t-Studenta* i jego dystrybuantę możemy stosować do oceny poziomu i przedziału ufności wielkości złożonej zasadniczo w trzech przypadkach.

1) *Wielkość złożoną można przedstawić w postaci liniowej funkcji niezależnych wielkości mierzonych bezpośrednio i mających rozkłady normalne.*

W tym przypadku ogólna metoda rozwiązywania zagadnienia oceny poziomu i przedziału ufności jest skomplikowana i z tego powodu nie będziemy jej przytaczali. Zainteresowany czytelnik znajdzie ją w [17] oraz w [18] (rozdział III). Ograniczymy się tu do podania jedynie wyrażenia pozwalającego obliczyć przybliżoną wartość zmiennej t podlegającej rozkładowi *t-Studenta* w przypadku, gdy wielkość złożoną z możemy rozwinąć na szereg Taylora i ograniczyć się do wyrazów liniowych. Odpowiada to przypadkowi 1A, omówionemu w rozdziale 5.2.3. Dla przejrzystości rozważań założymy, że $z = z(x, y)$, oraz że wielkość x zmierzylśmy N razy, a wielkość y została zmierzona M razy. Rozwińmy $z = z(x, y)$ na szereg Taylora w otoczeniu punktu $(\bar{x}, \bar{y})$ i założymy, że jest spełniony warunek (5.1.32), tzn. w całym otrzymanym z pomiarów przedziale zmienności wielkości x oraz y możemy ograniczyć się do wyrazów liniowych. Wtedy:

$$z(x, y) = z(\bar{x}, \bar{y}) + \left(\frac{\partial z}{\partial x}\right)_{\bar{x}, \bar{y}}(x - \bar{x}) + \left(\frac{\partial z}{\partial y}\right)_{\bar{x}, \bar{y}}(y - \bar{y}). \quad (6.2.5)$$

Obliczenia, których nie będziemy przytaczali, pokazują, że w tym przypadku wielkość t jest:

$$t = \frac{z(\bar{x}, \bar{y}) - z(x_0, y_0)}{\tilde{S}_{\bar{z}}}, \quad (6.2.6)$$

gdzie przybliżona wartość $\tilde{S}_{\bar{z}}$ spełnia równość:

$$\tilde{S}_{\bar{z}}^2 = \frac{n}{n-2} \left[\left(\frac{\partial z}{\partial x}\right)_{\bar{x}, \bar{y}}^2 \frac{1}{N} \tilde{S}_x^2 + \left(\frac{\partial z}{\partial y}\right)_{\bar{x}, \bar{y}}^2 \frac{1}{M} \tilde{S}_y^2 \right], \quad (6.2.7)$$

oraz

$$\tilde{S}_x^2 = \frac{1}{N} \sum_{i=1}^N (x_i - \bar{x})^2, \quad \tilde{S}_y^2 = \frac{1}{M} \sum_{i=1}^M (y_i - \bar{y})^2, \quad (6.2.8)$$

$n = N + M$, a liczba stopni swobody $k = n - 2$.

W przypadku ogólnym, gdy $z = z(x_1, x_2, \dots, x_r)$ oraz wszystkie wielkości $x_1, x_2, \dots, x_r$ mierzone bezpośrednio mają rozkład normalny i są niezależne oraz każda z nich została zmierzona N_i razy ($i = 1, 2, \dots, r$), wielkość t będzie:

$$t = \frac{z(\bar{x}_1, \bar{x}_2, \dots, \bar{x}_r) - z(x_{0,1}, x_{0,2}, \dots, x_{0,r})}{\tilde{S}_{\bar{z}}}, \quad (6.2.9)$$

zaś:

$$\tilde{S}_{\bar{z}}^2 = \frac{n}{n-r} \sum_{i=1}^r \left(\frac{\partial z}{\partial x_i} \right)_{\bar{x}_1, \bar{x}_2, \dots, \bar{x}_r}^2 \frac{1}{N_i} \tilde{S}_{x_i}^2, \quad (6.2.10)$$

$$\tilde{S}_{x_i}^2 = \frac{1}{N_i} \sum_{k=1}^{N_i} (x_{i,k} - \bar{x}_i)^2, \quad (6.2.11)$$

$$n = \sum_{i=1}^r N_i, \quad (6.2.12)$$

a r jest liczbą wielkości mierzonych bezpośrednio. W tym przypadku liczba stopni swobody $k = n - r$. Wielkość t określona wzorem (6.2.9) ma w przybliżeniu funkcję rozkładu $S(t, k) = S(t, n - r)$ daną wzorem (6.1.2). Mając wyznaczone $z(\bar{x}_1, \bar{x}_2, \dots, \bar{x}_r)$ oraz $\tilde{S}_{\bar{z}}$, przedział i poziom ufności znajdujemy postępując tak samo jak dla wielkości mierzonej bezpośrednio, tzn. znajdujemy prawdopodobieństwo α tego, że spełniona jest nierówność (6.2.4), która w tym przypadku przyjmie postać:

$$z(\bar{x}_1, \bar{x}_2, \dots, \bar{x}_r) - t_{\alpha} \tilde{S}_{\bar{z}} \leq z(x_{0,1}, x_{0,2}, \dots, x_{0,r}) \leq z(\bar{x}_1, \bar{x}_2, \dots, \bar{x}_r) + t_{\alpha} \tilde{S}_{\bar{z}}. \quad (6.2.13)$$

Współczynnik t_{α} dla danej wartości α i liczby stopni swobody $k = n - r$ odczytujemy z tablicy IV w Dodatku H. Jak łatwo pokazać, w przypadku gdy zachodzi równość $z = x$, tj. wielkość z jest mierzona bezpośrednio, ze wzoru (6.2.9) otrzymujemy (6.1.1).

2) *Wielkość złożona ma rozkład normalny, a wielkości mierzone bezpośrednio są niezależne.*

Mamy tu do czynienia z sytuacją przedstawioną w punkcie 1B, omówionym w rozdziale 5.2.3. Załóżmy, że wielkość złożona z jest funkcją dwóch wielkości x oraz y mierzonych bezpośrednio, oraz że wielkość x zmierzono N razy, a wielkość y zmierzono M razy. Zgodnie z tym, co powiedziano w rozdziale 6.1 wielkość t jest określona przez

$$t = \frac{\bar{z} - z_0}{S_{\bar{z}}}, \quad (6.2.14)$$

gdzie $\bar{z}$ jest dane wzorem (5.1.39), a $S_{\bar{z}}^2$ wzorem (5.2.28). Liczba stopni swobody $k = NM - 1$. Wielkość t , określona wzorem (6.2.14), ma funkcję rozkładu $S(t, k) = S(t, NM - 1)$ daną wzorem (6.1.2). Mając wyznaczone $\bar{z}$ oraz $S_{\bar{z}}$, oceny poziomu i przedziału ufności dokonujemy postępując tak samo jak w przypadku wielkości mierzonej bezpośrednio.

3) *Wielkość złożona ma rozkład normalny, a wielkości mierzone bezpośrednio są zależne.*

Przypadek ten jest identyczny z opisanym w punkcie 2B w rozdziale 5.2.3. W tym przypadku wielkość t dana jest wzorem (6.2.14), zaś $\bar{z}$ oraz $S_{\bar{z}}$ są dane odpowiednio wzorami (5.2.31) i (5.2.32). Liczba stopni swobody $k = N - 1$. W celu oszacowania poziomu i przedziału ufności postępujemy dalej tak jak w przypadku 2.

Rozkład *t-Studenta* i jego zastosowania zostały obszernie przedstawione w [18]. Książkę tę polecamy czytelnikom zainteresowanym tym zagadnieniem.

Przykład

Aby zilustrować zastosowanie rozkładu *t-Studenta* do oceny poziomu i przedziału ufności w przypadku pomiarów pośrednich, gdy zachodzi przypadek 1, wróćmy do przykładu z wyznaczaniem przyspieszenia ziemskiego g , przy użyciu wahadła matematycznego, podanego w rozdziale 3.4. Założyliśmy tam, że pomiar długości wahadła l wykonano $N = 7$ razy i pomiar okresu T wykonano $M = 10$ razy. Wobec tego całkowita liczba pomiarów bezpośrednich $n = N + M = 17$, a liczba stopni swobody $k = n - 2 = 15$. Wyniki obliczeń $\bar{l}$, $\bar{T}$, $\bar{g}$ oraz $S_{\bar{g}}$ przedstawione w przykładzie z rozdziału 3.4 są następujące:

$$\begin{aligned}\bar{l} &= 100.01428 \text{ cm} , & \bar{T} &= 2.007 \text{ s} , \\ \bar{g} &= 980.2205 \text{ cm/s}^2 , & S_{\bar{g}} &= 4.6690808 \text{ cm/s}^2 .\end{aligned}$$

Wykorzystując te same dane, obliczmy wielkości występujące we wzorze (6.2.10); otrzymujemy: $\tilde{S}_l = 0.1807017 \text{ cm}$ oraz $\tilde{S}_T = 0.0141774 \text{ s}$. Stąd ze wzoru (6.2.8) znajdujemy $\tilde{S}_{\bar{g}} = 5.020833 \text{ cm/s}^2$. Wyniki obliczeń podano bez zaokrągleń, tj. tak jak odczytano je z kalkulatora. Porównajmy teraz wyniki oceny przedziału ufności otrzymane przy użyciu rozkładu normalnego i rozkładu *t-Studenta*. Przyjmijmy, że poziom ufności wynosi $\alpha = 0.6827, 0.9545$ i 0.9973 , co w przypadku rozkładu normalnego odpowiada granicom przedziału ufności $\bar{g} \pm S_{\bar{g}}$, $\bar{g} \pm 2S_{\bar{g}}$, $\bar{g} \pm 3S_{\bar{g}}$. Wyniki oceny przedziału ufności przy użyciu tych rozkładów zebrano w tabeli 6.1.

Jak widać z tabeli 6.1, ocena przedziału ufności przy użyciu rozkładu normalnego jest bardziej optymistyczna niż przy użyciu rozkładu *t-Studenta*. Różnice nie są zbyt duże, ale liczba stopni swobody jest duża, a

Tabela 6.1: Porównanie oceny przedziału ufności przy użyciu rozkładu normalnego i rozkładu *t-Studenta*. Granice przedziału ufności podano w jednostkach $[\text{cm/s}^2]$, $\bar{g} = 980.2 \text{ cm/s}^2$

Granice przedziału ufności przy użyciu rozkładu	Poziom ufności α		
	0.6827	0.9543	0.9973
normalnego	$\bar{g} \pm 4.7$	$\bar{g} \pm 9.4$	$\bar{g} \pm 14.1$
<i>t-Studenta</i>	$\bar{g} \pm 5.2$	$\bar{g} \pm 11.0$	$\bar{g} \pm 18.0$

mianowicie $k = 15$. W przypadku mniejszej liczby stopni swobody różnice byłyby większe.

7. Rozkład dwumianowy (Bernoulliego) i rozkład Poissona

Dotychczas zajmowaliśmy się dwoma rozkładami zmiennej losowej ciągłej szczególnie ważnymi z punktu widzenia teorii błędów, a mianowicie rozkładem Gaussa (normalnym) i rozkładem *t-Studenta*, ponieważ opisywały one rozkład wyników pomiarów spowodowany występowaniem wielu czynników przypadkowych. Oprócz tych rozkładów istnieje wiele różnych rozkładów zmiennej losowej ciągłej i skokowej, ważnych zarówno ze względów teoretycznych, jak i praktycznych. W tym rozdziale zajmiemy się dwoma rozkładami zmiennej losowej skokowej, a mianowicie rozkładem dwumianowym nazywanym rozkładem Bernoulliego i rozkładem Poissona. Pierwszy z nich nie ma wprawdzie szczególnego znaczenia przy ocenie błędów, ale stanowi punkt wyjściowy do wyjaśnienia różnych własności rozkładów. W szczególności pozwala na uzasadnienie rozkładu normalnego (Gaussa), co zostało podane w Dodatku E. Rozkład Poissona wynikający z rozkładu dwumianowego jest bardzo ważny w tych wszystkich działach fizyki, w których w czasie pomiarów występuje zliczanie *kolejno pojawiających się przypadkowych zdarzeń* np. zliczanie cząstek emitowanych podczas rozpadu promieniotwórczego bądź fotonów emitowanych przez atomy wzbudzone.

7.1. Rozkład dwumianowy

Przed przystąpieniem do omówienia rozkładu dwumianowego musimy poznać, niezbędną do tego celu, terminologię. Wyobraźmy sobie, że przeprowadzamy N niezależnych **prób** (np. losowań kart, rzutów kośćmi itp.). Każdy wynik, którym jesteśmy zainteresowani, będziemy nazywali **sukcesem** (np. wylosowanie karty określonego koloru, wyrzucenia określonej ilości oczek kośćmi do gry). Każda niezależna próba jest nazywana **zdarzeniem**. Prawdopodobieństwo sukcesu w pojedynczej próbie (np. w jednym losowaniu) oznaczamy przez p , wtedy prawdopodobieństwo $1 - p = q$ jest prawdopodobieństwem uzyskania innego wyniku niż ten, którym jesteśmy zainteresowani, czyli **porażki**.

W celu wyprowadzenia rozkładu dwumianowego przeanalizujmy następujący przykład. Z 52 kartowej talii kart do gry losujemy kartę określonego koloru, np. kierowego. Za sukces będziemy uważali wylosowanie dowolnej karty tego koloru (np. asa, waleta, dwójki itp.), a za porażkę wylosowanie karty z każdego innego koloru. Po losowaniu kartę wkładamy z powrotem

do talii. Prawdopodobieństwo sukcesu w pojedynczej próbie $p = \frac{1}{4}$, gdyż jest ono równe stosunkowi liczby kart sprzyjających odniesieniu sukcesu, tj. liczbie kart danego koloru, która wynosi 13, do liczby wszystkich kart w talii, tj. 52. Prawdopodobieństwo porażki, tzn. wylosowania karty innego koloru niż kierowy, wynosi $1 - p = \frac{3}{4}$.

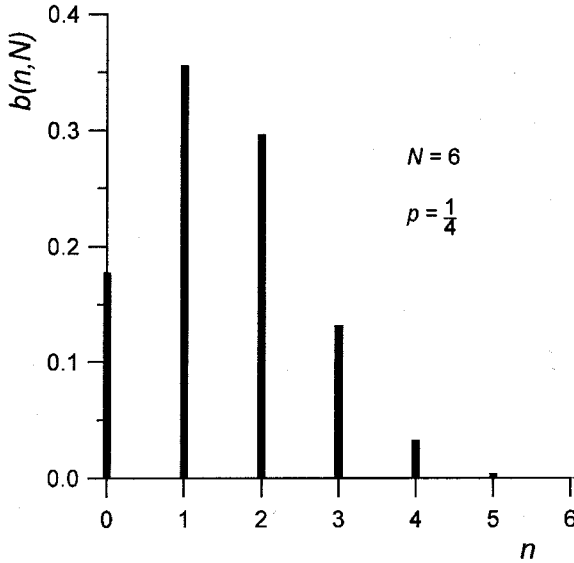
Będziemy teraz poszukiwali prawdopodobieństwa odniesienia n sukcesów (tzn. wylosowania n kart koloru kierowego), gdy wykonamy serię N niezależnych prób (tzn. N razy powtarzamy losowanie i po każdym losowaniu kartę wkładamy z powrotem do talii). Prawdopodobieństwo to będziemy oznaczali $b(n, N)$. Na przykład zapis $b(2, 5)$ będzie oznaczał prawdopodobieństwo dwóch sukcesów w pięciu próbach. Załóżmy teraz, że wykonujemy serię $N = 6$ kolejnych niezależnych od siebie prób. Prawdopodobieństwo tego, że uda się nam odnieść $n = N = 6$ sukcesów (tzn., że w każdej próbie wylosujemy kartę koloru kierowego), jest równe $p^n = p^6$. Wynik ten zapiszemy $b(6, 6) = p^6$. Prawdopodobieństwo odniesienia $n = N - 1 = 5$ sukcesów (np. pierwsze losowanie porażka i następne pięć - sukces) jest równe iloczynowi $p^5(1-p)^1$, czyli $p^n(1-p)^{N-n}$. Pamiętać jednak należy, że porażka może zdarzyć się nie tylko w pierwszej próbie, ale także w drugiej, bądź w trzeciej itd. Możliwych kombinacji, w których wystąpi 5 sukcesów i jedna porażka, jest, jak łatwo policzyć, 6, a więc $b(5, 6) = 6p^5(1-p)^1$. Prawdopodobieństwo 4 sukcesów i 2 porażek w jednej serii prób jest $p^4(1-p)^2$, czyli $p^4(1-p)^{N-4}$. W tym przypadku liczba możliwych kombinacji jest większa. Jak wynika z kombinatoryki, jest ona równa $\binom{6}{4} = 15$, wtedy $b(4, 6) = 15p^4(1-p)^2$. Wyniki powyższych obliczeń możemy przedstawić w postaci jednego wzoru, a mianowicie:

$$b(n, N) = \binom{N}{n} p^n (1-p)^{N-n}, \quad (7.1.1)$$

gdzie wprowadzono symbol Newtona:

$$\binom{N}{n} = \frac{N!}{n!(N-n)!}.$$

Wyrażenie (7.1.1) jest prawdopodobieństwem uzyskania n sukcesów w N niezależnych próbach, gdy prawdopodobieństwo uzyskania sukcesu w pojedynczej próbie jest p , a porażki $q = 1 - p$. Wyrażenie to nosi nazwę **rozkładu dwumianowego** lub **rozkładu Bernoulliego**. Rozkład dwumianowy jest zdefiniowany *tylko* dla całkowitych wartości n . Na rysunku 7.1, jako przykład rozkładu dwumianowego, przedstawiono rozkład prawdopodobieństwa n sukcesów w $N = 6$ próbach, gdy prawdopodobieństwo p sukcesu w pojedynczej próbie wynosi $1/4$.



Rys. 7.1: Rozkład $b(n, N)$ dla przykładu omówionego w tekście $N = 6, p = 1/4$

Rozkład dwumianowy jest unormowany do jedności, co oznacza, że prawdopodobieństwo otrzymania dowolnej wartości n (traktowanej jako sukces) wynosi 1, czyli że

$$\sum_{n=0}^N b(n, N) = \sum_{n=0}^N \binom{N}{n} p^n (1-p)^{N-n} = \sum_{n=0}^N \binom{N}{n} p^n q^{N-n} = 1. \quad (7.1.2)$$

Skorzystaliśmy tu z faktu, że ostatnia suma jest rozwinięciem na dwumian Newtona sumy $p + q$:

$$\sum_{n=0}^N \binom{N}{n} p^n q^{N-n} = (p + q)^N = 1. \quad (7.1.3)$$

Obliczmy teraz **średnią rozkładu** $\bar{n}_R$. Zgodnie ze wzorem (4.1.6) średnia rozkładu jest równa:

$$\bar{n}_R = \sum_{n=0}^N n b(n, N) = \sum_{n=0}^N n \frac{N!}{n!(N-n)!} p^n (1-p)^{N-n}. \quad (7.1.4)$$

Wynosząc Np przed znak sumy i skracając ułamek otrzymujemy:

$$\bar{n}_R = Np \sum_{n=1}^N \frac{(N-1)!}{(n-1)!((N-1)-(n-1))!} p^{n-1} (1-p)^{(N-1)-(n-1)}. \quad (7.1.5)$$

Dokonajmy zmiany wskaźnika sumowania oznaczając $m = n - 1$, wówczas wzór (7.1.5) przyjmie postać:

$$\bar{n}_R = Np \sum_{m=0}^{N-1} \frac{(N-1)!}{m!(N-1-m)!} p^m (1-p)^{N-1-m}.$$

Korzystając z tożsamości (7.1.3) otrzymamy:

$$\bar{n}_R = Np(p+1-p)^{N-1} = Np. \quad (7.1.6)$$

Otrzymany wynik jest zgodny z oczekiwaniami, bo przecież średnia liczba sukcesów powinna być równa iloczynowi liczby prób i prawdopodobieństwa uzyskania sukcesu w pojedynczej próbie. Zgodnie ze wzorem (4.1.8) **wariancja rozkładu** określona jest przez

$$\sigma^2 = \sum_{n=0}^N n^2 b(n, N) - (\bar{n}_R)^2. \quad (7.1.7)$$

W celu obliczenia sumy we wzorze (7.1.7), skorzystajmy z tożsamości $n^2 = n(n-1) + n$. Podstawiając ją do (7.1.7) i uwzględniając (7.1.4) otrzymujemy:

$$\sigma^2 = \sum_{n=0}^N n(n-1) \frac{N!}{n!(N-n)!} p^n (1-p)^{N-n} + \bar{n}_R - (\bar{n}_R)^2. \quad (7.1.8)$$

Wynosząc $N(N-1)p^2$ przed znak sumy i skracając ułamek dostajemy:

$$\begin{aligned} \sigma^2 = N(N-1)p^2 \sum_{n=2}^N \frac{(N-2)!}{(n-2)!((N-2)-(n-2))!} p^{n-2} (1-p)^{(N-2)-(n-2)} \\ + \bar{n}_R - (\bar{n}_R)^2. \end{aligned} \quad (7.1.9)$$

Dokonajmy zamiany wskaźnika sumowania oznaczając $m = n - 2$, wtedy

$$\sigma^2 = N(N-1)p^2 \sum_{m=0}^{N-2} \frac{(N-2)!}{m!(N-2-m)!} p^m (1-p)^{N-2-m} + \bar{n}_R - (\bar{n}_R)^2.$$

Korzystając z tożsamości (7.1.3) znajdujemy:

$$\sigma^2 = N(N-1)p^2 + \bar{n}_R - (\bar{n}_R)^2. \quad (7.1.10)$$

Wstawiając (7.1.6) do (7.1.10) otrzymujemy:

$$\sigma^2 = N^2 p^2 - Np^2 + Np - N^2 p^2,$$

tak więc

$$\sigma^2 = Np(1-p). \quad (7.1.11)$$

Rozkład dwumianowy, chociaż nie znajduje zastosowania przy ocenie błędu pomiarowego, jest szeroko stosowany, np. przy testowaniu hipotez (np. prognoz wyborczych), patrz np. [4, 13, 16]. Uogólnieniem rozkładu dwumianowego jest rozkład wielomianowy, którym nie będziemy się zajmowali, i który ma podobne zastosowania jak rozkład dwumianowy.

7.2. Rozkład Poissona i jego zastosowanie

7.2.1. Rozkład Poissona — zliczanie zdarzeń

W fizyce jądrowej często spotykamy się z sytuacją, gdy posługując się licznikiem cząstek naładowanych, zliczamy liczbę n cząstek emitowanych w przedziale czasu Δt (np. 1 s) w wyniku rozpadu promieniotwórczego jąder atomowych danej próbki. Jeżeli licznik jest sprawny i pomiar został wykonany prawidłowo, to nie będzie występował błąd pomiaru liczby cząstek n . Powtarzając jednak wielokrotnie pomiar w tym samym przedziale czasu Δt stwierdzimy występowanie rozrzutu otrzymanych wartości liczby zliczeń n . Rozrzut ten nie jest spowodowany błędem zliczania, ale statystycznym charakterem samego zjawiska promieniotwórczości. Ponieważ każde z N jąder atomowych, zawartych w badanej próbce, ma w przedziale czasu Δt jednakowe prawdopodobieństwo rozpadu p , możemy obliczyć średnią liczbę rozpadów $\bar{n} = Np$ w tym przedziale czasu. Nas interesuje jednak nie tylko średnia liczba zliczeń $\bar{n}$, ale również jaki jest ich rozkład. Jak wynika z rozważań przedstawionych w rozdziale 7.1, jeśli mamy N jąder i prawdopodobieństwo rozpadu promieniotwórczego pojedynczego jądra w dowolnym, ale ustalonym przedziale czasu Δt jest p , to prawdopodobieństwo zliczenia n rozpadów w tym przedziale czasu jest prawdopodobieństwem n „sukcesów” w N „próbach”. Wobec tego **rozkład liczby zliczeń** jest rozkładem dwumianowym $b(n, N)$. Jeżeli przedział czasu Δt jest na tyle mały, że można przyjąć, iż liczba jąder promieniotwórczych N w próbce praktycznie się nie zmienia w czasie jego trwania, to wtedy $n \ll N$ i $p \ll 1$. Wówczas rozkład liczby zliczeń n („sukcesów”) można opisać rozkładem otrzymanym z rozkładu dwumianowego w wyniku przejścia granicznego: $N \rightarrow \infty$ i $p \rightarrow 0$ w taki sposób, aby $Np = \lambda = \text{const.}$ Aby znaleźć ten rozkład, zapiszmy wzór (7.1.1) w innej równoważnej postaci:

$$b(n, N) = \frac{1}{n!} Np \dots (N-k)p \dots (N-n-1)p (1-p)^N (1-p)^{-n}. \quad (7.2.1)$$

W przypadku granicznym, gdy $\lambda = Np = \text{const}$, zachodzą związki:

$$\lim_{\substack{N \rightarrow \infty \\ p \rightarrow 0}} (N - k)p = \lambda \lim_{N \rightarrow \infty} \frac{N - k}{N} = \lambda,$$

$$\lim_{\substack{N \rightarrow \infty \\ p \rightarrow 0}} (1 - p)^N = \lim_{p \rightarrow 0} (1 - p)^{\frac{\lambda}{p}} = \left[\lim_{x \rightarrow \infty} \left(1 - \frac{1}{x}\right)^x \right]^\lambda = [e^{-1}]^\lambda = e^{-\lambda},$$

(w powyższym wyrażeniu $x = \frac{1}{p}$) oraz

$$\lim_{p \rightarrow 0} (1 - p)^{-n} = \left[\lim_{p \rightarrow 0} (1 - p) \right]^{-n} = 1.$$

Granice rozkładu $b(n, N)$, gdy $N \rightarrow \infty$, $p \rightarrow 0$ obliczamy korzystając z powyższych związków. Wynosi ona

$$\boxed{\lim_{\substack{N \rightarrow \infty \\ p \rightarrow 0}} b(n, N) = P_\lambda(n) = \frac{\lambda^n}{n!} e^{-\lambda}.} \quad (7.2.2)$$

Rozkład $P_\lambda(n)$ nosi nazwę **rozkładu Poissona**. Jest on, podobnie jak rozkład dwumianowy, unormowany do 1:

$$\sum_{n=0}^{\infty} P_\lambda(n) = e^{-\lambda} \sum_{n=0}^{\infty} \frac{\lambda^n}{n!} = e^{-\lambda} e^\lambda = 1. \quad (7.2.3)$$

Skorzystaliśmy tu z rozwinięcia funkcji e^λ na szereg Maclaurina, a mianowicie:

$$e^\lambda = \sum_{n=0}^{\infty} \frac{\lambda^n}{n!} = 1 + \lambda + \frac{1}{2!}\lambda^2 + \frac{1}{3!}\lambda^3 + \dots$$

Obliczmy teraz wartość średnią rozkładu $\bar{n}_R$. Korzystając z (4.1.6), znajdujemy:

$$\bar{n}_R = e^{-\lambda} \sum_{n=0}^{\infty} n \frac{\lambda^n}{n!} = e^{-\lambda} \sum_{n=1}^{\infty} n \frac{\lambda^n}{n!} = e^{-\lambda} \sum_{n=1}^{\infty} \lambda \frac{\lambda^{n-1}}{(n-1)!}.$$

Wynosząc λ przed znak sumy i dokonując zamiany wskaźnika sumowania na $m = n - 1$, otrzymujemy ważny związek, a mianowicie

$$\bar{n}_R = \lambda e^{-\lambda} \sum_{m=0}^{\infty} \frac{\lambda^m}{m!} = \lambda e^{-\lambda} e^\lambda = \lambda. \quad (7.2.4)$$

Oznacza to, że stała λ w rozkładzie Poissona jest równa średniej rozkładu, czyli

$$\boxed{\bar{n}_R = \lambda.} \quad (7.2.5)$$

W celu obliczenia wariancji rozkładu Poissona znajdziemy na początku $\overline{n^2}_R$:

$$\overline{n^2}_R = e^{-\lambda} \sum_{n=0}^{\infty} n^2 \frac{\lambda^n}{n!} = \lambda e^{-\lambda} \sum_{n=1}^{\infty} n \frac{\lambda^{n-1}}{(n-1)!}.$$

Dokonajmy następnie zamiany wskaźnika sumowania przez podstawienie $m = n - 1$. Mamy wówczas:

$$\overline{n^2}_R = \lambda e^{-\lambda} \sum_{m=0}^{\infty} (m+1) \frac{\lambda^m}{m!} = \lambda e^{-\lambda} \sum_{m=0}^{\infty} m \frac{\lambda^m}{m!} + \lambda e^{-\lambda} \sum_{m=0}^{\infty} \frac{\lambda^m}{m!}.$$

Pierwsza suma jest równa λe^λ , zaś druga e^λ , tak więc:

$$\overline{n^2}_R = \lambda^2 + \lambda. \quad (7.2.6)$$

Zgodnie z (4.1.8) **wariancja** rozkładu σ^2 wyraża się następująco:

$$\sigma^2 = \overline{n^2}_R - (\bar{n}_R)^2. \quad (7.2.7)$$

Podstawiając (7.2.5) i (7.2.6) do (7.2.7) otrzymujemy:

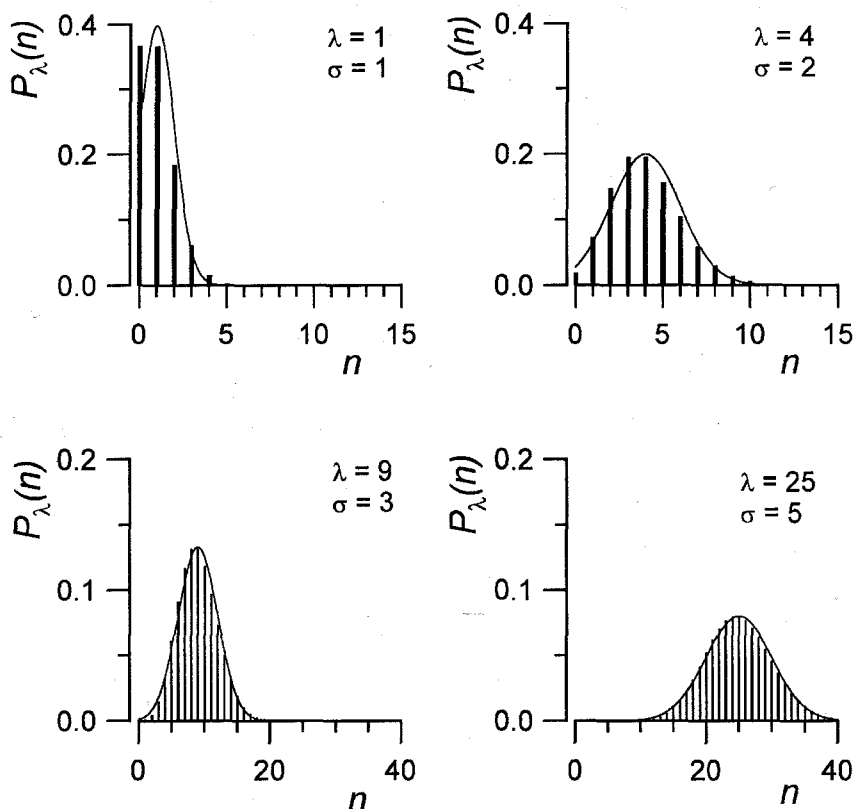
$$\sigma^2 = \lambda^2 + \lambda - \lambda^2,$$

czyli, że

$$\boxed{\sigma^2 = \lambda = \bar{n}_R.} \quad (7.2.8)$$

Oznacza to, że wariancja rozkładu Poissona jest równa średniej rozkładu. Rozkład Poissona jest więc scharakteryzowany jednym parametrem, podczas gdy rozkład normalny charakteryzują dwa parametry, a mianowicie: wariancja i średnia rozkładu. Jak wynika z rozważań przedstawionych w rozdziale 4.1, parametr λ rozkładu Poissona $P_\lambda(n)$ jest równy *średniej liczbie zliczeń dużej serii pomiarowej* (tzn. takiej, której liczebność zmierza do liczebności populacji generalnej) wykonanej w tych samych warunkach i w tym samym przedziale czasu Δt .

Rozkład Poissona $P_\lambda(n)$, podobnie jak **rozkład dwumianowy** $b(n, N)$, jest określony tylko dla całkowitych i nieujemnych n . Nie oznacza to jednak, że $\bar{n}$ i $\bar{n}_R$ muszą być liczbami całkowitymi. Rozkład $P_\lambda(n)$ jest *niesymetryczny*. Asymetria rozkładu maleje ze wzrostem $\bar{n}_R$ i dla dużych wartości $\bar{n}_R$ jest on prawie symetryczny. Na rysunku 7.2 przedstawiono



Rys. 7.2: Przykłady rozkładu Poissona $P_\lambda(n)$ dla $\lambda = 1, 4, 9, 25$. Linia ciągłą zaznaczono dla porównania rozkład normalny $N(\lambda, \sqrt{\lambda})$

przykłady rozkładu Poissona dla $\lambda = 1, 4, 9, 25$; liniami ciągłymi zaznaczono rozkład normalny o tej samej wariancji i średniej rozkładu. Jak widać z rysunku 7.2 dla $\lambda = 1$ rozkład jest niesymetryczny, dla $\lambda = 9$ rozkład jest prawie symetryczny i różnica między rozkładem Poissona i normalnym jest wyraźna, ale niezbyt duża. Dla $\lambda = 25$ rozkład Poissona jest praktycznie symetryczny i różnice pomiędzy nim a rozkładem normalnym są prawie niezauważalne. Oznacza to, że **rozkład Poissona** szybko zmierza do **rozkładu normalnego** i wyniki obliczeń poziomu jak i przedziału ufności przy użyciu tych rozkładów będą prawie jednakowe. Fakt ten ma, jak się przekonamy dalej, bardzo duże znaczenie praktyczne.

Ogólnie mówiąc, rozkład Poissona stosujemy zawsze, gdy mamy dużą liczbę zdarzeń, ale jedynie niewielką ich część posiada interesującą nas własność. Tak więc rozkład Poissona $P_\lambda(n)$ stosuje się *przede wszystkim*, gdy zliczamy:

- cząstki emitowane w różnych procesach zachodzących w jądrach atomowych;
- cząstki rozproszone w zderzeniach;
- emitowane lub absorbowane fotony;
- a także przy kontroli jakości produkcji, gdy liczba wadliwych wyrobów jest mała w porównaniu z liczbą wyrobów wyprodukowanych.

Tablice rozkładu Poissona dla małych wartości λ podano w Dodatku H.

7.2.2. Dystrybuanta rozkładu Poissona

Dystrybuanta $P(n < A)$, zgodnie z tym co powiedziano w rozdziale 4.1 (wzór (4.1.9)), wyraża się wzorem:

$$P(n < A) = \sum_{n=0}^{A-1} P_{\lambda}(n) = e^{-\lambda} \sum_{n=0}^{A-1} \frac{\lambda^n}{n!} \quad (7.2.9)$$

i jest prawdopodobieństwem tego, że liczba odniesionych sukcesów n będzie mniejsza niż A . **Prawdopodobieństwo** $P(A \leq n < B)$, zgodnie ze wzorem (4.1.10), wynosi:

$$P(A \leq n < B) = P(n < B) - P(n < A) = \sum_{n=A}^{B-1} P_{\lambda}(n) = e^{-\lambda} \sum_{n=A}^{B-1} \frac{\lambda^n}{n!}, \quad (7.2.10)$$

zaś prawdopodobieństwo $P(n \geq B)$ tego, że liczba odniesionych sukcesów będzie nie mniejsza niż B , wynosi

$$P(n \geq B) = \sum_{n=B}^{\infty} P_{\lambda}(n) = e^{-\lambda} \sum_{n=B}^{\infty} \frac{\lambda^n}{n!}. \quad (7.2.11)$$

Korzystanie ze wzorów (7.2.9) – (7.2.11) dla λ przekraczających 10 jest niewygodne ze względów rachunkowych, np. dla $\lambda = 100$ zachodzi konieczność obliczenia takich wielkości jak e^{-100} czy $100!$. Jak pokazano na rysunku 7.2, dla $\lambda = 25$ rozkład Poissona jest bardzo bliski rozkładowi normalnemu. A więc dla dużych $\bar{n}_R$ do oceny przedziału i poziomu ufności $\bar{n}_R$ możemy zastosować **rozkład normalny standaryzowany**. Zmienna standaryzowana w tym przypadku przyjmie postać:

$$u = \frac{\bar{n} - \bar{n}_R}{\sigma}. \quad (7.2.12)$$

We wzorze (7.2.12) σ^2 jest wariancją rozkładu Poissona i $\sigma^2 = \bar{n}_R$. Ponieważ nie znamy średniej rozkładu $\bar{n}_R$, musimy przybliżyć ją średnią arytmetyczną

liczby zliczeń ($\bar{n} = \frac{1}{N} \sum_{i=1}^N n_i$, gdzie n_i jest liczbą zliczeń w i -tym pomiarze, a N liczbą pomiarów). **Zmienna standaryzowana** u będzie więc miała postać:

$$u = \frac{\bar{n} - \bar{n}_R}{\sqrt{\bar{n}}}. \quad (7.2.13)$$

Powstaje pytanie, jak duże powinno być $\bar{n}$, aby można było korzystać ze wzoru (7.2.13)? W praktyce przyjmuje się $\bar{n} > 30$. Należy jednak zaznaczyć, że przyjęcie $\bar{n} = 30$ jest zbyt optymistyczne i wzór (7.2.13) zalecamy stosować dla $\bar{n} \geq 100$. Przykłady ilustrujące dokładność przybliżenia rozkładu Poissona $P_\lambda(n)$ rozkładem normalnym standaryzowanym $N(0, 1)$ można znaleźć w wielu opracowaniach, np. w [4]. Na zakończenie musimy odpowiedzieć na pytanie, co należy zrobić, jeżeli średnia liczba zliczeń $\bar{n}$ w czasie Δt jest zbyt mała, aby można było rozkład $P_\lambda(n)$ przybliżyć rozkładem $N(0, 1)$, a musimy oszacować poziom i przedział ufności średniej rozkładu $\bar{n}_R$. Odpowiedź jest następująca: możemy wydłużyć czas zliczania. Na przykład, jeśli w czasie 10 s średnia liczba zliczeń $\bar{n} = 20$, to możemy przedłużyć czas zliczania np. do 60 s.

8. Przedstawianie danych i graficzne oszacowanie błędu

Jednym z najpoważniejszych i najczęściej spotykanych błędów popełnianych przez niedoświadczonych eksperymentatorów, jak i numeryków jest chaotyczny zapis wyników pomiarów bądź obliczeń. Otrzymujemy wtedy nieopisany zbiór liczb, w którym łatwo jest się pogubić i po pewnym czasie zapomnieć, co która liczba przedstawia, oraz w jakich jednostkach jest wyrażony wynik pomiaru bądź obliczeń. Dlatego wszystkie dane należy *zawsze* przedstawiać w jakiś uporządkowany sposób zawierający również ich opis.

Z tego powodu, w tym rozdziale, omówimy podstawowe zasady przedstawiania danych dowolnego rodzaju w postaci uporządkowanych zbiorów i wykresów oraz zasady przedstawiania wielkości fizycznych i zapisu ich jednostek.

8.1. Zasady zapisu wielkości fizycznych i ich jednostek

Przygotowując dowolne opracowanie (np. opis ćwiczenia wykonanego w pracowni) musimy zwracać uwagę na symbole używane do oznaczania wielkości fizycznych i ich jednostek. Wszystkie symbole występujące w tekście opracowania, jak również w tabelach i na rysunkach muszą być wyraźne, aby nie było wątpliwości, którym symbolem oznaczono daną wielkość fizyczną. Ponieważ coraz częściej opracowania są przygotowywane przy użyciu komputerowych edytorów tekstu, podamy poniżej zasady używania liter jako symboli wielkości fizycznych. Zasady te są zalecane przez Międzynarodową Unię Fizyki Czystej i Stosowanej (International Union of Pure and Applied Physics, w skrócie IUPAP) i obowiązują w ramach układu SI [2]. Są one stosowane w tzw. dużej poligrafii i są następujące.

- 1) Liczby, symbole jednostek miar, symbole pierwiastków chemicznych piszemy drukiem prostym (antykwą).
- 2) Symbole skalarnych wielkości fizycznych piszemy drukiem pochylonym (kursywą, w anglojęzycznych edytorach: italic).
- 3) Symbole wielkości fizycznych wektorowych piszemy drukiem pochylonym półgrubym.

- 4) Symbole jednostek miar piszemy z jednym odstępem (spacją) po miarze wielkości fizycznej; symbol jednostki miary nie jest skrótem, lecz wielkością matematyczną i na jego oznaczenie używamy zazwyczaj pierwszej litery nazwy jednostki (są jednak wyjątki np. cd, mol, Hz).
- 5) Symbole jednostek miar piszemy małymi literami, chyba że nazwa jednostki pochodzi od nazwiska np. A, N, Hz, nazwy jednostek piszemy zawsze małą literą np. amper (A), niuton (N), sekunda (s).
- 6) Aby uniknąć pisania liczb z dużą liczbą zer dla określenia wartości znacznie większych i znacznie mniejszych niż jednostka podstawowa, wprowadza się ich wielokrotności lub podwielokrotności; są one jednostkami wtórnymi, wyraża się je przez dodanie do nazwy lub symbolu jednostki przedrostka lub jego symbolu wyrażającego odpowiedni mnożnik dziesiętny, przedrostek umieszcza się bezpośrednio przed nazwą jednostki (np. pikofarad), a jego symbol bezpośrednio przed symbolem jednostki (nie wolno stosować odstępów ani kresek) np. μs , pisze się je tak jak jednostki drukiem prostym. Nie używa się przedrostków złożonych, nie należy więc pisać: $\mu\mu\text{F}$ a pF, kW a GW. Wykładnik potęgi przy jednostce wtórnej odnosi się do całej jednostki wraz z przedrostkiem np. cm^2 oznacza $(0.01\text{m})^2$, μs^{-1} oznacza $(10^{-6}\text{s})^{-1}$.
- 7) Mnożenie jednostek miar zapisujemy umieszczając między nimi podniesioną kropkę lub pozostawiając między nimi jeden odstęp (spację), np. m·s lub m s. Dzielenie jednostek zapisujemy w jednej z trzech postaci, np. $\frac{\text{m}}{\text{s}}$, m/s i m s^{-1} . Nie wolno stosować podwójnej kreski ułamekowej, np. m/s/s należy zapisać m/s^2 . W przypadku gdy w mianowniku występuje więcej niż jedna jednostka, wskazane jest stosowanie nawiasów lub ujemnych wykładników potęg, np. prawidłowy jest zapis $\text{kg}/(\text{m}\cdot\text{s})$ lub $\text{kg m}^{-1}\text{s}^{-1}$. Powyższe zasady muszą być stosowane również przy tworzeniu zbiorów danych dowolnej postaci oraz przy ich graficznym przedstawianiu.

W przypadku gdy opracowanie jest przygotowane w postaci rękopisu, nie możemy stosować zasad wymienionych w punktach 1–3. Wielkości wektorowe oznaczamy wówczas pisząc nad ich symbolami strzałkę.

8.2. Tabelaryczne przedstawianie danych

Tabelaryczne przedstawianie danych polega na zestawieniu ich w postaci listy (np. wartości liczbowych) ułożonej w pewien uporządkowany sposób. Uporządkowanie to ma na celu przygotowanie danych do dalszej analizy numerycznej bądź graficznej, ułatwiającej znalezienie związków między nimi,

a także stworzenie możliwości ich wykorzystania. Wyniki pomiarów bądź obliczeń przedstawione w postaci uporządkowanej listy ich wartości liczbowych noszą nazwę tabel lub tablic.

8.2.1. Zasady sporządzania tabel

Przy sporządzaniu tabel należy postępować zgodnie z poniżej podanymi zasadami.

Każda tabela musi być opisana, tzn. zgodnie z przyjętymi zwyczajami *nad tabelą* zamieszczamy informację co ona przedstawia. Jest to nagłówek tabeli. W nagłówku tabeli należy wyjaśnić symbole występujących wielkości oraz podać ich jednostki. Jeżeli oznaczenia występujące w tabeli zostały wcześniej wyjaśnione w tekście, to nie musimy ich ponownie wyjaśniać, ale należy to zaznaczyć w nagłówku, zamieszczając odpowiednią uwagę, np. *Objaśnienia w tekście*. W przypadku gdy w opracowaniu występuje kilka tabel z takimi samymi oznaczeniami, to nie musimy ich opisywać przy każdej tabeli. Wystarczy jeśli zostały wyjaśnione w pierwszej tabeli, w której występują. W następnych można wyjaśnienia te opuścić, ale wtedy w nagłówku należy zamieścić uwagę: *Objaśnienia jak w tabeli...* W *wyjątkowych* przypadkach dodatkowe informacje można zamieszczać pod tabelą, należy jednak tego unikać.

Jeżeli w danym opracowaniu występuje więcej niż jedna tabela, to muszą one być ponumerowane. Sposób numeracji tabel jest zasadniczo dowolny, jednak raz przyjętego sposobu musimy przestrzegać w całym opracowaniu.

Przystępując do układania tabel musimy zwracać uwagę na ich przejrzystość. Jeżeli tabela zawiera dużo danych, w których jest się trudno zorientować, jej użyteczność jest mała. Dane zamieszczone w tabelach grupujemy w kolumnach. W jednej kolumnie zamieszczamy dane dotyczące jednej wielkości fizycznej, jednej próbki lub np. uzyskane w jednym laboratorium czy też jedną metodą. Niezależnie od nagłówka tabeli każda kolumna powinna mieć swój nagłówek. Jeżeli w nagłówku tabeli nie podajemy jednostek, to należy je podać w nagłówkach kolumn. Zazwyczaj w pierwszej kolumnie tabeli zamieszczamy liczbę porządkową, czyli numer wiersza. Jeżeli rezygnujemy z liczby porządkowej, to tę kolumnę po prostu opuszczamy. W następnej kolumnie zamieszczamy zmienną niezależną, a w następnych tabelaryzowane wielkości.

Jeżeli tabela nie mieści się na jednej stronie, to numerujemy poszczególne kolumny. Numery kolumn wpisujemy w wierszu poniżej nagłówków kolumn. Na następnych stronach nie zamieszczamy nagłówków kolumn, a jedynie ich numery.

W przypadku gdy danych jest niewiele, można sporządzać tabele w układzie poziomym, a nie w kolumnowym, np. tabela 3.1 w tym opracowaniu.

8.2.2. Rodzaje tabel

Tabele, w zależności od tego jakie związki między danymi chcemy w nich przedstawić, dzielimy na (patrz np. [19]):

- jakościowe,
- funkcyjne,
- statystyczne.

W tabelach jakościowych zamieszczamy zestawienia interesujących nas cech jakościowych. Przykładem może być tabela 8.1, w której zamieszczono definicje jednostek podstawowych Międzynarodowego Układu Jednostek SI.

Tabele funkcyjne przedstawiają zależność funkcyjną między wielkościami fizycznymi, np. zależność przebytej drogi od czasu. Czasami w takiej tabeli przedstawiamy wyniki kolejnych pomiarów (obliczeń) tej samej wielkości fizycznej. Przykładem może być np. tabela 4.2.

Tabele statystyczne podają zazwyczaj zależność między grupami danych, z których jedna jest ujęta jakościowo i traktowana jako zmienna niezależna, a pozostałe ilościowo. Tablice takie są zamieszczane np. w rocznikach statystycznych. Są one również stosowane do przedstawiania danych doświadczalnych, jak i wyników obliczeń. Jako przykład takiej tabeli zamieszczamy zestawienie wyników pomiarów stałej grawitacyjnej G otrzymanych w latach 1990–1996 (tabela 8.2).

W praktyce będziemy się spotykali z dwoma rodzajami tabel, a mianowicie funkcyjnymi i statystycznymi. Szczegółowe omówienie zasad sporządzania różnego rodzaju tablic można znaleźć np. w [19].

8.3. Graficzne przedstawianie danych

Przedstawianie wyników pomiarów w postaci wykresu stosuje się zazwyczaj wtedy, gdy mierzymy dwie lub więcej wielkości fizycznych, o których *wiemy* lub *sądzymy*, że są związane jakąś zależnością funkcyjną, np. $y = f(x)$. Wykres zależności $y = f(x)$, w przeciwieństwie do tabeli zawierającej wyniki pomiarów, stanowi najbardziej pogładowe przedstawienie zależności funkcyjnej, jest jej „obrazem”.

W badaniach fizycznych graficznym przedstawieniem wyników posługujemy się przede wszystkim, gdy:

- a) chcemy wykazać, że wielkości y i x są związane uprzednio założoną zależnością funkcyjną, tzn. chcemy porównać wyniki pomiarów z przewidywaniami teoretycznymi; występowanie systematycznych różnic oznacza niezgodność danych doświadczalnych z przewidywaniami teoretycznymi i może wskazywać na konieczność uściślenia teorii;

Tabela 8.1: Przykład tabeli jakościowej

Jednostki podstawowe układu SI [1]

Lp.	Wielkość	Jednostki miary		Definicja jednostki
		nazwa	symbol	
1	długość	metr	m	Metr jest to długość drogi przebytej w próżni przez światło w czasie $1/299\,792\,458$ sekundy.
2	masa	kilogram	kg	Kilogram jest masą międzynarodowego wzorca tej jednostki przechowywanego w Międzynarodowym Biurze Miar w Sèvres.
3	czas	sekunda	s	Sekunda jest to czas równy $9\,192\,631\,770$ okresom promieniowania odpowiadającego przejściu między dwoma nadsubtelnymi poziomami stanu podstawowego atomu ^{133}Cs (cezu 133).
4	natężenie prądu elektrycznego	amper	A	Amper jest to prąd elektryczny nie zmieniający się, który płynie w dwóch przewodach równoległych, prostoliniowych, nieskończenie długich, o przekroju kołowym znikomo małym, umieszczonych w próżni w odległości 1 m od siebie, wywołalby między tymi przewodami siłę $2 \cdot 10^{-7}$ N na każdy metr długości.
5	temperatura	kelwin	K	Kelwin jest to $1/273.16$ temperatury termodynamicznej punktu potrójnego wody; stosuje się do wyrażania temperatury termodynamicznej T i różnicy temperatur.
6	ilość materii	mol	mol	Mol jest to ilość materii, występująca gdy liczba cząstek jest równa liczbie atomów zawartych w masie 0.012 kg ^{12}C (węgla 12); przy stosowaniu mola należy określić rodzaj cząstek; mogą nim być: atomy, molekuly jony, elektrony itp., albo określone zespoły takich cząstek.
7	światłość	kandela	cd	Kandela jest to światłość, jaką ma w określonym kierunku źródło emitujące promieniowanie monochromatyczne o częstotliwości $540 \cdot 10^{12}$ Hz, i którego natężenie w tym kierunku jest równe $1/683$ W/sr.

Tabela 8.2: Przykład tabeli statystycznej

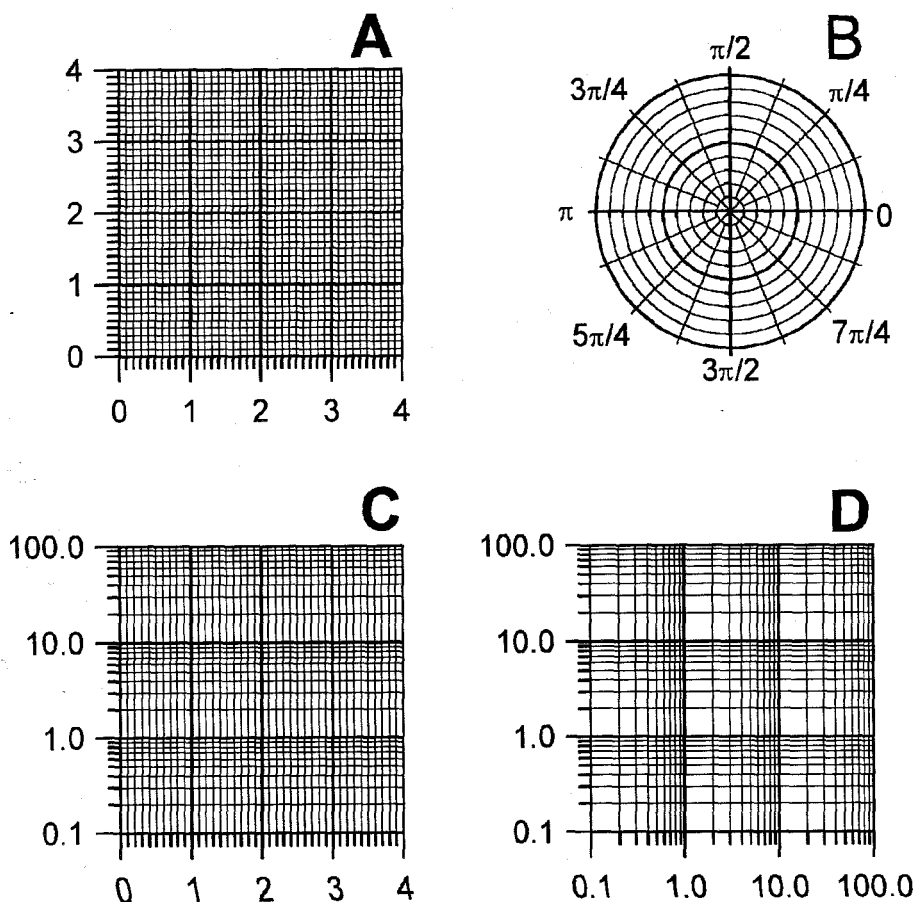
Zestawienie wyników pomiarów stałej grawitacyjnej G
otrzymanych w latach 1990–1996 [20]

Autorzy	$G [10^{-11} \text{ m}^3 \text{ s}^{-2} \text{ kg}^{-1}]$
Muller i inni (1990)	6.689 ± 0.027
Zumberge i inni (1991)	6.677 ± 0.013
Schurr i inni (1991)	6.66 ± 0.06
Yang i inni (1991)	6.672 ± 0.040
Schurr i inni (1992)	6.6613 ± 0.0011
Oldham i inni (1993)	6.671 ± 0.015
	6.703 ± 0.035
Walesch i inni (1994)	6.6724 ± 0.0015
Walesch i inni (1995)	6.6719 ± 0.0008
Fitzgerald i Armstrong (1995)	6.6656 ± 0.0006
Hubler i inni (1995)	6.678 ± 0.007
	6.669 ± 0.005
Meyer i inni (1995)	6.6685 ± 0.0007
Fitzgerald (1995)	6.6659 ± 0.0006
Michaelis i inni (1995/96)	6.71540 ± 0.00056
	6.7174 ± 0.0020
Bagley i Luther (1996)	6.6739 ± 0.0011
	6.6741 ± 0.0008

- b) w przypadku braku danych teoretycznych chcemy przewidzieć typ zależności funkcyjnej między mierzonymi wielkościami lub gdy chcemy ustalić zależność empiryczną między dwiema wielkościami, jest to przypadek występujący szczególnie często przy kalibracji przyrządów pomiarowych;
- c) chcemy graficznie wyznaczyć wartość jakiejś wielkości (współczynnika) występującej w zależności funkcyjnej (ten sposób w związku z rozwojem metod komputerowych jest coraz rzadziej stosowany);
- d) wyznaczonej zależności funkcyjnej nie można przybliżyć żadną prostą zależnością (np. pętla histerezy), w tym przypadku wykres stanowi ostateczną formę przedstawienia wyników;
- e) badamy rozkład wyników pomiarów w danej serii, sporządzamy wtedy histogramy omówione w rozdziale 4.

8.3.1. Najczęściej używane papiery funkcyjne

Wykresy wykonuje się na tzw. papierze funkcyjnym. Na papier funkcyjny jest zawsze naniesiona odpowiednia siatka. Siatkę stanowi układ nadrukowanych linii, między którymi zachodzą związki zależne od rodzaju siatki, np. siatka milimetrowa składa się z dwóch prostopadłych do siebie układów linii prostych równoległych o odległości między liniami 1 mm. Istnieje wiele rodzajów siatek, wybór określonego rodzaju siatki jest zależny od analizowanej zależności funkcyjnej. Szczegółowe omówienie różnych siatek, zasady tworzenia i posługiwania się nimi można znaleźć np. w [21] (patrz, także [13]).



Rys. 8.1: Przykłady papierów funkcyjnych do sporządzania wykresów. Siatka: A) liniowa, B) biegunowa, C) półlogarytmiczna, D) logarytmiczna

Najczęściej używane rodzaje papierów funkcyjnych przedstawiono na rysunku 8.1, należą do nich:

- papier z siatką milimetrową, nazywany potocznie papierem milimetrowym (rys. 8.1A);
- papier z siatką biegunową;
- papier z siatką półlogarytmiczną, nazywany potocznie papierem półlogarytmicznym lub papierem logarytmicznym (rys. 8.1C);
- papier z siatką logarytmiczną, nazywany potocznie papierem logarytmicznym, (rys. 8.1D).

W przypadku papierów z siatką półlogarytmiczną i logarytmiczną najczęściej przyjmuje się za podstawę logarytmu 10.

Papier z siatką milimetrową ma charakter uniwersalny i jest najczęściej stosowany. Używany jest do różnego rodzaju wykresów, przede wszystkim wówczas, gdy spodziewamy się liniowej zależności między mierzonymi wielkościami, oraz gdy mierzone wielkości zmieniają się w zakresie mniej więcej jednego rzędu.

Papier z siatką półlogarytmiczną jest używany przede wszystkim wtedy, gdy spodziewamy się zależności typu:

$$y = ab^{cx},$$

(przy czym a i c są stałe i mogą być nieznane), ponieważ $\log y = \log a + cx \log b$, a więc $\log y$ jest liniową funkcją x i wykres jest linią prostą. Stosuje się go również, gdy jedna z wielkości zmienia się znacznie (np. o kilka rzędów), a druga zmienia się parokrotnie.

Papier z siatką logarytmiczną stosuje się, gdy spodziewamy się zależności typu:

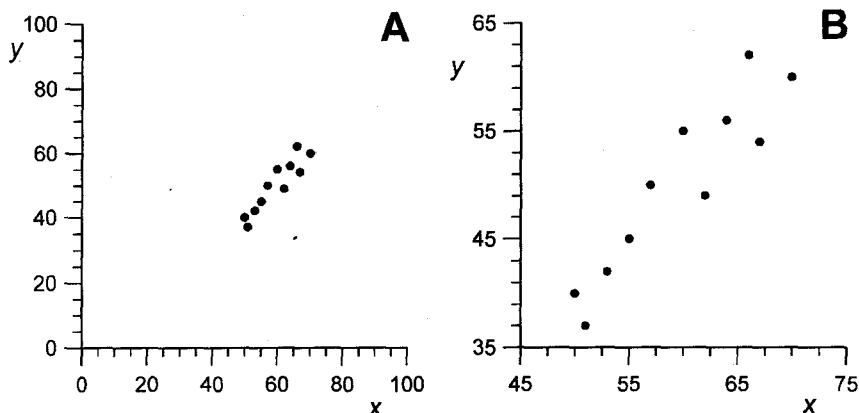
$$y = ax^b,$$

(stałe a i b mogą być nieznane), oraz gdy obie mierzone wielkości zmieniają się w zakresie paru rzędów.

8.3.2. Zasady sporządzania wykresów

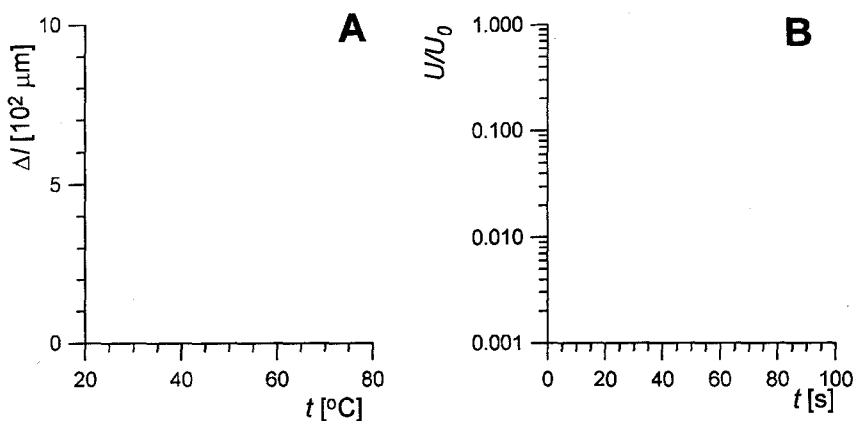
Przed przystąpieniem do wykonania wykresu należy dokonać wyboru papieru funkcyjnego z odpowiednią siatką, pamiętając o tym, że w fizyce na ogół na osi odciętych (poziomej) odkłada się zmienną niezależną, a na osi rzędnych (pionowej) zmienną zależną. Następną czynnością jest dobranie skali. Wybierając skalę należy przestrzegać następujących reguł:

- a) punkty doświadczalne nie mogą być zgrupowane na małym fragmencie rysunku (rys. 8.2A) i powinny pokrywać cały wykres (rys. 8.2B), a skala *nie musi* zaczynać się od zera;
- b) skala powinna być możliwie prosta.



Rys. 8.2: A) Wykres mało użyteczny — nieprawidłowy. B) Te same wyniki przedstawione w prawidłowy sposób

Na przykład przy siatce milimetrowej 1 cm powinien, jeśli tylko jest to możliwe, odpowiadać $1 \cdot 10^n$, $2 \cdot 10^n$ lub $5 \cdot 10^n$ (n dowolna liczba całkowita) jednostkom. Dopuszczalny jest oczywiście inny podział, ale proponowany powyżej jest najwygodniejszy przy odczytywaniu i nanoszeniu danych.

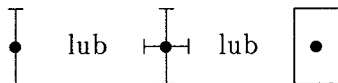


Rys. 8.3: Przykłady opisu osi: A) przy pomiarach rozszerzalności liniowej (Δl wydłużenie pręta); B) przy pomiarach napięcia na kondensatorze rozładowywanym przez opór

Kolejną czynnością jest opisanie obu osi układu współrzędnych za pomocą nazw lub symboli obu wielkości i ustalenie jednostek miar. Ustalać jednostki miar wygodnie jest wybrać mnożnik dziesiętny taki, aby działki skali opisywać liczbami 1, 2, 3, ..., czy też 10, 20, 30, ..., a nie 100000, 200000, ... lub 0.00001, 0.00002, ... itd. Na rysunku 8.3A przed-

stawiono przykład opisu osi na papierze funkcyjnym z siatką milimetrową, a na rysunku 8.3B z siatką półlogarytmiczną.

Na tak przygotowany papier funkcyjny nanosimy punkty doświadczalne. Punkty doświadczalne muszą być zaznaczone *wyraźnymi* symbolami (np. ●, ×, ○, *) tak, aby były dobrze widoczne. Jeżeli nanosimy na wykres punkty otrzymane w różnych warunkach, czy dla różnych próbek, to *muszą one różnić się symbolami*. Na wykresach powinny być zaznaczone niepewności pomiarowe. Można je zaznaczać na przykład za pomocą symboli:



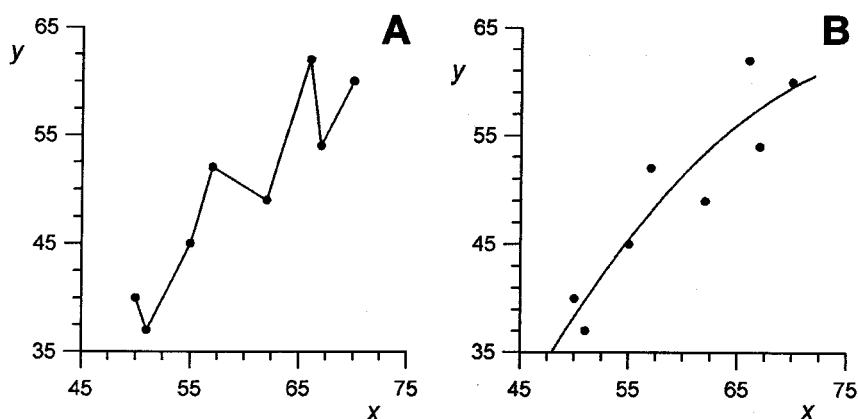
Prostokąt taki nazywamy **prostokątem błędów**.

Na ogół, gdy dla jednej wartości zmiennej niezależnej wykonujemy szeregi pomiarów zmiennej zależnej, wówczas na wykres nanosimy ich średnią arytmetyczną lub średnią ważoną. W niektórych przypadkach, gdy chcemy przedstawić rozrzut wyników pomiarów, nanosimy wszystkie zmierzone wartości. Na wykresach nanosimy zasadniczo błąd maksymalny. W przypadkach gdy nanosimy odchylenia standardowe serii lub średniej serii musimy to wyraźnie zaznaczyć w opisie rysunku.

Każdy rysunek musi być zaopatrzony podpisem, w którym podajemy, jaką zależność on przedstawia. Oprócz tego podpis pod rysunkiem musi zawierać objaśnienie stosowanych symboli. Jeśli są one wyjaśnione w tekście, do którego dołączony jest rysunek, to nie musimy ich powtarzać, ale należy to zaznaczyć w podpisie, zamieszczając odpowiednią formułę, np.: *Objaśnienia w tekście*. W przypadku gdy do tekstu jest dołączonych kilka rysunków z tymi samymi oznaczeniami, to należy objaśnić je na pierwszym rysunku, na którym one występują, a na następnym można zamieścić uwagę: *Oznaczenia jak na rys...*

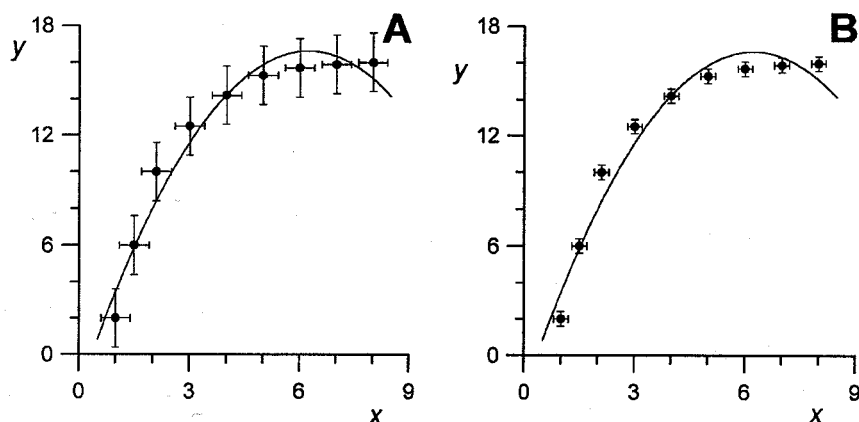
W celu pokazania przebiegu badanej zależności *między punktami doświadczalnymi prowadzimy krzywą gładką* — „na oko” — używając krzywika, albo jeśli znamy postać zależności funkcyjnej, to znajdujemy ją metodą najmniejszych kwadratów. Przeprowadzając krzywą przy użyciu krzywika powinniśmy się starać, aby na każdym odcinku krzywej liczba punktów po prawej i lewej stronie była możliwie jednakowa, oraz aby krzywa przechodziła jeśli nie przez wszystkie, to przez większość prostokątów błędów. Łączenie ze sobą poszczególnych punktów doświadczalnych jest *nieprawidłowe*. Ponieważ każdy pomiar jest obciążony jakimś błędem, rozrzut wyników pomiarów jest rzeczą naturalną i łączenie poszczególnych punktów może sugerować skokowy charakter zależności, co nie zawsze ma miejsce. Jeżeli jednak podejrzewamy skokowy charakter zależności, to musimy zagęścić punkty pomiarowe. Na rysunku 8.4A przedstawiono nieprawidłowy

wykres z linią łamaną, zaś na rysunku 8.4B wykres prawidłowy z linią gładką poprowadzoną przy użyciu krzywika.



Rys. 8.4: Prowadzenie linii na wykresie: A) nieprawidłowe, punkty pomiarowe połączone linią łamaną; B) prawidłowe

Często zdarza się, że na wykresie przedstawiamy porównanie przewidywań teoretycznych z wynikami pomiarów. Wtedy albo nie zaznaczamy punktów obliczonych i prowadzimy linię, albo muszą się one *bardzo wyraźnie odróżniać od punktów doświadczalnych*. Zazwyczaj nie nanosi się punktów obliczonych, a prowadzi się linię. Tak samo postępujemy w przypadku dopasowywania funkcji do danych doświadczalnych metodą najmniejszych kwadratów.



Rys. 8.5: A) Naniesienie granic błędów sugeruje zgodność wyników z poprowadzoną krzywą; B) to samo, ale błędy mniejsze, odchylenie od krzywej wskazuje na istotne różnice

Nanoszenie granic błędów punktów pomiarowych jest *konieczne*, gdy występują odstępstwa od przewidywań teoretycznych. Przypadek taki przedstawiono na rysunku 8.5. Zaznaczanie błędów pomiarowych, gdy wszystkie pomiary są obarczone tym samym błędem, może być ograniczone do jednego punktu, w przeciwnym razie należy zaznaczać błędy dla wszystkich punktów. W przypadkach gdy nanoszenie błędów nie wnosi żadnych informacji, nie nanosimy ich, ponieważ zaciemniałyby jedynie wykres. Zasady sporządzania wykresów zostały bardzo przystępnie przedstawione w [7] rozdział 11 (patrz także [10, 13, 19]).

8.4. Graficzne oszacowanie błędu

Nie należy do rzadkości odczytywanie danych z wykresu. Dane odczytujemy z wykresu wówczas, gdy nie możemy dopasować do wyników pomiarów zależności funkcyjnej metodą najmniejszych kwadratów, omówioną w rozdziale 3.5. Sytuacja taka może wystąpić na przykład przy kalibracji przyrządów, czy też wtedy, gdy błąd żadnej z mierzonych wielkości fizycznych nie może być pominięty (por. rozdział 3.5 i 5.4). W dalszym ciągu tego rozdziału zajmiemy się przypadkiem, gdy zależność między mierzonymi wielkościami fizycznymi jest dana jedynie w postaci wykresu. Ponieważ wielkości mierzone bezpośrednio są obarczone błędem pomiarowym, wobec tego i wykreślona krzywa też będzie obciążona błędem. Powoduje to, że wartość odczytana z wykresu będzie również obciążona błędem. Błąd ten można oszacować graficznie. Graficzne oszacowanie błędu polega na oszacowaniu błędu wykreślonej krzywej, a następnie na ocenie błędu wielkości odczytanej z wykresu.

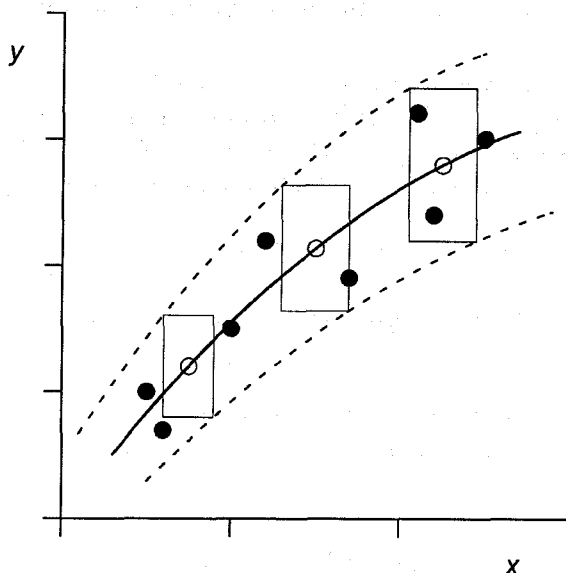
Przedstawione poniżej graficzne oszacowanie błędu jest niezależne od tego czy na wykresie zaznaczono błędy maksymalne, czy odchylenia standardowe; w opisie wykresu musi to być jednak wyraźnie zaznaczone. Należy jednak zaznaczyć, że na ogół graficznie oszacowujemy **błąd maksymalny**.

8.4.1. Oszacowanie błędu wykreślonej krzywej

Zasadę oszacowania błędu wykreślonej krzywej przedstawiono na rysunku 8.6. Poniżej omówimy sposób, w jaki dokonujemy tego oszacowania.

Założmy, że dokonaliśmy N pomiarów zależnych wielkości fizycznych x oraz y . W wyniku pomiarów otrzymaliśmy więc N par wartości (x_i, y_i) ($i = 1, 2, \dots, N$). Wartości (x_i, y_i) nanosimy na papier funkcyjny w układzie współrzędnych XY i w każdym punkcie zaznaczamy prostokąt błędów.

Na rysunku 8.6 punkty o współrzędnych (x_i, y_i) oznaczono czarnymi kropkami ($\bullet$); dla przejrzystości rysunku nie zaznaczono prostokątów błędów. Następnie, pamiętając o zasadach podanych w rozdziale 8.3, wykreślamy



Rys. 8.6: Przykład graficznego przedstawienia przedziału, w którym zawarta jest wykreślona krzywa. Linia ciągłą przedstawiono wykreśloną krzywą, a liniami przerywanymi granice przedziału. Pozostałe objaśnienia w tekście

przy użyciu krzywika krzywą (linia ciągła na rysunku 8.6). Po wykreśleniu krzywej *obieramy na niej kilka* (lub więcej w zależności od zakresu) punktów. Zazwyczaj pierwszy punkt obieramy na początku wykresu, a ostatni na końcu, pomiędzy nimi powinien być co najmniej jeden punkt. Na rysunku 8.6 obrane punkty oznaczono kółkiem (o). Wokół obranych punktów rysujemy prostokąty błędów i przez ich wierzchołki prowadzimy jedną krzywą powyżej, a drugą poniżej wykreślonej krzywej. Krzywe te na rysunku 8.6 wykreślone są linią przerywaną i czasem są nazywane **krzywymi błędu**.

Ponieważ błąd, jakim jest obarczony dowolny punkt leżący na wykreślonej krzywej $y = y(x)$ nie przekracza błędów, jakimi były obarczone punkty pomiarowe, to przyjmujemy, że rzeczywisty przebieg krzywej, w części przedstawionej na wykresie, jest zawarty w obszarze ograniczonym krzywymi błędu (linie przerywane na rysunku 8.6).

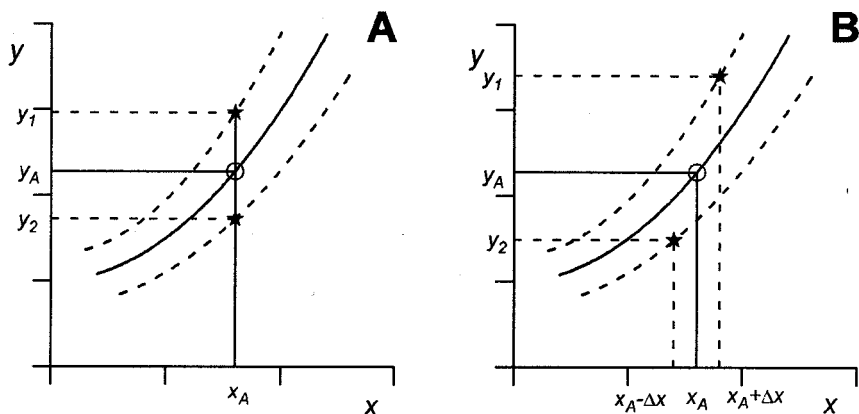
Pozostaje nam jeszcze wyjaśnić, jak znajdujemy prostokąt błędu w obranych punktach. Jeżeli prostokąty błędów wszystkich zmierzonych punktów (x_i, y_i) ($i = 1, 2, \dots, N$) są takie same, to prostokąty błędów w obranych punktach będą jednakowe i równe prostokątom błędów punktów zmierzonych. W przypadku gdy poszczególne pary (x_i, y_i) mają różne prostokąty błędów, to na końcach wykresu prostokąty błędu obranych punktów muszą być równe odpowiednim prostokątom błędu punktów skrajnych (x, y) ;

w punktach obranych między punktami skrajnymi, prostokątom błędów im najbliższych punktów.

Jeżeli znamy odchylenia standardowe (a nie błędy maksymalne), to zakładamy określony poziom ufności i korzystając z dystrybuanty rozkładu normalnego standaryzowanego bądź rozkładu *t-Studenta* znajdujemy przedziały ufności, a następnie wykreślamy odpowiadające im prostokąty błędów.

8.4.2. Oszacowanie błędu wartości odczytanej z wykresu

Przyjmijmy, że mamy zmierzoną zależność między wielkościami fizycznymi x i y . Wykres tej zależności jest przedstawiony na rysunkach 8.7 i 8.8. Aby nie komplikować rysunków, nie zostały zaznaczone na nich punkty



Rys. 8.7: Ocena błędu wartości odczytanej z wykresu zmiennej y . A) błąd zmiennej x jest do pominięcia; B) błędu zmiennej x nie można pominąć. Objasnienia w tekście

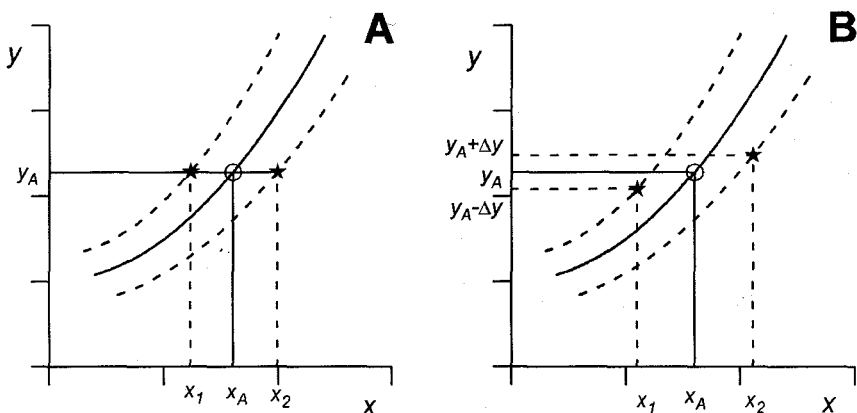
doświadczalne, a jedynie krzywa poprowadzona na ich podstawie (linia ciągła) przedstawiająca zależność $y(x)$ oraz krzywe błędów (linie przerywane). Mając ustaloną wartość x_A możemy z wykresu odczytać odpowiadającą jej wartość y_A , lub na odwrót, mając ustaloną wartość y_A odczytujemy x_A .

1) Odczytujemy z wykresu wartość y_A .

Zasadę oceny błędu wartości $y_A = y(x_A)$ przedstawiono na rysunku 8.7A. Z punktu x_A prowadzimy do przecięcia z górną krzywą błędów prostą równoległą do osi rzędnych (osi y). Rzędna punktu przecięcia z poprowadzoną krzywą (na rysunku 8.7A oznaczonego kółkiem $\circ$) jest szukaną wartością y_A . Rzędne punktów przecięcia tej prostej z krzywymi błędów, oznaczone na rysunku 8.7A gwiazdkami ($\star$), są granicami przedziału, w

którym jest zawarta szukana wartość y_A . Na rysunku 8.7A są to punkty y_1 i y_2 .

Bardziej złożona jest sytuacja, gdy wartość x_A jest obarczona błędem Δx . Aby odczytać z wykresu wartość y_A oraz przedział, w którym jest ona zawarta, postępujemy podobnie jak powyżej. Musimy zwrócić jednak uwagę na to, czy wykres przedstawia funkcję rosnącą, czy malejącą. Przyjmijmy, że rosnącą. Szukamy wtedy rzędnej y_2 punktu przecięcia prostej poprowadzonej z punktu $x_A - \Delta x$ i równoległej do osi y , z dolną krzywą błędu, a następnie rzędnej y_1 punktu przecięcia prostej, poprowadzonej z punktu $x_A + \Delta x$ i równoległej do osi y z górną krzywą błędu. Na rysunku 8.7B przedstawiono omawiany przypadek. Gdy funkcja jest malejąca, postępujemy odwrotnie. Wyznaczamy w ten sposób przedział, w którym jest zawarta szukana wartość y_A , w ogólności może on być niesymetryczny. W takim przypadku za granice **błędu odczytanej wartości** przyjmujemy większą z liczb $|y_1 - y_A|$ i $|y_2 - y_A|$.



Rys. 8.8: Ocena błędu wartości odczytanej z wykresu zmiennej x . A) błąd zmiennej y jest do pominięcia; B) błędu zmiennej y nie można pominąć

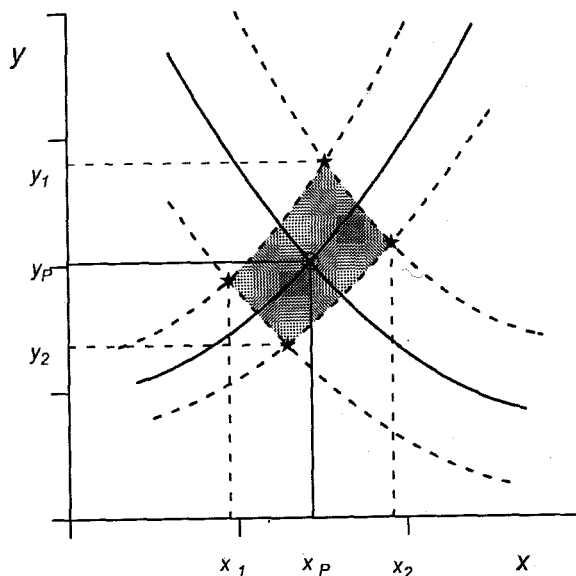
2) Odczytujemy z wykresu wartość x_A .

Postępowanie w tym przypadku jest takie samo, jak w poprzednim. Zasadę oceny błędu wartości x_A przedstawiono na rysunkach 8.8A i 8.8B.

3) Oszacowanie błędu współrzędnych punktu przecięcia się dwóch krzywych.

Nanosimy na papier funkcyjny wyniki pomiarów i przy użyciu krzywika wykreślamy obydwie krzywe. Odczytujemy z wykresu współrzędne (x_p, y_p) punktu przecięcia się krzywych. Dla każdej krzywej w pobliżu punktu przecięcia wykreślamy krzywe błędu. Przedstawia to rysunek 8.9, na którym punkt przecięcia oznaczono kółkiem ($\circ$), a krzywe błędu liniami przerywanymi. Jak widać z rysunku 8.9, w pobliżu punktu przecięcia obszary ogra-

niczone krzywymi błędów częściowo się pokrywają. Na rysunku 8.9 jest to obszar zakreskowany. Punkty skrajne zakreskowanego obszaru, oznaczone na rysunku 8.9 gwiazdkami (*), wyznaczają granice błędów współrzędnych punktu przecięcia.



Rys. 8.9: Ocena błędów współrzędnych punktu przecięcia krzywych. Objasnienia w tekście

9. Ocena błędu w przypadku, gdy błędy przypadkowe i systematyczne są porównywalne

Przypadek ten występuje zawsze tam, gdzie dokładność przyrządu pomiarowego jest porównywalna z rozrzutem wyników pomiarów, a w przypadku wielkości złożonej również, gdy któraś z wielkości mierzonych bezpośrednio została zmierzona jeden raz. Poniżej przedstawimy sposoby oszacowania błędu, gdy zachodzą wyżej wymienione przypadki oraz gdy błąd systematyczny jest spowodowany dokładnością przyrządu pomiarowego.

9.1. Ocena błędu maksymalnego

9.1.1. Pomiary bezpośrednie

Założmy, że wykonaliśmy N pomiarów bezpośrednich wielkości fizycznej x używając miernika o dokładności ε i otrzymaliśmy zbiór wartości $x_1, x_2, \dots, x_N$. Założmy ponadto, że rozrzut wyników pomiarów był porównywalny z dokładnością miernika, tzn. ε nie można pominąć w porównaniu z S_x ani $S_{\bar{x}}$ oraz odwrotnie, S_x ani $S_{\bar{x}}$ nie można pominąć w porównaniu z ε . Na podstawie serii pomiarów znajdujemy średnią arytmetyczną $\bar{x}$, która oprócz niepewności spowodowanej błędami przypadkowymi jest również obciążona błędem ε . Aby dokonać oszacowania **błędu maksymalnego**, należy rozpatrzyć dwa przypadki, a mianowicie:

- gdy liczba pomiarów jest tak duża, że do oceny błędów przypadkowych możemy stosować rozkład normalny;
- gdy liczba pomiarów jest mała i należy stosować rozkład *t-Studenta*.

W pierwszym przypadku, jak powiedziano w rozdziale 5.2.2, jako błąd maksymalny średniej arytmetycznej można przyjąć $3S_{\bar{x}}$. Wówczas błąd maksymalny wielkości mierzonej będzie

$$\Delta x = \varepsilon + 3S_{\bar{x}}, \quad (9.1.1)$$

a wynik pomiaru zapiszemy:

$$x = \bar{x} \pm \Delta x, \quad (9.1.2)$$

przestrzegając zasad podanych w rozdziale 1.1.4. W drugim przypadku, zgodnie z tym co powiedziano w rozdziale 6.2, jako **błąd maksymalny** średniej arytmetycznej można przyjąć: $t_{0.997}S_{\bar{x}}$, gdzie współczynnik $t_{0.997}$ jest zależny od liczby stopni swobody $k = N - 1$. W tym przypadku błąd maksymalny wielkości mierzonej będzie:

$$\Delta x = \varepsilon + t_{0.997}S_{\bar{x}}. \quad (9.1.3)$$

9.1.2. Pomiary pośrednie

W przypadku pomiarów pośrednich korzystamy z reguły przenoszenia błędów maksymalnych podanej w rozdziale 2.2, wzór (2.2.10). Jednak jeżeli błędy przypadkowe i systematyczne są porównywalne lub gdy jedno z wielkości mierzonych bezpośrednio są obarczone tylko błędami przypadkowymi, a inne systematycznymi, występuje różnorodność w oszacowaniu błędów maksymalnych, którą musimy wziąć pod uwagę.

Wyjaśnimy to na przykładzie. Niech wielkość złożona z będzie funkcją wielkości a, b, c, d, e , tzn. $z = f(a, b, c, d, e)$. Przyjmijmy, że:

- wielkość fizyczną a zmierzylismy jeden raz i otrzymaliśmy wartość a_m z dokładnością Δa ;
- wielkość b zmierzylismy N razy i okazało się, że błędy przypadkowe są na tyle duże, iż błąd systematyczny można pominąć; z pomiarów obliczyliśmy $\bar{b}$ i $S_{\bar{b}}$; za błąd maksymalny Δb możemy więc przyjąć $3S_{\bar{b}}$ lub $t_{0.997}S_{\bar{b}}$ w zależności od liczby pomiarów N ;
- wielkość c zmierzylismy M razy i błędy przypadkowe były porównywalne z błędem systematycznym; z pomiarów obliczyliśmy $\bar{c}$ i $S_{\bar{c}}$; błąd maksymalny Δc obliczamy ze wzoru (9.1.1) lub (9.1.3) w zależności od liczby pomiarów M ;
- wartość wielkości d i jej błąd maksymalny wyznaczyliśmy graficznie i otrzymaliśmy odpowiednio d_m i Δd ;
- z dopasowania metodą najmniejszych kwadratów wyznaczyliśmy wartość wielkości e , otrzymaliśmy wartość e_k i odchylenie standardowe S_e ; wartość $3S_e$, zgodnie z tym co powiedziano w rozdziale 9.1, możemy przyjąć za błąd maksymalny Δe .

Błąd maksymalny wielkości złożonej $z = f(a, b, c, d, e)$ będzie więc można przedstawić w postaci:

$$\begin{aligned} \Delta z = & \left| \left(\frac{\partial f}{\partial a} \right)_{a_m, \bar{b}, \bar{c}, d_m, e_k} \Delta a \right| + \left| \left(\frac{\partial f}{\partial b} \right)_{a_m, \bar{b}, \bar{c}, d_m, e_k} \Delta b \right| \\ & + \left| \left(\frac{\partial f}{\partial c} \right)_{a_m, \bar{b}, \bar{c}, d_m, e_k} \Delta c \right| + \left| \left(\frac{\partial f}{\partial d} \right)_{a_m, \bar{b}, \bar{c}, d_m, e_k} \Delta d \right| + \left| \left(\frac{\partial f}{\partial e} \right)_{a_m, \bar{b}, \bar{c}, d_m, e_k} \Delta e \right|. \end{aligned} \quad (9.1.4)$$

9.2. Statystyczna ocena błędu

9.2.1. Pomiary bezpośrednie

W przypadku serii pomiarów, jeśli rozrzut wyników jest porównywalny z dokładnością przyrządu pomiarowego, a jej liczebność N jest na tyle duża, że w nieobecności błędów systematycznych możemy do oceny błędu stosować rozkład Gaussa, to do oceny poziomu i przedziału ufności można stosować metody statystyczne [22].

Na początku założymy, że jedynym źródłem błędu jest **dokładność przyrządu pomiarowego** ε i jest to błąd maksymalny. Statystycznie przedział $[x_0 - \varepsilon, x_0 + \varepsilon]$ wyznaczony przez błąd maksymalny można interpretować jako przedział, w którym z prawdopodobieństwem $P = 1$ są zawarte wszystkie wyniki pomiarów. Ponieważ nie ma żadnego powodu by uważać, że jakaś wartość z tego przedziału może wystąpić częściej niż inna, możemy przyjąć, że funkcja rozkładu ma jednakową wartość dla wszystkich wartości x otrzymanych z pomiarów. Wynika z tego, że zawsze zachodzi nierówność $x_0 - \varepsilon \leq x \leq x_0 + \varepsilon$. Nierówność tą możemy zapisać $-\varepsilon \leq x - x_0 \leq \varepsilon$. Ponieważ $x - x_0 = \delta'$ jest błędem bezwzględnym, nierówność przyjmie postać $-\varepsilon \leq \delta' \leq \varepsilon$. Oznacza to, że błędy bezwzględne wyników pomiarów będą zawarte w przedziale określonym przez błąd maksymalny oraz ich funkcja rozkładu w tym przedziale będzie miała stałą wartość, a poza tym przedziałem będzie równa zero. Rozkład ten ma więc postać

$$\Theta(\delta') = \begin{cases} c & \text{dla } -\varepsilon \leq \delta' \leq \varepsilon \\ 0 & \text{dla pozostałych wartości } \delta'. \end{cases} \quad (9.2.1)$$

Warunek normalizacji (4.2.6) w tym przypadku przyjmie postać

$$\int_{-\infty}^{\infty} \Theta(\delta') d\delta' = \int_{-\varepsilon}^{\varepsilon} c d\delta' = 1. \quad (9.2.2)$$

Korzystając z warunku normalizacji znajdujemy wartość c , wynosi ona

$$c = \frac{1}{2\varepsilon}. \quad (9.2.3)$$

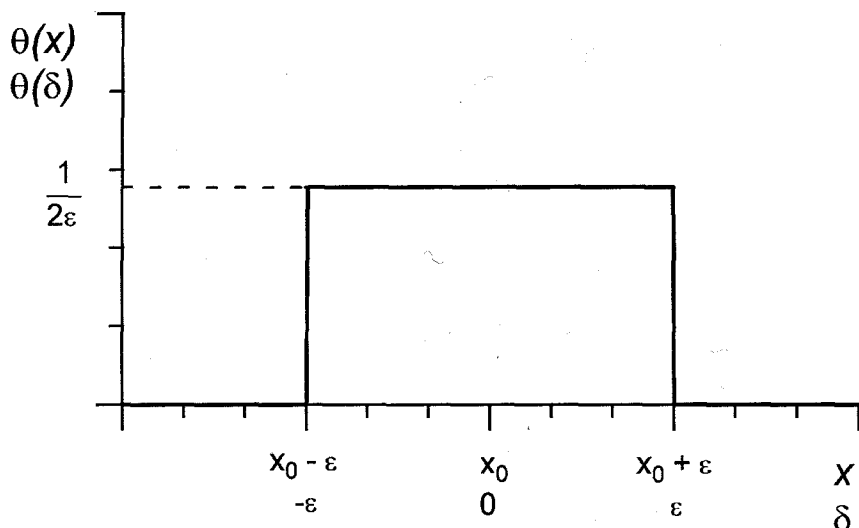
Wobec tego rozkład (9.2.1) zapiszemy:

$$\Theta(\delta') = \begin{cases} \frac{1}{2\varepsilon} & \text{gdy } -\varepsilon \leq \delta' \leq \varepsilon \\ 0 & \text{dla pozostałych wartości } \delta'. \end{cases}$$

(9.2.4)

Rozkład ten nosi nazwę **rozkładu prostokątnego** lub **jednostajnego**. Na rysunku 9.1 przedstawiono przykład rozkładu prostokątnego dla wyników

pomiarów $\Theta(x)$ i dla błędów bezwzględnych $\Theta(\delta)$. Jak widać z rysunku 9.1 oba te rozkłady są jednakowe, różnice występują jedynie na osi odciętych, ponieważ $\delta = x - x_0$.



Rys. 9.1: Przykładowy wykres rozkładu prostokątnego dla wyników pomiarów $\Theta(x)$ i dla błędów bezwzględnych $\Theta(\delta)$

Jeżeli wyniki pomiarów są obarczone tylko **błędami przypadkowymi**, to, zgodnie z tym co powiedziano w rozdziale 5.1, rozkład błędów przypadkowych będzie **rozkładem Gaussa**, tj.

$$\varphi(\delta'') = \frac{1}{\sigma\sqrt{2\pi}} e^{-\frac{(\delta'')^2}{2\sigma^2}}. \quad (9.2.5)$$

W przypadku gdy wynik pomiaru będzie obarczony błędem związanym z dokładnością przyrządu i błędem przypadkowym, całkowity błąd δ jest sumą obydwóch wymienionych błędów, tj. $\delta = \delta' + \delta''$. Ponieważ każdy z tych błędów jest zmienną losową i są one od siebie niezależne, to jak pokazano w twierdzeniu 1 Dodatku G funkcja rozkładu sumy tych dwóch błędów jest **plotem** rozkładu $\Theta(\delta')$ i $\varphi(\delta'')$, tj.

$$\xi(\delta) = \int_{-\infty}^{\infty} \Theta(\delta') \varphi(\delta - \delta') d\delta'. \quad (9.2.6)$$

Funkcja $\Theta(\delta')$ jest różna od zera, gdy $-\varepsilon \leq \delta' \leq \varepsilon$, a więc wzór (9.2.6)

przyjmie postać:

$$\xi(\delta) = \frac{1}{2\varepsilon} \int_{-\varepsilon}^{\varepsilon} \varphi(\delta - \delta') d\delta'. \quad (9.2.7)$$

Dokonajmy teraz zamiany zmiennych w całce (9.2.7) wprowadzając zmienną $\delta'' = \delta - \delta'$, wtedy $d\delta'' = -d\delta'$ i całka ta przyjmie postać:

$$\xi(\delta) = -\frac{1}{2\varepsilon} \int_{\delta+\varepsilon}^{\delta-\varepsilon} \varphi(\delta'') d\delta'' = \frac{1}{2\varepsilon} \int_{\delta-\varepsilon}^{\delta+\varepsilon} \varphi(\delta'') d\delta''. \quad (9.2.8)$$

Korzystając z własności całek oznaczonych wyrażenie to możemy zapisać jako

$$\xi(\delta) = \frac{1}{2\varepsilon} \left[\int_{-\infty}^{\delta+\varepsilon} \varphi(\delta'') d\delta'' - \int_{-\infty}^{\delta-\varepsilon} \varphi(\delta'') d\delta'' \right]. \quad (9.2.9)$$

Całki występujące w nawiasach są dystrybuantami Φ rozkładu normalnego (por. rozdział 5.1), wobec tego funkcja $\xi(\delta)$ przyjmuje postać

$$\xi(\delta) = \frac{1}{2\varepsilon} [\Phi(\delta + \varepsilon) - \Phi(\delta - \varepsilon)]. \quad (9.2.10)$$

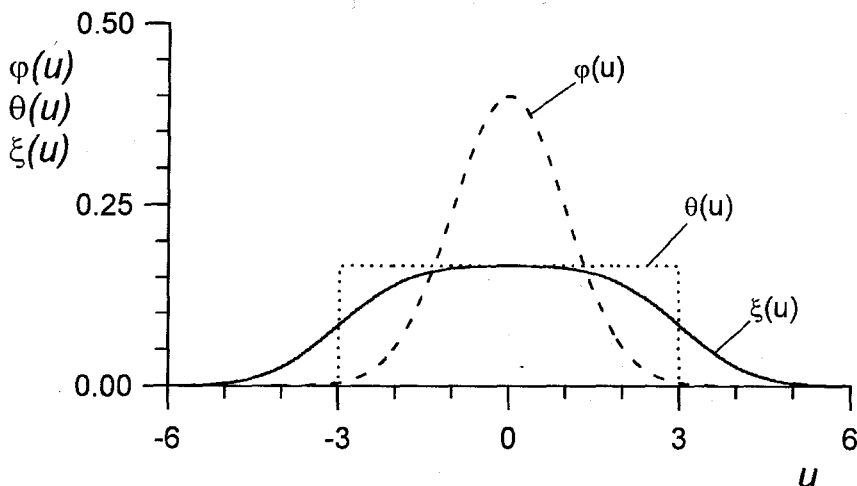
Tak więc **funkcja rozkładu błędów** $\xi(\delta)$, w przypadku gdy błąd maksymalny i błędy przypadkowe są porównywalne, wyraża się poprzez dystrybuanty rozkładu normalnego zależne od $\delta + \varepsilon$ i $\delta - \varepsilon$, czyli prawdopodobieństwo $dP(\delta)$ tego, że błąd pojedynczego pomiaru będzie zawarty między δ a $\delta + d\delta$ wynosi $dP(\delta) = \xi(\delta)d\delta$.

Korzystanie przy ocenie poziomu i przedziału ufności z funkcji rozkładu danej wzorem (9.2.10) i jej dystrybuanty jest niewygodne, ponieważ dla każdej wartości σ^2 i ε musimy ją tablicować. Z tego powodu przedstawimy funkcję rozkładu (9.2.10) korzystając z wprowadzonej w rozdziale 5.2 zmiennej losowej standaryzowanej. Zmienna standaryzowana $u = \delta/\sigma = (x - x_0)/\sigma$. Ponieważ $dP(\delta) = dP(u)$ oraz $d\delta = \sigma du$, to zachodzą związki $\xi(\delta)d\delta = \xi(\sigma u)\sigma du = \xi_u(u)du$ (por. rozdział 5.2.1), a górna i dolna granica całkowania we wzorze (9.2.9) wyniosą odpowiednio $u + \varepsilon/\sigma$ i $u - \varepsilon/\sigma$. Funkcja rozkładu (9.2.10) przyjmie więc postać:

$$\xi_u(u; \varepsilon/\sigma) = \frac{1}{2(\varepsilon/\sigma)} [\Phi_u(u + \varepsilon/\sigma) - \Phi_u(u - \varepsilon/\sigma)]. \quad (9.2.11)$$

Funkcja Φ_u jest dystrybuantą rozkładu normalnego standaryzowanego daną wzorem (5.2.13).

Na rysunku 9.2 przedstawiono przykładowy wykres funkcji rozkładu prostokątnego $\Theta(u)$, normalnego $\varphi(u)$ oraz ich splotu $\xi(u)$. Obliczenia wykonano przy założeniu $\sigma = 1$ oraz $\varepsilon = 3$, tzn. że odchylenie standardowe stanowi zaledwie 1/3 błędu maksymalnego. Szerokość splotu $\xi(u)$ jest większa niż którejkolwiek z funkcji $\Theta(u)$ i $\varphi(u)$.



Rys. 9.2: Przykład rozkładu prostokątnego $\Theta(u)$, rozkładu normalnego $\varphi(u)$ oraz ich splotu $\xi(u)$

Rozważania przedstawione powyżej dotyczą funkcji rozkładu pojedynczego pomiaru, jak i średniej arytmetycznej, ponieważ zmienna standaryzowana u jest zdefiniowana dla obu tych wielkości (por. rozdział 5.2.1). Zajmiemy się dalej rozkładem średniej arytmetycznej $\bar{x}$. W przypadku rozkładu średniej arytmetycznej, który jest również normalny (por. rozdział 5.1.4), zmienna standaryzowana u jest dana wzorem (5.2.9) i we wzorze (9.2.11) σ musimy zastąpić przez $\sigma_{\bar{x}}$.

Przejdziemy teraz do oceny poziomu i przedziału ufności dla wartości rzeczywistej x_0 . Zgodnie z tym co powiedziano w rozdziale 5.2, **zmienna standaryzowana** $u = (\bar{x} - x_0)/\sigma_{\bar{x}}$. Prawdopodobieństwo α (poziom ufności) tego, że zmienna standaryzowana u będzie zawarta w przedziale

$$-d_\alpha \leq u \leq d_\alpha \quad (9.2.12)$$

(gdzie $d_\alpha > 0$) jest

$$\alpha = P(-d_\alpha \leq u \leq d_\alpha) = \int_{-d_\alpha}^{d_\alpha} \xi_u(u, \varepsilon/\sigma_{\bar{x}}) du, \quad (9.2.13)$$

i zależy od stosunku $\varepsilon/\sigma_{\bar{x}}$. Rozwiązując nierówność (9.2.12) znajdujemy przedział ufności:

$$\boxed{\bar{x} - d_{\alpha}\sigma_{\bar{x}} \leq x_0 \leq \bar{x} + d_{\alpha}\sigma_{\bar{x}},} \quad (9.2.14)$$

przy **poziomie ufności** α . Ponieważ poziom ufności α dany wzorem (9.2.13) jest zależny od stosunku $\varepsilon/\sigma_{\bar{x}}$, to również współczynnik d_{α} wyznaczony dla ustalonego α jest zależny od stosunku $\varepsilon/\sigma_{\bar{x}}$.

Ponieważ prawdopodobieństwo α i współczynnik d_{α} są zależne od stosunku $\varepsilon/\sigma_{\bar{x}}$ (który jest parametrem rozkładu) nie tablicuje się całki (9.2.13), ale podaje się tablice współczynnika d_{α} dla różnych wartości α i stosunku $\varepsilon/\sigma_{\bar{x}}$. Tablica taka została zamieszczona w Dodatku H (tablica VI). Sposób posługiwania się tablicą VI jest taki sam jak tablicą współczynnika t_{α} i został omówiony w rozdziale 6.

Oceny poziomu i przedziału ufności można również dokonać w sposób przybliżony stosując **rozkład Gaussa**. Wariancja σ_{Θ}^2 rozkładu prostokątnego (9.2.4) zgodnie ze wzorem (4.2.12) wynosi

$$\sigma_{\Theta}^2 = \frac{1}{2\varepsilon} \int_{x_0-\varepsilon}^{x_0+\varepsilon} (x - x_0)^2 dx = \frac{\varepsilon^3}{3}. \quad (9.2.15)$$

Znając wariancję σ_{Θ}^2 rozkładu prostokątnego i wariancję $\sigma_{\bar{x}}^2$ rozkładu normalnego średniej arytmetycznej obliczamy *wariancję efektywną*

$$\tilde{\sigma}^2 = \sigma_{\bar{x}}^2 + \frac{\varepsilon^2}{3}. \quad (9.2.16)$$

Wariancję efektywną daną wzorem (9.2.16), traktujemy dalej jako wariancję rozkładu Gaussa. Oszacowania poziomu i przedziału ufności dokonujemy tak jak podano w rozdziale 5.2.2.

Czytelnikowi pozostawiamy sprawdzenie na ile oszacowanie poziomu i przedziału ufności przy użyciu rozkładu normalnego z wariancją efektywną $\tilde{\sigma}^2$ jest zgodne z oszacowaniem dokonany przy użyciu tablic współczynnika d_{α} .

Na zakończenie należy zaznaczyć, że wielkości $\sigma_{\bar{x}}^2$ nie znamy, ponieważ jest to wariancja dla populacji generalnej i w obliczeniach zastępujemy ją jej najlepszym przybliżeniem — to jest wariancją średniej serii $S_{\bar{x}}^2$. Wtedy wzór (9.2.16) przyjmie postać:

$$\tilde{\sigma}^2 = S_{\bar{x}}^2 + \frac{\varepsilon^2}{3}. \quad (9.2.17)$$

Przykład

W pewnym doświadczeniu mierzono czas spadania kulki stalowej w rurze wypełnionej olejem. Pomiary wykonano na tym samym odcinku rury i otrzymano:

$$t[s] = 3.6, 3.3, 3.7, 3.9, 3.2, 3.5, 3.8, 3.4, 3.7, 3.9, 3.3.$$

Dokładność pomiaru czasu: $\varepsilon = 0.1$ s. Znaleźć przedział ufności dla poziomu ufności $d = 0.9545$ (odpowiada to w przypadku zastosowania rozkładu normalnego przedziałowi $\bar{t} - 2S_{\bar{t}} \leq t_0 \leq \bar{t} + 2S_{\bar{t}}$) z uwzględnieniem dokładności pomiaru czasu.

Obliczamy: $\bar{t} = 3.573 \text{ s} \approx 3.57 \text{ s}$, $S_t^2 = 0.0622 \text{ s}^2$, $S_t = 0.24936 \text{ s} \approx 0.25 \text{ s}$, $S_{\bar{t}}^2 = 0.005653 \text{ s}^2$, $S_{\bar{t}} = 0.0752 \text{ s}$. Stosunek $\varepsilon/S_{\bar{t}} = 1.33$, a więc należy uwzględnić dokładność przyrządu pomiarowego. Musimy więc użyć dystrybuanty (9.2.11) rozkładu $\xi_u(u, \varepsilon/\sigma_{\bar{t}})$. Korzystając z tablicy VI w Dodatku H znajdujemy $d_\alpha = d_{0.9545} = 2.49$. Ze wzoru (9.2.14) otrzymujemy $(3.57 - 2.49 \cdot 0.0752) \text{ s} \leq t_0 \leq (3.57 + 2.49 \cdot 0.0752) \text{ s}$, czyli $3.38 \text{ s} \leq t_0 \leq 3.75 \text{ s}$ z prawdopodobieństwem 0.9545.

Obliczmy teraz to samo przy użyciu rozkładu Gaussa. Przy założeniu, że występują tylko błędy przypadkowe otrzymamy: $3.42 \text{ s} \leq t_0 \leq 3.72 \text{ s}$. Jak widać oszacowanie przedziału ufności przy użyciu rozkładu Gaussa jest bardziej optymistyczne.

Skorzystajmy teraz z przybliżonej metody. Ze wzoru (9.2.17) obliczamy efektywną wariancję tej serii pomiarowej. Wynosi ona: $\tilde{S}^2 = 0.008986 \text{ s}^2$, a efektywne odchylenie standardowe tej serii pomiarowej $\tilde{S} = 0.0948 \text{ s}$. Dalej postępujemy tak jak przy użyciu rozkładu Gaussa i znajdujemy: $3.38 \text{ s} \leq t_0 \leq 3.76 \text{ s}$. Otrzymaliśmy wynik bardzo zbliżony do otrzymanego przy użyciu dystrybuanty (9.2.13) funkcji $\xi_u(u; \varepsilon/\sigma_{\bar{t}})$.

9.2.2. Pomiary pośrednie

W przypadku pomiarów pośrednich oszacowanie poziomu i przedziału ufności, gdy występują jednocześnie błędy przypadkowe i systematyczne wielkości mierzonych bezpośrednio, jest znacznie bardziej skomplikowane. Nawet wtedy gdy można stosować rozwinięcie (5.1.31), rozkład każdej z wielkości mierzonych bezpośrednio jest dany wzorem (9.2.10) oraz wielkości te są niezależne, oszacowanie poziomu i przedziału ufności wymaga obliczania wielokrotnych spłotów rozkładów (9.2.10). Z tego powodu najwygodniej jest korzystać z metody przybliżonej.

Niech wielkość złożona z będzie funkcją r niezależnych wielkości mierzonych bezpośrednio $x_1, x_2, \dots, x_r$, tzn. $z = f(x_1, x_2, \dots, x_r)$. Oznaczmy ich błędy maksymalne przez ε_j ($j = 1, 2, \dots, r$), odchylenia standardowe średnich arytmetycznych każdej z wielkości x_j niech będą $\sigma_{\bar{x}_j}$ ($j = 1, 2, \dots, r$).

Wtedy efektywne odchylenia standardowe $\tilde{\sigma}_j$ mierzonych bezpośrednio wielkości fizycznych x_j będą dane wzorem (9.2.16), tj.

$$\tilde{\sigma}_j = \sqrt{\sigma_{\bar{x}_j}^2 + \frac{1}{3}\varepsilon_j^2}. \quad (9.2.18)$$

W tym przypadku *efektywną wariancję* $\tilde{\sigma}_{\bar{z}}^2$ znajdujemy korzystając z reguły przenoszenia błędów, gdy wielkości mierzone bezpośrednio są niezależne (3.4.8), wynosi ona

$$\tilde{\sigma}_{\bar{z}}^2 = \sum_{j=1}^r \left(\frac{\partial f}{\partial x_j} \right)_{\bar{x}_1, \dots, \bar{x}_r}^2 (\sigma_{\bar{x}_j}^2 + \frac{1}{3}\varepsilon_j^2). \quad (9.2.19)$$

Oczywiście w obliczeniach $\sigma_{\bar{x}_j}^2$ zastępujemy w tym wyrażeniu przez najlepsze przybliżenie, tj. $S_{\bar{x}_j}^2$. W celu znalezienia przedziału i poziomu ufności przy użyciu rozkładu normalnego postępujemy dalej tak jak zostało to przedstawione w rozdziale 5.2.

A. Prawdopodobieństwo

Jednym z najważniejszych pojęć rachunku prawdopodobieństwa jest zdarzenie losowe. **Zdarzeniem losowym** albo krótko **zdarzeniem** będziemy nazywali zjawisko, które może zajść albo nie. Takie określenie zdarzenia nie jest definicją. Pojęcie zdarzenia wyjaśnimy na przykładach. Zdarzeniami są na przykład: wynik pomiaru, otrzymanie reszki w rzucie monetą, otrzymanie w rzucie kostką do gry określonej liczby oczek (np. dwóch), wyciągnięcie z talii kart do gry karty określonego koloru (np. kierowego) itp.

Zdarzenia będziemy oznaczali literami $A, B, C, \dots$

Jeżeli zdarzenie nie może nastąpić, to nazywamy je **zdarzeniem niemożliwym**. Na przykład zdarzeniem niemożliwym będzie wyrzucenie 8 oczek kostką do gry, wyciągnięcie z talii kart pięciu króli itp.

Jeżeli zdarzenie zajdzie na pewno, to nazywamy je **zdarzeniem pewnym**. Na przykład w rzucie kostką do gry zdarzeniem pewnym będzie wyrzucenie nie więcej niż sześciu oczek.

Jeżeli A jest dowolnym zdarzeniem, to przez **zdarzenie przeciwne** będziemy rozumieli zdarzenie polegające na tym, że A nie nastąpi. Będziemy je oznaczali $\bar{A}$. Na przykład jeżeli zdarzenie A polega na wylosowaniu karty koloru kierowego lub karowego, to $\bar{A}$ polega na wylosowaniu karty koloru pikowego lub treflowego.

Zdarzenia A i B nazywamy **wykluczającymi się** lub **rozłącznymi**, jeśli zajście zdarzenia A wyklucza możliwość wystąpienia zdarzenia B . Na przykład podczas rzutu kostką do gry wyrzucenie trzech oczek (zdarzenie A) wyklucza możliwość wyrzucenia pozostałych (zdarzenie B). Z definicji tej wynika, że zdarzenia A i $\bar{A}$ zawsze się wykluczają.

Jeżeli jest dany zbiór zdarzeń rozłącznych $A_1, A_2, \dots, A_n$, to zdarzenia te nazywamy **elementarnymi**, jeżeli w wyniku każdej próby zajdzie jedno z nich. O zdarzeniach tych mówi się również, że tworzą pełną grupę zdarzeń. Na przykład gdy rzucamy kostką do gry, to wyrzucenie każdej z możliwych liczby oczek jest zdarzeniem elementarnym. Zdarzeń elementarnych jest w tym przypadku sześć i stanowi je wyrzucenie 1, 2, 3, 4, 5 i 6 oczek; tworzą one pełną grupę zdarzeń.

Przejdziemy teraz do podania definicji prawdopodobieństwa. Istnieje kilka definicji prawdopodobieństwa. Podamy trzy spośród nich:

Definicja aksjomatyczna prawdopodobieństwa

Niech będzie dany zbiór zdarzeń elementarnych $A_1, A_2, \dots, A_n$.

- a) Każdemu zdarzeniu A_i ze zbioru zdarzeń elementarnych przyporządkowujemy nieujemną liczbę rzeczywistą nazywaną jego prawdopodobieństwem

$$P(A_i) \geq 0.$$

- b) Prawdopodobieństwo tego, że zajdzie którekolwiek ze wszystkich możliwych zdarzeń elementarnych $A_1, A_2, \dots, A_n$

$$P(E) = 1.$$

- c) Jeżeli zdarzenia A_i oraz A_j są rozłączne, to prawdopodobieństwo zajścia zdarzenia A_i lub A_j

$$P(A_i + A_j) = P(A_i) + P(A_j).$$

Sformułowanie to określa, jakie własności musi mieć wielkość nazywana prawdopodobieństwem, nie określa jednak, w jaki sposób wyznaczać wartość prawdopodobieństwa. Dlatego w zastosowaniach praktycznych korzystamy z tzw. definicji operacyjnych, które, choć z matematycznego punktu widzenia są niezadowalające, pozwalają jednak na znalezienie wartości prawdopodobieństwa. Należy jednak traktować je jako dobre przybliżenia. Podamy dwie definicje operacyjne: klasyczną i częstościową nazywaną również statystyczną.

Definicja klasyczna prawdopodobieństwa

Niech będzie dany skończony zbiór zdarzeń $A_1, A_2, \dots, A_n$. O zdarzeniach tych zakładamy, że

- 1) wykluczają się parami, to znaczy dla każdej pary zdarzeń A_i oraz A_j ($i, j = 1, 2, \dots, n; i \neq j$) zajście zdarzenia A_i wyklucza możliwość zajścia zdarzenia A_j ,
- 2) tworzą zbiór zdarzeń elementarnych (pełną grupę zdarzeń), tzn. jedno z nich musi nastąpić,
- 3) są one jednakowo możliwe.

Jeżeli mamy zbiór zdarzeń elementarnych $A_1, A_2, \dots, A_n$ i jeżeli przy pewnych zdarzeniach należących do tego zbioru zachodzi określone zdarzenie A , to mówimy, że zdarzenia te *sprzyjają* zdarzeniu A .

Przykład

Rzucamy kostką do gry. W tym przypadku mamy, jak pokazano wyżej, 6 zdarzeń elementarnych. Jeżeli zdarzenie A polega na wyrzuceniu pięciu oczek, to liczba zdarzeń sprzyjających wynosi 1. Jeżeli zdarzenie A polega na wyrzuceniu liczby oczek mniejszej od 4, to zdarzeń sprzyjających mamy trzy, a mianowicie: wyrzucenie jednego, dwóch lub trzech oczek.

Niech spośród $A_1, A_2, \dots, A_n$ zdarzeń k sprzyja wystąpieniu zdarzenia A . Prawdopodobieństwem $P(A)$ zdarzenia A będziemy nazywali stosunek liczby **zdarzeń sprzyjających** k do liczby wszystkich zdarzeń n , tj.

$$P(A) = \frac{k}{n}. \quad (\text{A.1})$$

Klasyczną definicję prawdopodobieństwa i założenia stanowiące jej podstawę wyjaśnimy na przykładzie.

W urnie znajduje się k kul białych i l kul czarnych. Całkowita liczba kul wynosi $k + l = n$. Wyciągamy kule z urny. Każde wyciągnięcie kuli z urny jest osobnym zdarzeniem. Wyciągnięcie jednej kuli *wyklucza* wyciągnięcie jakiegokolwiek innej, a więc zdarzenia te są zdarzeniami wykluczającymi się. Zdarzenia polegające na wyciąganiu każdej kuli z osobna są zdarzeniami elementarnymi i tworzą pełną grupę zdarzeń. Jeżeli kule są identyczne (różnią się jedynie kolorem) i dokładnie wymieszane, to nie istnieje żadna przyczyna, która powodowałaby, że możliwość wyciągnięcia jednej kuli byłaby większa niż drugiej. Obliczmy teraz prawdopodobieństwo wylosowania kuli białej $P(B)$ i kuli czarnej $P(C)$. Zgodnie z wzorem (A.1) $P(B) = \frac{k}{n}$, zaś $P(C) = \frac{l}{n}$. Zwróćmy uwagę na fakt, że zdarzenie polegające na wyciągnięciu z urny dowolnej kuli, tj. białej lub czarnej, jest zdarzeniem pewnym. Obliczmy więc prawdopodobieństwo tego zdarzenia. Liczba zdarzeń sprzyjających będzie w tym przypadku $k + l$, a więc

$$P(B + C) = \frac{k + l}{n},$$

ale $k + l = n$ (patrz wyżej), a więc $P(B + C) = 1$. Wynika z tego ogólny wniosek, że prawdopodobieństwo zdarzenia pewnego jest równe 1. Jeżeli w urnie są tylko kule białe i czarne, to wyciągnięcie z urny innej kuli niż biała lub czarna będzie zdarzeniem niemożliwym. W tym przypadku liczba zdarzeń sprzyjających będzie oczywiście równa zero. Ze wzoru (A.1) wynika, że prawdopodobieństwo zdarzenia niemożliwego jest równe zero. Z tego co powiedziano wyżej wynika, że prawdopodobieństwo $P(A)$ zajścia zdarzenia A spełnia nierówność

$$0 \leq P(A) \leq 1. \quad (\text{A.2})$$

Definicja częstościowa (statystyczna) prawdopodobieństwa

Klasyczna definicja prawdopodobieństwa ma szereg wad. Przede wszystkim wymaga, aby zbiór zdarzeń elementarnych zawierał skończoną liczbę elementów (zdarzeń) n , ponieważ w przypadku gdy $n \rightarrow \infty$ definicja ta traci sens. Aby stosować definicję klasyczną musimy znać liczbę zdarzeń sprzyjających k . Warunki te można spełnić w szeregu przypadkach, np. gier hazardowych. Nie można ich jednak spełnić w przypadku badania zjawisk społecznych czy też przyrodniczych. Dlatego w zastosowaniach rachunku prawdopodobieństwa do badania zjawisk przyrodniczych i społecznych posługujemy się **definicją częstościową** nazywaną również **definicją statystyczną**.

Rozpatrzmy następujący przykład. Wiemy, że w urnie mamy jednakowe kule różniące się tylko kolorem — białe i czarne. Nie wiemy, ile wynosi całkowita liczba kul, ani ile jest kul białych, a ile czarnych. Chcemy jednak znaleźć prawdopodobieństwo $P(B)$ wyciągnięcia np. kuli białej. Prawdopodobieństwo to możemy określić eksperymentalnie. W tym celu wyciągamy n razy kulę z urny. Po każdorazowym wyciągnięciu kulę wkładamy do urny i mieszamy kule w urnie. Jeżeli n razy dokonaliśmy wyciągnięcia kuli z urny (mówimy, że dokonaliśmy losowania) i w k przypadkach wylosowaliśmy kulę białą, to stosunek $k/n = f$ nazywamy **częstością**. Okazuje się, że częstość f w miarę wzrostu liczby losowań n będzie zmierzała do pewnej stałej wartości. Możemy więc określić granicę, do której będzie zmierzała częstość $f = k/n$, gdy $n \rightarrow \infty$. Granicę tę

$$P(B) = \lim_{n \rightarrow \infty} \frac{k}{n} \quad (\text{A.3})$$

nazywamy **prawdopodobieństwem zdarzenia B** w sensie statystycznym. Z tej definicji będziemy korzystali w rozdziale 4.

Przejdźmy teraz do podania paru dalszych podstawowych własności prawdopodobieństwa wynikających z definicji (A.1).

Niech dwa zdarzenia A i B będą zdarzeniami *wykluczającymi się*. Chcemy obliczyć prawdopodobieństwo tego, że zajdzie zdarzenie A lub B . Prawdopodobieństwo to oznaczmy przez $P(A + B)$. Niech liczba zdarzeń sprzyjających zdarzeniu A będzie równa k , a zdarzeniu B niech będzie równa l . Liczba zdarzeń, w których może wystąpić A lub B , wynosi więc $k + l$. Niech liczba wszystkich zdarzeń wynosi n . Zgodnie ze wzorem (A.1) prawdopodobieństwo $P(A + B)$ jest równe:

$$P(A + B) = \frac{k + l}{n} = \frac{k}{n} + \frac{l}{n}. \quad (\text{A.4})$$

Ponieważ $\frac{k}{n} = P(A)$, a $\frac{l}{n} = P(B)$ wzór (A.4) przyjmie postać:

$$P(A + B) = P(A) + P(B). \quad (\text{A.5})$$

W przypadku gdy mamy s zdarzeń wykluczających się $A_1, A_2, \dots, A_s$, to

$$P(A_1 + A_2 + \dots + A_s) = \sum_{i=1}^s P(A_i). \quad (\text{A.6})$$

Jest to tzw. prawo dodawania prawdopodobieństw zdarzeń wykluczających się.

Przykład

W urnie mamy 10 kul białych, 8 kul czarnych i 5 kul zielonych. Prawdopodobieństwo wyciągnięcia kuli białej oznaczmy przez $P(A)$, kuli czarnej przez $P(B)$, a zielonej przez $P(C)$. Zgodnie ze wzorem (A.1) wynoszą one:

$$P(A) = \frac{10}{23}, \quad P(B) = \frac{8}{23}, \quad \text{oraz} \quad P(C) = \frac{5}{23}.$$

Chcemy obliczyć prawdopodobieństwo $P(A + C)$, tj. prawdopodobieństwo wyciągnięcia kuli białej lub zielonej. Zgodnie ze wzorem (A.5) wynosi ono

$$P(A + C) = \frac{10 + 5}{23} = \frac{10}{23} + \frac{5}{23} = \frac{15}{23}.$$

Zdarzenia przeciwne A i $\bar{A}$ są zdarzeniami wykluczającymi się, więc prawdopodobieństwo $P(A + \bar{A})$ zajścia zdarzenia A lub $\bar{A}$ jest równe:

$$P(A + \bar{A}) = P(A) + P(\bar{A}),$$

co więcej, zajście zdarzenia A lub $\bar{A}$ jest zdarzeniem pewnym, a więc

$$P(A) + P(\bar{A}) = 1. \quad (\text{A.7})$$

Dotychczas była mowa o prawdopodobieństwie zajścia jakiegoś zdarzenia, gdy nie były podawane *żadne dodatkowe warunki*. Takie prawdopodobieństwo nazywamy **prawdopodobieństwem bezwarunkowym**. Często jednak pojawia się sytuacja, gdy zdarzenie A może zajść wówczas, gdy zajdzie zdarzenie B . Wtedy prawdopodobieństwo zajścia zdarzenia A pod warunkiem, że zaszło zdarzenie B , nazywamy **prawdopodobieństwem warunkowym** zdarzenia A i oznaczamy je $P(A|B)$. Niech całkowita liczba zdarzeń będzie równa n , liczba zdarzeń sprzyjających zdarzeniu B niech

będzie równa k , a liczba zdarzeń sprzyjających zdarzeniu A , gdy zaszło zdarzenie B , niech będzie równe l . Wtedy

$$P(A|B) = \frac{l}{k}. \quad (\text{A.8})$$

Wzór (A.8) możemy przekształcić do postaci

$$P(A|B) = \frac{l}{n} \cdot \frac{n}{k} = \frac{P(A, B)}{P(B)}, \quad (\text{A.9})$$

gdzie

$$P(A, B) = \frac{l}{n} \quad (\text{A.10})$$

i oznacza prawdopodobieństwo *jednoczesnego* zajścia zdarzenia A i zdarzenia B . Dzieje się tak dlatego, że l przypadków sprzyja zajściu zarówno zdarzenia A , jak i zdarzenia B .

Z wzoru (A.9) wynika, że:

$$P(A, B) = P(B) \cdot P(A|B). \quad (\text{A.11})$$

Jest to tzw. prawo mnożenia prawdopodobieństw. Możemy je sformułować następująco: Prawdopodobieństwo jednoczesnego zajścia dwóch zdarzeń jest równe iloczynowi prawdopodobieństwa jednego z nich i prawdopodobieństwa warunkowego drugiego zdarzenia przy założeniu, że zaszło pierwsze zdarzenie.

Ponieważ jest obojętne, które zdarzenie nazywamy pierwszym, a które drugim, to wzór (A.11) możemy zapisać w postaci

$$P(A, B) = P(A)P(B|A). \quad (\text{A.12})$$

O zdarzeniach mówimy że są **niezależne**, gdy prawdopodobieństwo zajścia jednego z nich nie zależy od tego czy zaszło drugie zdarzenie, czy też nie. W przypadku dwóch zdarzeń A i B oznacza to, że jeżeli zdarzenie A nie zależy od zdarzenia B , to wtedy $P(A|B) = P(A)$. W przypadku zdarzeń niezależnych wzór (A.11) przyjmie postać

$$P(A, B) = P(B) \cdot P(A). \quad (\text{A.13})$$

W przypadku s **zdarzeń niezależnych** $A_1, A_2, \dots, A_s$ prawo mnożenia prawdopodobieństw możemy zapisać:

$$P(A_1, A_2, \dots, A_s) = \prod_{j=1}^s P(A_j). \quad (\text{A.14})$$

Z wzoru (A.14) korzystaliśmy w rozdziale 5.

Przykłady

1) Mamy trzy stopery. Prawdopodobieństwo tego, że: pierwszy z nich nie ulegnie zepsuciu w ciągu dnia pracy (zdarzenie A) wynosi 0.98, drugi (zdarzenie B) wynosi 0.94 i trzeci (zdarzenie C) wynosi 0.92. Obliczamy prawdopodobieństwo tego, że żaden ze stoperów nie ulegnie zepsuciu w ciągu dnia pracy. Zdarzenia A , B i C są niezależne, ponieważ zepsucie się jednego stopera nie ma wpływu na pracę pozostałych. Wobec tego

$$P(A, B, C) = P(A) \cdot P(B) \cdot P(C).$$

2) Obliczmy prawdopodobieństwo wylosowania z talii kart damy lub waleta kier. Oznaczmy je przez $P(D + W, \text{kier})$. Prawdopodobieństwo wylosowania dowolnej *karty koloru kierowego*

$$P(\text{kier}) = \frac{\text{liczba kart koloru kierowego}}{\text{liczba kart w talii}} = \frac{13}{52}.$$

Zaś wylosowania *damy lub waleta*

$$P(D + W) = \frac{\text{liczba dam i waletów w talii}}{\text{liczba kart w talii}} = \frac{8}{52}.$$

Wobec tego

$$P(D + W, \text{kier}) = P(\text{kier}) \cdot P(D + W) = \frac{1}{4} \cdot \frac{8}{52} = \frac{2}{52}.$$

Na zakończenie należy jeszcze raz podkreślić, że wzór (A.6) możemy stosować *tylko wtedy*, gdy zdarzenia wykluczają się, a wzór (A.14) *tylko wtedy*, gdy zdarzenia są niezależne.

B. Klasy mierników elektrycznych

Klasą miernika nazywamy maksymalny błąd względny miernika pomnożony przez 100%. Klasa miernika jest więc jego maksymalnym błędem procentowym.

$$\text{klasa miernika} = \left| \frac{\delta_{\max}}{x_0} \right| \cdot 100\%$$

Na przykład miernik klasy 0.5 nie powinien wykazywać większego niż 0.5% błędu względnego pomiaru, oczywiście wówczas, gdy nie występują inne błędy systematyczne.

Obowiązujące normy przewidują siedem klas dokładności, a mianowicie: 0.1, 0.2, 0.5, 1, 1.5, 2.5, 5. Klasa miernika jest zawsze podawana przez producenta i umieszczona na tabliczce ze skalą przyrządu. Umieszczenie na tabliczce ze skalą np. liczby 0.5 oznacza, że miernik jest klasy 0.5.

C. Kowariancja wielkości złożonych

W Dodatku tym znajdziemy wyrażenie podające kowariancję dwóch wielkości fizycznych a i b zależnych od tych samych niezależnych zmiennych losowych x oraz y mierzonych bezpośrednio, tzn. że $a = a(x, y)$ i $b = b(x, y)$. Załóżmy, że wielkości fizyczne x oraz y zmierzylśmy odpowiednio N i M razy, jako wynik tych pomiarów otrzymaliśmy zbiory ich wartości $x_1, x_2, \dots, x_N$ oraz $y_1, y_2, \dots, y_M$. Tak więc w wyniku pomiarów mamy NM par (x_i, y_j) ($i = 1, 2, \dots, N$; $j = 1, 2, \dots, M$) i dwa NM elementowe zbiory odpowiadających im wartości $a_{ij} = a(x_i, y_j)$ oraz $b_{ij} = b(x_i, y_j)$. Ponieważ wielkości fizyczne $a(x, y)$ i $b(x, y)$ są funkcjami tych samych zmiennych losowych nie będą od siebie niezależne. Ich współzależność jest scharakteryzowana kowariancją $C_{a,b}$ lub zależnym od niej współczynnikiem korelacji $\rho_{a,b}$. Zgodnie z tym co powiedziano w rozdziale 3.4.2 kowariancja wielkości fizycznych $a(x, y)$ i $b(x, y)$ wyrazi się wzorem:

$$C_{a,b} = \frac{1}{NM} \sum_{i=1}^N \sum_{j=1}^M (a_{ij} - a_0)(b_{ij} - b_0), \quad (\text{C.1})$$

gdzie $a_0 = a(x_0, y_0)$ oraz $b_0 = b(x_0, y_0)$ są wartościami rzeczywistymi wielkości a oraz b , zaś x_0 i y_0 oznaczają odpowiednio wartości rzeczywiste niezależnych wielkości x oraz y mierzonych bezpośrednio.

Przyjmijmy, że warunki stosowalności rozwinięcia (3.3.3) są spełnione i obliczmy wartości $a(x_0, y_0)$, $a(x_i, y_j)$, $b(x_0, y_0)$, $b(x_i, y_j)$ rozwijając funkcje $a(x, y)$ i $b(x, y)$ na szereg Taylora funkcji dwóch zmiennych w otoczeniu punktu $(\bar{x}, \bar{y})$. Otrzymamy wówczas:

$$a(x_i, y_j) = a(\bar{x}, \bar{y}) + \left(\frac{\partial a}{\partial x}\right)_{\bar{x}, \bar{y}} (x_i - \bar{x}) + \left(\frac{\partial a}{\partial y}\right)_{\bar{x}, \bar{y}} (y_j - \bar{y}), \quad (\text{C.2})$$

$$a(x_0, y_0) = a(\bar{x}, \bar{y}) + \left(\frac{\partial a}{\partial x}\right)_{\bar{x}, \bar{y}} (x_0 - \bar{x}) + \left(\frac{\partial a}{\partial y}\right)_{\bar{x}, \bar{y}} (y_0 - \bar{y}), \quad (\text{C.3})$$

$$b(x_i, y_j) = b(\bar{x}, \bar{y}) + \left(\frac{\partial b}{\partial x}\right)_{\bar{x}, \bar{y}} (x_i - \bar{x}) + \left(\frac{\partial b}{\partial y}\right)_{\bar{x}, \bar{y}} (y_j - \bar{y}) \quad (\text{C.4})$$

oraz

$$b(x_0, y_0) = b(\bar{x}, \bar{y}) + \left(\frac{\partial b}{\partial x}\right)_{\bar{x}, \bar{y}} (x_0 - \bar{x}) + \left(\frac{\partial b}{\partial y}\right)_{\bar{x}, \bar{y}} (y_0 - \bar{y}), \quad (\text{C.5})$$

Podstawiając (C.2) – (C.5) do (C.1) znajdujemy:

$$C_{a,b} = \frac{1}{NM} \sum_{i=1}^N \sum_{j=1}^M \left[\left(\frac{\partial a}{\partial x} \right)_{\bar{x}, \bar{y}} (x_i - x_0) + \left(\frac{\partial a}{\partial y} \right)_{\bar{x}, \bar{y}} (y_j - y_0) \right] \times \\ \times \left[\left(\frac{\partial b}{\partial x} \right)_{\bar{x}, \bar{y}} (x_i - x_0) + \left(\frac{\partial b}{\partial y} \right)_{\bar{x}, \bar{y}} (y_j - y_0) \right]. \quad (C.6)$$

Wyrażenie (C.6) można zapisać w postaci:

$$C_{a,b} = \left(\frac{\partial a}{\partial x} \right)_{\bar{x}, \bar{y}} \left(\frac{\partial b}{\partial x} \right)_{\bar{x}, \bar{y}} \frac{1}{N} \sum_{i=1}^N (x_i - x_0)^2 \\ + \left(\frac{\partial a}{\partial y} \right)_{\bar{x}, \bar{y}} \left(\frac{\partial b}{\partial y} \right)_{\bar{x}, \bar{y}} \frac{1}{M} \sum_{j=1}^M (y_j - y_0)^2 \\ + \left[\left(\frac{\partial a}{\partial x} \right)_{\bar{x}, \bar{y}} \left(\frac{\partial b}{\partial y} \right)_{\bar{x}, \bar{y}} + \left(\frac{\partial a}{\partial y} \right)_{\bar{x}, \bar{y}} \left(\frac{\partial b}{\partial x} \right)_{\bar{x}, \bar{y}} \right] \frac{1}{NM} \sum_{i=1}^N \sum_{j=1}^M (x_i - x_0)(y_j - y_0). \quad (C.7)$$

Ostatni składnik w powyższym wyrażeniu, zgodnie z tym, co powiedziano w rozdziale 3.4.1, może być pominięty. Jeżeli dodatkowo skorzystamy z definicji odchylenia standardowego serii pomiarów bezpośrednich (3.2.3) to wzór (C.7) zapiszemy:

$$C_{a,b} = \left(\frac{\partial a}{\partial x} \right)_{\bar{x}, \bar{y}} \left(\frac{\partial b}{\partial x} \right)_{\bar{x}, \bar{y}} S_x^2 + \left(\frac{\partial a}{\partial y} \right)_{\bar{x}, \bar{y}} \left(\frac{\partial b}{\partial y} \right)_{\bar{x}, \bar{y}} S_y^2. \quad (C.8)$$

Jeżeli mamy Q wielkości a_i ($i = 1, 2, \dots, Q$) zależnych od r niezależnych wielkości mierzonych bezpośrednio $x_1, x_2, \dots, x_r$, to znaczy $a_i = a_i(x_1, x_2, \dots, x_r)$ ($i = 1, 2, \dots, Q$), wtedy wzór (C.8) przyjmie postać:

$$C_{a_i, a_j} = \sum_{k=1}^r \left(\frac{\partial a_i}{\partial x_k} \right) \left(\frac{\partial a_j}{\partial x_k} \right)_{\bar{x}_1, \bar{x}_2, \dots, \bar{x}_r} S_{x_k}^2. \quad (C.9)$$

D. Wyprowadzenie wzoru (3.5.19)

Średnia kwadratów odchyleń punktów od prostej $y = ax + b$ wynosi

$$S_y^2 = \frac{1}{N} \sum_{i=1}^N (y_i - ax_i - b)^2 = \frac{1}{N} \sum_{i=1}^N \delta_i^2, \quad (\text{D.1})$$

gdzie

$$\delta_i = y_i - ax_i - b. \quad (\text{D.2})$$

Ponieważ wielkości δ_i są nieznane, możemy jedynie znaleźć wielkości

$$d_i = y_i - \hat{a}x_i - \hat{b}. \quad (\text{D.3})$$

Z tego powodu w równaniu (D.1) należy zastąpić δ_i przez d_i . Należy zatem znaleźć związek między tymi wielkościami. Odejmując stronami (D.3) od (D.2) otrzymujemy:

$$\delta_i - d_i = (\hat{a} - a)x_i + (\hat{b} - b) = \Delta d_i, \quad (\text{D.4})$$

lub

$$\delta_i = d_i + \Delta d_i. \quad (\text{D.5})$$

Podnieśmy obustronnie równanie (D.5) do kwadratu i wysumujmy po wszystkich N pomiarach. W wyniku tego otrzymamy:

$$\sum_{i=1}^N \delta_i^2 = \sum_{i=1}^N d_i^2 + 2 \sum_{i=1}^N d_i \Delta d_i + \sum_{i=1}^N (\Delta d_i)^2. \quad (\text{D.6})$$

Obliczenie $\sum_{i=1}^N d_i \Delta d_i$.

Korzystając z zależności (D.4) znajdujemy, że

$$\sum_{i=1}^N d_i \Delta d_i = (\hat{a} - a) \sum_{i=1}^N d_i x_i - (\hat{b} - b) \sum_{i=1}^N d_i. \quad (\text{D.7})$$

Korzystając z (D.3) układ równań (3.5.11) możemy zapisać w postaci:

$$\begin{aligned} \sum_{i=1}^N x_i d_i &= 0, \\ \sum_{i=1}^N d_i &= 0. \end{aligned} \quad (\text{D.8})$$

Podstawiając (D.8) do (D.7) znajdujemy:

$$\sum_{i=1}^N d_i \Delta d_i = 0. \quad (\text{D.9})$$

Obliczenie $\sum_{i=1}^N (\Delta d_i)^2$

Ze wzoru (D.4) wynika, że

$$(\Delta d_i)^2 = \Delta d_i \cdot \Delta d_i = (\hat{a} - a)x_i \Delta d_i + (\hat{b} - b)\Delta d_i,$$

wobec tego

$$\sum_{i=1}^N (\Delta d_i)^2 = (\hat{a} - a) \sum_{i=1}^N x_i \Delta d_i + (\hat{b} - b) \sum_{i=1}^N \Delta d_i. \quad (\text{D.10})$$

Korzystając z (D.5) i (D.8) znajdujemy, że

$$\sum_{i=1}^N x_i \Delta d_i = \sum_{i=1}^N \delta_i x_i - \sum_{i=1}^N d_i x_i = \sum_{i=1}^N \delta_i x_i$$

oraz

$$\sum_{i=1}^N \Delta d_i = \sum_{i=1}^N \delta_i - \sum_{i=1}^N d_i = \sum_{i=1}^N \delta_i.$$

Wstawiając powyższe związki do (D.10) otrzymujemy:

$$\sum_{i=1}^N (\Delta d_i)^2 = (\hat{a} - a) \sum_{i=1}^N \delta_i x_i + (\hat{b} - b) \sum_{i=1}^N \delta_i. \quad (\text{D.11})$$

Aby obliczyć $\sum_{i=1}^N (\Delta d_i)^2$, musimy wyznaczyć różnice $\hat{a} - a$ oraz $\hat{b} - b$. Współczynniki $\hat{a}$ i $\hat{b}$, dane odpowiednio wzorami (3.5.13) i (3.5.14), możemy przedstawić w postaci:

$$\hat{a} = \sum_{k=1}^N A_k y_k, \quad (\text{D.12})$$

$$\hat{b} = \sum_{k=1}^N B_k y_k, \quad (\text{D.13})$$

gdzie

$$A_k = \frac{\sum_{i=1}^N x_i - N x_k}{\left(\sum_{i=1}^N x_i\right)^2 - N \sum_{i=1}^N x_i^2}, \quad (\text{D.14})$$

$$B_k = \frac{x_k \sum_{i=1}^N x_i - \sum_{i=1}^N x_i^2}{\left(\sum_{i=1}^N x_i\right)^2 - N \sum_{i=1}^N x_i^2}. \quad (\text{D.15})$$

Ze wzoru (D.14) wynika, że

$$\sum_{k=1}^N A_k = 0, \quad \sum_{n=1}^N x_k A_k = 1, \quad (\text{D.16})$$

a ze wzoru (D.15), że

$$\sum_{k=1}^N B_k = 1, \quad \sum_{n=1}^N x_k B_k = 0. \quad (\text{D.17})$$

Zmierzonej wartości x_i odpowiada wartość rzeczywista y_{0i} , która jest równa:

$$y_{0i} = ax_i + b. \quad (\text{D.18})$$

Z wzorów (D.2) i (D.18) otrzymujemy:

$$y_{0i} = y_i - \delta_i. \quad (\text{D.19})$$

Różnicę $\hat{a} - a$ możemy teraz zapisać:

$$\hat{a} - a = \sum_{k=1}^N A_k y_k - a.$$

Podstawiając y_k obliczone ze wzoru (D.19) oraz y_{0k} ze wzoru (D.18) dostajemy:

$$\begin{aligned} \hat{a} - a &= \sum_{k=1}^N A_k y_{0k} + \sum_{k=1}^N A_k \delta_k - a \\ &= a \sum_{k=1}^N A_k x_k + b \sum_{k=1}^N A_k + \sum_{k=1}^N A_k \delta_k - a. \end{aligned}$$

Biorąc pod uwagę (D.16) znajdujemy, że

$$\hat{a} - a = \sum_{k=1}^N A_k \delta_k. \quad (\text{D.20})$$

W ten sam sposób obliczamy różnicę $\hat{b} - b$. Jest ona równa:

$$\hat{b} - b = \sum_{k=1}^N B_k \delta_k. \quad (\text{D.21})$$

Podstawiając (D.20) i (D.21) do (D.11) otrzymujemy:

$$\sum_{i=1}^N (\Delta d_i)^2 = \left(\sum_{i=1}^N x_i \delta_i \right) \left(\sum_{k=1}^N A_k \delta_k \right) + \left(\sum_{i=1}^N \delta_i \right) \left(\sum_{k=1}^N B_k \delta_k \right). \quad (\text{D.22})$$

Aby obliczyć $\sum (\Delta d_i)^2$ musimy dokonać kilku przybliżeń. Rozpatrzmy na początku pierwszy składnik prawej strony równania (D.22). Jest on równy:

$$\begin{aligned} \left(\sum_{i=1}^N x_i \delta_i \right) \left(\sum_{k=1}^N A_k \delta_k \right) &= \\ &= (x_1 \delta_1 + x_2 \delta_2 + \dots + x_N \delta_N) (A_1 \delta_1 + A_2 \delta_2 + \dots + A_N \delta_N) \\ &= x_1 A_1 \delta_1^2 + x_2 A_2 \delta_2^2 + \dots + x_N A_N \delta_N^2 + x_1 A_2 \delta_1 \delta_2 + x_2 A_1 \delta_1 \delta_2 + \dots \\ &= \sum_{i=1}^N x_i A_i \delta_i^2 + \sum_{\substack{j,k=1 \\ j \neq k}}^N (x_j A_k + x_k A_j) \delta_j \delta_k. \end{aligned}$$

Ponieważ błędy przypadkowe mają różne znaki i różne wartości, druga suma w tym wyrażeniu będzie w przybliżeniu równa zero, czyli

$$\left(\sum_{i=1}^N x_i \delta_i \right) \left(\sum_{k=1}^N A_k \delta_k \right) \approx \sum_{i=1}^N x_i A_i \delta_i^2. \quad (\text{D.23})$$

Postępując podobnie drugi człon prawej strony równania (D.22) możemy przedstawić w postaci:

$$\begin{aligned} \left(\sum_{i=1}^N \delta_i \right) \left(\sum_{k=1}^N B_k \delta_k \right) &= (\delta_1 + \delta_2 + \dots + \delta_N) (B_1 \delta_1 + B_2 \delta_2 + \dots + B_N \delta_N) \\ &= \sum_{i=1}^N B_i \delta_i^2 + \sum_{\substack{j,k=1 \\ j \neq k}}^N B_j \delta_j \delta_k. \end{aligned}$$

Z powodów podanych wyżej, druga suma jest w przybliżeniu równa zero, wobec tego:

$$\left(\sum_{i=1}^N \delta_i \right) \left(\sum_{k=1}^N B_k \delta_k \right) \approx \sum_{i=1}^N B_i \delta_i^2. \quad (\text{D.24})$$

Podstawiając (D.23) i (D.24) do (D.22) znajdujemy:

$$\sum_{i=1}^N (\Delta d_i)^2 \approx \sum_{i=1}^N A_i x_i \delta_i^2 + \sum_{i=1}^N B_i \delta_i^2. \quad (\text{D.25})$$

Dokonyjmy teraz kolejnego przybliżenia. W wyrażeniach $x_i A_i \delta_i^2$ i $B_i \delta_i^2$ przybliżmy kwadraty odchyłeń δ_i^2 przez ich wartości średnie $\frac{1}{N} \sum_{i=1}^N \delta_i^2$, wówczas wzór (D.25) przyjmie postać:

$$\sum_{i=1}^N (\Delta d_i)^2 = \left(\frac{1}{N} \sum_{i=1}^N \delta_i^2 \right) \left(\sum_{i=1}^N A_i x_i \right) + \left(\frac{1}{N} \sum_{i=1}^N \delta_i^2 \right) \left(\sum_{i=1}^N B_i \right). \quad (\text{D.26})$$

Ponieważ, zgodnie z (D.16) i (D.17), $\sum_{i=1}^N A_i x_i = 1$ oraz $\sum_{i=1}^N B_i = 1$, otrzymamy:

$$\sum_{i=1}^N (\Delta d_i)^2 = \frac{2}{N} \sum_{i=1}^N \delta_i^2. \quad (\text{D.27})$$

Wstawiając (D.27) i (D.9) do (D.6) dostajemy:

$$\sum_{i=1}^N \delta_i^2 = \frac{N}{N-2} \sum_{i=1}^N d_i^2. \quad (\text{D.28})$$

Podstawiając otrzymaną wartość $\sum \delta_i^2$ do (D.1) znajdujemy:

$$S_y^2 = \frac{1}{N-2} \sum_{i=1}^N d_i^2 = \frac{1}{N-2} \sum_{i=1}^N (y_i - \hat{a}x_i - \hat{b})^2. \quad (\text{D.29})$$

Otrzymaliśmy więc wzór (3.5.19).

E. Model błędów przypadkowych Laplace'a i uzasadnienie rozkładu normalnego

Laplace w roku 1783 [23] podał prosty model tłumaczący powstawanie błędów przypadkowych. W modelu Laplace'a zakłada się, że:

- mierzymy wielkość fizyczną x o wartości rzeczywistej x_0 ;
- każdy jej pomiar obarczony jest błędem przypadkowym, spowodowanym działaniem dużej liczby N niezależnych czynników (źródeł);
- każdy z tych czynników powoduje powstanie takiego samego błędu $|\epsilon|$;
- prawdopodobieństwo wystąpienia błędu $-\epsilon$, jak i $+\epsilon$, jest jednakowe i wynosi $1/2$.

Obliczenia prowadzące do uzasadnienia rozkładu normalnego, które przedstawimy poniżej, zaczerpnięto z pracy [24].

Przy założeniach modelu Laplace'a błąd pomiaru jest sumą błędów spowodowanych przez każdy z tych N czynników. Prawdopodobieństwo tego, że k spośród N czynników spowoduje wystąpienie błędu $+\epsilon$, zgodnie z rozważaniami przedstawionymi w rozdziale 7.1, jest dane rozkładem dwumianowym:

$$b(k, N) = \binom{N}{k} p^k q^{N-k}. \quad (\text{E.1})$$

Ponieważ przyjęto powyżej, że błędy $-\epsilon$ i $+\epsilon$ są jednakowo prawdopodobne, to $p = q = 1/2$ i wzór (E.1), po rozwinięciu symbolu Newtona, przyjmie postać

$$b(k, N) = \frac{N!}{k!(N-k)!} \left(\frac{1}{2}\right)^N. \quad (\text{E.2})$$

W przypadku, gdy każdy z k czynników powoduje błąd $+\epsilon$, to każdy z $N-k$ czynników powoduje błąd $-\epsilon$. Całkowity błąd bezwzględny jest wówczas:

$$\delta = m\epsilon, \quad (\text{E.3})$$

gdzie:

$$m = k - (N - k) = 2k - N. \quad (\text{E.4})$$

Ze wzoru (E.4) otrzymujemy:

$$k = \frac{1}{2}(m + N). \quad (\text{E.5})$$

Podstawiając (E.5) do (E.2) znajdujemy prawdopodobieństwo $p(m, N)$, że w wyniku działania *wszystkich* N czynników wystąpi błąd $\delta = m\epsilon$.

$$p(m, N) = \frac{N!}{\left[\frac{1}{2}(m+N)\right]! \left[\frac{1}{2}(N-m)\right]!} \left(\frac{1}{2}\right)^N. \quad (\text{E.6})$$

Wzór (E.6) można znacznie uprościć, gdy liczba niezależnych czynników N jest bardzo duża. Skorzystamy w tym celu ze wzoru Stirlinga, który jest przybliżeniem $n!$ dla dużych n :

$$n! \approx \sqrt{2\pi n} n^{(n+\frac{1}{2})} e^{-n}. \quad (\text{E.7})$$

Logarytmując (E.7) otrzymujemy:

$$\ln n! \approx \left(n + \frac{1}{2}\right) \ln n - n + \ln \sqrt{2\pi}, \quad (\text{E.8})$$

a logarytmując (E.6) dostajemy:

$$\ln p(m, N) = \ln N! - \ln \left[\frac{1}{2}(N+m)\right]! - \ln \left[\frac{1}{2}(N-m)\right]! - N \ln 2. \quad (\text{E.9})$$

Jeżeli wyrażenia $N!$, $\left[\frac{1}{2}(N+m)\right]!$ oraz $\left[\frac{1}{2}(N-m)\right]!$ przybliżymy wzorem Stirlinga, a następnie zlogarytmujemy i podstawimy do wzoru (E.9), to przyjmiemy on postać:

$$\begin{aligned} \ln p(m, N) &\approx \left(N + \frac{1}{2}\right) \ln N - N + \ln \sqrt{2\pi} \\ &- \frac{1}{2}(N+m+1) \ln \left[\frac{1}{2}(N+m)\right] + \frac{1}{2}(N+m) - \ln \sqrt{2\pi} \\ &- \frac{1}{2}(N-m+1) \ln \left[\frac{1}{2}(N-m)\right] - \frac{1}{2}(N-m) - \ln \sqrt{2\pi} - N \ln 2. \end{aligned} \quad (\text{E.10})$$

Występujące we wzorze (E.10) logarytmy wyrażenia $\frac{1}{2}(N \pm m)$ możemy przekształcić, a mianowicie:

$$\ln \left[\frac{1}{2}(N \pm m)\right] = \ln \left[\frac{1}{2}N \left(1 \pm \frac{m}{N}\right)\right] = -\ln 2 + \ln N + \ln \left(1 \pm \frac{m}{N}\right). \quad (\text{E.11})$$

Podstawiając (E.11) do (E.10) i wykonując proste obliczenia, otrzymujemy:

$$\begin{aligned} \ln p(m, N) &= \frac{1}{2} \ln \left(\frac{2}{\pi N}\right) - \frac{1}{2}(N+m+1) \ln \left(1 + \frac{m}{N}\right) \\ &- \frac{1}{2}(N-m+1) \ln \left(1 - \frac{m}{N}\right). \end{aligned} \quad (\text{E.12})$$

Kolejnego przybliżenia możemy dokonać korzystając z założenia, że liczba N jest bardzo duża. W tym przypadku logarytmy we wzorze (E.12) możemy

rozwinąć na szereg Taylora w otoczeniu zera i ograniczyć się do pierwszych dwóch wyrazów rozwinięcia, tj.

$$\ln \left(1 \pm \frac{m}{N} \right) \approx \pm \frac{m}{N} - \frac{m^2}{2N^2}. \quad (\text{E.13})$$

Podstawiając (E.13) do (E.12) otrzymujemy:

$$\ln p(m, N) \approx \ln \sqrt{\frac{2}{\pi N}} - \frac{N-1}{N} \frac{m^2}{2N}, \quad (\text{E.14})$$

ale $(N-1)/N \rightarrow 1$, gdy $N \rightarrow \infty$, a więc wzór (E.14) przyjmie postać:

$$\ln p(m, N) = \ln \sqrt{\frac{2}{\pi N}} - \frac{m^2}{2N}, \quad (\text{E.15})$$

a zatem szukane prawdopodobieństwo, że działanie N czynników spowoduje wystąpienie błędu $\delta = m\epsilon$ jest:

$$p(m, N) = \sqrt{\frac{2}{\pi N}} e^{-\frac{m^2}{2N}}. \quad (\text{E.16})$$

Otrzymaliśmy proste wyrażenie, które jest przybliżeniem wzoru (E.6) dla dużych N . Chociaż zostało ono otrzymane przy założeniu, że $N \rightarrow \infty$, jest bardzo dobrym przybliżeniem wzoru (E.6), nawet gdy $N = 10$.

Zastosujmy teraz wzór (E.16) w celu znalezienia funkcji rozkładu prawdopodobieństwa błędów bezwzględnych $\varphi(\delta)$. Prawdopodobieństwo tego, że błąd bezwzględny δ będzie zawarty między δ a $\delta + \Delta\delta$ jest w przybliżeniu równe $\varphi(\delta)\Delta\delta$, dla dostatecznie małych $\Delta\delta$. Traktujemy tu δ jako zmienną ciągłą, natomiast m występujące w równaniu (E.3) jest zmienną dyskretną. Należy teraz powiązać ze sobą $\varphi(\delta)\Delta\delta$, prawdopodobieństwo wystąpienia wartości δ z danego przedziału z prawdopodobieństwem $p(m, N)$ otrzymania odpowiadającej jej wartości dyskretnej m . Aby można było porównać z dobrym przybliżeniem oba te prawdopodobieństwa, należy tak dobrać szerokość przedziału $\Delta\delta$, aby odpowiadała odstępowi między sąsiednimi wartościami m , czyli $\Delta m = 2$, odpowiadającego $\Delta k = 1$ (por. (E.4)). Wtedy

$$\varphi(\delta)\Delta\delta = p(m, N), \quad (\text{E.17})$$

gdzie

$$\Delta\delta = 2\epsilon. \quad (\text{E.18})$$

Wstawiając (E.18) i (E.16) do wzoru (E.17) znajdujemy:

$$\varphi(\delta) = \frac{1}{\epsilon\sqrt{2\pi N}} e^{-\frac{m^2}{2N}}. \quad (\text{E.19})$$

Następnie ze wzoru (E.3) obliczamy $m = \delta/\epsilon$ i wstawiamy do (E.19); otrzymujemy:

$$\varphi(\delta) = \frac{1}{\epsilon\sqrt{2\pi N}} e^{-\frac{\delta^2}{2N\epsilon^2}}.$$

Oznaczmy $N\epsilon = \sigma^2$, wówczas

$$\boxed{\varphi(\delta) = \frac{1}{\sigma\sqrt{2\pi}} e^{-\frac{\delta^2}{2\sigma^2}}.} \quad (\text{E.20})$$

Tak więc korzystając z modelu błędów przypadkowych Laplace'a i rozkładu Bernoulliego otrzymaliśmy w przypadku granicznym rozkład Gaussa.

F. Obliczenie niektórych całek

1. Obliczenie całki $\int_{-\infty}^{+\infty} e^{-x^2} dx$.

Rozpatrzmy całkę

$$I = \int_{-a}^a e^{-x^2} dx = \int_{-a}^a e^{-y^2} dy.$$

Równość ta zachodzi, ponieważ wartość całki nie zależy od oznaczenia zmiennej i wyboru osi w układzie współrzędnych. Wobec tego

$$I^2 = \int_{-a}^a e^{-x^2} dx \cdot \int_{-a}^a e^{-y^2} dy = \int_{-a}^a \int_{-a}^a e^{-(x^2+y^2)} dx dy,$$

i jest całką z funkcji $\exp[-(x^2 + y^2)]$ po powierzchni kwadratu o boku $2a$ (rysunek F.1). Przejdźmy teraz do współrzędnych biegunowych (r, Θ) . Wtedy element powierzchni $dx dy = r dr d\Theta$ oraz $x^2 + y^2 = r^2$. Rozpatrzmy całkę po powierzchni koła o promieniu b

$$B = \int_0^b \int_0^{2\pi} e^{-r^2} r dr d\Theta = 2\pi \int_0^b e^{-r^2} r dr = \pi (1 - e^{-b^2}).$$

Jeżeli $b = a$, to $I^2 > B$ (okrąg 1); jeżeli $b = \sqrt{2}a$, to $I^2 < B$ (okrąg 2); wobec tego

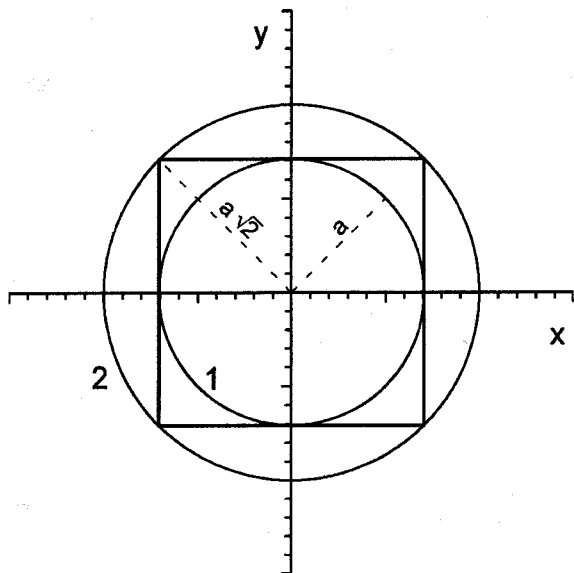
$$\pi (1 - e^{-a^2}) < I^2 < \pi (1 - e^{-2a^2}).$$

Gdy $a \rightarrow \infty$, to oba te ograniczenia zbiegają do wspólnej granicy równej π . A więc

$$\lim_{a \rightarrow \infty} I^2 = \lim_{a \rightarrow \infty} \int_{-a}^a e^{-(x^2+y^2)} dx dy = I_0^2 = \pi,$$

wobec tego

$$I_0 = \int_{-\infty}^{+\infty} e^{-x^2} dx = \sqrt{\pi}.$$



Rys. F.1: Rysunek pomocniczny do obliczenia całki z funkcji Gaussa [7]

2. Obliczenie całki $\int_{-\infty}^{+\infty} x e^{-x^2} dx$.

Obliczmy na początku całkę nieoznaczoną $\int x e^{-x^2} dx$. Stosując podstawienie $x^2 = y$, po zróżniczkowaniu otrzymujemy $2x dx = dy$ i

$$I_1 = \int x e^{-x^2} dx = \frac{1}{2} \int e^{-y} dy = \frac{1}{2} e^{-y} + C = \frac{1}{2} e^{-x^2} + C.$$

Wobec tego całka oznaczona

$$\int_{-a}^a x e^{-x^2} dx = 0,$$

czyli

$$I_1 = \int_{-\infty}^{+\infty} x e^{-x^2} dx = 0.$$

3. Obliczenie całki $\int_{-\infty}^{+\infty} x^2 e^{-x^2} dx$.

Obliczamy ją całkując przez części, a mianowicie

$$I_2 = \int_{-\infty}^{+\infty} x^2 e^{-x^2} dx = \int_{-\infty}^{+\infty} x (x e^{-x^2}) dx = \int_{-\infty}^{+\infty} x \left(\frac{d}{dx} \left(-\frac{e^{-x^2}}{2} \right) \right) dx,$$

a więc

$$I_2 = -\frac{xe^{-x^2}}{2} \Big|_{-\infty}^{+\infty} - \int_{-\infty}^{+\infty} \left(-\frac{e^{-x^2}}{2} \right) dx = \frac{1}{2} \int_{-\infty}^{+\infty} e^{-x^2} dx.$$

Stąd ostatecznie mamy

$$I_2 = \int_{-\infty}^{+\infty} x^2 e^{-x^2} dx = \frac{1}{2} \sqrt{\pi}.$$

G. Funkcja rozkładu średniej arytmetycznej

Na początku udowodnimy kilka twierdzeń, na podstawie których znajdziemy funkcję rozkładu średniej arytmetycznej.

TWIERDZENIE 1.

Jeżeli $\varphi_1(x_1)$ i $\varphi_2(x_2)$ są odpowiednio funkcjami rozkładu niezależnych ciągłych zmiennych losowych x_1 i x_2 , to zmienna losowa $y = x_1 + x_2$ będzie miała funkcję rozkładu $\varphi(y)$, która wyraża się następująco:

$$\varphi(y) = \int_{-\infty}^{+\infty} \varphi_1(x_1) \varphi_2(y - x_1) dx_1. \quad (G.1)$$

Całkę tę nazywamy **splotem** funkcji φ_1 i φ_2 .

Dowód.

Prawdopodobieństwo wystąpienia zmiennej losowej x_1 z przedziału $[x_1, x_1 + dx_1]$ jest $\varphi_1(x_1)dx_1$, a zmiennej losowej x_2 z przedziału $[x_2, x_2 + dx_2]$ jest $\varphi_2(x_2)dx_2$. Wobec tego, że zmienne losowe x_1 i x_2 są niezależne, prawdopodobieństwo jednoczesnego wystąpienia zmiennej x_1 z przedziału $[x_1, x_1 + dx_1]$ i zmiennej x_2 z przedziału $[x_2, x_2 + dx_2]$ jest równe iloczynowi prawdopodobieństw wystąpienia każdej z tych zmiennych w odpowiednim przedziale, czyli

$$\varphi_1(x_1) \varphi_2(x_2) dx_1 dx_2. \quad (G.2)$$

Ponieważ, jak założono, zmienna losowa y jest sumą zmiennych losowych x_1 i x_2 ,

$$y = x_1 + x_2, \quad (G.3)$$

to w celu znalezienia prawdopodobieństwa $\varphi(y)dy$, że zmienna losowa y jest zawarta między y a $y + dy$, musimy wysumować prawdopodobieństwa (G.2) po wszystkich parach x_1 i x_2 , dla których zmienna losowa y , będąca ich sumą, ma tę samą wartość. Dla zmiennych losowych ciągłych możliwości realizacji warunku (G.3) mamy nieskończenie wiele i sumowanie musimy zastąpić całkowaniem. Aby obliczyć funkcję rozkładu zmiennej losowej ciągłej y dokonujemy podstawienia $x_2 = y - x_1$, wtedy $dx_2 = dy$. Prawdopodo-

bieństwo wystąpienia zmiennej losowej y z przedziału $[y, y + dy]$ będzie:

$$\varphi(y)dy = \int_{-\infty}^{+\infty} \varphi_1(x_1)\varphi_2(y-x_1)dx_1dy,$$

a więc

$$\varphi(y) = \int_{-\infty}^{+\infty} \varphi_1(x_1)\varphi_2(y-x_1)dx_1, \quad (\text{G.4})$$

a tego należało dowieść.

W dalszym ciągu tego Dodatku będziemy zajmowali się rozkładem normalnym. Ponieważ wynik pomiaru t_i ($i = 1, 2, \dots, N$) wielkości fizycznej t jest zmienną losową, to również różnica $t_i - t_0$ (ogólnie $t_0 = \text{const}$, w naszym przypadku jest to wartość rzeczywista) jest zmienną losową o tej samej funkcji rozkładu, ale maksimum jej jest przesunięte o t_0 i występuje dla zera. Wobec tego, nie zmniejszając ogólności rozważań, możemy dla przejrzystości obliczeń przyjąć, że rozpatrujemy zmienne losowe $x = t - t_0$. Otrzymamy wówczas wzory, w których zmienną losową x należy zastąpić przez powyższą różnicę i wtedy będą one identyczne ze wzorami rozdziału 5.

TWIERDZENIE 2.

Splot dwóch rozkładów Gaussa, których wariancje wynoszą σ_1^2 i σ_2^2 , jest rozkładem Gaussa o wariancji $\sigma^2 = \sigma_1^2 + \sigma_2^2$.

Dowód.

W tym przypadku całka (G.1) będzie splotem dwóch rozkładów Gaussa i przyjmie postać:

$$\varphi(y) = \int_{-\infty}^{+\infty} \frac{1}{\sigma_1\sqrt{2\pi}} e^{-\frac{x_1^2}{2\sigma_1^2}} \frac{1}{\sigma_2\sqrt{2\pi}} e^{-\frac{(y-x_1)^2}{2\sigma_2^2}} dx_1. \quad (\text{G.5})$$

Możemy ją zapisać:

$$\varphi(y) = \frac{1}{2\pi\sigma_1\sigma_2} \int_{-\infty}^{+\infty} \exp \left[-\frac{y^2}{2\sigma_2^2} + 2\frac{yx_1}{2\sigma^2} - \frac{2(\sigma_1^2 + \sigma_2^2)}{4\sigma_1^2\sigma_2^2} x_1^2 \right] dx_1. \quad (\text{G.6})$$

Aby obliczyć całkę (G.6) musimy przekształcić wyrażenie w wykładniku potęgi. Jeżeli dodamy i odejmiemy

$$\frac{2y^2\sigma_1^2}{4\sigma_2^2(\sigma_1^2 + \sigma_2^2)},$$

oraz skorzystamy ze wzoru $a^2 - 2ab + b^2 = (a - b)^2$, to wykładnik potęgi przyjmie postać:

$$-\frac{y^2}{2(\sigma_1^2 + \sigma_2^2)} - \left(\frac{\sqrt{\sigma_1^2 + \sigma_2^2}}{\sqrt{2}\sigma_1\sigma_2}x_1 - \frac{\sigma_1}{\sqrt{2}\sigma_2\sqrt{\sigma_1^2 + \sigma_2^2}}y \right)^2. \quad (\text{G.7})$$

Oznaczmy:

$$s = \frac{\sqrt{\sigma_1^2 + \sigma_2^2}}{\sqrt{2}\sigma_1\sigma_2}x_1 - \frac{\sigma_1}{\sqrt{2}\sigma_2\sqrt{\sigma_1^2 + \sigma_2^2}}y.$$

Wtedy wyrażenie (G.7) przyjmie postać:

$$-\frac{y^2}{2(\sigma_1^2 + \sigma_2^2)} - s^2, \quad (\text{G.8})$$

oraz

$$dx_1 = \frac{\sqrt{2}\sigma_1\sigma_2}{\sqrt{\sigma_1^2 + \sigma_2^2}}ds. \quad (\text{G.9})$$

Podstawiając (G.8) i (G.9) do (G.6) otrzymamy:

$$\varphi(y) = \frac{1}{\pi\sqrt{2}\sqrt{\sigma_1^2 + \sigma_2^2}} e^{-\frac{y^2}{2(\sigma_1^2 + \sigma_2^2)}} \int_{-\infty}^{+\infty} e^{-s^2} ds. \quad (\text{G.10})$$

Ponieważ występująca w tym wyrażeniu całka jest równa $\sqrt{\pi}$ (obliczenie podano w Dodatku F), wzór (G.10) przyjmie postać:

$$\varphi(y) = \frac{1}{\sqrt{2\pi(\sigma_1^2 + \sigma_2^2)}} e^{-\frac{y^2}{2(\sigma_1^2 + \sigma_2^2)}}.$$

Suma $\sigma_1^2 + \sigma_2^2 = \sigma^2$ jest wariancją rozkładu $\varphi(y)$ będącego splotem dwóch rozkładów Gaussa, wobec tego:

$$\varphi(y) = \frac{1}{\sigma\sqrt{2\pi}} e^{-\frac{y^2}{2\sigma^2}}, \quad (\text{G.11})$$

czego należało dowieść.

TWIERDZENIE 3.

Jeżeli każda z niezależnych zmiennych losowych $x_1, x_2, \dots, x_N$ podlega rozkładowi Gaussa o wariancji odpowiednio $\sigma_1^2, \sigma_2^2, \dots, \sigma_N^2$, to zmienna losowa:

$$y = x_1 + x_2 + \dots + x_N \quad (\text{G.12})$$

również podlega rozkładowi Gaussa o wariancji:

$$\sigma^2 = \sigma_1^2 + \sigma_2^2 + \dots + \sigma_N^2. \quad (\text{G.13})$$

Dowód.

Wynika z twierdzenia 1 oraz 2 i jest rekurencyjny.

Twierdzenie 4.

Jeżeli zmienna losowa x podlega rozkładowi Gaussa z wariancją σ_x^2 , to zmienna losowa:

$$y = ax, \quad (\text{G.14})$$

($a > 0$) podlega również rozkładowi Gaussa o wariancji:

$$\sigma_y^2 = a^2 \sigma_x^2. \quad (\text{G.15})$$

Dowód.

Polega na wykonaniu prostych przekształceń, a mianowicie:

$$\varphi(x)dx = \frac{1}{\sigma_x \sqrt{2\pi}} \exp \left[-\frac{x^2}{2\sigma_x^2} \right] dx = \frac{1}{a\sigma_x \sqrt{2\pi}} \exp \left[-\frac{(ax)^2}{2(\sigma_x a)^2} \right] d(ax),$$

ale $ax = y$ i wzór powyższy przyjmie postać:

$$\varphi(x)dx = \frac{1}{a\sigma_x \sqrt{2\pi}} \exp \left[-\frac{y^2}{2(\sigma_x a)^2} \right] dy = \varphi(y)dy.$$

Stąd wynika, że

$$\varphi(y) = \frac{1}{\sigma_y \sqrt{2\pi}} e^{-\frac{y^2}{2\sigma_y^2}}, \quad (\text{G.16})$$

gdzie $\sigma_y = a\sigma_x$.

Przejdziemy teraz do znalezienia na podstawie twierdzeń 1–4 funkcji rozkładu prawdopodobieństwa średniej arytmetycznej serii N niezależnych pomiarów, gdy poszczególne pomiary podlegają rozkładowi Gaussa. Z twierdzenia 2 i 3 wynika, że wielkość:

$$y = \sum_{i=1}^N x_i \quad (\text{G.17})$$

ma rozkład Gaussa o wariancji:

$$\sigma_y^2 = \sum_{i=1}^N \sigma_{x_i}^2. \quad (\text{G.18})$$

Ponieważ wszystkie zmienne losowe x_i ($i = 1, 2, \dots, N$) zostały otrzymane w jednej serii N pomiarów, $\sigma_{x_i}^2$ jest jednakowe dla wszystkich pomiarów i wynosi σ_x^2 , a więc wzór (G.18) przyjmie postać:

$$\sigma_y^2 = N\sigma_x^2. \quad (\text{G.19})$$

Z twierdzenia 4 wynika, że zmienna losowa:

$$\bar{x} = \frac{y}{N} \quad (\text{G.20})$$

ma rozkład Gaussa o wariancji:

$$\sigma_{\bar{x}}^2 = \frac{1}{N^2}\sigma_y^2. \quad (\text{G.21})$$

Podstawiając (G.19) do (G.21) dostajemy:

$$\sigma_{\bar{x}}^2 = \frac{1}{N}\sigma_x^2. \quad (\text{G.22})$$

Otrzymaliśmy więc na innej drodze wzór (3.4.11). Sposób jego otrzymania uzasadnia nazwę odchylenie standardowe średniej arytmetycznej.

Funkcja rozkładu prawdopodobieństwa średniej arytmetycznej $\varphi(\bar{x})$ będzie więc równa:

$$\varphi(\bar{x}) = \frac{1}{\sigma_{\bar{x}}\sqrt{2\pi}} \exp \left[-\frac{\bar{x}^2}{2\sigma_{\bar{x}}^2} \right]. \quad (\text{G.23})$$

Zgodnie z tym co powiedziano na początku tego Dodatku (po udowodnieniu twierdzenia 1) wzór (G.23) przyjmie postać

$$\varphi(\bar{t}) = \frac{1}{\sigma_{\bar{t}}\sqrt{2\pi}} \exp \left[-\frac{(\bar{t} - t_0)^2}{2\sigma_{\bar{t}}^2} \right]. \quad (\text{G.24})$$

H. Tablice statystyczne

W Dodatku zostały zamieszczone tablice niektórych funkcji ważnych w teorii błędów. Ograniczyliśmy się do podania tablic funkcji występujących w tym opracowaniu. Są to tablice:

- funkcji rozkładu normalnego standaryzowanego $\varphi_u(u)$,
- dystrybuanty rozkładu normalnego standaryzowanego $\Phi_u(a)$,
- całki błędu $\Phi_u(-a, a)$,
- współczynników t_α ,
- rozkładu Poissona dla małych wartości średnich,
- współczynników d_α :

Tablice mają na celu umożliwienie oceny poziomu i przedziału ufności na poziomie podstawowym. Zostały skonstruowane tak, aby pomiędzy podanymi wartościami można było stosować interpolację liniową.

W celu szybkiego i łatwego porównania oceny przedziału i poziomu ufności przy użyciu rozkładu normalnego standaryzowanego i rozkładu *t-Studenta* w tablicy współczynników t_α podano ich wartości dla poziomu ufności $\alpha = 0.6827, 0.9545$ oraz 0.9973 ; ponieważ odpowiadają one granicom przedziału ufności $\bar{x} \pm \sigma_{\bar{x}}, \bar{x} \pm 2\sigma_{\bar{x}}$ i $\bar{x} \pm 3\sigma_{\bar{x}}$ dla rozkładu $N(0, 1)$. W ten sam sposób ułożono tablicę współczynników d_α .

Uwaga:

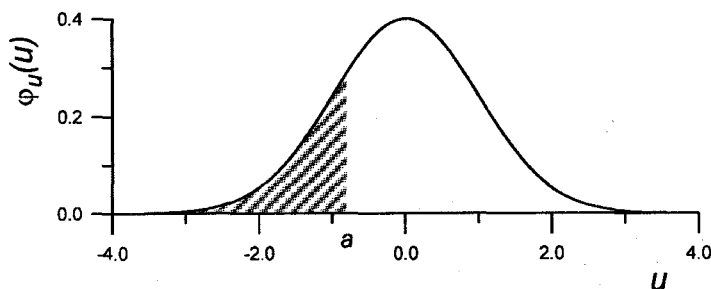
zapis 0.0_31140 oznacza 0.0001140 ,

zapis 0.9_4128 oznacza 0.9999128 .

TABLICA I: Wartości funkcji rozkładu normalnego standaryzowanego $N(0, 1)$. Tablica podaje wartości funkcji $\varphi_u(u) = (2\pi)^{-1/2} \exp(-u^2/2)$

u	0.00	0.02	0.04	0.05	0.06	0.08
0.0	0.3989	0.3989	0.3986	0.3984	0.3982	0.3977
0.1	0.3970	0.3961	0.3951	0.3945	0.3939	0.3925
0.2	0.3910	0.3894	0.3876	0.3867	0.3857	0.3836
0.3	0.3814	0.3790	0.3765	0.3752	0.3739	0.3712
0.4	0.3683	0.3653	0.3621	0.3605	0.3589	0.3555
0.5	0.3521	0.3485	0.3448	0.3429	0.3410	0.3372
0.6	0.3332	0.3292	0.3251	0.3230	0.3209	0.3166
0.7	0.3123	0.3079	0.3034	0.3011	0.2989	0.2943
0.8	0.2897	0.2850	0.2803	0.2780	0.2756	0.2709
0.9	0.2661	0.2613	0.2565	0.2541	0.2516	0.2468
1.0	0.2420	0.2371	0.2323	0.2299	0.2275	0.2227
1.1	0.2179	0.2131	0.2083	0.2059	0.2036	0.1989
1.2	0.1942	0.1895	0.1849	0.1826	0.1804	0.1758
1.3	0.1714	0.1669	0.1626	0.1604	0.1582	0.1539
1.4	0.1497	0.1456	0.1415	0.1394	0.1374	0.1334
1.5	0.1295	0.1257	0.1219	0.1200	0.1182	0.1145
1.6	0.1109	0.1074	0.1040	0.1023	0.1006	0.09728
1.7	0.09405	0.09089	0.08780	0.08628	0.08478	0.08183
1.8	0.07895	0.07614	0.07341	0.07206	0.07074	0.06814
1.9	0.06562	0.06316	0.06077	0.05959	0.05844	0.05618
2.0	0.05399	0.05186	0.04980	0.04879	0.04780	0.04586
2.1	0.04398	0.04217	0.04041	0.03955	0.03871	0.03706
2.2	0.03547	0.03394	0.03246	0.03174	0.03103	0.02965
2.3	0.02833	0.02705	0.02582	0.02522	0.02463	0.02349
2.4	0.02239	0.02134	0.02033	0.01984	0.01936	0.01842
2.5	0.01753	0.01667	0.01585	0.01545	0.01506	0.01431
2.6	0.01358	0.01289	0.01223	0.01191	0.01160	0.01100
2.7	0.01042	0.009871	0.009347	0.009094	0.008846	0.008370
2.8	0.007915	0.007483	0.007071	0.006873	0.006679	0.006307
2.9	0.005953	0.005616	0.005296	0.005143	0.004993	0.004705
3.0	0.004432	0.004173	0.003928	0.003810	0.003695	0.003475
3.1	0.003267	0.003070	0.002884	0.002794	0.002707	0.002541
3.2	0.002384	0.002236	0.002096	0.002029	0.001964	0.001840
3.3	0.001723	0.001612	0.001508	0.001459	0.001411	0.001319
3.4	0.001232	0.001151	0.001075	0.001038	0.001003	0.0009358
3.5	0.0008727	0.0008135	0.0007581	0.0007317	0.0007061	0.0006575
3.6	0.0006119	0.0005693	0.0005294	0.0005105	0.0004921	0.0004573
3.7	0.0004248	0.0003944	0.0003661	0.0003526	0.0003396	0.0003149
3.8	0.0002919	0.0002705	0.0002506	0.0002411	0.0002320	0.0002147
3.9	0.0001987	0.0001837	0.0001698	0.0001633	0.0001569	0.0001449
4.0	0.0001338	0.0001235	0.0001140	0.0001094	0.0001051	0.00009687

TABLICA II: Wartości dystrybuanty rozkładu normalnego standaryzowanego. Tablica podaje wartości funkcji $\phi_u(a) = (2\pi)^{-1/2} \int_{-\infty}^a \exp(-u^2/2) du$



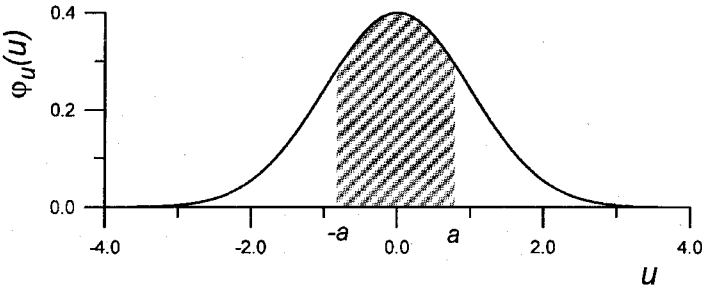
a	0.00	-0.02	-0.04	-0.05	-0.06	-0.08
-4.0	0.043167	0.042910	0.042673	0.042561	0.042454	0.042252
-3.9	0.044810	0.044427	0.044074	0.043908	0.043747	0.043446
-3.8	0.047235	0.046673	0.046152	0.045906	0.045669	0.045223
-3.7	0.051078	0.049961	0.049201	0.048842	0.048496	0.047841
-3.6	0.055191	0.053147	0.051363	0.050311	0.049261	0.048166
-3.5	0.059526	0.057158	0.054501	0.052601	0.050554	0.048318
-3.4	0.064093	0.061313	0.058299	0.055983	0.053471	0.050707
-3.3	0.068894	0.065645	0.062189	0.059441	0.056357	0.052982
-3.2	0.073927	0.070241	0.066476	0.063377	0.060001	0.056386
-3.1	0.079192	0.075073	0.070799	0.067221	0.063396	0.059271
-3.0	0.084697	0.080126	0.075491	0.071544	0.067341	0.062926
-2.9	0.090442	0.085473	0.080399	0.075981	0.071281	0.066346
-2.8	0.096437	0.091073	0.085599	0.080881	0.075881	0.070556
-2.7	0.102692	0.096923	0.090949	0.085641	0.080081	0.074226
-2.6	0.109217	0.103043	0.096769	0.090261	0.083501	0.076456
-2.5	0.116022	0.109443	0.102769	0.095981	0.088961	0.081686
-2.4	0.123107	0.116123	0.109049	0.101861	0.094541	0.087056
-2.3	0.130472	0.123183	0.115799	0.108301	0.100681	0.092906
-2.2	0.138117	0.130523	0.122849	0.115061	0.107141	0.099156
-2.1	0.146042	0.138143	0.130169	0.122081	0.113861	0.105556
-2.0	0.154247	0.145943	0.137569	0.129081	0.120461	0.111706
-1.9	0.162732	0.154023	0.145449	0.136801	0.128061	0.119156
-1.8	0.171497	0.162383	0.153409	0.144461	0.135421	0.126256
-1.7	0.180532	0.171023	0.162149	0.152801	0.143461	0.133956
-1.6	0.189837	0.180023	0.170749	0.161401	0.151961	0.142356
-1.5	0.199412	0.189383	0.179709	0.170081	0.160421	0.150656
-1.4	0.209257	0.198923	0.189049	0.179161	0.169261	0.159256
-1.3	0.219372	0.208723	0.198549	0.188301	0.177961	0.167456
-1.2	0.229757	0.218823	0.208349	0.197801	0.187161	0.176456
-1.1	0.240412	0.229183	0.218309	0.207461	0.196521	0.185456
-1.0	0.251337	0.239823	0.228749	0.217601	0.206361	0.195056
-0.9	0.262532	0.250723	0.239449	0.228081	0.216621	0.205156
-0.8	0.274007	0.261923	0.250349	0.238681	0.226921	0.215056
-0.7	0.285752	0.273423	0.261549	0.249681	0.237721	0.225756
-0.6	0.297767	0.285123	0.272849	0.260601	0.248361	0.235956
-0.5	0.310042	0.297123	0.284649	0.272201	0.259761	0.247256
-0.4	0.322577	0.309323	0.296749	0.284201	0.271661	0.259056
-0.3	0.335372	0.321823	0.309049	0.296301	0.283561	0.270756
-0.2	0.348427	0.334623	0.321549	0.308601	0.295661	0.282656
-0.1	0.361742	0.347623	0.334349	0.321001	0.307561	0.294056
0.0	0.500000	0.492020	0.484040	0.480100	0.476100	0.468100

Ciąg dalszy TABLICY II.

<i>a</i>	0.00	0.02	0.04	0.05	0.06	0.08
0.0	0.5000	0.5080	0.5160	0.5199	0.5239	0.5319
0.1	0.5398	0.5478	0.5557	0.5596	0.5636	0.5714
0.2	0.5793	0.5871	0.5948	0.5987	0.6026	0.6103
0.3	0.6179	0.6255	0.6331	0.6368	0.6406	0.6480
0.4	0.6554	0.6628	0.6700	0.6736	0.6772	0.6844
0.5	0.6915	0.6985	0.7054	0.7088	0.7123	0.7190
0.6	0.7257	0.7324	0.7389	0.7422	0.7454	0.7517
0.7	0.7580	0.7642	0.7704	0.7734	0.7764	0.7823
0.8	0.7881	0.7939	0.7995	0.8023	0.8051	0.8106
0.9	0.8159	0.8212	0.8264	0.8289	0.8315	0.8365
1.0	0.8413	0.8461	0.8508	0.8531	0.8554	0.8599
1.1	0.8643	0.8686	0.8729	0.8749	0.8770	0.8810
1.2	0.8849	0.8888	0.8925	0.8944	0.8962	0.8997
1.3	0.90320	0.90658	0.90988	0.91149	0.91309	0.91621
1.4	0.91924	0.92220	0.92507	0.92647	0.92785	0.93056
1.5	0.93319	0.93574	0.93822	0.93943	0.94062	0.94295
1.6	0.94520	0.94738	0.94950	0.95053	0.95154	0.95352
1.7	0.95543	0.95728	0.95907	0.95994	0.96080	0.96246
1.8	0.96407	0.96562	0.96712	0.96784	0.96856	0.96995
1.9	0.97128	0.97257	0.97381	0.97441	0.97500	0.97615
2.0	0.97725	0.97831	0.97932	0.97982	0.98030	0.98124
2.1	0.98214	0.98300	0.98382	0.98422	0.98461	0.98537
2.2	0.98610	0.98679	0.98745	0.98778	0.98809	0.98870
2.3	0.98928	0.98983	0.920358	0.920613	0.920863	0.921344
2.4	0.921802	0.922240	0.922656	0.922857	0.923053	0.923431
2.5	0.923790	0.924132	0.924457	0.924614	0.924766	0.925060
2.6	0.925339	0.925604	0.925855	0.925975	0.926093	0.926319
2.7	0.926533	0.926736	0.926928	0.927020	0.927110	0.927282
2.8	0.927445	0.927599	0.927744	0.927814	0.927882	0.928012
2.9	0.928134	0.928250	0.928359	0.928411	0.928462	0.928559
3.0	0.928650	0.928736	0.928817	0.928856	0.928893	0.928965
3.1	0.930324	0.930957	0.931553	0.931836	0.932112	0.932636
3.2	0.933129	0.933590	0.934024	0.934230	0.934429	0.934810
3.3	0.935166	0.935499	0.935811	0.935959	0.936103	0.936376
3.4	0.936631	0.936869	0.937091	0.937197	0.937299	0.937493
3.5	0.937674	0.937842	0.937999	0.938074	0.938146	0.938282
3.6	0.938409	0.938527	0.938637	0.938689	0.938739	0.938834
3.7	0.938922	0.940039	0.940799	0.941158	0.941504	0.942159
3.8	0.942765	0.943327	0.943848	0.944094	0.944331	0.944777
3.9	0.945190	0.945573	0.945926	0.946092	0.946253	0.946554
4.0	0.946833	0.947090	0.947327	0.947439	0.947546	0.947748

TABLICA III: Wartości całki błędu. Tablica podaje wartości całki:

$$\Phi_u(-a, a) = (2/\pi)^{1/2} \int_{-a}^a \exp(-u^2/2) du.$$



a	0.00	0.02	0.04	0.05	0.06	0.08
0.0	0.0	0.01596	0.03191	0.03988	0.04784	0.06376
0.1	0.07966	0.09552	0.1113	0.1192	0.1271	0.1428
0.2	0.1585	0.1741	0.1897	0.1974	0.2051	0.2205
0.3	0.2358	0.2510	0.2661	0.2737	0.2812	0.2961
0.4	0.3108	0.3255	0.3401	0.3473	0.3545	0.3688
0.5	0.3829	0.3969	0.4108	0.4177	0.4245	0.4381
0.6	0.4515	0.4647	0.4778	0.4843	0.4907	0.5035
0.7	0.5161	0.5285	0.5407	0.5467	0.5527	0.5646
0.8	0.5763	0.5878	0.5991	0.6047	0.6102	0.6211
0.9	0.6319	0.6424	0.6528	0.6579	0.6629	0.6729
1.0	0.6827	0.6923	0.7017	0.7063	0.7109	0.7199
1.1	0.7287	0.7373	0.7457	0.7499	0.7540	0.7620
1.2	0.7699	0.7775	0.7850	0.7887	0.7923	0.7995
1.3	0.8064	0.8132	0.8198	0.8230	0.8262	0.8324
1.4	0.8385	0.8444	0.8501	0.8529	0.8557	0.8611
1.5	0.8664	0.8715	0.8764	0.8789	0.8812	0.8859
1.6	0.8904	0.8948	0.8990	0.90106	0.90309	0.90704
1.7	0.91087	0.91457	0.91814	0.91988	0.92159	0.92492
1.8	0.92814	0.93124	0.93423	0.93569	0.93711	0.93989
1.9	0.94257	0.94514	0.94762	0.94882	0.95000	0.95230
2.0	0.95450	0.95662	0.95865	0.95964	0.96060	0.96247
2.1	0.96427	0.96599	0.96765	0.96844	0.96923	0.97074
2.2	0.97219	0.97358	0.97491	0.97555	0.97618	0.97739
2.3	0.97855	0.97966	0.98072	0.98123	0.98173	0.98269
2.4	0.98360	0.98448	0.98531	0.98571	0.98611	0.98686
2.5	0.98758	0.98826	0.98891	0.98923	0.98953	0.920120
2.6	0.920678	0.921207	0.921709	0.921951	0.922186	0.922638
2.7	0.923066	0.923472	0.923856	0.924040	0.924220	0.924564
2.8	0.924890	0.925198	0.925489	0.925628	0.925764	0.926023
2.9	0.926268	0.926500	0.926718	0.926822	0.926924	0.927118
3.0	0.927300	0.927472	0.927634	0.927712	0.927787	0.927930
3.1	0.928065	0.928191	0.928311	0.928367	0.928422	0.928527
3.2	0.928626	0.928718	0.928805	0.928846	0.928886	0.928962
3.3	0.930332	0.930998	0.931622	0.931919	0.932206	0.932751
3.4	0.933261	0.933738	0.934183	0.934394	0.934598	0.934986
3.5	0.935347	0.935685	0.935999	0.936148	0.936291	0.936564
3.6	0.936818	0.937054	0.937274	0.937378	0.937478	0.937668
3.7	0.937844	0.938008	0.938160	0.938232	0.938301	0.938432
3.8	0.938553	0.938665	0.938770	0.938819	0.938866	0.938955
3.9	0.940381	0.941145	0.941852	0.942185	0.942505	0.943108
4.0	0.943666	0.944180	0.944655	0.944878	0.945093	0.945496

TABLICA IV: Wartości współczynników t_{α} -Studenta. Tablica podaje wartości współczynników $t_{\alpha,k}$ spełniających równości:

$$\alpha = P(-t_{\alpha,k} \leq t \leq t_{\alpha,k}) = 2 \int_0^{t_{\alpha,k}} S(t,k) dt.$$

$\alpha \backslash k$	0.50	0.60	0.70	0.80	0.90	0.99	0.6827	0.9545	0.9973	$\alpha \backslash k$
1	1.000	1.376	1.963	3.078	6.314	63.657	1.837	13.968	235.784	1
2	0.816	1.061	1.386	1.886	2.920	9.925	1.321	4.527	19.206	2
3	0.765	0.978	1.250	1.638	2.353	5.841	1.197	3.307	9.219	3
4	0.741	0.941	1.190	1.533	2.132	4.604	1.142	2.869	6.620	4
5	0.727	0.920	1.156	1.476	2.015	4.032	1.111	2.649	5.507	5
6	0.718	0.906	1.134	1.440	1.943	3.707	1.091	2.517	4.904	6
7	0.711	0.896	1.119	1.415	1.895	3.499	1.077	2.429	4.530	7
8	0.706	0.889	1.108	1.397	1.860	3.355	1.067	2.366	4.277	8
9	0.703	0.883	1.100	1.383	1.833	3.250	1.059	2.320	4.094	9
10	0.700	0.879	1.093	1.372	1.812	3.169	1.053	2.284	3.957	10
11	0.697	0.876	1.088	1.363	1.796	3.106	1.048	2.255	3.850	11
12	0.695	0.873	1.083	1.356	1.782	3.055	1.043	2.231	3.764	12
13	0.694	0.870	1.079	1.350	1.771	3.012	1.040	2.212	3.694	13
14	0.692	0.868	1.076	1.345	1.761	2.977	1.037	2.195	3.636	14
15	0.691	0.866	1.074	1.341	1.753	2.947	1.034	2.181	3.586	15
16	0.690	0.865	1.071	1.337	1.746	2.921	1.032	2.169	3.544	16
17	0.689	0.863	1.069	1.333	1.740	2.898	1.030	2.158	3.507	17
18	0.688	0.862	1.067	1.330	1.734	2.878	1.029	2.149	3.475	18
19	0.688	0.861	1.066	1.328	1.729	2.861	1.027	2.140	3.447	19
20	0.687	0.860	1.064	1.325	1.725	2.845	1.026	2.133	3.422	20
21	0.686	0.859	1.063	1.323	1.721	2.831	1.024	2.126	3.400	21
22	0.686	0.858	1.061	1.321	1.717	2.819	1.023	2.120	3.380	22
23	0.685	0.858	1.060	1.319	1.714	2.807	1.022	2.115	3.361	23
24	0.685	0.857	1.059	1.318	1.711	2.797	1.021	2.110	3.345	24
25	0.684	0.856	1.058	1.316	1.708	2.787	1.020	2.105	3.330	25
26	0.684	0.856	1.058	1.315	1.706	2.779	1.020	2.101	3.316	26
27	0.684	0.855	1.057	1.314	1.703	2.771	1.019	2.097	3.303	27
28	0.683	0.855	1.056	1.313	1.701	2.763	1.018	2.093	3.291	28
29	0.683	0.854	1.055	1.311	1.699	2.756	1.018	2.090	3.280	29
30	0.683	0.854	1.055	1.310	1.697	2.750	1.017	2.087	3.270	30
40	0.681	0.851	1.050	1.303	1.684	2.704	1.013	2.064	3.199	40
50	0.679	0.849	1.047	1.299	1.676	2.678	1.010	2.051	3.157	50
60	0.679	0.848	1.045	1.296	1.671	2.660	1.008	2.043	3.130	60
70	0.678	0.847	1.044	1.294	1.667	2.648	1.007	2.036	3.111	70
80	0.678	0.846	1.043	1.292	1.664	2.639	1.006	2.032	3.096	80
90	0.677	0.846	1.042	1.291	1.662	2.632	1.006	2.028	3.085	90
100	0.677	0.845	1.042	1.290	1.660	2.626	1.005	2.025	3.077	100
∞	0.674	0.842	1.036	1.282	1.645	2.576	1.000	2.000	3.000	∞

TABLICA VI: Wartości współczynników d_{α} . Tablica podaje wartości współczynników $d_{\alpha, \varepsilon/\sigma}$ spełniających równości:

$$\alpha = P(-d_{\alpha, \varepsilon/\sigma} \leq u \leq d_{\alpha, \varepsilon/\sigma}) = 2 \int_0^{d_{\alpha, \varepsilon/\sigma}} \xi(u; \varepsilon/\sigma) du.$$

α ε/σ	0.50	0.60	0.70	0.80	0.90	0.99	0.6827	0.9545	0.9973	α ε/σ
0	0.674	0.842	1.036	1.282	1.645	2.576	1.000	2.000	3.000	0
0.10	0.676	0.843	1.038	1.284	1.647	2.579	1.002	2.003	3.003	0.10
0.20	0.679	0.847	1.043	1.290	1.656	2.592	1.007	2.013	3.018	0.20
0.30	0.685	0.854	1.052	1.301	1.669	2.613	1.015	2.030	3.042	0.30
0.40	0.693	0.864	1.064	1.315	1.688	2.642	1.027	2.052	3.075	0.40
0.50	0.703	0.877	1.079	1.335	1.712	2.678	1.042	2.081	3.116	0.50
0.60	0.715	0.892	1.098	1.357	1.741	2.719	1.060	2.115	3.162	0.60
0.70	0.730	0.910	1.120	1.384	1.774	2.768	1.081	2.155	3.216	0.70
0.80	0.746	0.931	1.145	1.415	1.812	2.820	1.105	2.199	3.273	0.80
0.90	0.765	0.954	1.174	1.449	1.854	2.877	1.133	2.247	3.335	0.90
1.00	0.786	0.980	1.205	1.486	1.899	2.938	1.163	2.299	3.401	1.00
1.10	0.810	1.009	1.239	1.527	1.948	3.001	1.196	2.354	3.469	1.10
1.20	0.835	1.039	1.276	1.570	2.000	3.068	1.232	2.413	3.540	1.20
1.30	0.862	1.073	1.315	1.616	2.055	3.137	1.270	2.474	3.612	1.30
1.40	0.891	1.108	1.357	1.665	2.112	3.208	1.311	2.537	3.689	1.40
1.50	0.922	1.145	1.401	1.716	2.171	3.280	1.353	2.603	3.765	1.50
1.60	0.955	1.184	1.447	1.769	2.232	3.354	1.398	2.670	3.842	1.60
1.70	0.989	1.225	1.494	1.824	2.295	3.430	1.445	2.739	3.920	1.70
1.80	1.025	1.268	1.544	1.880	2.360	3.506	1.493	2.809	4.001	1.80
1.90	1.062	1.312	1.595	1.939	2.426	3.584	1.542	2.881	4.082	1.90
2.00	1.100	1.357	1.647	1.998	2.493	3.663	1.594	2.953	4.164	2.00
2.20	1.180	1.451	1.755	2.120	2.631	3.823	1.700	3.102	4.330	2.20
2.40	1.264	1.550	1.868	2.247	2.773	3.986	1.810	3.254	4.498	2.40
2.60	1.351	1.652	1.984	2.378	2.918	4.152	1.924	3.409	4.669	2.60
2.80	1.440	1.757	2.104	2.511	3.066	4.320	2.041	3.567	4.843	2.80
3.00	1.531	1.864	2.226	2.647	3.216	4.490	2.160	3.727	5.017	3.00
3.50	1.767	2.140	2.540	2.997	3.602	4.921	2.468	4.135	5.459	3.50
4.00	2.009	2.425	2.864	3.357	3.998	5.360	2.785	4.553	5.908	4.00
4.50	2.254	2.715	3.195	3.726	4.402	5.805	3.110	4.978	6.363	4.50
5.00	2.502	3.009	3.531	4.101	4.812	6.255	3.439	5.409	6.822	5.00
5.50	2.751	3.305	3.872	4.480	5.227	6.710	3.772	5.844	7.284	5.50
6.00	3.000	3.603	4.215	4.864	5.647	7.167	4.107	6.284	7.749	6.00
7.00	3.500	4.201	4.907	5.640	6.497	8.090	4.784	7.173	8.686	7.00
8.00	4.000	4.800	5.603	6.425	7.357	9.021	5.463	8.072	9.631	8.00
9.00	4.500	5.400	6.301	7.215	8.226	9.959	6.145	8.979	10.581	9.00
10.00	5.000	6.000	7.000	8.009	9.101	10.902	6.827	9.893	11.536	10.00

Literatura

- [1] J. M. Masalski, J. Studnicki, *Legalne jednostki miar i stałe fizyczne*, PWN, Warszawa 1988.
- [2] J. M. Massalski, *Postępy fizyki*, **48**, 227(1997).
- [3] *Słownik języka polskiego*, wyd. X, Warszawa 1995, t. III, s. 867.
- [4] J. R. Taylor, *Wstęp do analizy błędu pomiarowego*, Wydawnictwo Naukowe PWN, Warszawa 1995.
- [5] E. R. Cohen, B. N. Taylor, *Physics Today* **49**, (1996), suplement s. BG9.
- [6] W. Rubinowicz, W. Królikowski, *Mechanika teoretyczna*, Wydawnictwo Naukowe PWN, Warszawa 1995.
- [7] G. L. Squires, *Praktyczna fizyka*, Wydawnictwo Naukowe PWN, Warszawa 1992.
- [8] F. Leja, *Rachunek różniczkowy i całkowy*, PWN, Warszawa 1954.
- [9] L. Górniewicz, R. S. Ingarden, *Analiza matematyczna dla fizyków*, PWN, Warszawa 1981, t. I.
- [10] H. Szydłowski, *Pracownia fizyczna*, PWN, Warszawa 1979.
- [11] T. Dryński, *Ćwiczenia laboratoryjne z fizyki*, PWN, Warszawa 1980.
- [12] A. Zawadzki, H. Hofmokr, *Laboratorium fizyczne*, PWN, Warszawa 1968.
- [13] *Teoria pomiarów*, praca zb. pod red. H. Szydłowskiego, PWN, Warszawa 1981.
- [14] T. Czechowski, *Elementarny wykład rachunku prawdopodobieństwa*, PWN, Warszawa 1968, s. 126.
- [15] C. R. Rao, *Modele liniowe statystyki matematycznej*, PWN, Warszawa 1982.
- [16] E. B. Wilson jr, *Wstęp do badań naukowych*, PWN, Warszawa 1968.
- [17] M. G. Kendall, A. Stuard, *The Advanced Theory o Statistics*, 2nd ed., vol. 2, Charles Griffin, London 1967.
- [18] W. I. Romanowski, *Podstawowe zagadnienia teorii błędów*, PWN, Warszawa 1955.
- [19] A. Strzałkowski, A. Śliżyński, *Matematyczne metody opracowania wyników pomiarów*, PWN, Warszawa 1978.
- [20] G. T. Gillies, *Reports on Progress in Physics* **60**, 151 (1997).

- [21] J. Łukaszewicz, M. Warmus, *Metody numeryczne i graficzne*, PWN, Warszawa 1956.
- [22] H. Szydłowski, *Wstęp do pracowni fizycznej*, Wydawnictwo Naukowe UAM, Poznań 1996.
- [23] S. Brandt, *Metody statystyczne i obliczeniowe analizy danych*, PWN, Warszawa 1976.
- [24] S. Chandrasekhar, *Rev. Mod. Phys.* **15**, 1 (1943).

Skorowidz

aproksymacja metodą najmniejszych kwadratów 67

błąd 11

bezwzględny 12, 27, 29, 41, 42, 86, 88, 105, 163, 164

gruby 16, 21

maksymalny 12, 24, 26, 27, 29, 30, 31, 35, 102, 130, 154, 156, 158, 161, 163, 166, 168

średniej arytmetycznej 113, 161, 162

wielkości złożonej 162

względny 25, 30, 31

odczytanej wartości 159

pomiarowy 156

pomiaru 10, 11, 13, 51, 105, 112, 131

pośredniego 27

procentowy 12, 13

przypadkowy 16, 18, 26, 35, 38, 48, 50, 88, 89, 113, 121, 162, 164, 165, 168

systematyczny 16, 18, 20, 24, 26, 27, 30, 35, 92, 161–163, 168

wielkości odczytanej z wykresu 156

wykreślonej krzywej 156

względny 12, 13

centralne twierdzenie graniczne 119–121, 123

cyfra pewna 13

cyfra znacząca 13

częstość 70, 78–80, 86

długość przedziału klasowego 78, 80
dokładność

miernika 76

odczytu 24

przyrządu 18, 24, 27, 161, 163, 164

dopasowanie

funkcji liniowych 56

innych krzywych 65

prostej 58

dystribuanta 75, 76, 82, 84, 121, 143
rozkładu 74

normalnego 93, 165

normalnego

standaryzowanego 111, 116, 118–120, 122, 158

Poissona 143

t-Studenta 129, 131

zmiennej losowej ciągłej 82

funkcja błędu 112

funkcja rozkładu 81, 82, 91, 93, 110, 121, 163

średniej arytmetycznej 98, 99, 110, 122

błędów bezwzględnych 89

prawdopodobieństwa 76, 81, 88

wielkości złożonej 103, 107

gęstość prawdopodobieństwa 81

graficzne przedstawienie wyników 148

granica średniej arytmetycznej serii 86

histogram 69–72, 75, 76, 79, 80, 86, 88, 150

komórkowy 79, 87

słupkowy 71, 77, 87

jednostka fizyczna 145

jednostka miary 9

jednostki podwielokrotne 15, 16

jednostki wielokrotne 15, 16

klasa przyrządu 24

korelacja 52

kowariancja 51, 57

krzywa błędu 157–159

liczba stopni swobody 43, 57, 60, 128–130, 132, 133

liczebność 69, 141

przedziału 78

klasowego 79

- serii 86
- występowania wyniku 70
- metoda
 - najmniejszych kwadratów 39, 54, 65, 119, 123, 125, 126, 154-156
 - największej wiarygodności 96, 123, 126
 - pochodnej logarytmicznej 31
 - różniczki zupełnej 27, 30
- minimalizacja 39, 125, 127
- mnożenie jednostek miar 146
- model błędów przypadkowych Laplace'a 90
- nagłówek kolumny 147
- nagłówek tabeli 147
- niepewność 11
 - przypadkowa 18
- niezależne
 - próby 135, 136
 - serie pomiarowe 100
 - wielkości 168
 - wielkości fizyczne 105
 - zmiennie losowe 103, 106, 107
- normalna całka błędu 112
- normalny rozkład błędów 88, 90
- ocena błędu 35
- odchylenie standardowe 47, 57, 66, 73, 86, 97, 110, 119, 154, 156, 158, 162, 166
- efektywne 169
- pojedynczego pomiaru serii 41
- serii 41, 43, 48, 86, 100, 118
 - pomiarów bezpośrednich 40, 50
 - pomiarów pośrednich 46, 48, 49
- średniej 101, 120
 - arytmetycznej 43, 48, 168
 - serii 100, 102, 119
 - ważonej 102
- papier
 - funkcyjny 151, 154, 156, 159
 - logarytmiczny 152
 - milimetrowy 152
 - półlogarytmiczny 152
 - z siatką biegunową 152
 - z siatką logarytmiczną 152
 - z siatką milimetrową 152
 - z siatką półlogarytmiczną 152
- parametr funkcji 56
- parametr rozkładu 107
- podpis pod rysunkiem 154
- podział błędów 16
- pomiar 9-11
 - bezpośredni 10, 24, 26, 35, 113, 116, 123, 161
 - dokładny 21
 - pośredni 10, 27, 35, 43, 110, 116, 162
 - precyzyjny 21
- pomiary
 - niezależne 43, 45, 46, 49, 117
 - o niejednakowej precyzji 100, 126
 - pośrednie 119, 123
 - sprzeczne 21
 - wielkości złożonej 10
 - zależne 45, 49, 50, 119
 - zgodne 21
- populacja 68
 - generalna, 68, 89, 141, 167
- porażka 135
- poziom istotności 113
- poziom ufności 113, 114, 116-120, 122, 128, 130, 132, 142, 143, 158, 163, 165-169
- próba 68
 - losowa 68, 72, 98
- prawo błędów Gaussa 88-90
- precyzja pomiaru 41
- prosta regresji zmiennych 63
- prostokąt błędów 154, 157, 158
- przedstawianie wyników 35
- przedział
 - klasowy 78, 80, 85-88
 - ufności 113, 114, 116-119, 122, 128, 130-132, 142, 143, 158, 163, 165, 167-169
 - zmienności 35, 69, 77, 78, 131
- przenoszenie błędów kwadratowe 47

- przenoszenie błędów liniowe 47
- regresja liniowa 63, 120
- reguła przenoszenia błędów 30, 43, 47, 48
- maksymalnych 162
- przypadkowych 57
- rodzaje tabel 148
- rozkład 81
- Bernoulliego 135, 136
- dwumianowy 90, 135, 136, 137, 139, 141
- Gaussa 88, 90, 92, 93, 98, 103, 108, 113, 123, 126, 135, 163, 164, 167
- dla wyników pomiarów 91
- jednostajny 163
- liczby zliczeń 139
- normalny 88, 90, 99, 105–107, 110, 113, 116, 118, 120, 122–124, 126, 128, 130–132, 141, 142, 161, 165, 166
- standaryzowany 108, 109, 111, 122, 129, 143
- Poissona 135, 139–141
- prawdopodobieństwa 72
- zmiennej losowej skokowej 73
- prostokątny 163–166
- wielkości złożonej 116
- wyników pomiarów 70
- zmiennej losowej 106
- ciągłej 75
- skokowej 69, 72
- t-Studenta* 88, 128–130, 131, 135, 158, 161
- rozrzut wyników pomiaru 44, 161
- rozstęp 78
- rysunek 145
- rzeczywista zależność funkcyjna 67
- seria pomiarowa 11, 35, 41, 43, 69, 98
- siatka 151
- biegunowa 151
- liniowa 151
- logarytmiczna 151
- milimetrowa 151, 154
- półlogarytmiczna 151, 154
- skończona dokładność pomiarów 10
- splot rozkładów 164, 166
- zmiennych losowych 106
- sukces 135
- symbol jednostki miar 146
- symbol wielkości fizycznej skalarnej 145
- symbol wielkości fizycznej wektorowej 145
- szereg rozdzielczy zmiennej losowej ciągłej 77
- średnia arytmetyczna 38–41, 43–45, 48, 70, 86, 95, 97, 98, 100, 101, 107, 110, 118, 119, 122, 128, 143, 154, 161, 166
- serii 85
- pomiarów pośrednich 44, 45
- wielkości złożonej 117
- średnia
- geometryczna 38
- harmoniczna 38
- kwadratowa 38, 39, 41
- rozkładu 73, 86, 91, 92, 122, 137, 141, 143
- zmiennej losowej 73
- ważona 40, 100–102, 154
- średni błąd kwadratowy pojedynczego pomiaru 41
- tabela 145, 147
- funkcyjna 148
- jakościowa 148
- statystyczna 148
- tabelaryczne przedstawianie danych 146
- tablica 147
- waga pomiaru 40, 101, 127
- wariancja 73, 84, 93, 97, 98, 107, 121–124, 126, 128, 138, 141
- średniej serii 167
- efektywna 167, 169
- rozkładu 86, 92, 95
- Gaussa 92
- normalnego 167
- Poissona 141–143
- prostokątnego 167
- serii 41, 86, 95

pomiarów pośrednich 49
 wartość
 dokładna 11
 najbardziej wiarygodna 96
 oczekiwana zmiennej losowej 73
 odczytana z wykresu 156, 158
 prawdziwa 11
 rzeczywista 12, 27, 38, 41, 88,
 92-95, 97, 98, 100, 101, 103,
 113, 118, 123, 124, 126, 166
 rzeczywista wielkości złożonej
 46, 103, 118
 średnia
 rozkładu 84
 serii pomiarów bezpośrednich
 35
 serii pomiarów pośrednich 43
 wielkości złożonej 45
 z różnych serii pomiarowych
 39
 wyznaczona 27
 zmierzona 12
 warunek normalizacyjny 70, 80-82,
 89
 warunek zgodności 21
 wielkość
 fizyczna 9-11, 13, 30, 43, 98, 113,
 122, 123, 126, 130, 145
 mierzona 116
 bezpośrednio 10, 43, 45, 48,
 107, 116, 118, 132
 wyznaczona 10
 złożona 10, 27, 43, 44, 103, 107,
 110, 116, 118, 131, 132
 wielkości
 niezależne 43, 47, 49, 120
 skorelowane 52
 zależne mierzone bezpośrednio
 49, 52
 współczynnik
 istotności 113
 korelacji liniowej 51
 ufności 113
 wyniki sprzeczne 21, 102
 wyniki zgodne 21, 102

zależne wielkości mierzone bezpo-
 średnio 49
 zależność empiryczna 67
 zaokrąglanie błędów 14
 zaokrąglanie wyniku pomiaru 13, 14
 zdarzenie 135
 zdarzenia niezależne 96
 zgodność wyników 21, 102
 zliczanie zdarzeń, 139
 zmienna 71
 ciągła 82
 losowa 71, 75, 79, 88, 98, 103,
 106, 108, 110, 164
 ciągła 75-77, 84, 85, 88, 123,
 135
 dyskretna 72
 skokowa 72, 76, 77, 79, 86, 135
 standaryzowana 108, 110-112,
 117, 118, 165
 zmienna standaryzowana 110, 143,
 144, 166



zależne wielkości fizyczne 156

BIBLIOTEKA
Instytutu Fizyki
UMK w Toruniu

M. 2718

mfn 15666/98



ISBN 83-231-0910-9