

Analiza niepewności pomiarowych

Andrzej Majhofer

1 Plan wykładu

- Wprowadzenie. Polecana literatura
- Charakterystyki zbiorów danych liczbowych: średnia, średnie odchylenie kwadratowe (dyspersja rozkładu). Graficzna prezentacja i analiza danych (histogram).
- Definicje: pomiar, błąd pomiaru, niepewność pomiaru. Klasyfikacja: błąd grubo, błąd przypadkowy, błąd systematyczny.
- Elementy rachunku prawdopodobieństwa i statystyki matematycznej.
- Składowa przypadkowa niepewności pomiaru (błąd przypadkowy). Ocena parametrów rozkładu błędów przypadkowych. Uwzględnienie dokładności przyrządów. Wynik końcowy i niepewność bezpośredniego pomiaru pojedynczej wielkości fizycznej.
- Pomiar pośredni i propagacja niepewności.
- Dopasowanie zależności funkcyjnej do danych pomiarowych.
- Statystyczne metody testowania hipotez.
- Eliminacja wpływu efektów systematycznych i wprowadzanie poprawek.
- Dobra praktyka laboratoryjna.
- Publikacja wyników: publikacja w czasopiśmie naukowym, raport laboratoryjny...

2 Literatura

1. J. R. Taylor, **Wstęp do analizy błędu pomiarowego**, Wydawnictwo Naukowe PWN, Warszawa, 1995.
2. G. L. Squires, **Praktyczna fizyka**, Wydawnictwo Naukowe PWN, Warszawa, 1992.
3. H. Abramowicz, **Jak analizować wyniki pomiarów**, Wydawnictwo Naukowe PWN, Warszawa, 1992.
4. A. Zięba, **Analiza danych w naukach ścisłych i technice**, Wydawnictwo Naukowe PWN, Warszawa, 2013.

3 Literatura uzupełniająca

1. S. Brandt, **Analiza danych**, Wydawnictwo Naukowe PWN, Warszawa, 1998.
2. W. T. Eadie, D. Drijard, F. E. James, M. Roos i B. Sadoulet, **Metody statystyczne w fizyce doświadczalnej**, PWN, Warszawa, 1989.
3. J. J. Jakubowski i R. Sztencel, **Wstęp do teorii prawdopodobieństwa**, SCRIPT, Warszawa, 2001.
4. W. Feller, **Wstęp do rachunku prawdopodobieństwa**, PWN, Warszawa, 1977.
5. R. Nowak, **Statystyka dla fizyków**, Wydawnictwo Naukowe PWN, Warszawa, 2002.

4 Charakterystyki zbiorów danych liczbowych

Dla zbioru danych (pełna populacja):

$$i \rightarrow x_i, i = 1, 2, \dots, N$$

definiujemy

średnią:

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

inaczej:

$$\bar{x} = \frac{1}{N}(x_1 + x_2 + x_3 + \dots + x_N);$$

średnie odchylenie kwadratowe:

$$\sigma^2 = \frac{1}{N} \sum_{i=1}^{i=N} (x_i - \bar{x})^2;$$

wielkość $\sigma := \sqrt{\sigma^2}$ nazywana jest też **dyspersją rozkładu**;

medianę:

Mediana jest to liczba dzieląca *uporządkowany* zbiór wartości x_i na dwa równoliczne podzbiory (na ogół nie jest określona jednoznacznie). *Gdy liczba N danych wartości x_i jest parzysta, często przyjmuje się dodatkowo, że mediana równa jest średniej arytmetycznej $x_{N/2}$ i $x_{1+N/2}$.*

5 Bezpośredni pomiar pojedynczej wielkości fizycznej

Pomiar $\rightarrow$ porównanie z wzorcem

Źródła odchylenia wyniku od wartości dokładnej:

- metoda pomiaru;
- wystąpienie błędów grubych i systematycznych;
- sposób postępowania $\rightarrow$ ustalenie i kontrolowanie warunków pomiaru;
- jakość przyrządów $\rightarrow$ „odtwarzania wzorca”.

Definicja:

błąd pomiaru = wartość zmierzona – wartość dokładna

Główne typy odchyłeń od wartości dokładnej (błędy):

- **błąd gruby:**

pomyłka zapisu, źle odczytany zakres miernika, zmierzenie nie tej wielkości co trzeba, awaria aparatury (np. przerwy w zasilaniu....);

unikanie i eliminowanie błędów grubych: staranność postępowania i szczegółowe dokumentowanie przebiegu pomiaru;

- **błąd systematyczny:**

odchylenie wyniku od wartości dokładnej mające tę samą wartość przy powtarzaniu pomiaru w tych samych warunkach (np. temperatura otoczenia różna od temperatury kalibracji przyrządów, błąd wskazań miernika....);

ocena wielkości błędu systematycznego: możliwa tylko poprzez wykonywanie pomiarów różnymi metodami lub poprzez badanie zależności wyniku od zmian warunków wykonywania pomiarów (→ wprowadzenie poprawek);

- **błąd przypadkowy:**

odchylenie wyniku pomiaru od wartości dokładnej przyjmujące różne wartości podczas powtarzania pomiarów w tych samych warunkach;

błędu przypadkowego nie można całkowicie wyeliminować, ale można ocenić parametry rozkładu pojawiających się wartości odchyłeń z nim związanych → model matematyczny.

Założenia matematycznego modelu błędu pomiaru:

Przyjmijmy oznaczenia:

- $\mu =$ *wartość dokładna mierzonej wielkości*
(zakładamy, że istnieje i jest jednoznacznie określona,
mianowaną liczbą rzeczywistą);
- $x =$ *wynik pomiaru*
 $x = \mu + \varepsilon$
 $\varepsilon =$ *błąd pomiaru*

Nasze rozważania rozpoczniemy od opisu błędów przypadkowych (zakładamy, że wyeliminowane są błędy grube i systematyczne):

- $\varepsilon =$ *błąd przypadkowy* powstaje w wyniku złożenia działania *bardzo wielu* czynników *niezależnych*, a każdy z tych czynników daje *bardzo mały* wkład do odchylenia wyniku pomiaru od wartości dokładnej (będziemy poszukiwali granicy, gdy liczba czynników dąży do nieskończoności, a wkład każdego z nich dąży do zera);
- do opisu błędów przypadkowych posłużymy się rachunkiem prawdopodobieństwa.

Prawdopodobieństwo w języku potocznym rozumiane bywa jako:

- miara naszej niewiedzy
- stopień przekonania
- średnia częstość występowania zjawiska w długiej serii powtórzeń.

Posługiwanie się „intuicyjnym” rozumieniem pojęcia prawdopodobieństwa prowadzi do paradoksów, gdy próbujemy wyznaczyć prawdopodobieństwo zdarzeń złożonych.

Dla uniknięcia takich trudności posłużymy się aksjomatycznym sformuowaniem teorii (rachunku) prawdopodobieństwa.

Analogia:

Rachunek prawdopodobieństwa jest aksjomatyczną teorią *zdarzeń przypadkowych (losowych)* tak, jak

geometria euklidesowa jest aksjomatyczną teorią *figur geometrycznych*.

Obie są niezależne od możliwych zastosowań, chociaż mogą być bardzo użyteczne.

6 Przypomnienie podstaw rachunku prawdopodobieństwa

Niech:

Ω będzie przestrzenią zdarzeń

Prawdopodobieństwo:

Prawdopodobieństwem nazywamy odwzorowanie:

$$\Omega \supset A \rightarrow p(A) \in [0, 1] ,$$

spełniające warunki:

1. A – zdarzenie pewne $\rightarrow p(A) = 1$;
2. A – zdarzenie niemożliwe $\rightarrow p(A) = 0$;
3. $\Omega \supset A, B$ oraz $A \cap B = \emptyset$, to:

$$p(A \cup B) = p(A) + p(B).$$

Definiuje się też:

$A' = \Omega \setminus A$, zdarzenie przeciwne do A :
 $p(A') = 1 - p(A)$.

W szczególności: $p(\Omega) = 1$, $p(\Omega') = 0$.

Prawdopodobieństwo warunkowe $p(A|B)$

Prawdopodobieństwo zajścia zdarzenia A pod warunkiem zajścia zdarzenia B , $p(B) \neq 0$:

$$p(A|B) = \frac{p(A \cap B)}{p(B)}$$

Zdarzenia A i B są **niezależne**, gdy:

$$p(A|B) = p(A) \rightarrow p(A \cap B) = p(A)p(B).$$

Uwaga: Gdy zdefiniowana jest przestrzeń zdarzeń i określone zostało na niej prawdopodobieństwo $p(\cdot)$, to stwierdzenie, że dane dwa zdarzenia są niezależne następuje na podstawie wyniku obliczeń.

Przykład 1.

Rzucamy trzy razy symetryczną monetą. Każdy z możliwych wyników (ciąg reszek, R , i orłów, O) jest jednakowo prawdopodobny. Czy wyrzucenie orła w drugim rzucie (zdarzenie A) jest zdarzeniem niezależnym od wyrzucenia orła w rzucie pierwszym (zdarzenie B)?

- Zaczynamy od zbudowania przestrzeni zdarzeń:

$$\Omega = \{(O, O, O), (O, O, R), (O, R, O), (R, O, O), \\ (O, R, R), (R, O, R), (R, R, O), (R, R, R)\}.$$

- Każde ze zdarzeń wymienionych w definicji Ω jest jednakowo prawdopodobne. Zdarzeń jest 8, a więc prawdopodobieństwo każdego z nich wynosi $1/8$.
- Określamy prawdopodobieństwo zdarzenia *w pierwszym rzucie orzeł*. Zdarzenie to jest sumą (teoriomnogościową) czterech różnych zdarzeń wymienionych w definicji Ω .

Mamy więc:

$$P((O, \cdot, \cdot)) = 4 \cdot \frac{1}{8} = \frac{1}{2} =: P(B).$$

Analogicznie dla zdarzenia *orzeł w drugim rzucie*:

$$P((\cdot, O, \cdot)) = 4 \cdot \frac{1}{8} = \frac{1}{2} =: P(A).$$

Iloczyn zdarzeń A i B to zdarzenie wyrzucenia orła w pierwszym i drugim rzucie – odpowiadają mu dwa zdarzenia w definicji Ω . Otrzymujemy więc:

$$P(A \cap B) = 2 \cdot \frac{1}{8} = \frac{1}{4}.$$

Ostatecznie ($P(B) \neq 0$):

$$P(A|B) = \frac{P(A \cap B)}{P(B)} = \frac{1}{2} = P(A).$$

Wniosek: zdarzenia A i B są niezależne.

Uwaga: Można też było postąpić odwrotnie – przyjąć, że kolejne rzuty są zdarzeniami niezależnymi, a w każdym rzucie prawdopodobieństwa wyrzucenia orła i reszki są sobie równe i wynoszą $1/2$. Takie założenie pozwala wyznaczyć prawdopodobieństwo uzyskania każdego wyniku serii o dowolnej, skończonej długości.

Przykład 2. Próby Bernoulliego i rozkład dwumianowy.

Uogólnijmy zagadnienie z Przykładu 1. Wykonujemy serię **niezależnych prób**. W każdej z nich uzyskujemy sukces z prawdopodobieństwem p lub ponosimy porażkę z prawdopodobieństwem q – oczywiście mamy $p + q = 1$.

Jakie jest prawdopodobieństwo uzyskania **dokładnie** k sukcesów w ciągu n takich prób?

- Prawdopodobieństwo uzyskania zadanego ciągu sukcesów i porażek równe jest (na podstawie założenia o niezależności prób) iloczynowi, w którym p występuje tyle razy ile jest w tym ciągu sukcesów, a q tyle razy ile porażek.
- Każde ze zdarzeń sprzyjających wynikowi, o który pytamy w zadaniu (dokładnie k sukcesów w n próbach) pojawia się więc z prawdopodobieństwem:

$$p^k q^{n-k}$$

- Liczba takich zdarzeń wynosi ($n \geq k$):

$$\frac{n!}{k!(n-k)!},$$

przy czym: $n! := 1 \cdot 2 \cdot 3 \cdots (n-1) \cdot n$, a dodatkowo przyjmuje się $0! = 1$.

- Ostatecznie otrzymujemy (tzw. dwumianowy rozkład prawdopodobieństwa):

$$P(k; n, p) = \frac{n!}{k!(n-k)!} p^k q^{n-k}$$

7 Zmienna losowa, wartość oczekiwana

Zmienna losowa

Funkcja liczbowa określona na przestrzeni zdarzeń nazywana jest zmienną losową. Inaczej: zmienna losowa przyporządkowuje liczby zdarzeniom losowym.

Zmienna losowa, to odwzorowanie $x : \Omega \rightarrow \mathbf{R}$
(ogólniej: $x : \Omega \rightarrow \mathbf{R}^n$)

Przykład 3.

Każdemu wynikowi rzutu sześcienną kostką (zdarzenie losowe) przyporządkowujemy liczbę oczek na górnej ścianie kostki.

Wartość oczekiwana zmiennej losowej:

Wartością oczekiwaną, $\mathcal{E}(x)$, zmiennej losowej x nazywamy:

$$\mathcal{E}(x) = \sum_i x_i p_i,$$

dla dyskretnej zmiennej losowej (i numeruje możliwe wartości zmiennej) lub

$$\mathcal{E}(x) = \int x f(x) dx,$$

dla zmiennej losowej ciągłej, przy czym sumowanie (całkowanie) przebiega cały zakres zmienności x , a $f(x)$ oznacza gęstość rozkładu prawdopodobieństwa otrzymania wartości x .

Wariancja zmiennej losowej nazywamy:

$$Var(x) = \mathcal{E}((x - \mathcal{E}(x))^2).$$

Kowariancja zmiennych losowych x i y nazywamy:

$$Cov(x, y) = \mathcal{E}((x - \mathcal{E}(x))(y - \mathcal{E}(y))).$$

Zmienne losowe $x : \Omega \rightarrow \mathbf{R}^n$ i $y : \Omega \rightarrow \mathbf{R}^m$ są **niezależne**,

gdy dla każdego dwóch zbiorów otwartych

$B_1 \subset \mathbf{R}^n, B_2 \subset \mathbf{R}^m$, zdarzenia $x \in B_1$ i $y \in B_2$ są niezależne.

Wówczas, dla ciągłych zmiennych losowych, gęstość prawdopodobieństwa $g(x, y) = f_1(x)f_2(y)$, gdzie $f_1(x)$ i $f_2(y)$ oznaczają gęstości prawdopodobieństwa zmiennych x i y .

Przykład 4.

Wynik pomiaru jest zmienną losową. Interesuje nas wartość oczekiwana tej zmiennej oraz wartość oczekiwana kwadratu odchylenia tej zmiennej od jej wartości oczekiwanej.

W przypadku, gdy wyeliminowane są błędy grube i systematyczne, **wartość oczekiwaną** zmiennej **wynik pomiaru** utożsamiamy z **wartością mierzoną** (poszukiwaną). Zadanie do rozwiązania: jak wyznaczyć wartość oczekiwaną na podstawie skończonej (na ogół niewielkiej) liczby prób?

Przykład 5.

Wartość oczekiwana zmiennej losowej k dla rozkładu dwumianowego:

$$\mathcal{E}(k) = np$$

Wartość oczekiwana kwadratu odchylenia zmiennej losowej k od jej wartości oczekiwanej:

$$\mathcal{E}((k - \mathcal{E}(k))^2) = npq = np(1 - p)$$

Przykład 6.

Wyznaczyć rozkład, do którego dąży rozkład dwumianowy, gdy $n \rightarrow \infty$, a oczekiwana liczba sukcesów jest skończona i stała, $\lambda = np$. Wynikiem jest tzw. **rozkład Poissona**.

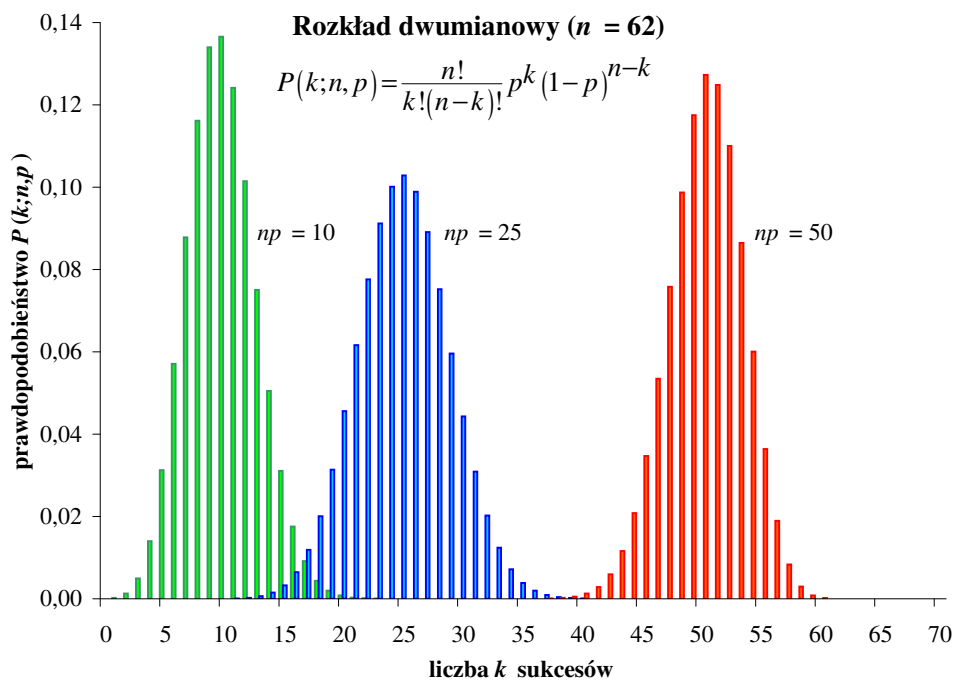


Figure 1: Histogramy prawdopodobieństw liczby sukcesów k opisanych rozkładem dwumianowym w przypadku $n = 62$ prób i różnych wartości p , prawdopodobieństwa sukcesu w pojedynczej próbie

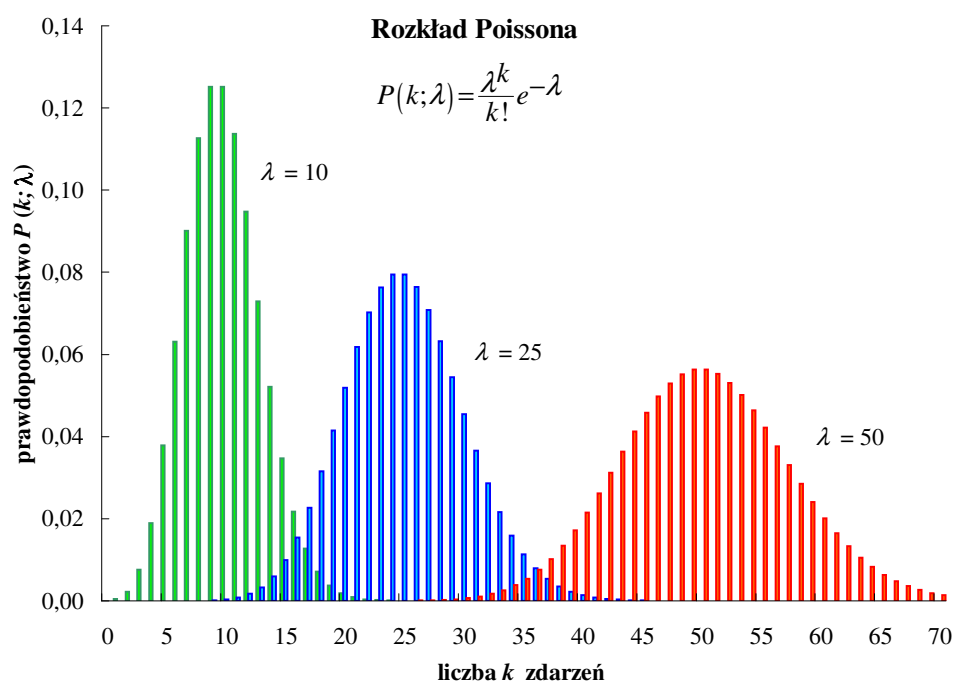


Figure 2: Histogramy rozkładu Poissona dla różnych λ – oczekiwanej liczby przypadków.

8 Wzór Stirlinga

Obliczanie $n!$ jest bardzo kłopotliwe. W rachunkach bardzo przydatny okazuje się wzór przybliżony, tzw. wzór Stirlinga:

$$\begin{aligned} n! &= n^n e^{-n} \sqrt{2\pi n} e^{\frac{\Theta}{12n}} \\ &= S(n) e^{\frac{\Theta}{12n}}, \end{aligned}$$

gdzie $0 < \Theta < 1$.

Poniższa tabela ilustruje charakter przybliżenia $n! \simeq S(n)$ (ilustracja wzorowana na: Jerzy Neyman, *Zasady Rachunku Prawdopodobieństwa i statystyki matematycznej*, PWN, Warszawa, 1969):

Tabela 1. Porównanie dokładnych wartości $n!$ z obliczonymi za pomocą wzoru Stirlinga.

n	$n!$	$S(n)$	$n!/S(n)$	$n! - S(n)$
2	2	1,9	1,042	0,1
4	24	23,5	1,021	0,5
6	720	710,1	1,014	9,9
8	40 320	39 902,4	1,010	417,6
10	3 628 800	3 598 699,6	1,008	30 100,4
11	39 916 800	39 615 625,2	1,008	301 174,8
12	479 001 600	475 687 487,7	1,007	3 314 112,3

9 Rozkład Gaussa – Graniczna postać rozkładu dwumianowego ($n \rightarrow \infty$)

Dla rozkładu dwumianowego:

$$P(k; n, p) = \frac{n!}{k!(n-k)!} p^k (1-p)^{n-k},$$

dla dużych n otrzymujemy:

$$P(k; n, p) \simeq \frac{1}{\sqrt{2\pi np(1-p)}} \exp\left(-\frac{(k-np)^2}{2np(1-p)}\right).$$

Wprowadźmy oznaczenia $\mu := np$ oraz $\sigma := \sqrt{np(1-p)}$ i rozważmy ciągły rozkład gęstości prawdopodobieństwa:

$$f(x; \mu, \sigma) = \frac{1}{\sqrt{2\pi}\sigma} \exp\left(-\frac{(x-\mu)^2}{2\sigma^2}\right)$$

nazywany **rozkładem normalnym** lub **rozkładem Gaussa**.

Dla $x_1 < x_2$ prawdopodobieństwo tego, że zmienna losowa opisana rozkładem normalnym przyjmuje wartość pomiędzy x_1 i x_2 równe jest polu pod krzywą $f(x; \mu, \sigma)$ zawartemu pomiędzy x_1 i x_2 .

Przykład 7.

Dla rozkładu normalnego (Gaussa):

$$f(x; \mu, \sigma) = \frac{1}{\sqrt{2\pi}\sigma} \exp\left(-\frac{(x-\mu)^2}{2\sigma^2}\right)$$

otrzymujemy

- $P(\mu - 1\sigma \leq x \leq \mu + 1\sigma) \simeq 0,682689$
- $P(\mu - 2\sigma \leq x \leq \mu + 2\sigma) \simeq 0,954500$
- $P(\mu - 3\sigma \leq x \leq \mu + 3\sigma) \simeq 0,997300$
- $P(\mu - 4\sigma \leq x \leq \mu + 4\sigma) \simeq 0,999937$

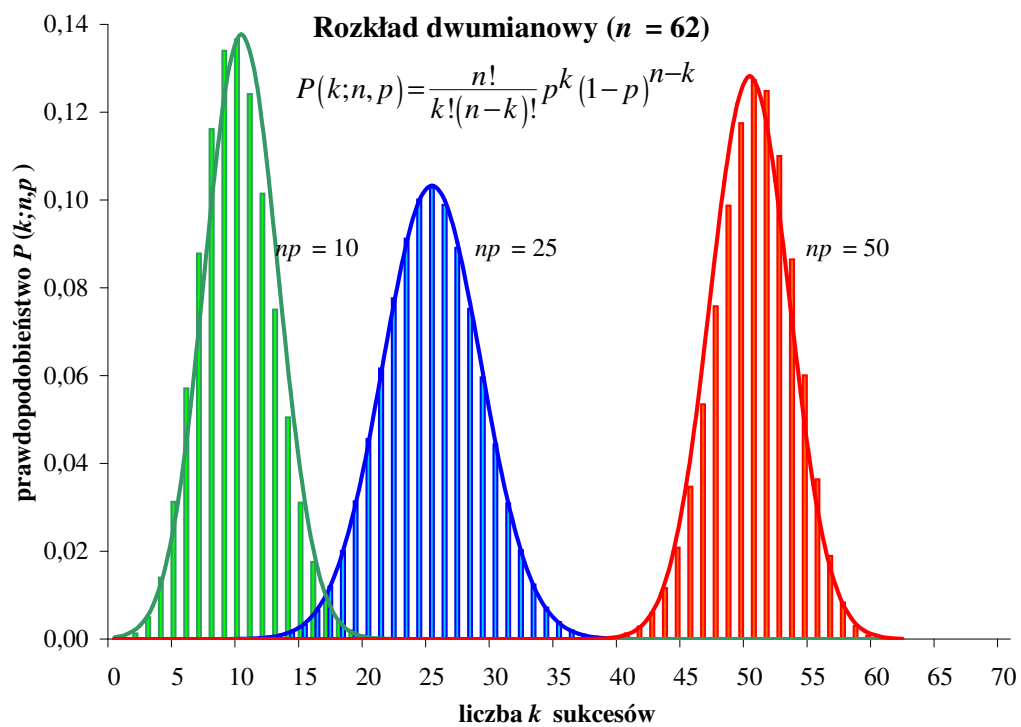


Figure 3: Porównanie dokładnych histogramów rozkładu dwumianowego z wyprowadzonym wzorem przybliżonym – linie ciągłe.

Przykłady rozkładów Gaussa

$$f(x; \mu, \sigma) = \frac{1}{\sqrt{2\pi}\sigma} \exp\left(-\frac{(x-\mu)^2}{2\sigma^2}\right)$$

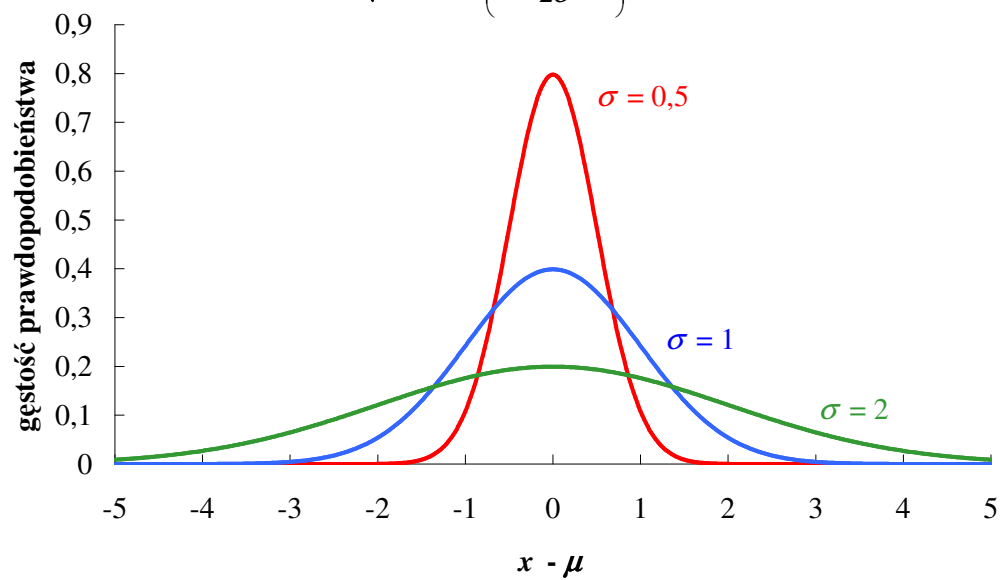


Figure 4: Rozkład Gaussa: gęstości rozkładu dla różnych wartości dyspersji σ .

10 Twierdzenie graniczne i prawo wielkich liczb

Jeżeli dany jest ciąg niezależnych zmiennych losowych pochodzących z dowolnego rozkładu, który ma skończoną wartość oczekiwaną μ i skończoną dyspersję σ , to rozkład zmiennej losowej:

$$z_n := \frac{\bar{x}_n - \mu}{\sigma/\sqrt{n}},$$

gdzie

$$\bar{x}_n := \frac{1}{n} \sum_{i=1}^n x_i,$$

dla $n \rightarrow \infty$ dąży do standardowego rozkładu Gaussa:

$$N(z; 0, 1) = \frac{1}{\sqrt{2\pi}} \exp\left(-\frac{z^2}{2}\right)$$

Twierdzenie graniczne pokazuje, że rozkład Gaussa (normalny) jest naturalnym modelem rozkładu zmiennej losowej będącej wynikiem złożenia wielu zmiennych losowych posiadających skończoną wartość oczekiwaną i skończoną wartość średniego odchylenia kwadratowego.

Prawo wielkich liczb Bernoulliego

Niech, k będzie zmienną losową o rozkładzie prawdopodobieństwa:

$$P(k; n, p) = \frac{n!}{k!(n-k)!} p^k (1-p)^{n-k}.$$

Tworzymy nową zmienną losową:

$$x_n := \frac{k}{n} - p$$

wówczas:

$$\bigwedge_{\varepsilon > 0} \lim_{n \rightarrow \infty} P(|x_n| > \varepsilon) = 0,$$

(w granicy $n \rightarrow \infty$ prawdopodobieństwo, że uzyskamy x_n większe co do wartości bezwzględnej od wybranej, dowolnie małej liczby dodatniej, dąży do zera).

Uwaga: Zbieżność opisanana w powyższym twierdzeniu nazywana jest **zbieżnością stochastyczną**. Moglibyśmy więc powyżej powiedzieć: *Ciąg x_n jest stochastycznie zbieżny do zera, gdy n dąży do nieskończoności.*

Istnieje wiele różnych postaci twierdzenia wielkich liczb.

Twierdzenia te opisują związek ścisłego pojęcia prawdopodobieństwa z jego (intuicyjną) interpretacją „częstościową”.

11 Model błędu przypadkowego: złożenie wielu małych odchyłeń

Błąd przypadkowy opiszemy jako odchylenie wyniku pomiaru x od wartości dokładnej μ powstające na skutek złożenia bardzo wielu niezależnych czynników przypadkowych, z których każdy powoduje odchylenie wyniku o $+\varepsilon$ lub $-\varepsilon$ z prawdopodobieństwem:

$$p = q = \frac{1}{2}.$$

Dla n czynników zaburzających mamy:

$$\begin{aligned} x &= \mu + k\varepsilon - (n - k)\varepsilon \\ &= \mu + (2k - n)\varepsilon. \end{aligned}$$

Wartości $x - \mu$ odpowiada więc:

$$\begin{aligned} k &= \frac{x - \mu}{2\varepsilon} + \frac{n}{2} = \frac{x - \mu + n\varepsilon}{2\varepsilon} \\ \mathcal{E}(k) &= np = \frac{n}{2} \end{aligned}$$

Stosując przybliżone wyrażenie na prawdopodobieństwo w rozkładzie dwumianowym dla dużych n , otrzymujemy prawdopodobieństwo uzyskania w wyniku pomiaru danej wartości $x - \mu$ równe:

$$\begin{aligned} P(x - \mu; n, \frac{1}{2}) &= \frac{1}{\sqrt{2\pi n \frac{1}{4}}} \exp\left(-\frac{(x - \mu)^2}{2n \frac{1}{4}(2\varepsilon)^2}\right) \\ &= \frac{2\varepsilon}{\sqrt{2\pi n \varepsilon^2}} \exp\left(-\frac{(x - \mu)^2}{2n \varepsilon^2}\right). \end{aligned}$$

Dokonajmy przejścia granicznego $n \rightarrow \infty$ (liczba losowych czynników zaburzających pomiar dąży do nieskończoności) oraz $\varepsilon \rightarrow 0$ z dodatkowym warunkiem: $n\varepsilon^2 = \sigma^2 = \text{const}$.

Dla prawdopodobieństwa dzielonego przez elementarną zmianę wyniku pomiaru (zmiana k o 1 zmienia wynik o 2ε) otrzymujemy:

$$\lim_{n \rightarrow \infty, \varepsilon \rightarrow 0; n\varepsilon^2 = \sigma^2} \left\{ \frac{1}{2\varepsilon} P(x - \mu; n, \frac{1}{2}) \right\} = \frac{1}{\sqrt{2\pi\sigma^2}} \exp\left(-\frac{(x - \mu)^2}{2\sigma^2}\right).$$

Wyrażenie po prawej stronie należy przy tym rozumieć jako gęstość prawdopodobieństwa dla ciągłej zmiennej losowej x .

Interpretacja parametrów rozkładu:

Jeżeli otrzymany rozkład służy jako model rozkładu wyników pomiarów (po wyeliminowaniu błędów grubych i systematycznych), to parametry μ i σ interpretujemy jako:

- μ – nieznana dokładna wartość wielkości mierzonej,
- σ – *szerokość* rozkładu określająca dokładność, z jaką potrafimy kontrolować stałość warunków wykonywania pomiaru, inaczej mówiąc, σ charakteryzuje proces pomiaru.

Uwaga: Rozkład Gaussa jest powszechnie stosowanym modelem rozkładu prawdopodobieństwa wartości błędu przypadkowego.

Przeprowadzając pomiary zwykle zakładamy jedynie, że rozkład jakiego podlegają wyniki ma dobrze określone (i skończone):

wartość oczekiwaną i średnie odchylenie kwadratowe. Nasze zadanie polega wówczas na skonstruowaniu z serii wyników pomiarów wielkości:

- $\hat{\mu}$ – najlepszej oceny wartości oczekiwanej oraz
- $u(\hat{\mu})$ – niepewności tej oceny,

w taki sposób, aby:

$$\begin{aligned} \mathcal{E}(\hat{\mu}) &= \mu \\ \mathcal{E}(u^2(\hat{\mu})) &= \mathcal{E}((\hat{\mu} - \mu)^2). \end{aligned}$$

12 Pomiar pojedynczej wielkości i jego niepewność

Wykonujemy serię niezależnych pomiarów tej samej wielkości (tą samą metodą i w tych samych warunkach). Otrzymujemy serię wyników:

$$x_i, i = 1, 2, \dots, n.$$

Zakładamy, że:

- wyeliminowane zostały błędy grube i systematyczne;
- wartość oczekiwana każdego z wyników pomiaru x_i równa jest dokładnej wartości wielkości mierzonej μ :

$$\mathcal{E}(x_i) = \mu;$$

- średnie odchylenie kwadratowe od wartości dokładnej (wariancja rozkładu), dla każdego x_i wynosi σ^2 :

$$\mathcal{E}((x_i - \mu)^2) = \sigma^2;$$

- wartości oczekiwane μ i σ^2 istnieją i są skończone.

Pytanie:

Jak wyznaczyć μ i σ na podstawie wyników serii niezależnych pomiarów?

Tworzymy wielkości:

$$\begin{aligned}\bar{x} &:= \frac{1}{n} \sum_{i=1}^n x_i, \\ s^2 &:= \frac{1}{n-1} \sum_{i=1}^n (x_i - \bar{x})^2.\end{aligned}$$

Można udowodnić, że:

$$\begin{aligned}\mathcal{E}(\bar{x}) &= \mu, \\ \mathcal{E}((\bar{x} - \mu)^2) &= \frac{\sigma^2}{n} =: \sigma_{\bar{x}}^2, \\ \mathcal{E}(s^2) &= \sigma^2, \\ \mathcal{E}((x_i - \bar{x})^2) &= \frac{n-1}{n} \sigma^2, \\ \mathcal{E}((s^2 - \sigma^2)^2) &= \frac{1}{n} \mathcal{E}((x_i - \mu)^4) - \frac{n-3}{n(n-1)} \sigma^4.\end{aligned}$$

Jeśli przyjmiemy, że każda z wielkości x_i podlega rozkładowi Gaussa o wartości oczekiwanej μ i dyspersji σ , to:

$$\mathcal{E}((x_i - \mu)^4) = 3\sigma^4 \Rightarrow \mathcal{E}((s^2 - \sigma^2)^2) = \frac{2}{n-1} \sigma^4.$$

Wniosek: Na podstawie wyników serii pomiarów, najlepszą oceną wartości μ i σ^2 są wielkości $\hat{\mu} := \bar{x}$ i $\widehat{\sigma^2} := s^2$ zdefiniowane jak wyżej. Wielkość $s := \sqrt{s^2}$ określa naszą ocenę dokładności pomiaru, inaczej, **niepewność** pojedynczego wyniku (x_i dla $i = 1, 2, \dots, n$).

13 Wyznaczanie oceny kowariancji zmiennych x i y na podstawie serii pomiarów

Wykonujemy serię niezależnych pomiarów par wielkości:

$$(x_i, y_i), i = 1, 2, \dots, n.$$

Założenia: Wyniki pomiarów dla $i \neq j$ są niezależne;

$$\mathcal{E}(x_i) = \mu$$

$$\mathcal{E}(y_i) = \nu$$

$$\mathcal{E}((x_i - \mu)(y_i - \nu)) = Cov(x, y)$$

Tworzymy wielkości:

$$\bar{x} := \frac{1}{n} \sum_{i=1}^n x_i$$

$$\bar{y} := \frac{1}{n} \sum_{i=1}^n y_i$$

$$\hat{C}_{x,y} := \frac{1}{n-1} \sum_{i=1}^n (x_i - \bar{x})(y_i - \bar{y}).$$

Wówczas:

$$\mathcal{E}(\hat{C}_{x,y}) = Cov(x, y)$$

oraz

$$Cov(\bar{x}, \bar{y}) := \mathcal{E}((\bar{x} - \mu)(\bar{y} - \nu))$$

$$\mathcal{E}((\bar{x} - \mu)(\bar{y} - \nu)) = \frac{1}{n} Cov(x, y)$$

$$\mathcal{E}((\bar{x} - \mu)(\bar{y} - \nu)) = \frac{1}{n(n-1)} \mathcal{E}\left(\sum_{i=1}^n (x_i - \bar{x})(y_i - \bar{y})\right)$$

$$\mathcal{E}\left(\frac{1}{n} \hat{C}(x, y)\right) = Cov(\bar{x}, \bar{y})$$

14 Wynik pomiaru i jego zapis

Zakładamy, że błędy grube i systematyczne zostały wyeliminowane lub są zanedbywalnie małe w porównaniu z błędami przypadkowymi. Wówczas, na podstawie znajomości wyników serii **bezpośrednich, niezależnych, równoważnych pomiarów**,

$$x_i, i = 1, 2, \dots, n,$$

wyznaczamy:

- **wynik pomiaru** określony jako wielkość:

$$\bar{x} := \frac{1}{n} \sum_{i=1}^n x_i;$$

- **niepewność pomiaru** $u(\bar{x})$, gdy wyeliminowano błędy grube i systematyczne, określoną wzorem:

$$u^2(\bar{x}) = s_{\bar{x}}^2 := \frac{1}{n(n-1)} \sum_{i=1}^n (x_i - \bar{x})^2;$$

- $u(\bar{x})$ zaokrąglamy do dwóch **cyfr znaczących**,
- $\bar{x}$ – zaokrąglamy tak, aby ostatnia cyfra znacząca $\bar{x}$ była na tym samym miejscu dziesiętnym, co ostatnia cyfra znacząca $u(\bar{x})$.
- wynik pomiaru podajemy w postaci:
 $x =$ **obliczona wartość** $\bar{x}$, $u(x) =$ **obliczona wartość** $u(\bar{x})$
– obie wartości odpowiednio (patrz wyżej) zaokrąglone.

Uwaga 1:

Stosowane są też inne sposoby podawania wyniku (przykłady poniżej):

niech odpowiednio zaokrąglone wartości wynoszą: $\bar{x} = 1,0238$ g oraz $u(\bar{x}) = 0,0036$ g, wówczas wynik możemy podać w postaci:

- $x = 1,0238$ g, $u(x) = 0,0036$ g
zapis według podanej wyżej reguły;
- $x = 1,0238(0,0036)$ g;
- $x = 1,0238(36)$ g;
- $x = (1,0238 \pm 0,0036)$ g; taki zapis jest zgodny z przyjętymi normami. Należy jednak pamiętać, że może on prowadzić do nieporozumień. Tu i w poprzednich przykładach $u(\bar{x})$ oznacza tzw. *niepewność standardową* – wartość oceny *odchylenia standardowego*, tj. pierwiastka z oceny *średniego odchylenia kwadratowego*.

Ta informacja powinna zawsze być podana tuż obok wyniku, bo w zapisie $\hat{x} \pm U$ często U oznacza tzw. *niepewność rozszerzoną*: $U = ku$, najczęściej z $k = 2$. Takie rozumienie zapisu $x \pm U$ jest zwyczajowo stosowane w naukach technicznych – jest to pozostałość zwyczaju podawania zakresu, w którym wartość mierzonej wielkości mieści się *niemal na pewno* (dla błędów podlegających rozkładowi Gaussa, z prawdopodobieństwem 95%).

Uwaga 2: Wielkości s^2 i $s_{\bar{x}}^2$ są funkcjami zmiennych losowych (wyników pomiarów) i w związku z tym są też zmiennymi losowymi i możemy pytać o oceny ich wariancji. Niepewność wyznaczenia wartości s i $s_{\bar{x}}$ oceniamy korzystając z przybliżenia (zakładamy, że wyniki pomiarów podlegają rozkładowi Gaussa):

$$\mathcal{E}((s^2 - \sigma^2)^2) = \frac{2}{n-1}\sigma^4 \Rightarrow \frac{\sigma_s}{\sigma} = \frac{\sigma_{s_{\bar{x}}}}{\sigma_{\bar{x}}} \simeq \frac{1}{\sqrt{2n-2}}$$

Tabela 2. Zależność względnej dokładności oceny s od długości serii n .

n	$\sigma_s/s \leq$
4	1/2
6	1/3
10	1/4
14	1/5
20	1/6
52	1/10
202	1/20
1 252	1/50
5 002	1/100
20 002	1/200

Stąd spotykane w starych podręcznikach zalecenie, żeby wykonywać $n > 5$ pomiarów, bo dopiero dla serii $n \geq 6$ pomiarów możemy ocenić wartość niepewności wynikającej z występowania błędów przypadkowych (tzw. niepewność statystyczną) z dokładnością lepszą niż 30%.

Powyższa tabela uzasadnia także zaokrąglanie niepewności do dwóch cyfr znaczących – dla serii krótszych niż ok. 1300 pomiarów już druga cyfra znana jest niedokładnie.

15 Uwzględnienie dokładności przyrządu

Dotychczasowe rozważania zakładały, że wynik pomiaru może być dowolną liczbą rzeczywistą. Wiemy jednak, że odczyt wyniku z przyrządu dokonywany jest z dokładnością do niewielu cyfr. Ogranicza to możliwość zwiększenia dokładności wyniku poprzez wydłużenie serii pomiarów. Musimy ten fakt uwzględnić w naszym modelu.

Zakładamy, że: wyeliminowane są błędy grube i systematyczne (poza błędem wynikającym z ograniczonej dokładności odczytu).

Rozważmy przyrząd cyfrowy, którego wskazanie zmienia się co $\Delta = 1, 2, \dots, 5$ na ostatnim miejscu odczytu wyniku.

Rysunek 5. pokazuje porównanie rzeczywistego sygnału na wejściu (cienka, niebieska linia prosta) z odczytem zdyskretyzowanym (poziome, zielone odcinki). Przełączenie wskazań odbywa się tak, żeby ich średnie odchylenie od wartości dokładnej, dla danej dokładności odczytu, było jak najmniejsze – optymalna dyskretyzacja odczytu.

Błąd ξ związany z odczytem definiujemy jako:

$\xi := Y - X$, gdzie X oznacza sygnał na wejściu, a Y odczyt wskazania przyrządu.

Rysunek 6. przedstawia ξ jako funkcję sygnału X w jednostkach rozdzielczości przyrządu Δ . Ten błąd jest **błędem systematycznym**, bo powtarzając pomiary dla tej samej wartości sygnału wejściowego będziemy popełniali stale błąd ξ o tej samej wartości.

Jego wartości jednak nie znamy, ale jak widać na Rysunku 6., w przedziale $-0,5\Delta \leq \xi \leq 0,5\Delta$ każda jego wartość ξ jest możliwa i tak samo prawdopodobna.

Model pomiaru przyrządem cyfrowym

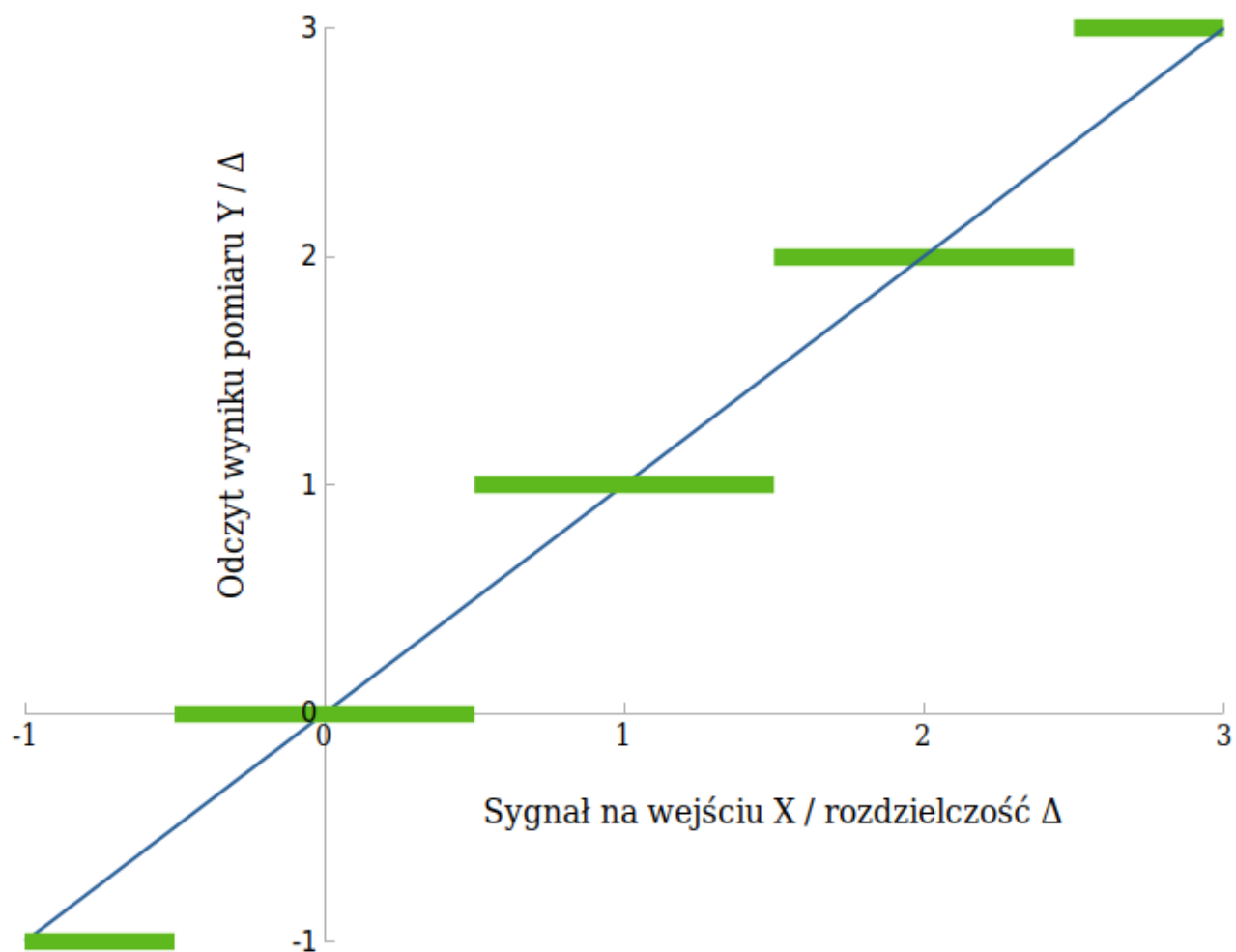


Figure 5: Porównanie sygnału na wejściu przyrządu cyfrowego w jednostkach Δ – cienka niebieska prosta, z odczytem cyfrowym – poziome, zielone odcinki.

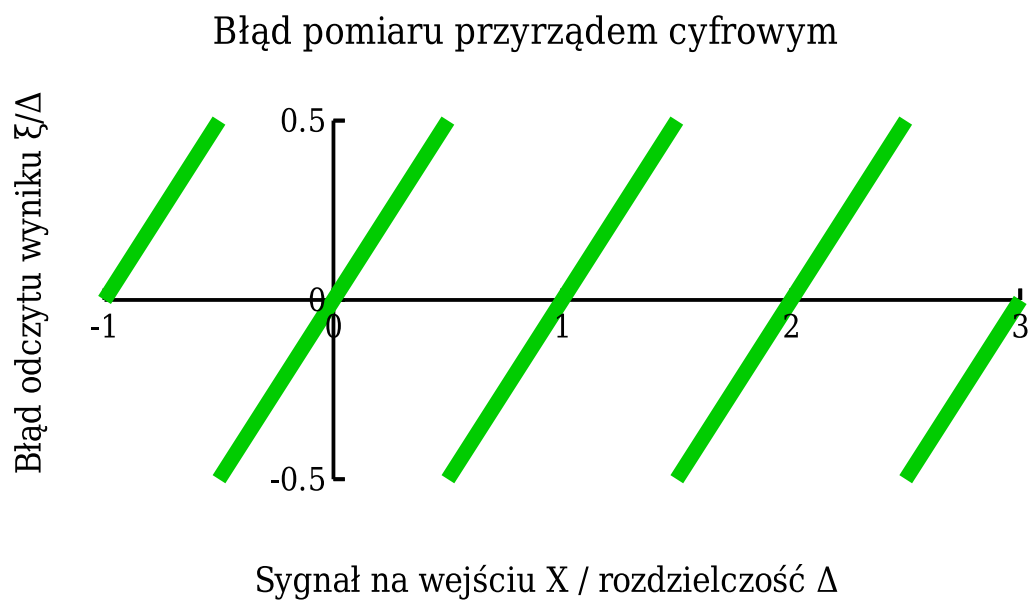


Figure 6: Błąd dyskretyzacji wyniku pomiaru przyrządem cyfrowym – zielone odcinki.

Rozkład gęstości prawdopodobieństwa $g(\xi)$ zmiennej losowej opisującej błąd dyskretyzacji (kwantowania), jest więc rozkładem jednostajnym.

Rozkład jednostajny zmiennej losowej ξ :

w przedziale $\xi \in [-a, a]$:

$$g(\xi) = \frac{1}{2a}$$

oraz $g(\xi) = 0$ poza tym przedziałem.

Wartość oczekiwana tak określonej zmiennej losowej:

$$\mathcal{E}(\xi) = \int_{-a}^a \frac{\xi}{2a} d\xi = 0,$$

a wartość jej wariancji ξ ($\mathcal{E}(\xi) = 0$):

$$\mathcal{E}(\xi^2) = \int_{-a}^a \frac{\xi^2}{2a} d\xi = \frac{a^2}{3}.$$

Z dotychczasowych rozważań wynika, że w modelu błędu związanego ze skończoną rozdzielczością przyrządu powinniśmy przyjąć $a = \Delta/2$, dla *bezpieczeństwa*, w obliczeniach przyjmujemy $a = \Delta$ i z tym założeniem będziemy ten model stosowali także w przypadku przyrządów analogowych (np. linijki), gdzie Δ oznacza najmniejszą działkę skali analogowej.

Uwaga: W zależności od sposobu działania przyrządu i naszej o tym wiedzy możemy zbudować inny model rozkładu wskazań dokładniej opisujący wykonywane pomiary. Tutaj ograniczamy się do najprostszego modelu uwzględniającego rozdzielczość wskazań przyrządu pomiarowego.

16 Niepewność wyniku bezpośredniego pomiaru pojedynczej wielkości fizycznej – wzór końcowy

Czynniki prowadzące do powstania błędu przypadkowego podczas pomiaru są niezależne od *przetwarzania* sygnału przez miernik.

Przyjmujemy więc, że oba rodzaje błędów opisywane są przez **niezależne zmienne losowe**. Wariancja sumy niezależnych zmiennych losowych równa jest sumie wariancji tych zmiennych.

W naszym przypadku: wariancji średniej z serii pomiarów $\sigma_{\bar{x}}^2$ i wariancji błędu wskazań $\mathcal{E}(\xi^2)$. Obie te wielkości nie są nam znane $\Rightarrow$ do obliczeń bierzemy ich oceny: $s_{\bar{x}}^2$ i $\Delta^2/3$ wyznaczone na podstawie serii pomiarów i rozdzielczości wskazań miernika.

Ostatecznie: wynik bezpośredniego pomiaru pojedynczej wielkości fizycznej otrzymany na podstawie serii równoważnych pomiarów podajemy jako:

$$x = \bar{x}, \quad u(\bar{x}) = \sqrt{s_{\bar{x}}^2 + \frac{\Delta^2}{3}}.$$

Uwagi:

- Zasady zaokrąglania stosujemy, jak podano poprzednio. Informujemy, jak wyznaczono niepewność $u(\bar{x})$.
- Podany wzór ogranicza sensowną liczbę pomiarów w serii: wydłużanie serii ma sens dokąd $s_{\bar{x}}^2 \geq \Delta^2/3$ (przypominamy, że $s_{\bar{x}}^2 \propto 1/n$, gdzie n jest długością serii).
- „W życiu” często wyraźnie dominuje jedna ze składowych niepewności: (a) rozrzut kolejnych wyników serii jest znacznie większy od Δ – dominuje błąd przypadkowy; (b) kolejne pomiary serii dają identyczne wyniki – dominuje błąd związany z dokładnością przyrządu $\rightarrow$ nie należy przedłużać serii.

17 Test „ 3σ ”

Przykład 8.

Rzucono 8 razy monetą. Otrzymano 7 orłów. Czy moneta jest „uczciwa”?

Odpowiedź:

Rachunek prawdopodobieństwa pozwala nam tylko na obliczenie prawdopodobieństwa otrzymania uzyskanego wyniku dla danej wartości p – prawdopodobieństwa wyrzucenia *orła* w pojedynczym rzucie. Dla „uczciwej” monety $p = \frac{1}{2}$.

Przyjmując $p = \frac{1}{2}$ otrzymujemy:

$$\begin{aligned}P(k \geq 7; 8, \frac{1}{2}) &= \frac{9}{256} < 0,0352, \\P(k \geq 6; 8, \frac{1}{2}) &= \frac{37}{256} < 0,145.\end{aligned}$$

Jeśli zdecydujemy, że odrzucamy hipotezę *moneta jest uczciwa* ($p = \frac{1}{2}$), gdy $P(\text{otrzymanego wyniku}) < 0,04$, to uznajemy, że otrzymane wyniki pozwalają nam uznać, że *moneta nie jest uczciwa*.

Powyższy przykład pomoże nam sformułować metodę porównywania wyników pomiarów z przewidywaniami teoretycznymi lub z innymi wynikami pomiarów.

Założmy, że wykonana została seria pomiarów wielkości x i jako jej wynik otrzymano (stosujemy notację jak dla przypadku dominacji błędu przypadkowego):

$$x = \hat{x}_1 \pm s_1$$

Często zadawane pytania:

1. Czy otrzymany wynik jest zgodny z wartością

$$x = \mu_0$$

przewidywaną przez teorię?

2. Czy wynik ten jest zgodny z wynikiem

$$x = \hat{x}_2 \pm s_2$$

otrzymanym w innej serii pomiarów, inną metodą lub innymi przyrządami? Inaczej mówiąc: czy oba wyniki doświadczalne są ze sobą zgodne?

Założmy dalej, że wspomniane wyniki doświadczalne uzyskaliśmy dla bardzo długich serii pomiarów, a błędy systematyczne są pomijalnie małe – możemy więc przyjąć, że:

$$s_i \simeq \sigma_i.$$

Dla zmiennej losowej x podlegającej rozkładowi Gaussa o średniej μ i dyspersji σ mamy:

$$P(|x - \mu| \geq 3\sigma) < 0,0028.$$

Oznacza to, że rzadziej niż raz na ok. 357 losowań zmiennej opisanej rozkładem Gaussa otrzymamy wynik różniący się od wartości oczekiwanej o więcej niż 3σ . Pozwala to zdefiniować sposób testowania wyników pomiarów zapewniający, że odrzucimy hipotezę prawdziwą (tzn. popełnimy tzw. **błąd pierwszego rodzaju**) rzadziej niż raz na 357 testowanych przypadków.

Test 3σ :

1. Jeżeli porównujemy wynik pomiaru $\hat{x}_1$ z wartością μ_0 przewidywaną przez teorię:
 $|\hat{x}_1 - \mu_0| > 3s_1 \Rightarrow$ odrzucamy hipotezę $x = \mu_0$.
2. Dla dwóch niezależnych pomiarów o wynikach $\hat{x}_1$ i $\hat{x}_2$:
 $|\hat{x}_1 - \hat{x}_2| > 3\sqrt{s_1^2 + s_2^2} \Rightarrow$ odrzucamy hipotezę, że w obu pomiarach mierzona była ta sama wielkość.

Uwagi:

- Jeżeli (jak to w praktyce najczęściej bywa i co sugeruje użyta powyżej notacja) zamiast dokładnych wartości dyspersji, σ , odpowiednich rozkładów Gaussa posługujemy się ich ocenami, s , to nasza decyzja staje się mniej kategoryczna – prawdopodobieństwo popełnienia błędu pierwszego rodzaju może być większe niż to wynika z rozkładu Gaussa (niż przyjęta wartość krytyczna tego prawdopodobieństwa).
- W jeszcze większym stopniu zastrzeżenie to odnosi się do przypadku, gdy ocena niepewności dotyczy składowej błędu nie podlegającej rozkładowi Gaussa – np. zawiera składową związaną z dokładnością przyrządu. Często i w takich przypadkach stosowany bywa test „ 3σ ”. Podjętej decyzji nie można jednak wówczas przypisać prawdopodobieństwa błędu I rodzaju. Wystarczy wiara, że czynnik 3 jest *wystarczająco* duży, aby decyzja była *wystarczająco* bezpieczna w codziennej praktyce.
- W wielu przypadkach, nawet w bardzo skomplikowanych sytuacjach, tworzone są ściśle modele matematyczne rozkładów prawdopodobieństwa wielkości podlegających testom, a akceptowane prawdopodobieństwo błędu I rodzaju wynika z oceny kosztów popełnienia tego błędu.

18 Pomiary pośrednie i „propagacja małych błędów”.

Nie zawsze możliwy jest bezpośredni pomiar interesującej nas wielkości y , albo też pomiar taki nie byłby wystarczająco dokładny. W takiej sytuacji często posługujemy się znanymi związkami funkcyjnymi pomiędzy wielkością y i innymi wielkościami dostępnymi bezpośrednim pomiarom:

$$y = f(x_1, x_2, \dots, x_k),$$

gdzie $f(x_1, x_2, \dots, x_k)$ jest znaną funkcją, a wielkości $x_1, x_2, \dots, x_k$ potrafimy mierzyć bezpośrednio.

Zakładamy, że:

- dokładna wartość wielkości y wynosi $\eta = f(\mu_1, \dots, \mu_k)$, gdzie $\mu_1, \dots, \mu_k$ są dokładnymi wartościami wielkości $x_1, \dots, x_k$;
- wielkości $x_i, i = 1, \dots, k$ są mierzone **niezależnie** i w wyniku pomiarów otrzymano wartości $\hat{x}_1, \dots, \hat{x}_k$, przy czym:

$$\mathcal{E}(\hat{x}_i) = \mu_i;$$

- dla kolejnych $\hat{x}_i, i = 1, \dots, k$ zachodzi także

$$\mathcal{E}((\hat{x}_i - \mu_i)^2) = \sigma_{\hat{x}_i}^2;$$

- w wyniku pomiarów wyznaczono te $s_{\hat{x}_1}, \dots, s_{\hat{x}_k}$ – oceny niepewności $\sigma_{\hat{x}_1}, \dots, \sigma_{\hat{x}_k}$ w taki sposób, że:

$$\mathcal{E}(s_{\hat{x}_i}^2) = \sigma_{\hat{x}_i}^2.$$

Twierdzenie:

Jeśli spełnione są powyższe założenia oraz

$$\eta = \sum_{i=1}^k a_i \mu_i,$$

gdzie $a_i, i = 1, \dots, k$ są znanymi stałymi, to

$$\begin{aligned}\mathcal{E}(\hat{y}) &= \mathcal{E}\left(\sum_{i=1}^k a_i \hat{x}_i\right) \\ &= \sum_{i=1}^k a_i \mu_i = \eta, \\ \mathcal{E}\left(\left(\sum_{i=1}^k a_i \hat{x}_i - \eta\right)^2\right) &= \sum_{i=1}^k a_i^2 \sigma_{\hat{x}_i}^2.\end{aligned}$$

Wniosek: Wynik pomiaru pośredniego wynosi $y = \hat{y} \pm s_{\hat{y}}$, gdzie

$$\begin{aligned}\hat{y} &= \sum_{i=1}^k a_i \hat{x}_i \\ s_{\hat{y}}^2 &= \sum_{i=1}^k a_i^2 s_{\hat{x}_i}^2\end{aligned}$$

Uwaga: Jeśli mierzone wielkości nie są niezależne i dla $i \neq j$

$$\mathcal{E}((\hat{x}_i - \mu_i)(\hat{x}_j - \mu_j)) = C_{ij} \neq 0$$

to:

$$\mathcal{E}\left(\left(\sum_{i=1}^k a_i \hat{x}_i - \eta\right)^2\right) = \sum_{i=1}^k a_i^2 \sigma_{\hat{x}_i}^2 + \sum_{i \neq j} a_i a_j \hat{C}_{ij}$$

Propagacja małych niepewności

W przypadku, gdy $\eta = f(\mu_1, \dots, \mu_k)$, dla *odpowiednio* małych niepewności $\sigma_{\hat{x}_i}$ i funkcji $f(x_1, \dots, x_k)$ różniczkowalnej w sposób ciągły (gadkiej) otrzymujemy przybliżenie:

$$\eta := \mathcal{E}(y) \simeq \hat{y} := f(\hat{x}_1, \hat{x}_2, \dots, \hat{x}_k)$$

oraz

$$\mathcal{E}((\hat{y} - \eta)^2) \simeq \sum_{i=1}^k \left(\frac{\partial f}{\partial x_i} \Big|_{\mu_1, \dots, \mu_k} \right)^2 \sigma_{\hat{x}_i}^2,$$

gdzie

$$\frac{\partial f}{\partial x_i} \Big|_{\mu_1, \dots, \mu_k} := \lim_{h \rightarrow 0} \frac{f(\mu_1, \dots, \mu_i + h, \dots, \mu_k) - f(\mu_1, \dots, \mu_i, \dots, \mu_k)}{h}.$$

Na podstawie pomiaru znamy jednak tylko wartości $s_{\hat{x}_i}$, to jest oceny wartości $\sigma_{\hat{x}_i}$. Definiujemy więc wielkość:

$$s_{\hat{y}}^2 := \sum_{i=1}^k \left(\frac{\partial f}{\partial x_i} \Big|_{\hat{x}_1, \dots, \hat{x}_k} \right)^2 s_{\hat{x}_i}^2,$$

a **wynik końcowy** podajemy w postaci:

$$y = \hat{y} \pm s_{\hat{y}}.$$

Uwagi:

- Niepewności $\sigma_{\hat{x}_i} \simeq s_{\hat{x}_i}$ są *odpowiednio* małe, gdy pochodne

$$\frac{\partial f}{\partial x_i} \Big|_{\hat{x}_1, \dots, \hat{x}_k}$$

są z dobrym przybliżeniem stałe dla $x_i \in [\hat{x}_i - s_{\hat{x}_i}, \hat{x}_i + s_{\hat{x}_i}]$.

- Jeśli pomiary argumentów f nie są niezależne, $\hat{C}_{\hat{x}_i, \hat{x}_j} \neq 0$, to:

$$s_{\hat{y}}^2 := \sum_{i=1}^k \left(\frac{\partial f}{\partial x_i} \Big|_{\hat{x}_1, \dots, \hat{x}_k} \right)^2 s_{\hat{x}_i}^2 + \sum_{i \neq j} \frac{\partial f}{\partial x_i} \Big|_{\hat{x}_1, \dots, \hat{x}_k} \frac{\partial f}{\partial x_j} \Big|_{\hat{x}_1, \dots, \hat{x}_k} \hat{C}_{\hat{x}_i, \hat{x}_j}.$$

19 Metoda najmniejszych kwadratów

- Wiemy (zakładamy), że wielkości fizyczne x i y wiąże zależność: $y = f(x; a_1, a_2, \dots, a_k)$, gdzie a_i są nieznanymi nam parametrami (np. charakteryzującymi właściwości badanego materiału).
- Dla $N > k$ różnych wartości wielkości x mierzymy odpowiadające im wartości y : $x_i \rightarrow \hat{y}_i, i = 1, \dots, N$.
- Zakładamy, że x_i znane są *dokładnie*, a każdy wynik $\hat{y}_i$ jest zmienną losową opisaną rozkładem Gaussa o średniej $f(x_i; a_1, \dots, a_k)$ i dyspersji σ_i .

Pytanie:

Jak na podstawie opisanych wyżej wyników pomiarów wyznaczyć wartości parametrów a_j i ich niepewności?

Idea postępowania: Przyjmujemy, że to, co obserwujemy reprezentuje sytuację typową, to jest bliską sytuacji najbardziej prawdopodobnej.

Wniosek:

Będziemy szukali takiego zbioru wartości $a_j, j = 1, \dots, k$, dla którego otrzymany wynik odpowiada **maksimum gęstości prawdopodobieństwa**. Poszukiwanie maksimum iloczynu funkcji Gaussa (jako modelu gęstości prawdopodobieństwa otrzymanych wyników) prowadzi do warunku:

$$\mathcal{R} := \sum_{i=1}^N \frac{(\hat{y}_i - f(x_i; a_1, \dots, a_k))^2}{\sigma_i^2} \rightarrow \text{minimum}.$$

Poszukujemy takich wartości $a_j, j = 1, \dots, k$, dla których $\mathcal{R}$ przyjmuje wartość minimalną.

Minimum wielkości $\mathcal{R}$ zdefiniowanej powyżej pozwala znajdować dobre oceny wartości parametrów $a_j, j = 1, \dots, k$ także w przypadkach, gdy wyniki pomiarów, $\hat{y}_i$, podlegają rozkładowi innemu niż rozkład Gaussa, jeśli tylko dla każdego i spełnione są warunki:

$$\begin{aligned}\mathcal{E}(\hat{y}_i) &= f(x_i; a_1, \dots, a_k), \\ \mathcal{E}((\hat{y}_i - f(x_i; a_1, \dots, a_k))^2) &= \sigma_i^2,\end{aligned}$$

a wielkości występujące po prawych stronach obu równości istnieją i są skończone.

Metoda znajdowania wartości parametrów $a_j, j = 1, \dots, k$ poprzez minimalizację wielkości:

$$\mathcal{R} := \sum_{i=1}^N \frac{(\hat{y}_i - f(x_i; a_1, \dots, a_k))^2}{\sigma_i^2} \rightarrow \textit{minimum}$$

nazywana jest **metodą najmniejszych kwadratów**.

Uwaga:

Wzory służące do wyznaczania parametrów wyprowadzamy jakby wartości niepewności σ_i były znane, a obliczenia wykonujemy podstawiając zamiast σ_i oceny ich wartości – s_i (ogólnie: oceny niepewności $u(\hat{y}_i)$).

20 Pomiar o różnych dokładnościach – średnia ważona

Interesuje nas wartość wielkości x . Znamy wyniki N **niezależnych** pomiarów: $x = \hat{x}_i \pm s_i, i = 1, \dots, N$.

Pytanie: Jak na podstawie tych wyników najlepiej wyznaczyć wartość x ?

Zakładamy, że każda z wartości $\hat{x}_i$ jest zmienną losową opisaną rozkładem o średniej μ i dyspersji σ_i (znamy tylko przybliżone wartości $s_i \simeq \sigma_i$ oraz $\hat{x}_i \simeq \mu$).

Poszukujemy $\hat{x}$ – wielkości lepiej przybliżającej μ niż każda z wielkości $\hat{x}_i$. W tym celu konstruujemy wielkość:

$$\mathcal{R} := \sum_{i=1}^N \frac{(\hat{x}_i - \hat{x})^2}{\sigma_i^2}$$

i poszukujemy minimum tak zdefiniowanego $\mathcal{R}$ względem $\hat{x}$.

Wynik:

$$\hat{x} = \frac{\sum_{i=1}^N (\hat{x}_i / \sigma_i^2)}{\sum_{i=1}^N (1 / \sigma_i^2)}.$$

Jak łatwo sprawdzić dla $\hat{x}$ spełnione są związki:

$$\begin{aligned} \mathcal{E}(\hat{x}) &= \mu, \\ \mathcal{E}((\hat{x} - \mu)^2) &= \frac{1}{\sum_{i=1}^N (1 / \sigma_i^2)}. \end{aligned}$$

Jeżeli w powyższych wzorach zastąpimy wszystkie σ_i znanymi nam ich przybliżeniami, s_i , to otrzymujemy praktyczny przepis:

Odpowiedź 1: $x = \hat{x}_w \pm s_{int}$, gdzie oceny wartości *średniej ważonej*, $\hat{x}_w$, i *wewnętrznej* oceny dyspersji, s_{int} , zdefiniowane są wzorami:

$$\hat{x}_w := \frac{\sum_{i=1}^N (\hat{x}_i / s_i^2)}{\sum_{i=1}^N (1 / s_i^2)},$$

$$s_{int}^2 := \frac{1}{\sum_{i=1}^N (1 / s_i^2)}.$$

Model alternatywny

Nasze dane możemy też rozumieć następująco: wynik $\hat{x}_i \pm s_{xi}$ otrzymano na podstawie serii n_i bezpośrednich pomiarów wielkości x . W pomiarach bezpośrednich uzyskano wyniki x_{ij} przy czym: $i = 1, \dots, N$, a dla każdego i , $j = 1, \dots, n_i$.

Przy takiej interpretacji, podane wartości $\hat{x}_i, s_{xi}$ otrzymane byłyby więc w następujący sposób:

$$\hat{x}_i = \frac{1}{n_i} \sum_{j=1}^{n_i} x_{ij},$$

$$s_i^2 = \frac{1}{n_i(n_i - 1)} \sum_{j=1}^{n_i} (x_{ij} - \hat{x}_i)^2.$$

Jeśli wszystkie pomiary wykonano w tych samych warunkach, z użyciem tych samych przyrządów, to:

$$\mathcal{E}(s_i^2) = \frac{\sigma^2}{n_i} \Rightarrow n_i \simeq \frac{\sigma^2}{s_i^2},$$

gdzie σ jest nieznaną, dokładną wartością dyspersji (niepewności) pojedynczego pomiaru.

Najlepszą ocenę wielkości x będzie więc średnia arytmetyczna **wszystkich** pojedynczych pomiarów:

$$\begin{aligned}\hat{x} &= \frac{1}{\sum_{i=1}^N n_i} \sum_{i=1}^N \sum_{j=1}^{n_i} x_{ij} \\ &= \frac{1}{\sum_{i=1}^N n_i} \sum_{i=1}^N n_i \hat{x}_i \\ &\simeq \frac{1}{\sum_{i=1}^N (\sigma^2/s_i^2)} \sum_{i=1}^N \frac{\sigma^2 \hat{x}_i}{s_i^2}.\end{aligned}$$

Po skróceniu licznika i mianownika przez σ^2 otrzymujemy $\hat{x} = \hat{x}_w$, gdzie $\hat{x}_w$ oznacza średnią ważoną zdefiniowaną poprzednio (minimalizacja $\mathcal{R}$).

W celu wyznaczenia oceny niepewności $\hat{x}_w$ zbadajmy wielkość:

$$A := \sum_{i=1}^N \sum_{j=1}^{n_i} (x_{ij} - \hat{x})^2 \Rightarrow \mathcal{E}(A) = \left(\sum_{i=1}^N n_i - 1 \right) \sigma^2.$$

Mamy także

$$\begin{aligned} \sum_{i=1}^N \sum_{j=1}^{n_i} (x_{ij} - \hat{x})^2 &= \sum_{ij} (x_{ij} - \hat{x}_i + \hat{x}_i - \hat{x})^2 \\ &= \sum_{i=1}^N \left[\sum_{j=1}^{n_i} (x_{ij} - \hat{x}_i)^2 + n_i (\hat{x}_i - \hat{x})^2 - 2(\hat{x} - \hat{x}_i) \left(\sum_{j=1}^{n_i} x_{ij} - n_i \hat{x}_i \right) \right]. \end{aligned}$$

Obliczając wartość oczekiwaną obu stron równości dostajemy:

$$\begin{aligned} \left(\sum_{i=1}^N n_i - 1 \right) \sigma^2 &= \sum_{i=1}^N (n_i - 1) \sigma^2 + \mathcal{E} \left(\sum_{i=1}^N n_i (\hat{x}_i - \hat{x})^2 \right) \\ (N - 1) \sigma^2 &= \mathcal{E} \left(\sum_{i=1}^N n_i (\hat{x}_i - \hat{x})^2 \right). \end{aligned}$$

Wniosek:

$$\begin{aligned} \sigma^2 &= \frac{1}{N - 1} \mathcal{E} \left(\sum_{i=1}^N n_i (\hat{x}_i - \hat{x})^2 \right) \\ \sigma_{\hat{x}}^2 &= \frac{\sigma^2}{\sum_{i=1}^N n_i}, \end{aligned}$$

co po podstawieniu $n_i \simeq \sigma^2 / s_i^2$ pozwala na uzyskanie zewnętrznej oceny $\sigma_{\hat{x}}^2$:

$$s_{ext}^2 := \frac{\sum_{i=1}^N [(\hat{x}_i - \hat{x}_w)^2 / s_i^2]}{(N - 1) \sum_{i=1}^N (1 / s_i^2)}.$$

Odpowiedź 2: $x = \hat{x}_w \pm s_{ext}$.

Odpowiedź ostateczna: $x = \hat{x}_w \pm \text{Max}\{s_{int}, s_{ext}\}$.

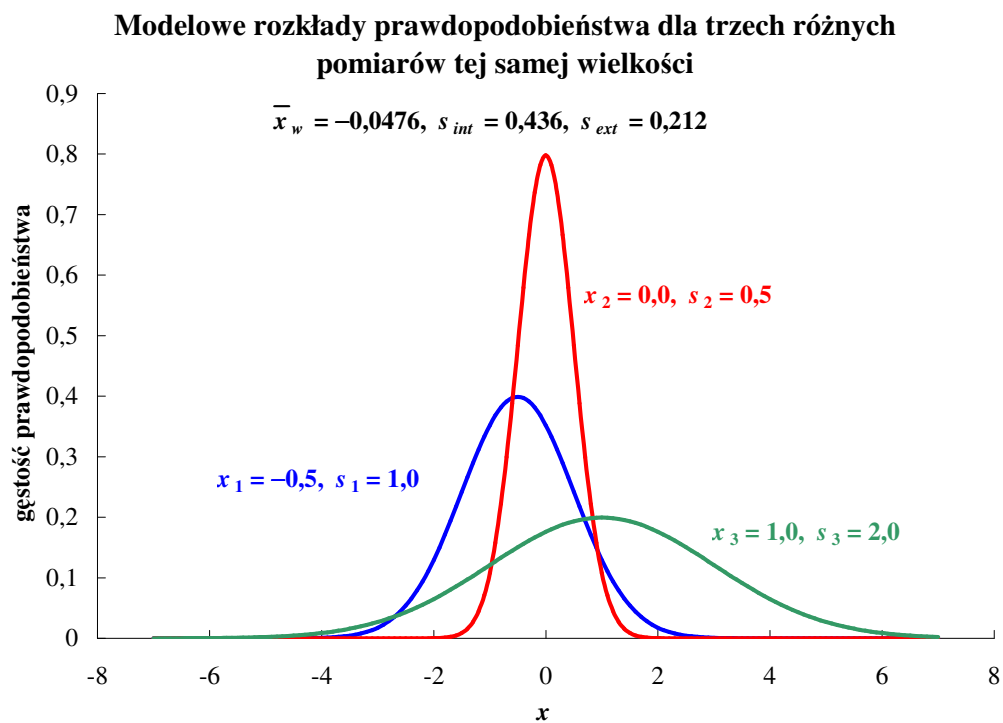


Figure 7: Średnia ważona pomiarów o różnych dokładnościach: $s_{int} > s_{ext}$ (ilustracja wzorowana na: W. H. Heini Gränicher, *Messung beendet – was nun?*, B. G. Teubner Stuttgart, 1994).

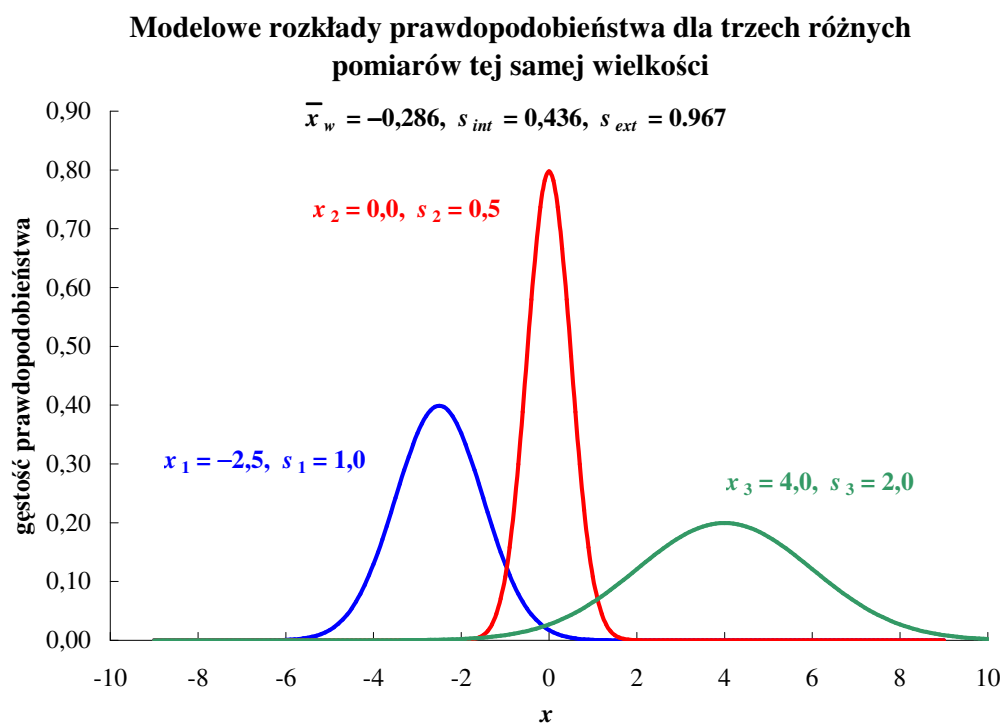


Figure 8: Średnia ważona pomiarów o różnych dokładnościach: $s_{int} < s_{ext}$ (ilustracja wzorowana na: W. H. Heini Gränicher, *Messung beendet – was nun?*, B. G. Teubner Stuttgart, 1994).

21 Wyznaczanie współczynnika proporcjonalności: $y = ax$

Dla serii **dokładnie** znanych wartości $x_i, i = 1, \dots, N$ znamy wyniki pomiarów odpowiadających im wartości $y_i = \hat{y}_i \pm s_i$.

Wiemy (zakładamy), że wielkości x i y spełniają zależność: $y = ax$, gdzie a jest pewnym nieznanym nam parametrem.

Jak wyznaczyć najlepszą ocenę wartości a na podstawie znanych nam wyników pomiarów: $x_i \rightarrow \hat{y}_i \pm s_i$?

Posłużymy się *metodą najmniejszych kwadratów*. Zakładamy, że każdy z wyników pomiarów $\hat{y}_i$ jest zmienną losową opisaną rozkładem o wartości oczekiwanej ax_i i dyspersji σ_i . Znane nam niepewności pomiarów $\hat{y}_i$ dają nam ocenę wartości $\sigma_i \rightarrow s_i \simeq \sigma_i$. Konstruujemy wielkość:

$$\mathcal{R} = \sum_{i=1}^N \frac{(\hat{y}_i - ax_i)^2}{\sigma_i^2}$$

i poszukujemy minimum $\mathcal{R}$ ze względu na a . Otrzymujemy:

$$\hat{a} = \frac{\sum_{i=1}^N (\hat{y}_i x_i / \sigma_i^2)}{\sum_{i=1}^N (x_i^2 / \sigma_i^2)}$$

Jak łatwo sprawdzić:

$$\begin{aligned} \mathcal{E}(\hat{a}) &= a, \\ \mathcal{E}((\hat{a} - a)^2) &= \frac{1}{\sum_{i=1}^N (x_i^2 / \sigma_i^2)}. \end{aligned}$$

W obliczeniach podstawiamy w miejsce nieznanych nam *dokładnych* wartości dyspersji σ_i znane nam oceny przybliżone s_i (ogólniej: oceny niepewności u_i).

W przypadku, gdy nie są znane niepewności σ_i ani ich oceny s_i korzystamy z założenia, że wszystkie one są sobie równe, $\sigma_i = \sigma$, a wartość σ wyznaczymy w przybliżeniu korzystając z twierdzenia:

Jeśli $f(x, a_1, \dots, a_k)$ jest liniową funkcją parametrów $a_1, \dots, a_k$, to:

$$\mathcal{E}(\mathcal{R}) = \mathcal{E} \left(\sum_{i=1}^N \frac{(\hat{y}_i - f(x_i; \hat{a}_1, \dots, \hat{a}_k))^2}{\sigma_i^2} \right) = N - k,$$

gdzie N – liczba punktów pomiarowych (danych), a k – liczba wyznaczanych parametrów zależności $y = f(x; a_1, \dots, a_k)$.

Mamy więc ostatecznie dla $y = ax$:

- Dane: $x_i \rightarrow \hat{y}_i \pm s_i$, to $a = \hat{a} \pm s_{\hat{a}}$,

$$\hat{a} = \frac{\sum_{i=1}^N (\hat{y}_i x_i / s_i^2)}{\sum_{i=1}^N (x_i^2 / s_i^2)},$$

$$s_{\hat{a}}^2 = \frac{1}{\sum_{i=1}^N (x_i^2 / s_i^2)}$$

- Dane: $x_i \rightarrow \hat{y}_i$, to $a = \hat{a} \pm s_{\hat{a}}$,

$$\hat{a} = \frac{\sum_{i=1}^N \hat{y}_i x_i}{\sum_{i=1}^N x_i^2},$$

$$s_{\hat{a}}^2 = \frac{\sum_{i=1}^N (\hat{y}_i - \hat{a} x_i)^2}{(N - 1) \sum_{i=1}^N x_i^2}$$

22 Wyznaczanie parametrów prostej: $y = ax + b$

Dla dokładnie znanych wartości x_i znamy wyniki pomiarów wartości $\hat{y}_i$. Znamy (zakładamy) zależność $y = ax + b$. Analogicznie jak poprzednio metoda najmniejszych kwadratów pozwala nam wyznaczyć najlepsze oceny wartości parametrów a i b .

Uwaga: wyznaczone wartości $\hat{a}$ i $\hat{b}$ nie są już niezależne, to znaczy, że *kowariancja* $C_{\hat{a}\hat{b}} := \mathcal{E}((\hat{a} - a)(\hat{b} - b)) \neq 0$

Oznaczenie: $w_i := 1/s_i^2$

- Dane: $x_i \rightarrow \hat{y}_i \pm s_i$ (znane niepewności $\hat{y}_i$), to
 $a = \hat{a} \pm s_{\hat{a}}, b = \hat{b} \pm s_{\hat{b}},$

$$\begin{aligned} D &= \sum_{i=1}^N w_i \sum_{j=1}^N w_j x_j^2 - \left(\sum_{i=1}^N x_i w_i \right)^2 \\ \hat{a} &= \frac{1}{D} \left(\sum_{i=1}^N w_i \sum_{j=1}^N w_j x_j \hat{y}_j - \sum_{i=1}^N w_i x_i \sum_{j=1}^N w_j \hat{y}_j \right) \\ \hat{b} &= \frac{1}{D} \left(\sum_{i=1}^N w_i x_i^2 \sum_{j=1}^N w_j \hat{y}_j - \sum_{i=1}^N w_i x_i \sum_{j=1}^N w_j x_j \hat{y}_j \right) \\ s_{\hat{a}}^2 &= \frac{1}{D} \sum_{i=1}^N w_i \\ s_{\hat{b}}^2 &= \frac{1}{D} \sum_{i=1}^N w_i x_i^2 \\ C_{\hat{a}\hat{b}} &= -\frac{1}{D} \sum_{i=1}^N w_i x_i, \end{aligned}$$

gdzie $C_{\hat{a}\hat{b}}$ oznacza kowariancję ocen wartości parametrów a i b .

Rysunek 9. ilustruje wpływ wielkości niepewności poszczególnych y_i na wartości parametrów i ich niepewności. Wartości y_i na wszystkich wykresach są te same, a niepewności oznaczone pionowymi kreskami.

- Dane: $x_i \rightarrow \hat{y}_i$ (nie są znane niepewności $\hat{y}_i$), to
 $a = \hat{a} \pm s_{\hat{a}}, b = \hat{b} \pm s_{\hat{b}},$

$$\begin{aligned}
 D &= N \sum_{i=1}^N x_i^2 - \left(\sum_{i=1}^N x_i \right)^2 \\
 \hat{a} &= \frac{1}{D} \left(N \sum_{i=1}^N x_i \hat{y}_i - \sum_{i=1}^N x_i \sum_{j=1}^N \hat{y}_j \right) \\
 \hat{b} &= \frac{1}{D} \left(\sum_{i=1}^N x_i^2 \sum_{j=1}^N \hat{y}_j - \sum_{i=1}^N x_i \sum_{j=1}^N x_j \hat{y}_j \right) \\
 s^2 &= \frac{1}{N-2} \sum_{i=1}^N (\hat{y}_i - \hat{a}x_i - \hat{b})^2 \\
 s_{\hat{a}}^2 &= \frac{Ns^2}{D} \\
 s_{\hat{b}}^2 &= \frac{s^2 \sum_{i=1}^N x_i^2}{D} \\
 C_{\hat{a}\hat{b}} &= -\frac{s^2 \sum_{i=1}^N x_i}{D}
 \end{aligned}$$

- W zasadzie **zawsze obie** zmienne obarczone są niepewnościami pomiarowymi. Za zmienną niezależną x przyjmujemy wówczas wielkość mierzona **dokładniej**, tzn. z **mniejszą** niepewnością. Niech σ_y i σ_x oznaczają, odpowiednio niepewności pomiaru zmiennych x i y . Wówczas zmienna x **jest mierzona dokładniej**, jeśli $\sigma_y > \hat{a}\sigma_x$, gdzie $\hat{a}$ oznacza ocenę wartości współczynnika kierunkowego prostej. Takiej oceny można dokonać w sposób przybliżony na podstawie wykresu, przed przystąpieniem do obliczeń.

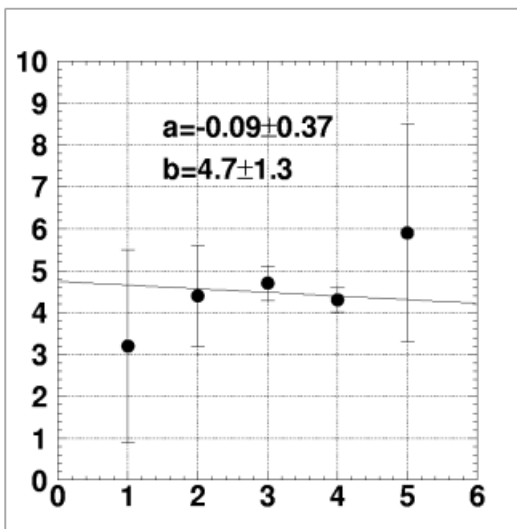
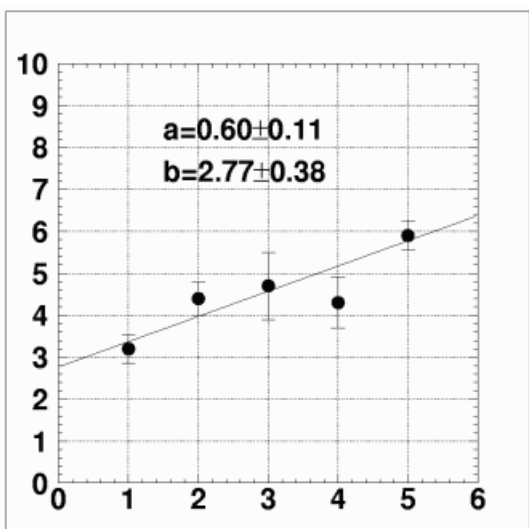
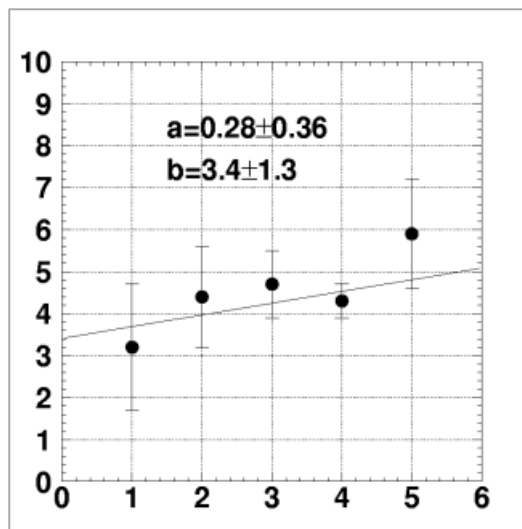
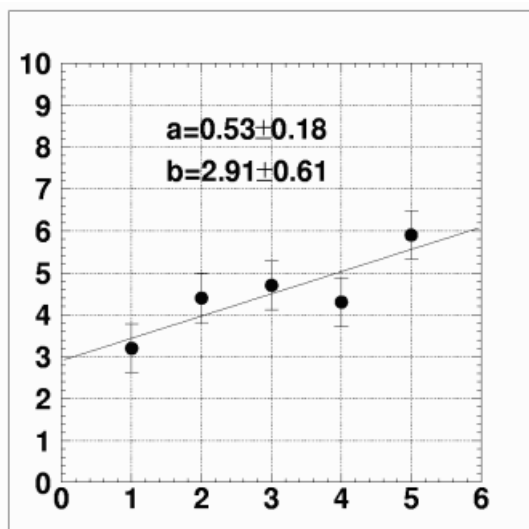


Figure 9: Ilustracja wpływu rozkładu niepewności pomiarowych na wartości parametrów prostej wyznaczonych metodą najmniejszych kwadratów.

- Wyznaczone metodą najmniejszych kwadratów parametry funkcji mogą posłużyć do wyznaczania przybliżonej wartości y w punktach, dla których nie wykonywano pomiarów – oznaczmy tę wartość jako $\tilde{y}$. W przypadku omawianego tu dopasowania parametrów prostej, dla **dokładnie znanej** wartości x mamy:

$$\begin{aligned}\tilde{y} &= \hat{a}x + \hat{b} \\ u^2(\tilde{y}) &= x^2\sigma_{\hat{a}}^2 + 2xC_{\hat{a},\hat{b}} + \sigma_{\hat{b}}^2,\end{aligned}$$

gdzie za $\sigma_{\hat{a}}^2, C_{\hat{a},\hat{b}}, \sigma_{\hat{b}}^2$ podstawiamy, odpowiednio, wyznaczone poprzednio ich oceny $s_{\hat{a}}^2, \hat{C}_{\hat{a},\hat{b}}, s_{\hat{b}}^2$.

23 Błąd systematyczny i dokładność przyrządów

Błąd systematyczny przy wielokrotnym powtarzaniu pomiarów tej samej wielkości w warunkach praktycznie niezmiennych pozostaje stały lub zmienia się według określonego prawa wraz ze zmianą warunków.

Praktycznie niezmiennie warunki:

- pomiary przeprowadzane są tą samą metodą,
- pomiary przeprowadza się tym samym przyrządem pomiarowym,
- pomiary przeprowadza ten sam obserwator,
- pomiary przeprowadzane są w tym samym miejscu,
- pomiary powtarzane są w krótkich odstępach czasu,
- podczas wykonywania pomiarów panują te same warunki zewnętrzne (temperatura, ciśnienie, wilgotność, oświetlenie,...).

(Przytoczona tu analiza zaczerpnięta została z:

W. Jakubiec i J. Malinowski *Metrologia wielkości geometrycznych*, Wydawnictwa Naukowo-Techniczne, Warszawa, 1999).

Eliminacja błędu systematycznego:

- usunięcie źródeł błędu;
- wprowadzenie poprawek do wyniku pomiaru:
 - obliczonych,
 - wyznaczonych doświadczalnie (pomiary różnymi metodami).

Rodzaje błędów systematycznych (przykłady):

- błąd temperaturowy;
- błąd odkształceń sprężystych (np. przy pomiarze długości);
- błąd metody (np. pomiar oporu metodą jednoczesnego pomiaru prądu i spadku potencjału $\rightarrow$ *układ dokładnego pomiaru prądu* albo *układ dokładnego pomiaru spadku potencjału*);
- błąd odczytu przyrządu (paralaksa, interpolacja, błąd kwantowania);
- błąd histerezy (spowodowany np. tarciem, albo luzami części ruchomych przyrządów);
- **dokładność przyrządu.**

Przykład 9. Eliminacja błędu temperaturowego.

(Przykład zaczerpnięty z podręcznika W. Jakubca i J. Malinowskiego *Metrologia wielkości geometrycznych*, WNT, Warszawa, 1999.)

Przyjmuje się, że podawane wymiary przedmiotów odpowiadają temperaturze $t_0 = 20^\circ\text{C}$.

Długość mierzono w temperaturze $t \neq 20^\circ\text{C} \rightarrow$ otrzymano wynik L .

Wartość poprawiona $\rightarrow L' = L - \delta_t$,

$$\delta_t = L [\alpha_p(t_p - t_0) - \alpha_n(t_n - t_0)],$$

gdzie α_p to współczynnik rozszerzalności temperaturowej mierzonego przedmiotu, α_n – współczynnik rozszerzalności temperaturowej przyrządu pomiarowego (linijki, taśmy mierniczej..), t_p – temperatura przedmiotu, t_n temperatura przyrządu.

Niepewność poprawionego wyniku obliczamy korzystając z prawa *propagacji małych niepewności*:

$$\begin{aligned}\varepsilon_{L'}^2 &= \varepsilon_L^2 + \varepsilon_{\delta t}^2, \\ \varepsilon_{\delta t}^2 &= L^2 [\alpha_p^2 \varepsilon_{t_p}^2 + \alpha_n^2 \varepsilon_{t_n}^2 + (t_p - t_0)^2 \varepsilon_{\alpha_p}^2 + (t_n - t_0)^2 \varepsilon_{\alpha_n}^2],\end{aligned}$$

gdzie $\varepsilon_{(.)}$ oznacza niepewność odpowiedniej wielkości.

Uwzględnienie skończonej dokładności przyrządów – cd.

Po wprowadzeniu poprawek uwzględniających wszystkie rozpoznane źródła błędów systematycznych pozostaje jeszcze błąd związany z kalibracją przyrządów. Kalibracja przyrządów pomiarowych nigdy nie jest idealna. Błędy kalibracji (np. naniesienia skali) przejawiają się jako błędy systematyczne. Moglibyśmy je wyeliminować, gdybyśmy powtarzali serie pomiarowe, w każdej z nich stosując inny egzemplarz przyrządu danego typu i uśredniając wyniki wszystkich tak otrzymanych serii. Zamiast tego posłużymy się modelem statystycznym (tzw. *metoda randomizacji i centryzacji błędu systematycznego*).

Zakładamy, że:

- używany przyrząd jest losowo wybranym przedstawicielem klasy przyrządów;
- niedokładność każdego z przyrządów tej klasy przejawia się jako błąd systematyczny ξ ;
- wewnątrz danej klasy wartość ξ jest zmienną losową opisaną pewnym rozkładem prawdopodobieństwa, spełniającym warunki:

$$\begin{aligned}\mathcal{E}(\xi) &= 0, \\ \mathcal{E}(\xi^2) &= \sigma^2(\xi),\end{aligned}$$

przy czym ξ jest zmienną losową **niezależną** od innych czynników losowych (błędów losowych) prowadzących do rozrzutu wyników wewnątrz jednej serii pomiarowej.

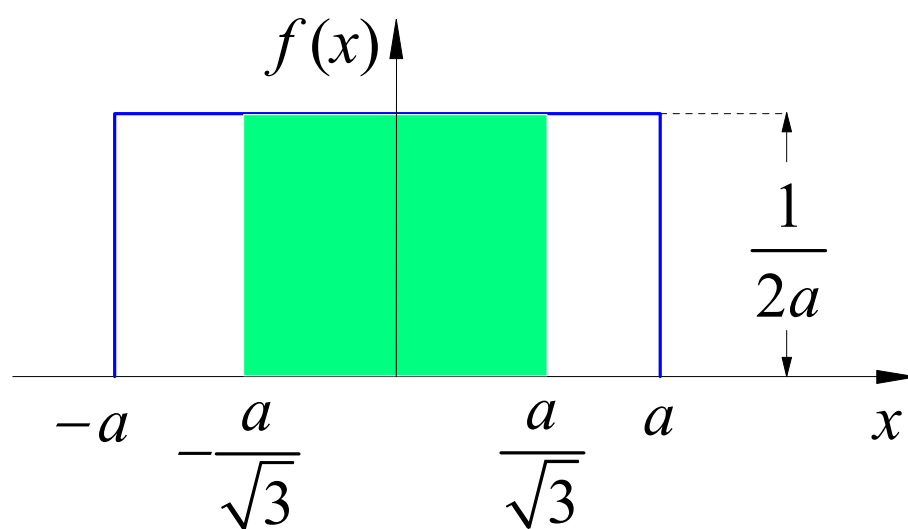
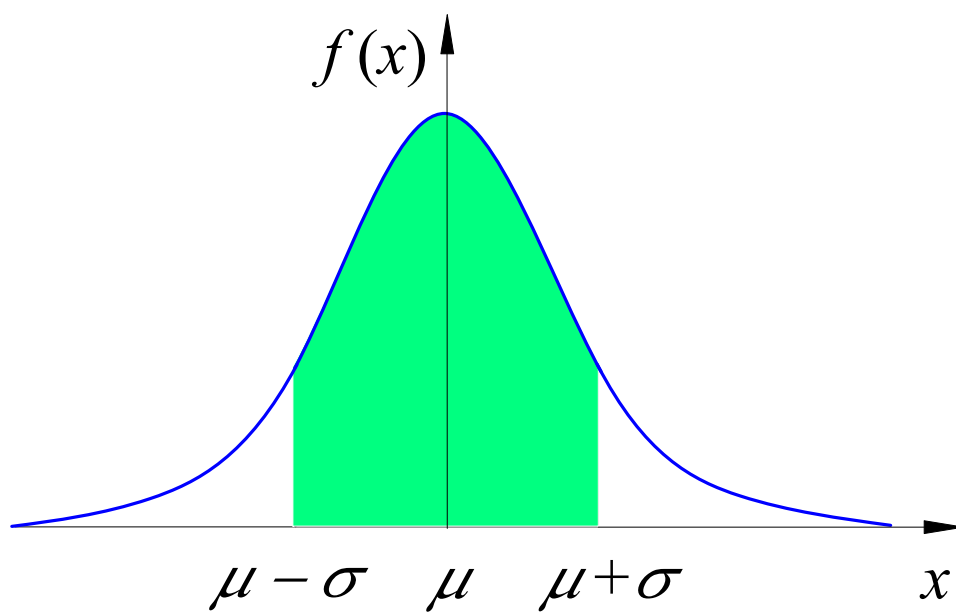


Figure 10: Przykłady modelowych rozkładów prawdopodobieństwa błędu wskazań przyrządu: rozkład Gaussa – górny wykres, rozkład jednostajny – dolny wykres.

Wartość $\sigma^2(\xi)$ wyznaczamy na podstawie:

- szczegółowych informacji podanych przez producenta,
- modelowego rozkładu tej części błędu systematycznego.

Najczęściej stosowane modele rozkładu prawdopodobieństwa dla *losowej składowej błędu systematycznego* (tzw. *randomizacja i centryzacja błędu systematycznego*):

- rozkład Gaussa (normalny) o dyspersji $\sigma = \Delta/3$;
- rozkład jednostajny o $a = \Delta$ (oznaczenia jak na wykresie); wówczas $\sigma^2(\xi) = \Delta^2/3$.

Wartość Δ określamy na podstawie:

- klasy przyrządu, albo
- jako *najmniejszą działkę skali*.

Uwagi:

- wielokrotne (n krotne) „przykładanie” przyrządu w celu zmierzenia jednej wielkości (np. nasza linijka jest krótsza od mierzonego odcinka) prowadzi do oceny:

$$\sigma^2(\xi) = n\sigma_{prz\acute{y}rz}^2;$$

- pomiar wielokrotności (n) interesującej nas wielkości (np. pomiar czasu trwania n okresów wahadła, pomiar grubości n kartek papieru...) prowadzi do oceny:

$$\sigma^2(\xi) = \frac{1}{n^2}\sigma_{prz\acute{y}rz}^2.$$

Wynik podajemy w postaci:

$$\begin{aligned}x &= \bar{x} \pm u(\bar{x}), \\ u^2(\bar{x}) &= s_{\bar{x}}^2 + \sigma^2(\xi),\end{aligned}$$

gdzie $\bar{x}$ oznacza średnią serii pomiarów poprawioną ze względu na wszystkie rozpoznane źródła błędów systematycznych, a $s_{\bar{x}}$ oznacza niepewność średniej (uwzględniającą niepewności wprowadzanych poprawek poprzez prawo propagacji małych błędów).

24 Rozkład χ^2

Niech $x_1, \dots, x_N$ będą **niezależnymi** zmiennymi losowymi, każda opisana rozkładem Gaussa o wartości średniej $\mu = 0$ i dyspersji $\sigma = 1$. Wówczas wielkość:

$$x := \chi^2 = \sum_{i=1}^N x_i^2$$

podlega **rozkładowi** χ^2 o gęstości prawdopodobieństwa zadanej wzorem:

$$f(x; N) = \frac{1}{2\Gamma(\frac{N}{2})} \left(\frac{x}{2}\right)^{\frac{N}{2}-1} \exp(-x/2),$$

gdzie N nazywane jest liczbą stopni swobody,

$$\Gamma(\tfrac{1}{2}) = \sqrt{\pi}, \quad \Gamma(1) = 1, \quad \Gamma(z+1) = z\Gamma(z),$$

a dla z równego liczbie naturalnej n , $\Gamma(n+1) = n!$.

$$\begin{aligned}\mathcal{E}(x) &= N \\ \mathcal{E}((x-N)^2) &= 2N\end{aligned}$$

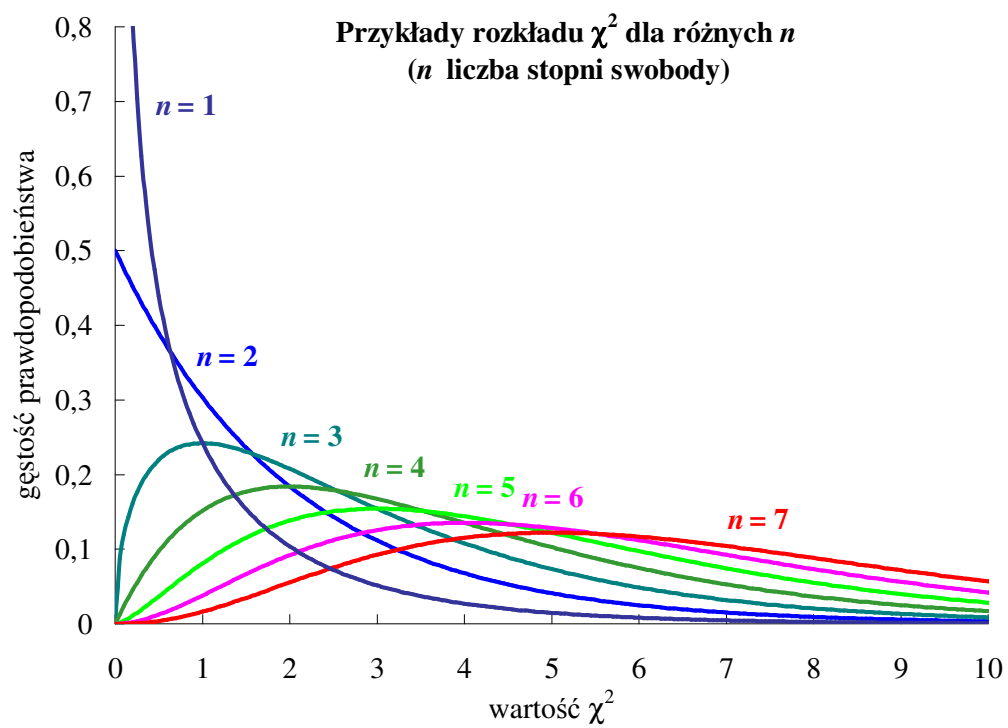


Figure 11: Wykresy gęstości rozkładu χ^2 dla małych liczb stopni swobody n .

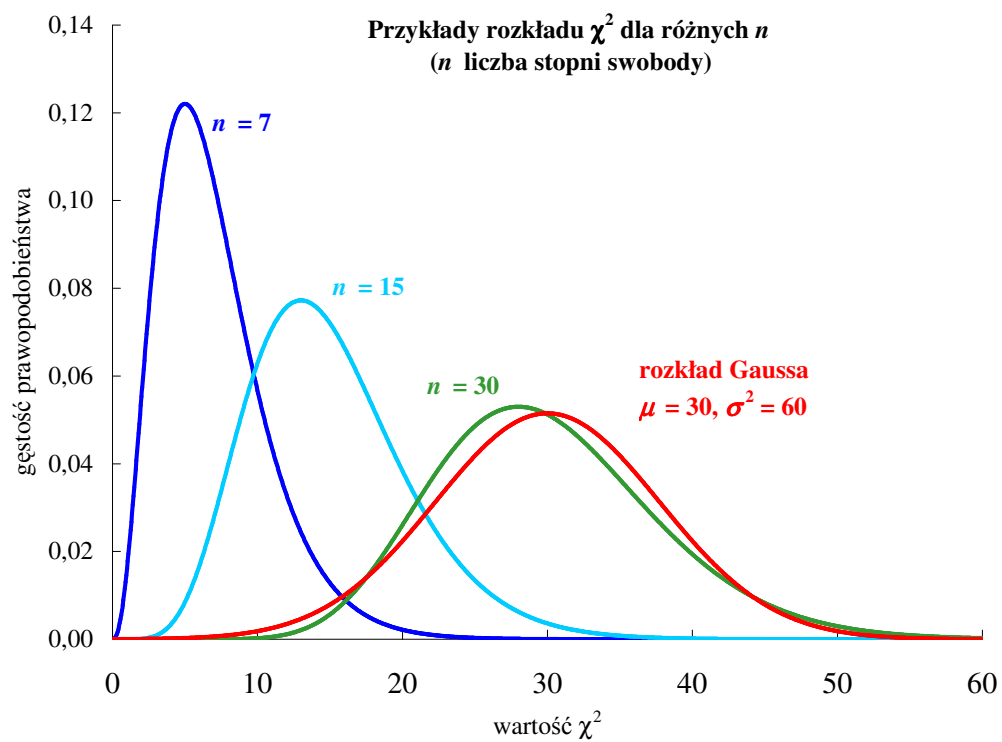


Figure 12: Wykresy gęstości rozkładu χ^2 dla dużych liczb stopni swobody n .

25 Test zgodności dopasowania – test χ^2

Sprawdzenie dopasowania zależności funkcyjnej do wyników pomiarów.

Hipoteza:

$$y = f(x; a_1, a_2, \dots, a_k).$$

Pomiary:

$$x_i \rightarrow y_i \pm \sigma_i, i = 1, \dots, N > k,$$

gdzie x_i i σ_i znane są dokładnie (założenie).

Tworzymy wielkość:

$$\mathcal{R} = \sum_{i=1}^N \frac{(y_i - f(x_i; a_1, \dots, a_k))^2}{\sigma_i^2}.$$

Wówczas:

1. Jeśli parametry $f(x; a_1, \dots, a_k)$ znane są dokładnie (tzn. ich wartości liczbowe są częścią testowanej hipotezy) to:

$$\mathcal{E}(\mathcal{R}) = N.$$

Jeśli dodatkowo, dla każdego i , y_i podlega rozkładowi Gaussa o wartości oczekiwanej $f(x_i; a_1, \dots, a_k)$ i dyspersji σ_i to $\mathcal{R}$ podlega rozkładowi χ^2 o N stopniach swobody.

2. Jeśli $f(x; a_1, \dots, a_k)$ jest liniową funkcją parametrów $a_1, \dots, a_k$, a wartości tych parametrów wyznaczone zostały na podstawie danych doświadczalnych (metoda najmniejszych kwadratów $\rightarrow$ minimalizacja wartości $\mathcal{R} \rightarrow \hat{a}_1, \dots, \hat{a}_k$), to

$$\mathcal{E}(\mathcal{R}) = N - k,$$

gdy wartość $\mathcal{R}$ obliczna jest po podstawieniu $\hat{a}_1, \dots, \hat{a}_k$.

Jeśli dodatkowo, dla każdego i , y_i podlega rozkładowi Gaussa o wartości oczekiwanej $f(x_i; a_1, \dots, a_k)$ i dyspersji σ_i , to wartość $\mathcal{R}$ obliczona dla wyznaczonych za pomocą minimalizacji $\mathcal{R}$ wartości $\hat{a}_i$ podlega rozkładowi χ^2 o $N - k$ stopniach swobody.

Test:

Hipotezę odrzucamy, jeżeli wynikające z rozkładu χ^2 o odpowiadającej naszemu zagadnieniu liczbie stopni swobody prawdopodobieństwo uzyskania otrzymanej wartości $\mathcal{R}$ lub większej jest mniejsze od przyjętej przez nas wartości krytycznej (np. 0,05) – test jednostronny. Przyjęta wartość krytyczna równa jest akceptowanemu przez nas prawdopodobieństwu popełnienia **błędu I rodzaju**.

W przypadku 1. odrzucenie hipotezy oznacza, że podane wartości parametrów lub podany kształt zależności, to znaczy postać funkcji $f(x, a_1, \dots, a_k)$, nie opisują otrzymanych wyników pomiarów.

W przypadku 2. odrzucenie hipotezy oznacza, że podany kształt zależności, to znaczy postać funkcji $f(x, a_1, \dots, a_k)$, nie opisuje otrzymanych wyników pomiarów.

Uwagi:

- Jeśli otrzymana wartość $\mathcal{R}$ jest zbyt mała (wyraźnie mniejsza od liczby stopni swobody), to najprawdopodobniej:
 - pomiary nie były niezależne lub
 - przeceniliśmy wartości niepewności pomiarowych.
- W praktyce używamy zamiast dokładnych wartości niepewności jedynie dostępnych nam ocen tych wartości. Nie sprawdzamy też przeważnie zasadności założenia o gaussowskim rozkładzie wyników pomiarów. Opisany powyżej test jest też często stosowany w przypadkach, gdy $f(x; a_1, \dots, a_k)$ **nie jest** liniową funkcją parametrów $a_1, \dots, a_k$. Wszystkie te czynniki zmniejszają konkluzywność przeprowadzonego testu (tzn. w niekontrolowany sposób zmieniają prawdopodobieństwo błędu I rodzaju w stosunku do przyjętej wartości krytycznej).

26 Sprawdzenie zgodności obserwowanego rozkładu wyników z zadaniem (przewidywanym) rozkładem prawdopodobieństwa.

Hipoteza:

Prawdopodobieństwo otrzymania wartości $x_i, i = 1, \dots, M$
(otrzymania wartości x z przedziału o numerze i) $\rightarrow p_i$

$$\sum_{i=1}^M p_i = 1.$$

Pomiary:

W przedziale i otrzymaliśmy n_i przypadków: $i \rightarrow n_i$,

$$\sum_{i=1}^M n_i = N > M.$$

Na podstawie testowanej hipotezy przewidywana liczba przypadków $i \rightarrow p_i N$.

Założenie:

Wartości $p_i N$ są duże (ten warunek musimy uwzględnić planując wykonanie testu). Warunek ten oznacza, że obserwowanym liczbom przypadków n_i można przypisać niepewności $\sigma_i \simeq \sqrt{p_i N}$ lub, inaczej, że n_i podlegają w przybliżeniu rozkładowi Poissona.

Tworzymy wielkość:

$$\mathcal{R} = \sum_{i=1}^M \frac{(n_i - p_i N)^2}{p_i N}.$$

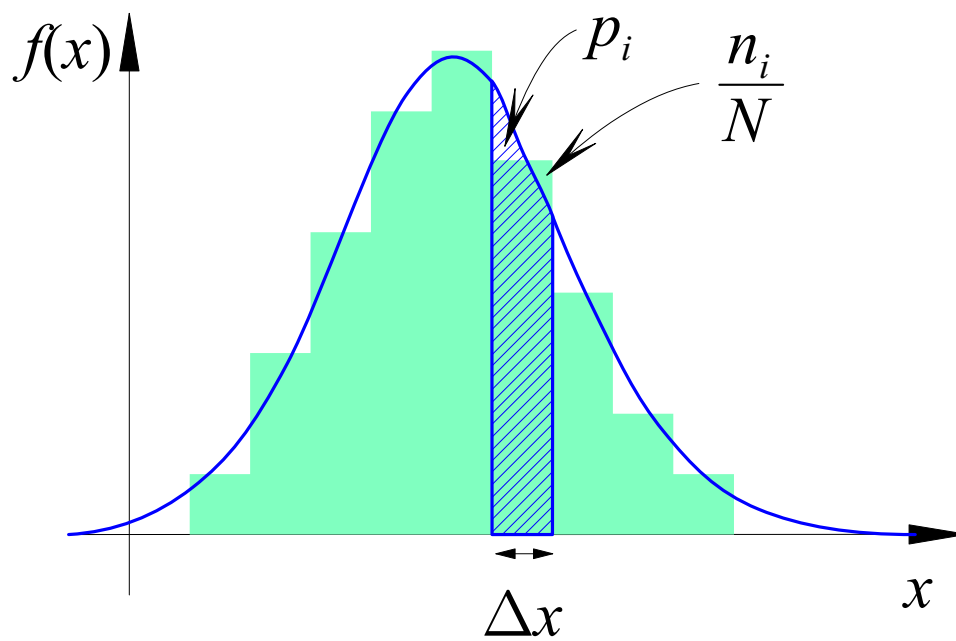


Figure 13: Test zgodności rozkładu obserwowanego – histogram wartości n_i/N – z modelową gęstością prawdopodobieństwa – p_i obliczane jako pole pod krzywą.

Wówczas:

$\mathcal{R}$ podlega rozkładowi χ^2 o $M - 1$ stopniach swobody.

Każdy dodatkowy warunek nałożony na wartości n_i poza warunkiem

$$\sum_{i=1}^M n_i = N$$

zmniejsza liczbę stopni swobody o 1, to znaczy:

k dodatkowych warunków $\rightarrow$ liczba stopni swobody $= M - k - 1$.

Test: Porównanie obliczonej wartości $\mathcal{R}$ z wartością krytyczną rozkładu χ^2 .

Praktyka: Często zalecane jest, by test stosować, gdy $p_i N \geq 5$, $M \geq 6 \Rightarrow N \geq 30$.

Uwaga: Można dopasować szerokości przedziałów do wartości $p_i N$.

27 Jak piszemy pracę naukową?

Przed rozpoczęciem pisania tekstu należy jasno odpowiedzieć sobie na dwa pytania:

- Co jest **najważniejszą informacją**, jaką chcemy przekazać?
- Dla **kogo** piszemy naszą pracę – kto ją będzie czytał?

Układ tekstu pracy:

- tytuł → informuje o zawartości;
- autor (autorzy) i ich adresy (uczelnia, instytut badawczy...);
- streszczenie → rozszerza tytuł;
- wstęp → dlaczego zajęliśmy się opisywanym zagadnieniem (dotychczasowy stan badań...);
- opis używanych metod: podstawy teoretyczne, używane wzory, stosowane przybliżenia, definicje i oznaczenia podstawowych wielkości...;
- techniczna realizacja pomiaru: schemat układu pomiarowego, użyta aparatura (typ i producent), metody kontroli warunków pomiaru (np. kalibrowania przyrządów);
- uzyskane wyniki: (a) syntetyczna prezentacja wyników pomiarów, (b) dyskusja źródeł niepewności i ocena jej wartości;
- wnioski i podsumowanie;
- podziękowania (współpracownicy, którzy nie są współautorami, sponsorzy,);
- dodatki (nie zawsze) → wyprowadzenie mało znanych zależności, szczegóły techniczne procedur pomocniczych,...;
- cytowana literatura.

Uwagi dodatkowe:

- Używany język powinien być prosty i jasny, a zdania krótkie.
- Wzory matematyczne to też proza: $a = b + c$, to pełne zdanie.
- Wszystkie oznaczenia muszą być zdefiniowane **przed** ich pierwszym użyciem (lub tuż po).
- Używane wzory powinny być wyprowadzone lub powinien być podany odnośnik do pracy (podręcznika), gdzie takie wyprowadzenie można znaleźć.
- Rysunki, tabele, odnośniki muszą być ponumerowane według kolejności ich omawiania w tekście, przy czym, **oddzielna numeracja** dotyczy rysunków, tabel i odnośników.
- W miarę możliwości należy cytować prace źródłowe, a jeśli nie, to powszechnie dostępne prace przeglądowe lub podręczniki.
- O ile to możliwe, stosujemy standardowe oznaczenia występujących wielkości, np.: m – masa, $\vec{r}$ – położenie, $\vec{v}$ – prędkość, R – opór elektryczny, t – czas, T – temperatura w skali Kelvina.
- Wykres przedstawianej przez nas zależności powinien wypełniać całe przeznaczone dla niego pole; nachylenie prostej powinno być bliskie 45° .
- Osie wykresów muszą być wyraźnie opisane z podaniem jednostek, w jakich mierzone są przedstawiane wielkości.
- Tabele powinny zawiera nazwy prezentowanych wielkości i jednostki, w których są one mierzone.
- **Główna informacja** powinna wynikać z tytułu, streszczenia, wykresów (rysunków) i tabel.