



Fundusze Europejskie
dla Wielkopolski

Dofinansowane przez
Unię Europejską



SAMORZĄD
WOJEWÓDZTWA
WIELKOPOLSKIEGO



Automatyka i robotyka

**Skrypt dla studentów
studiów niestacjonarnych**

Rok I

redakcja

Robert Cieślak

Konin 2024

Tytuł

Automatyka i robotyka
Skrypt dla studentów studiów niestacjonarnych
Rok I

Redakcja naukowa

Robert Cieślak

Autorzy

Robert Cieślak
Bohdana Hładysz
Monika Muszyńska
Paweł Sobczak
Piotr Świta

Projekt pn.
„Rozwój studiów o profilu praktycznym i form kształcenia ustawicznego
dostosowanych do potrzeb Wielkopolski Wschodniej”,
realizowany przez Akademię Nauk Stosowanych w Koninie,
jest współfinansowany przez Unię Europejską
ze środków Funduszu na rzecz Sprawiedliwej Transformacji
w ramach programu Fundusze Europejskie dla Wielkopolski 2021-2027



Fundusze Europejskie
dla Wielkopolski

Dofinansowane przez
Unię Europejską



SAMORZĄD
WOJEWÓDZTWA
WIELKOPOLSKIEGO



Wydawca

Akademia Nauk Stosowanych w Koninie
ul. Przyjaźni 1, 62-510 Konin



Spis treści

1. Matematyka w mechanice i budowie maszyn (<i>Bohdana Hładysz</i>)	5
1.1. Wstęp	5
1.2. Funkcje cyklotryczne i ich własności. Liczba Eulera. Logarytm naturalny	5
1.3. Liczby zespolone i ich własności	12
1.4. Rachunek różniczkowy funkcji jednej zmiennej i jego zastosowanie techniczne	23
1.5. Rachunek całkowy funkcji jednej zmiennej i jego zastosowanie techniczne	38
1.6. Funkcje dwu lub więcej zmiennych	57
1.7. Działanie na macierzach. Obliczanie wyznacznika i macierzy odwrotnej. Operacje elementarne. Rozwiązywanie układów równań liniowych	70
1.8. Równania różniczkowe zwyczajne	88
1.9. Obliczanie całek podwójnych, potrójnych, krzywoliniowych i powierzchniowych	101
1.10. Obliczanie całek oraz rozwiązywanie równań nieliniowych przy pomocy metod numerycznych	104
1.11. Zakończenie	109
1.12. Bibliografia	109
2. Fizyka w budowie maszyn (<i>Paweł Sobczak</i>)	111
2.1. Wstęp	111
2.2. Wybrane zagadnienia fizyczne	111
2.2.1. Wielkości fizyczne, jednostki	111
2.2.2. Praca, moc energia	113
2.2.3. Wybrane zagadnienia z kinematyki	115
2.2.4. Wybrane zagadnienia z dynamiki	116
2.3. Wprowadzenie do pracowni fizycznej	122
2.3.1. Niepewności pomiarowe	122
2.3.2. Przykładowy protokół z pracowni fizycznej	124
2.4. Zadania	132
2.4.1. Przykładowe rozwiązanie zadania	132
2.4.2. Zadania	134
2.5. Zakończenie	135
2.6. Bibliografia	135
3. Wytrzymałość w mechanice i budowie maszyn (<i>Piotr Świta</i>)	136
3.1. Wstęp	136
3.2. Schematy statyczne i modele materiałowe	138
3.2.1. Idealizacja konstrukcji	138
3.2.2. Idealizacja obciążeń	139
3.2.3. Idealizacja materiału	140
3.3. Siły wewnętrzne	144
3.4. Naprężenia, odkształcenia, przemieszczenia	148
3.5. Założenia wytrzymałości materiałów	149
3.6. Badania doświadczalne parametrów wytrzymałościowych materiałów	150
3.7. Warunek wytrzymałości	154
3.7.1. Ściskanie lub rozciąganie	154
3.7.2. Zginanie i ścinanie	154
3.7.3. Skręcanie	157
3.8. Warunek sztywności	165



3.9. Warunek stateczności	169
3.9.1. Wstęp do stateczności	169
3.9.2. Badania doświadczalne słupów ściskanych osiowo	170
3.9.3. Podstawowe zagadnienia teoretyczne stateczności sprężystej pręta	173
3.10. Bibliografia	181
4. Metaloznawstwo i obróbka cieplna (Monika Muszyńska)	183
4.1. Wstęp	183
4.2. Budowa metali	184
4.2.1. Wiązania metaliczne	185
4.2.2. Defekty w strukturze krystalicznej	186
4.2.3. Właściwości mechaniczne i fizyczne metali	186
4.3. Produkcja metali	187
4.3.1. Wydobycie rud metali	187
4.3.2. Procesy wzbogacania rud	187
4.3.3. Redukcja rud	188
4.3.4. Rafinacja i oczyszczanie	189
4.3.5. Formowanie i obróbka	189
4.4. Podział metali	192
4.4.1. Metale żelazne	192
4.4.2. Metale nieżelazne	193
4.4.3. Metale ziem rzadkich	194
4.4.4. Stopy metali	195
4.5. Wykres żelazo – węgiel	195
4.6. Podział stali	198
4.6.1. Podział ze względu na zawartość węgla	198
4.6.2. Podział według składu chemicznego	198
4.6.3. Podział według struktury	199
4.6.4. Stale szybko tnące	199
4.6.5. Stale żarowytrzymałe	199
4.6.6. Podział według zastosowania stali	200
4.7. Żeliwo	200
4.7.1. Podział żeliwa	201
4.7.2. Właściwości i zastosowanie żeliwa	202
4.8. Obróbka cieplna i cieplno-chemiczna metali	203
4.8.1. Procesy obróbki cieplnej	204
4.8.2. Obróbka cieplno-chemiczna	205
4.8.3. Zastosowanie obróbki cieplnej i cieplno-chemicznej	206
4.9. Badania metalograficzne	207
4.9.1. Mikroskop metalograficzny, zasada działania	208
4.10. Badania metalograficzne – przykłady zadań	211
4.11. Bibliografia	217
5. Grafika inżynierska w mechanice i budowie maszyn (Robert Cieślak)	219
5.1. Zagadnienia teoretyczne z przykładami	219
5.2. Formaty rysunkowe i linie rysunkowe	220
5.3. Rzutowanie prostokątne, przekroje i wymiarowanie	221
5.4. Zadania	238
5.5. Bibliografia	243



Bohdana Hładysz

1. MATEMATYKA W MECHNICE I BUDOWIE MASZYN**1.1. WSTĘP**

*Tylko państwa, które pielęgnują matematykę,
mogą być silne i potężne.*

(Stefan Banach)

Skrypt ten przeznaczony dla studentów kierunków technicznych wyższych uczelni, w tym dla kierunków związanych z mechaniką i budową maszyn.

Materiał przedstawiono bez nadmiernego formalizmu matematycznego, pomijając dowody twierdzeń oraz zachowując kolejność tematów. Każdy rozdział składa się z teorii i przykładów. Ostatnie uszeregowano od prostego (najlepiej ilustrują teorię) do złożonego, co jest jedną z najważniejszych zasad w nauczaniu matematyki. Również jest wiele przykładów poświęconych zastosowaniu różnych dziedzin matematyki, które demonstrują o ile trudne bądź łatwe ono może być.

Zawartość poniższych rozdziałów nie jest wyczerpującym źródłem wiedzy zarówno jak teoretycznej, tak i praktycznej. Za pomocą podręczników [1-15] oraz różnych stron internetowych, np. takich jak Wikipedia Wolna encyklopedia, Khan Academy lub chatGPT (w jakości kierunkowskazu dla samokształcenia) można pogłębić wiedzę w zakresie matematyki wyższej i jej zastosowania.

Tekst zasadniczy

**1.2. FUNKCJE CYKLOMETRYCZNE I ICH WŁASNOŚCI. LICZBA EULERA.
LOGARTYM NATURALNY**

Ciąg liczbowy. Granica ciągu liczbowego. Liczba Eulera

Definicja (Nieskończonym) *ciągami liczbowymi* nazywa się odwzorowanie $a: \mathbb{N} \rightarrow \mathbb{R}$ ze zbioru liczb naturalnych $\mathbb{N}$ w zbiór liczb rzeczywistych $\mathbb{R}$, mianowicie

$$\begin{array}{ccccccc} n & 1 & 2 & 3 & \dots & n & \dots \\ a: \downarrow & \downarrow & \downarrow & \downarrow & \dots & \downarrow & \dots \\ a_n & a_1 & a_2 & a_3 & \dots & a_n & \dots \end{array}$$

Ciąg oznaczamy przez $\{a_n\}_{n=1}^{\infty}$ lub $\{a_n\}$. Element a_n nazywamy *n-tym elementem ciągu* liczbowego $\{a_n\}_{n=1}^{\infty}$.

Przykład *Ciągiem Fibonacciego* nazywa się ciąg liczbowy, który jest określony w następujący sposób: $f_1 = f_2 = 1, f_{n+2} = f_n + f_{n+1}, n \in \mathbb{N}$. Słownie ciąg Fibonacciego możemy określić tak: pierwsze dwa elementy ciągu wynoszą 1, wtedy jak każdy następny element jest sumą dwóch poprzednich. Bernoulli wskazał na związek liczb Fibonacciego ze



złotym stosunkiem [1], tak więc wartości $\frac{f_n}{f_{n-1}}$ liczby f_n do jej poprzedniczki f_{n-1} bardzo szybko zbliżają się do liczby $\varphi = \frac{1+\sqrt{5}}{2} \approx 1,61803399$, która opisuje złoty stosunek. Również ciąg Fibonacciego można spotkać w trochę innej postaci, kiedy $f_1 = 0$, $f_2 = 1$, $f_{n+2} = f_n + f_{n+1}$, $n \in \mathbb{N}$, ale po eliminacji pierwszego elementu otrzymamy ten sam ciąg liczbowy, co wcześniej.

Definicja Liczba $g \in \mathbb{R}$ nazywa się *granica* (według Cauchy'ego) *ciągu liczbowego* $\{a_n\}$, jeśli

$$\forall \varepsilon > 0 \quad \exists N \in \mathbb{N} \quad \forall n \geq N: |a_n - g| < \varepsilon$$

Innymi słowy, (1) dla dowolnej małej liczby ε istnieje numer N , zaczynając od którego odległość między elementami ciągu a_n a liczbą g jest mniejsza od ε lub (2) przy zwiększeniu numerów elementów ciągu, zmniejsza się odległość między elementami ciągu a liczbą g .

Mówimy, że ciąg $\{a_n\}$ *zbieżny do liczby* g . Oznaczamy przez $\lim_{n \rightarrow \infty} a_n = g$ lub $a_n \rightarrow g, n \rightarrow \infty$.

Definicja Ciąg liczbowy nazywa się *rozbieżny do $+\infty$ ($-\infty$)*, jeśli

$$\forall \alpha \in \mathbb{R} \quad \exists N \in \mathbb{N} \quad \forall n \geq N: a_n > \alpha \quad (a_n < \alpha)$$

Innymi słowy, dla dowolnej liczby $\alpha \in \mathbb{R}$ istnieje numer N , zaczynając od którego wszystkie elementy ciągu a_n są większe (mniejsze) od wskazanej liczby α . Odpowiednia granica istnieje, ale nie jest liczbą $\lim_{n \rightarrow \infty} a_n = \pm\infty$ (tzw. granica niewłaściwa).

Definicja Ciąg liczbowy, który nie jest zbieżnym oraz nie jest rozbieżnym do $+\infty$ lub $-\infty$, nazywa się *rozbieżnym*. Mówimy, że odpowiednia granica nie istnieje $\nexists \lim_{n \rightarrow \infty} a_n$.

Twierdzenie Zbieżny ciąg liczbowy ma dokładnie jedną granicę oraz jest ograniczony. Rozbieżnym do $+\infty$ lub $-\infty$ ma dokładnie jedną granicę.

Twierdzenie Jeżeli ciągi liczbowe $\{a_n\}$, $\{b_n\}$ są rozbieżne do $+\infty$ (innymi słowy $\lim_{n \rightarrow \infty} a_n = \lim_{n \rightarrow \infty} b_n = +\infty$), to

$$\bullet \quad \forall C \in \mathbb{R}: \lim_{n \rightarrow \infty} (a_n \pm C) = +\infty$$

$$[[\infty \pm C]] = \infty$$

$$\bullet \quad \lim_{n \rightarrow \infty} (a_n + b_n) = +\infty$$

$$[[\infty + \infty]] = \infty$$

$$\bullet \quad \forall C \neq 0: \lim_{n \rightarrow \infty} \frac{C}{a_n} = 0$$

$$[[C/\infty, C \neq 0]] = 0$$

$$\bullet \quad \forall C > 1: \lim_{n \rightarrow \infty} C^{a_n} = +\infty$$

$$[[C^\infty, C > 1]] = \infty$$

$$\bullet \quad \forall C, -1 < C < 1: \lim_{n \rightarrow \infty} C^{a_n} = 0$$

$$[[C^\infty, -1 < C < 1]] = 0$$

$$\bullet \quad \forall C > 0: \lim_{n \rightarrow \infty} (a_n)^C = +\infty$$

$$[[\infty^C, C > 0]] = \infty$$

$$\bullet \quad \forall C < 0: \lim_{n \rightarrow \infty} (a_n)^C = 0$$

$$[[\infty^C, C < 0]] = 0$$



$$\bullet \quad \forall C > 0: \lim_{n \rightarrow \infty} (C \cdot a_n) = +\infty$$

$$[[C \cdot \infty, C > 0]] = \infty$$

$$\bullet \quad \forall C < 0: \lim_{n \rightarrow \infty} (C \cdot a_n) = -\infty$$

$$[[C \cdot \infty, C < 0]] = -\infty$$

$$\bullet \quad \lim_{n \rightarrow \infty} (a_n \cdot b_n) = +\infty$$

$$[[\infty \cdot \infty]] = \infty$$

$$\bullet \quad \lim_{n \rightarrow \infty} (a_n)^{b_n} = +\infty$$

$$[[\infty^\infty]] = \infty$$

Symbolami nieoznaczonymi nazywamy następujące wyraży, wynikające pod symbolem granicy: $[[\infty - \infty]]$, $[[\frac{0}{0}]]$, $[[\frac{\infty}{\infty}]]$, $[[0 \cdot \infty]]$, $[[1^\infty]]$, $[[\infty^0]]$, $[[0^0]]$. Zaznaczone symbole są nieoznaczone, dlatego że w zależności od szczególnego przykładu mogą być różne odpowiedzi, a jak wiemy ciąg zbieżny oraz ciąg rozbieżny do $+\infty$ lub $-\infty$ ma tylko jedną granicę.

Definicja Liczba Euler'a (zwaną również liczbą Neper'a) nazywa się granica zbieżnego ciągu liczbowego $\{(1 + \frac{1}{n})^n\}$. Granica w przybliżeniu wynosi 2,718281828459 i oznaczamy symbolem e . Są również inne sposoby dla określenia liczby Euler'a.

$$\lim_{n \rightarrow \infty} \left(1 + \frac{1}{n}\right)^n = e \approx 2,7$$

Liczba Eulera jest jedną z najważniejszych matematycznych stałych, ponieważ pojawia się w każdej dziedzinie matematyki (rachunek różniczkowy i całkowy, równania różniczkowe, liczby zespolone, teoria prawdopodobieństwo, statystyka itd.). Również liczba e jest często spotykana w ekonomii i finansach, bo jest wykorzystywana w zadaniach związanych z jakimkolwiek wzrostem.

Funkcja logarytmiczna (w tym logarytm naturalny)

$$y = \log_a x, a > 0, a \neq 1$$

Definicja Logarytmem $\log_a b$ o podstawie a od liczby b nazywa się wykładnik potęgi, do której należy podnieść podstawę a aby otrzymać liczbę logarytmową b .

Własności funkcji logarytmicznej $y = \log_a x, a > 0, a \neq 1$:

- dziedzina: $D_y = (0; +\infty)$
- zbiór wartości: $ZW_y = \mathbb{R}$
- miejsca zerowe: $x = 1$
- parzystość, nieparzystość: brak
- monotoniczność:
 - $0 < a < 1$: $y \searrow$
 - $a > 1$: $y \nearrow$
- wklęsłość, wypukłość:
 - $0 < a < 1$: y wypukła
 - $a > 1$: y wklęsła



- funkcja logarytmiczna jest odwrotną do funkcji wykładniczej
- własności logarytmów: niech podstawy logarytmów $a, b > 0$ i $\neq 1$
 - $\log_a u + \log_a v = \log_a(uv), u > 0, v > 0$
 - $\log_a u - \log_a v = \log_a\left(\frac{u}{v}\right), u > 0, v > 0$
 - $\log_a u^n = n \log_a u, u > 0, n \in \mathbb{R}$
 - $\log_a^m u = \frac{1}{m} \log_a u, m \in \mathbb{R}$
 - $\log_a u = \frac{\log_b u}{\log_b a}$
 - $\log_a b = \frac{1}{\log_b a}$
 - $a^{\log_a u} = u$

Definicja Logarytmem dziesiętnym od liczby b nazywa się logarytm o podstawie 10 od tej liczby, oznaczamy przez $\lg b$ lub $\log b$.

Definicja Logarytmem naturalnym od liczby b nazywa się logarytm o podstawie e od tej liczby, oznaczamy przez $\ln b$ ($e \approx 2,7$ tj. liczba Eulera).

Przykład Obliczyć

$$\log_a 1, a > 0, a \neq 1; \log_5 125; \log_{125} 5; \log_2 64; \log_8 64; \log 10; \log 100;$$

$$\ln e^5; \ln \sqrt{e^7}; 10^{\log b}, b > 0; 2 \ln 3e - \ln 9; e^{\ln 2}.$$

Rozwiązanie

- $\log_a 1 = \log_a a^0 = 0$
- $\log_5 125 = \log_5 5^3 = 3$
- $\log_{125} 5 = \log_{5^3} 5 = \frac{1}{3} \log_5 5 = \frac{1}{3}$
- $\log_2 64 = \log_2 2^6 = 6$
- $\log_8 64 = \log_{2^3} 2^6 = \frac{6}{3} \log_2 2 = \frac{6}{3} = 2$
- $\log 10 = \log_{10} 10^1 = 1$
- $\log 100 = \log_{10} 10^2 = 2$
- $\ln e^5 = 5$
- $\ln \sqrt{e^7} = \ln e^{1/7} = \frac{1}{7}$
- $10^{\log b} = 10^{\log_{10} b} = b$
- $2 \ln 3e - \ln 9 = \ln(3e)^2 - \ln 9 = \ln \frac{9e^2}{9} = \ln e = 1$
- $e^{\ln 2} = e^{\log_e 2} = 2$

Przykład Znaleźć y z wymienionych niżej równości



- $5y + 2 - 8 \ln x = 4$
- $2^x - 3^y + 7 = 0$
- $e^y = 2x^2 + 5 - \log x$

Rozwiązanie

- $5y + 2 - 8 \ln x = 4$

$$5y = 8 \ln x + 2$$

$$y = \frac{8}{5} \ln x + \frac{2}{5}$$

- $2^x - 3^y + 7 = 0$

$$3^y = 2^x + 7$$

$$\log_3 3^y = \log_3 (2^x + 7)$$

$$y = \log_3 (2^x + 7)$$

- $e^y = 2x^2 + 5 - \log x$

$$\ln e^y = \ln(2x^2 + 5 - \log x)$$

$$y = \ln(2x^2 + 5 - \log x)$$

Funkcje elementarne

(w tym funkcje cyklotometryczne i ich własności)

Definicja Niech $f: X \rightarrow Y$, $g: Y \rightarrow Z$, wtedy funkcja $h(x) = g(f(x)): X \rightarrow Z$ nazywa się *złożeniem lub superpozycją funkcji f i g* . Funkcja f nazywa się *funkcją wewnętrzną*, a g – *zewnętrzną*. Funkcja h nazywa się *funkcją złożoną*.

Definicja Funkcja $f: X \rightarrow Y$ przyjmująca jako swoje wartości wszystkie elementy zbioru Y , nazywa się *surjekcją* lub *funkcją „na”*, tzn.

$$\forall y \in Y \quad \exists x \in X: y = f(x)$$

Funkcja $f: X \rightarrow Y$, która różnym argumentom przyporządkuje różne wartości, nazywa się *injekcją* lub *funkcją różnowartościową*, tzn.

$$\forall x_1, x_2 \in X, x_1 \neq x_2: f(x_1) \neq f(x_2)$$

Funkcja nazywa się *bijekcją* lub *wzajemnie jednoznaczna*, gdy ona jest surjekcją oraz injekcją.

Definicja Funkcja $f^{-1}: Y \rightarrow X$ nazywa się *funkcją odwrotną* do funkcji f , jeżeli są spełnione poniższe warunki

$$\forall x \in X: f^{-1}(f(x)) = x$$

$$\forall y \in Y: f(f^{-1}(y)) = y$$

W przypadku, gdy taka funkcja f^{-1} istnieje, funkcja f nazywa się *odwracalna*.

Twierdzenie Funkcja jest odwracalna wtedy i tylko wtedy, gdy ona jest bijekcją.



Definicja *Podstawowymi funkcjami elementarnymi* nazywają się takie funkcje: potęgowa (w tym stała), wykładnicza, logarytmiczna, trygonometryczne i cyklometryczne, tzn.

- funkcja potęgowa: $y = x^a, a \in \mathbb{R}$;
- funkcja wykładnicza: $y = a^x, a > 0, a \neq 1$;
- funkcja logarytmiczna: $y = \log_a x, a > 0, a \neq 1$;
- funkcje trygonometryczne: $y = \sin x, y = \cos x, y = \operatorname{tg} x, y = \operatorname{ctg} x$;
- funkcje cyklometryczne: $y = \arcsin x, y = \arccos x, y = \operatorname{arctg} x, y = \operatorname{arcctg} x$.

Definicja *Funkcjami elementarnymi* nazywają się funkcje, które można otrzymać z podstawowych funkcji elementarnych za pomocą skończonej liczby takich działań jak suma, różnica, iloczyn, iloraz oraz złożenie.

Funkcje cyklometryczne: $y = \arcsin x, y = \arccos x, y = \operatorname{arctg} x, y = \operatorname{arcctg} x$

Definicja *Funkcje cyklometryczne* to są funkcje odwrotne do funkcji trygonometrycznych ograniczonych do pewnych przedziałów

- $y = \arcsin x$: funkcja $\sin x, x \in \langle -\frac{\pi}{2}; \frac{\pi}{2} \rangle$ jest wzajemnie jednoznaczna, z czego wynika, że istnieje funkcja odwrotna w zaznaczonym przedziale, którą oznaczamy przez $y = \arcsin x$. Własności:
 - dziedzina $D_{\arcsin x} = \langle -1; 1 \rangle$
 - zbiór wartości $ZW_{\arcsin x} = \langle -\frac{\pi}{2}; \frac{\pi}{2} \rangle$
 - parzystość, nieparzystość: funkcja $y = \arcsin x$ nieparzysta
 - monotoniczność: rosnąca
 - wklęsłość, wypukłość:
 - $x \in \langle -1; 0 \rangle$: funkcja $y = \arcsin x$ wklęsła
 - $(0; 0)$ tj. punkt przegięcia
 - $x \in (0; 1)$: funkcja $y = \arcsin x$ wypukła
- $y = \arccos x$: funkcja $\cos x, x \in \langle 0; \pi \rangle$ jest wzajemnie jednoznaczna, z czego wynika, że istnieje funkcja odwrotna w zaznaczonym przedziale, którą oznaczamy przez $y = \arccos x$. Własności:
 - dziedzina $D_{\arccos x} = \langle -1; 1 \rangle$
 - zbiór wartości $ZW_{\arccos x} = \langle 0; \pi \rangle$
 - parzystość, nieparzystość: brak
 - monotoniczność: malejąca
 - wklęsłość, wypukłość:



- $x \in \langle -1; 0 \rangle$: funkcja $y = \arccos x$ wypukła
- $\left(0; \frac{\pi}{2}\right)$ tj. punkt przegięcia
- $x \in (0; 1)$: funkcja $y = \arccos x$ wklęsła
- $y = \operatorname{arctg} x$: funkcja $\operatorname{tg} x, x \in \left(-\frac{\pi}{2}; \frac{\pi}{2}\right)$ jest wzajemnie jednoznaczna, z czego wynika, że istnieje funkcja odwrotna w zaznaczonym przedziale, którą oznaczamy przez $y = \operatorname{arctg} x$. Własności:
 - dziedzina $D_{\operatorname{arctg} x} = \mathbb{R}$
 - zbiór wartości $ZW_{\operatorname{arctg} x} = \left(-\frac{\pi}{2}; \frac{\pi}{2}\right)$
 - parzystość, nieparzystość: funkcja $y = \operatorname{arctg} x$ nieparzysta
 - monotoniczność: rosnąca
 - wklęsłość, wypukłość:
 - $x \in (-\infty; 0)$: funkcja $y = \operatorname{arctg} x$ wypukła
 - $(0; 0)$ tj. punkt przegięcia
 - $x \in (0; +\infty)$: funkcja $y = \operatorname{arctg} x$ wklęsła
- $y = \operatorname{arcctg} x$: funkcja $\operatorname{ctg} x, x \in (0; \pi)$ jest wzajemnie jednoznaczna, z czego wynika, że istnieje funkcja odwrotna w zaznaczonym przedziale, którą oznaczamy przez $y = \operatorname{arcctg} x$. Własności:
 - dziedzina $D_{\operatorname{arcctg} x} = \mathbb{R}$
 - zbiór wartości $ZW_{\operatorname{arcctg} x} = (0; \pi)$
 - parzystość, nieparzystość: brak
 - monotoniczność: malejąca
 - wklęsłość, wypukłość:
 - $x \in (-\infty; 0)$: funkcja $y = \operatorname{arcctg} x$ wklęsła
 - $\left(0; \frac{\pi}{2}\right)$ tj. punkt przegięcia
 - $x \in (0; +\infty)$: funkcja $y = \operatorname{arcctg} x$ wypukła

Inne własności funkcji cyklometrycznych:

$$\begin{aligned}\arcsin(-x) &= -\arcsin x \\ \arccos(-x) &= \pi - \arccos x \\ \operatorname{arctg}(-x) &= -\operatorname{arctg} x \\ \operatorname{arcctg}(-x) &= \pi - \operatorname{arcctg} x\end{aligned}$$

$$\begin{aligned}\arcsin(\sin x) &= x, x \in \left\langle -\frac{\pi}{2}; \frac{\pi}{2} \right\rangle \\ \sin(\arcsin x) &= x, x \in \langle -1; 1 \rangle \\ \arccos(\cos x) &= x, x \in \langle 0; \pi \rangle \\ \cos(\arccos x) &= x, x \in \langle -1; 1 \rangle\end{aligned}$$



$$\begin{aligned}\operatorname{arctg}(\operatorname{tg} x) &= x, x \in \left(-\frac{\pi}{2}; \frac{\pi}{2}\right) \\ \operatorname{tg}(\operatorname{arctg} x) &= x, x \in \mathbb{R}\end{aligned}$$

$$\begin{aligned}\operatorname{arcctg}(\operatorname{ctg} x) &= x, x \in (0; \pi) \\ \operatorname{ctg}(\operatorname{arcctg} x) &= x, x \in \mathbb{R}\end{aligned}$$

1.3. LICZBY ZESPOLONE I ICH WŁASNOŚCI

Definicje podstawowe

Zbiory liczb (poniższe symbole wykorzystują się tylko i wyłącznie dla wymienionych dalej zbiorów liczb):

- naturalnych $\mathbb{N} = \{1; 2; 3; 4; \dots\}$
- naturalnych z zerem $\mathbb{N}_0 = \{0; 1; 2; 3; 4; \dots\}$
- całkowitych $\mathbb{Z} = \{\dots; -4; -3; -2; -1; 0; 1; 2; 3; 4; \dots\}$
- wymiernych $\mathbb{Q} = \left\{\frac{m}{n} \mid m \in \mathbb{Z}, n \in \mathbb{N}\right\}$
- rzeczywistych $\mathbb{R} = (-\infty; +\infty)$
- niewymiernych $\mathbb{R} \setminus \mathbb{Q} = \{x \in \mathbb{R} \mid x \notin \mathbb{Q}\}$
- zespolonych $\mathbb{C} = \{a + bi \mid a, b \in \mathbb{R}\}$

W zbiorze liczb rzeczywistych nie wszystkiego da się znaleźć. Na przykład, nie możemy obliczyć pierwiastków parzystego stopnia lub logarytmów od liczb ujemnych, często wielomian n -go stopnia ma mniej niż n miejsc zerowych.

Definicja Jednostką urojoną nazywa się taka liczba i , że $i^2 = -1$. Liczba $a + bi$, gdzie $a, b \in \mathbb{R}$ nazywa się *liczbą zespoloną*. Zbiór liczb zespolonych oznaczamy przez $\mathbb{C}$. Dla liczby zespolonej $z = a + bi \in \mathbb{C}$ liczbę $\operatorname{Re} z = a$ nazywamy *częścią rzeczywistą*, a liczbę $\operatorname{Im} z = b$ *częścią urojoną*. Liczba zespolona $z = a + bi$, która ma zerową część rzeczywistą (tzn. $a = 0$), nazywa się *czysto urojona*. Natomiast liczba zespolona, dla której część urojona jest zerowa (tzn. $b = 0$), będzie również liczbą rzeczywistą. Dwie liczby zespolone nazywamy *równymi*, gdy są równe ich części rzeczywiste i urojone.

Tak więc, poprawnym jest zapis relacji zbiorów liczb $\mathbb{N} \subset \mathbb{N}_0 \subset \mathbb{Z} \subset \mathbb{Q} \subset \mathbb{R} \subset \mathbb{C}$.

Równoważną jest definicja liczb zespolonych jako zbiór par uporządkowanych $(a; b)$, gdzie $a, b \in \mathbb{R}$, tzn. $\mathbb{C} = \{(a; b) \mid a, b \in \mathbb{R}\}$.

Potęgowanie jednostki urojonej i :

$$\begin{cases} i^{4k} = 1, k \in \mathbb{Z} \\ i^{4k+1} = i, k \in \mathbb{Z} \\ i^{4k+2} = -1, k \in \mathbb{Z} \\ i^{4k+3} = -i, k \in \mathbb{Z} \end{cases}$$

Dlatego $i^0 = 1$, $i^1 = i$, $i^2 = -1$, $i^3 = -i$, $i^4 = 1$, $i^5 = i$, $i^6 = -1$, $i^7 = -i$, $i^8 = 1$, itd.



Przykład Obliczyć $i^{10}, i^{12}, i^{25}, i^{30}, i^{41}, i^{65}, i^{83}$.

Rozwiązanie

$$i^{10} = (i^2)^5 = (-1)^5 = -1$$

$$i^{12} = (i^2)^6 = (-1)^6 = 1$$

$$i^{25} = (i^2)^{10} \cdot i = (-1)^{10} \cdot i = i$$

$$i^{30} = (i^{10})^3 = (-1)^3 = -1$$

$$i^{41} = (i^{10})^4 \cdot i = (-1)^4 \cdot i = i$$

$$i^{65} = (i^2)^{32} \cdot i = (-1)^{32} \cdot i = i$$

$$i^{83} = (i^2)^{40} \cdot i^3 = 1 \cdot (-i) = -i$$

Interpretacja geometryczna liczb zespolonych [2]

Interpretując liczbę zespoloną jako parę uporządkowaną, otrzymamy punkt na tzw. płaszczyźnie zespolonej, w której oś pozioma odpowiada za część rzeczywistą (oznaczamy przez $\text{Re } z$ i nazywamy dalej *osią rzeczywistą*), a oś pionowa – za część urojoną (oznaczamy przez $\text{Im } z$ i nazywamy dalej *osią urojoną*). Każda niezerowa liczba zespolona $z \in \mathbb{C}, z \neq 0$ określa i również może być określona za pomocą takich wartości:

- *argument główny* $\text{Arg } z = \varphi$ – miara kąta skierowanego $\varphi \in \langle 0; 2\pi \rangle$ pomiędzy dodatnim kierunkiem osi rzeczywistej a odcinkiem, łączącym punkty $O(0; 0)$ i naszą liczbę zespoloną $z = (a; b)$;
- *moduł* $|z|$ – odległość od punktu $O(0; 0)$ do naszej liczby zespolonej $z = (a; b)$.

Argument główny nie może być określony dla zerowej liczby zespolonej.

Generalnie argument liczby zespolonej może być wyznaczony z dokładnością do 2π . Inaczej mówiąc, *argumentami* liczby zespolonej z , które oznaczamy przez $\arg z$, nazywają się kąty $\arg z = \text{Arg } z + 2\pi k, k \in \mathbb{Z}$. Tak więc, każda niezerowa liczba zespolona ma nieskończenie wiele argumentów $\arg z$ i tylko jeden argument główny $\text{Arg } z$.

Pojęcia modułu określone zarówno jak dla liczby niezerowej, tak i dla zerowej liczby zespolonej (i wynosi 0).

Z definicji modułu (zgodnie z twierdzeniem Pitagorasa) wynika poniższy wzór:

$$z = (a; b) \Rightarrow |z| = \sqrt{a^2 + b^2}$$

Działanie na liczbach zespolonych: niech $z_1 = a_1 + b_1 i, z_2 = a_2 + b_2 i, k \in \mathbb{R}$

- dodawanie/odejmowanie liczb zespolonych: sumą liczb zespolonych z_1 i z_2 jest liczba zespolona

$$z_1 \pm z_2 = (a_1 \pm a_2) + (b_1 \pm b_2)i$$



- mnożenie liczb zespolonych przez stałą rzeczywistą: iloczynem liczby zespolonej z_1 a stałej $k \in \mathbb{R}$ jest liczba zespolona

$$k \cdot z_1 = ka_1 + kb_1i$$

- mnożenie liczb zespolonych: iloczynem dwóch liczb zespolonych z_1 i z_2 jest liczba zespolona

$$z_1 \cdot z_2 = (a_1a_2 - b_1b_2) + (a_1b_2 + a_2b_1)i$$

$$\text{dlatego że } z_1 \cdot z_2 = (a_1 + b_1i)(a_2 + b_2i) = a_1a_2 + a_1b_2i + a_2b_1i + b_1b_2i^2 = \\ \llbracket i^2 = -1 \rrbracket = a_1a_2 + a_1b_2i + a_2b_1i - b_1b_2 = (a_1a_2 - b_1b_2) + (a_1b_2 + a_2b_1)i$$

- sprzężenie liczb zespolonych: *sprzężeniem* liczby zespolonej $z = a + bi$ nazywa się liczba zespolona

$$\bar{z} = a - bi$$

- dzielenie liczb zespolonych przez zespolonych: ilorazem dwóch liczb zespolonych jest liczba zespolona

$$\frac{z_1}{z_2} = \frac{a_1a_2 + b_1b_2}{a_2^2 + b_2^2} + \frac{-a_1b_2 + a_2b_1}{a_2^2 + b_2^2}i, \quad z_2 \neq 0$$

$$\text{dlatego że } \frac{z_1}{z_2} = \frac{z_1 \cdot \bar{z}_2}{z_2 \cdot \bar{z}_2} = \frac{a_1 + b_1i}{a_2 + b_2i} = \frac{(a_1 + b_1i)(a_2 - b_2i)}{(a_2 + b_2i)(a_2 - b_2i)} = \frac{(a_1a_2 + b_1b_2) + (-a_1b_2 + a_2b_1)i}{a_2^2 + b_2^2} = \\ = \frac{a_1a_2 + b_1b_2}{a_2^2 + b_2^2} + \frac{-a_1b_2 + a_2b_1}{a_2^2 + b_2^2}i$$

Przykład Dla podanych liczb zespolonych $z_1 = (1; -2)$ i $z_2 = (-3; \sqrt{5})$ obliczyć

$$3z_1 + z_2, \bar{z}_1, z_1 \cdot \bar{z}_1, |z_1|, z_1 \cdot z_2, \frac{z_1}{z_2}.$$

Rozwiązanie

$$3z_1 + z_2 = 3(1; -2) + (-3; \sqrt{5}) = (3; -6) + (-3; \sqrt{5}) = (0; \sqrt{5} - 6)$$

lub

$$z_1 = (1; -2) = 1 - 2i$$

$$z_2 = (-3; \sqrt{5}) = -3 + \sqrt{5}i$$

$$3z_1 + z_2 = 3(1 - 2i) - 3 + \sqrt{5}i = (\sqrt{5} - 6)i = (0; \sqrt{5} - 6)$$

$$\bar{z}_1 = 1 + 2i = (1; 2)$$

$$z_1 \cdot \bar{z}_1 = (1 - 2i)(1 + 2i) = 1 + 2i - 2i - 4i^2 = 1 - 4(-1) = 5 = (5; 0)$$

$$|z_1| = \sqrt{1^2 + (-2)^2} = \sqrt{5}$$

$$z_1 \cdot z_2 = (1 - 2i)(-3 + \sqrt{5}i) = -3 + \sqrt{5}i + 6i - 2\sqrt{5}i^2 = -3 + (\sqrt{5} + 6)i - 2\sqrt{5}(-1) \\ = 2\sqrt{5} - 3 + (\sqrt{5} + 6)i = (2\sqrt{5} - 3; \sqrt{5} + 6)$$



$$\begin{aligned}\frac{z_1}{z_2} &= \frac{1-2i}{-3+\sqrt{5}i} = \frac{(1-2i)(-3-\sqrt{5}i)}{(-3+\sqrt{5}i)(-3-\sqrt{5}i)} = \frac{-3-\sqrt{5}i+6i+2\sqrt{5}i^2}{9+5} \\ &= \frac{-3-2\sqrt{5}+(6-\sqrt{5})i}{14} = \frac{-3-2\sqrt{5}}{14} + \frac{6-\sqrt{5}}{14}i = \left(\frac{-3-2\sqrt{5}}{14}; \frac{6-\sqrt{5}}{14}\right)\end{aligned}$$

Własności modułu liczby zespolonej: niech $z = (a; b)$, $z_1, z_2 \in \mathbb{C}$, $\varphi = \arg z$

- $|z_1 \cdot z_2| = |z_1| \cdot |z_2|$
- $|z_1 + z_2| \leq |z_1| + |z_2|$, w tym $|z_1 + z_2| = |z_1| + |z_2| \Leftrightarrow z_1 = z_2$
- $\cos \varphi = \frac{a}{|z|} = \frac{a}{\sqrt{a^2+b^2}}$
- $\sin \varphi = \frac{b}{|z|} = \frac{b}{\sqrt{a^2+b^2}}$

Własności sprzężenia liczby zespolonej: niech $z = a + bi \in \mathbb{C}$

- $z = (a; b) \Rightarrow \bar{z} = (a; -b)$
- $\overline{\bar{z}} = z$
- $|\bar{z}| = |z| = \sqrt{a^2+b^2}$

Postać trygonometryczna liczby zespolonej: niech $z = a + bi \in \mathbb{C}$, $\varphi = \text{Arg } z$, wtedy z własności modułu wynika, że

$$\begin{aligned}a &= |z| \cos \varphi \\ b &= |z| \sin \varphi\end{aligned} \Rightarrow z = |z| \cos \varphi + i|z| \sin \varphi = |z|(\cos \varphi + i \sin \varphi)$$

ostatnia postać nazywa się *postacią trygonometryczną* liczby zespolonej z .

Przykład Przedstawić liczby zespolone w postaci trygonometrycznej

- $z_1 = 5\sqrt{2} - 5\sqrt{2}i$
- $z_2 = 1 + \sqrt{3}i$
- $z_3 = 7i$

Rozwiązanie

- $z_1 = 5\sqrt{2} - 5\sqrt{2}i$

$$|z_1| = \sqrt{(5\sqrt{2})^2 + (-5\sqrt{2})^2} = \sqrt{50+50} = 10$$

$$z_1 = a + bi \Rightarrow a = 5\sqrt{2}, b = -5\sqrt{2}$$

$$\begin{aligned}\cos \varphi &= \frac{a}{|z_1|} = \frac{5\sqrt{2}}{10} = \frac{\sqrt{2}}{2} \\ \sin \varphi &= \frac{b}{|z_1|} = \frac{-5\sqrt{2}}{10} = -\frac{\sqrt{2}}{2}\end{aligned} \Rightarrow \varphi = \text{Arg } z_1 = \frac{7\pi}{4} \in (0; 2\pi)$$

$$z_1 = 10 \left(\cos \frac{7\pi}{4} + i \sin \frac{7\pi}{4} \right)$$

- $z_2 = 1 + \sqrt{3}i$

$$|z_2| = \sqrt{1^2 + (\sqrt{3})^2} = \sqrt{4} = 2$$



$$z_2 = a + bi \Rightarrow a = 1, b = \sqrt{3}$$

$$\cos \varphi = \frac{a}{|z_2|} = \frac{1}{2} \Rightarrow \varphi = \text{Arg } z_2 = \frac{\pi}{3} \in \langle 0; 2\pi \rangle$$

$$\sin \varphi = \frac{b}{|z_2|} = \frac{\sqrt{3}}{2}$$

$$z_2 = 2 \left(\cos \frac{\pi}{3} + i \sin \frac{\pi}{3} \right)$$

$$\bullet \quad z_3 = 7i$$

$$|z_3| = \sqrt{0^2 + 7^2} = \sqrt{49} = 7$$

$$z_3 = a + bi \Rightarrow a = 0, b = 7$$

$$\cos \varphi = \frac{a}{|z_3|} = \frac{0}{7} = 0 \Rightarrow \varphi = \text{Arg } z_2 = \frac{\pi}{2} \in \langle 0; 2\pi \rangle$$

$$\sin \varphi = \frac{b}{|z_3|} = \frac{7}{7} = 1$$

$$z_3 = 7 \left(\cos \frac{\pi}{2} + i \sin \frac{\pi}{2} \right)$$

Postać wykładnicza liczby zespolonej: niech $z = a + bi \in \mathbb{C}, \varphi = \text{Arg } z$, wtedy ze

wzorów Eulera $\sin \varphi = \frac{e^{i\varphi} - e^{-i\varphi}}{2i}$ wynika tzw. *postać wykładnicza* liczby zespolonej, która

$$\cos \varphi = \frac{e^{i\varphi} + e^{-i\varphi}}{2}$$

wygląda następująco

$$z = |z|e^{i\varphi}$$

Przykład Przedstawić liczbę zespoloną w postaci wykładniczej

$$\bullet \quad z_1 = 5\sqrt{2} - 5\sqrt{2}i \quad \bullet \quad z_2 = 1 + \sqrt{3}i \quad \bullet \quad z_3 = 7i$$

Rozwiązanie

$$\bullet \quad z_1 = 5\sqrt{2} - 5\sqrt{2}i$$

$$\varphi = \text{Arg } z_1 = \frac{7\pi}{4}$$

$$|z_1| = 10$$

$$z_1 = |z_1|e^{i\varphi} = 10 \exp\left(\frac{7\pi i}{4}\right)$$

$$\bullet \quad z_2 = 1 + \sqrt{3}i$$

$$\varphi = \text{Arg } z_2 = \frac{\pi}{3}$$

$$|z_2| = 2$$

$$z_2 = |z_2|e^{i\varphi} = 2 \exp\left(\frac{\pi i}{3}\right)$$



- $z_3 = 7i$

$$\varphi = \operatorname{Arg} z_3 = \frac{\pi}{2}$$

$$|z_3| = 7$$

$$z_3 = |z_3|e^{i\varphi} = 7 \exp\left(\frac{\pi i}{2}\right)$$

Działanie na liczbach zespolonych: dla potęgowania i pierwiastkowania liczb zespolonych obowiązują wzory de Moivre'a, które po połączeniu ze wzorami Eulera wyglądają następująco:

- potęgowanie: niech $n \in \mathbb{Z}$, wtedy

$$z = |z|(\cos \varphi + i \sin \varphi) \Rightarrow z^n = |z|^n(\cos n\varphi + i \sin n\varphi)$$

$$z = |z|e^{i\varphi} \Rightarrow z^n = |z|^n e^{in\varphi}$$

- pierwiastkowanie: niech $n \in \mathbb{N}$, wtedy pierwiastków n -go stopnia z liczby zespolonej z istnieje dokładnie n sztuk $\{z_k | k = \overline{0, n-1}\}$, mianowicie

$$z = |z|(\cos \varphi + i \sin \varphi) \Rightarrow z_k = \sqrt[n]{|z|} \left(\cos \frac{\varphi + 2\pi k}{n} + i \sin \frac{\varphi + 2\pi k}{n} \right),$$

$$k = \overline{0, n-1}$$

$$z = |z|e^{i\varphi} \Rightarrow z_k = \sqrt[n]{|z|} \exp\left(i \frac{\varphi + 2\pi k}{n}\right), k = \overline{0, n-1}$$

Przykład Obliczyć zaznaczone potęgi liczb zespolonych

- $z_1 = (1; 2), z_1^2, z_1^3$
- $z_2 = 3e^{\pi i}, z_2^5, z_2^{2n}, z_2^{2n+1}, n \in \mathbb{N}$
- $z_3 = \cos \frac{\pi}{3} + i \sin \frac{\pi}{3}, z_3^2, z_3^3$

Rozwiązanie

- $z_1 = (1; 2), z_1^2, z_1^3$

$$z_1 = (1; 2) = 1 + 2i$$

$$z_1^2 = (1 + 2i)^2 = 1 + 4i + 4i^2 = -3 + 4i = (-3; 4)$$

$$z_1^3 = z_1^2 z_1 = (-3 + 4i)(1 + 2i) = -3 - 6i + 4i + 8i^2 = -11 - 2i$$

$$= (-11; -2)$$

- $z_2 = 3e^{\pi i}, z_2^5, z_2^{2n}, z_2^{2n+1}, n \in \mathbb{N}$

$$z_2^5 = 3^5 e^{5\pi i} = 243(\cos 5\pi + i \sin 5\pi) = 243(-1 + i \cdot 0) = -243$$

$$z_2^{2n} = 3^{2n} e^{2\pi n i} = 9^n(\cos 2\pi n + i \sin 2\pi n) = 9^n(1 + i \cdot 0) = 9^n$$



$$z_2^{2n+1} = 3^{2n+1} e^{(2n+1)\pi i} = 3^{2n+1} (\cos(2n+1)\pi + i \sin(2n+1)\pi) \\ = 3^{2n+1} (-1 + i \cdot 0) = -3^{2n+1} = -3 \cdot 9^n$$

$$\bullet \quad z_3 = \cos \frac{\pi}{3} + i \sin \frac{\pi}{3}, z_3^2, z_3^3$$

$$z_3^2 = \cos \frac{2\pi}{3} + i \sin \frac{2\pi}{3} = \cos \frac{\pi}{3} + i \sin \frac{\pi}{3} = -\frac{1}{2} + i \frac{\sqrt{3}}{2} = \left(-\frac{1}{2}; \frac{\sqrt{3}}{2}\right)$$

$$z_3^3 = \cos \frac{3\pi}{3} + i \sin \frac{3\pi}{3} = \cos \pi + i \sin \pi = -1 = (-1; 0)$$

Przykład Znaleźć wszystkie pierwiastki zaznaczonych stopni liczb zespolonych

$$\bullet \quad z_1 = 8e^{\pi i}, \sqrt[3]{z_1}$$

$$\bullet \quad z_2 = 1, \sqrt[4]{z_2}$$

Rozwiązanie

$$\bullet \quad z_1 = 8e^{\pi i}, \sqrt[3]{z_1}$$

$$z_1 = 8e^{\pi i} = 8(\cos \pi + i \sin \pi) = 8(-1 + 0) = -8$$

$$z_1 = 8e^{\pi i} \Rightarrow |z_1| = 8, \varphi = \pi$$

$$\sqrt[3]{z_1} = \sqrt[3]{8} \exp\left(i \frac{\pi + 2\pi k}{3}\right), k = \overline{0, 2}$$

$$k = 0 \Rightarrow \sqrt[3]{z_1} = \sqrt[3]{8} \exp\left(i \frac{\pi}{3}\right) = 2 \exp\left(\frac{\pi i}{3}\right) = 2\left(\frac{1}{2} + \frac{\sqrt{3}}{2}i\right) = 1 + i\sqrt{3}$$

$$k = 1 \Rightarrow \sqrt[3]{z_1} = \sqrt[3]{8} \exp\left(i \frac{\pi + 2\pi}{3}\right) = 2 \exp(\pi i) = 2(-1 + 0 \cdot i) = -2$$

$$k = 2 \Rightarrow \sqrt[3]{z_1} = \sqrt[3]{8} \exp\left(i \frac{\pi + 4\pi}{3}\right) = 2 \exp\left(\frac{5\pi i}{3}\right) = 2\left(\frac{1}{2} - \frac{\sqrt{3}}{2}i\right) = 1 - i\sqrt{3}$$

$$\bullet \quad z_2 = 1, \sqrt[4]{z_2}$$

$$z_2 = 1 = 1(\cos \pi + i \sin \pi) \Rightarrow |z_2| = 1, \varphi = 0$$

$$\sqrt[4]{z_2} = \sqrt[4]{1} \left(\cos \frac{0 + 2\pi k}{4} + i \sin \frac{0 + 2\pi k}{4} \right), k = \overline{0, 3}$$

$$k = 0 \Rightarrow \sqrt[4]{z_2} = \cos 0 + i \sin 0 = 1$$

$$k = 1 \Rightarrow \sqrt[4]{z_2} = \cos \frac{\pi}{2} + i \sin \frac{\pi}{2} = i$$

$$k = 2 \Rightarrow \sqrt[4]{z_2} = \cos \pi + i \sin \pi = -1$$

$$k = 3 \Rightarrow \sqrt[4]{z_2} = \cos \frac{3\pi}{2} + i \sin \frac{3\pi}{2} = -i$$

Wielomiany zespolone i ich miejsca zerowe

Definicja Wielomianem o zespolonych współczynnikach n -go stopnia ($n \in \mathbb{N}$) zmiennej z nazywa się wielomian postaci $w_n(z) = a_n z^n + a_{n-1} z^{n-1} + \dots + a_2 z^2 + a_1 z + a_0$, gdzie $a_i \in \mathbb{C}, i = \overline{0, n}$ oraz $a_n \neq 0$. Liczby $a_i \in \mathbb{C}, i = \overline{1, n}$ zwane współczynnikami, $a_0 \in \mathbb{C}$ –



wyrazem wolnym. Liczba zespolona z_0 nazywa się *miejszem zerowym wielomianu zespolonego* $w_n(z)$ jeżeli $w_n(z_0) = 0$.

Jeśli współczynniki wielomianu zespolonego $a \in \mathbb{R}$ są liczbami rzeczywistymi, to otrzymamy wielomian o współczynnikach rzeczywistych, który (gdy nie będzie zaznaczono innego) nadal będzie nazywał się wielomianem zespolonym, co oznacza, że jest on wielomianem zmiennej zespolonej. Innymi słowy,

- wielomian zespolony (wielomian zmiennej zespolonej) – tj. o współczynnikach zespolonych lub rzeczywistych;
- wielomian rzeczywisty (wielomian zmiennej rzeczywistej) – tj. o współczynnikach rzeczywistych.

Twierdzenie Jeśli liczba zespolona $z = a + ib$ jest miejscem zerowym wielomianu o współczynnikach rzeczywistych, to jej sprzężenie $\bar{z} = a - ib$ również będzie miejscem zerowym tego wielomianu.

Innymi słowy, miejsca zerowe wielomianu o współczynnikach rzeczywistych, części urojone których nie są zerowe, tworzą tzw. pary $z_{1,2} = a \pm ib$.

Twierdzenie Wielomian zespolony n -go stopnia

$$w_n(x) = a_n z^n + a_{n-1} z^{n-1} + \dots + a_2 z^2 + a_1 z + a_0$$

ma dokładnie n miejsc zerowych $z_1, z_2, z_3, \dots, z_n$ spośród liczb zespolonych oraz może być przepisany w postaci

$$w_n(z) = a_n(z - z_1)(z - z_2)(z - z_3) \dots (z - z_n)$$

W przypadku wielomianów zmiennej rzeczywistej w odpowiednim przedstawieniu mogą być nie tylko czynniki liniowe, a również kwadratowe.

Rozwiązywać równania $w_n(z) = 0$ w zbiorze liczb zespolonych można na dwa sposoby: (1) pamiętając, że z liczby zespolonej istnieje n pierwiastków n -go stopnia; (2) pamiętając, że $-1 = i^2$, liczba rozwiązań równania $w_n(z) = 0$ wynosi dokładnie n oraz w przypadku wielomianu o współczynnikach rzeczywistych zespolone miejsca zerowe tworzą tzw. pary $z, \bar{z}$.

Przykład Rozwiązać równania w zbiorach liczb $\mathbb{N}, \mathbb{N}_0, \mathbb{Z}, \mathbb{Q}, \mathbb{R}, \mathbb{C}$ (w tym z zbiorze liczb zespolonych)

- | | |
|-----------------------|-------------------------------------|
| • $z^2 - 1 = 0$ | • $z^2 - 2z + 4 = 0$ |
| • $z^2 + 1 = 0$ | • $z(z^2 - 2)(2z - 1)(z^2 + 9) = 0$ |
| • $z^2 - 2z + 10 = 0$ | • $z^2 + iz + 6 = 0$ |
| • $z^4 - 16 = 0$ | • $z^2 - 3iz + 10 = 0$ |

Rozwiązanie



- $z^2 - 1 = 0$

$$z^2 = 1$$

$$z = \pm 1$$

$$\mathbb{N}, \mathbb{N}_0: 1$$

$$\mathbb{Z}, \mathbb{Q}, \mathbb{R}, \mathbb{C}: \pm 1$$

- $z^2 + 1 = 0$

$$z^2 = -1$$

$$z = \sqrt{-1} = \sqrt{\cos \pi + i \sin \pi} = \cos \frac{\pi + 2\pi k}{2} + i \sin \frac{\pi + 2\pi k}{2}, k = \overline{0, 1}$$

$$k = 0 \Rightarrow z = \cos \frac{\pi}{2} + i \sin \frac{\pi}{2} = i$$

$$k = 1 \Rightarrow z = \cos \frac{3\pi}{2} + i \sin \frac{3\pi}{2} = -i$$

$$z = \pm i$$

$$\mathbb{N}, \mathbb{N}_0, \mathbb{Z}, \mathbb{Q}, \mathbb{R}: \emptyset$$

$$\mathbb{C}: \pm i$$

Inaczej, $z = i$ jest miejscem zerowym wielomianu $\Rightarrow \bar{z} = -i$ również będzie miejscem zerowym tego wielomianu, czyli w zbiorze liczb zespolonych otrzymujemy rozwiązania równania $z = \pm i$.

- $z^2 - 2z + 10 = 0$

$$z^2 - 2z + 10 = 0$$

$$\Delta = -36 = (6i)^2$$

$$z_{1,2} = \frac{2 \pm \sqrt{\Delta}}{2} = \frac{2 \pm 6i}{2} = 1 \pm 3i$$

$$\mathbb{N}, \mathbb{N}_0, \mathbb{Z}, \mathbb{Q}, \mathbb{R}: \emptyset$$

$$\mathbb{C}: 1 \pm 3i$$

- $z^4 - 16 = 0$

$$z^4 - 16 = 0$$

$$z^4 = 16$$

$$(z^2)^2 = 16$$

$$z^2 = \pm 4$$

$$\begin{array}{l} z^2 = 4 \\ z = \pm 2 \end{array} \vee \begin{array}{l} z^2 = -4 \\ z^2 = (2i)^2 \\ z = \pm 2i \end{array}$$

$$\mathbb{N}, \mathbb{N}_0: 2$$



$$\mathbb{Z}, \mathbb{Q}, \mathbb{R}: \pm 2$$

$$\mathbb{C}: \pm 2; \pm 2i$$

- $z^2 - 2z + 4 = 0$

$$z^2 - 2z + 4 = 0$$

$$\Delta = -12 = (2\sqrt{3}i)^2$$

$$z = \frac{2 \pm 2\sqrt{3}i}{2} = 1 \pm \sqrt{3}i$$

$$\mathbb{N}, \mathbb{N}_0, \mathbb{Z}, \mathbb{Q}, \mathbb{R}: \emptyset$$

$$\mathbb{C}: 1 \pm \sqrt{3}i$$

- $z(z^2 - 2)(2z - 1)(z^2 + 9) = 0$

$$z(z^2 - 2)(2z - 1)(z^2 + 9) = 0$$

$$z = 0 \quad \vee \quad \begin{array}{l} z^2 - 2 = 0 \\ z^2 = 2 \\ z = \pm\sqrt{2} \end{array} \quad \vee \quad \begin{array}{l} 2z - 1 = 0 \\ z = \frac{1}{2} \end{array} \quad \vee \quad \begin{array}{l} z^2 + 9 = 0 \\ z^2 = -9 \\ z^2 = (3i)^2 \\ z = \pm 3i \end{array}$$

$$\mathbb{N}: \emptyset$$

$$\mathbb{N}_0: 0$$

$$\mathbb{Z}: 0$$

$$\mathbb{Q}: 0, \frac{1}{2}$$

$$\mathbb{R}: 0, \frac{1}{2}, \pm\sqrt{2}$$

$$\mathbb{C}: 0; \frac{1}{2}; \pm\sqrt{2}; \pm 3i$$

- $z^2 + iz + 6 = 0$

$$z^2 + iz + 6 = 0$$

$$\Delta = i^2 - 24 = -25 = (5i)^2$$

$$z = \frac{-i \pm 5i}{2}$$

$$z = \frac{-i - 5i}{2} \quad \vee \quad z = \frac{-i + 5i}{2}$$

$$z = -3i \quad \quad \quad z = 2i$$

$$\mathbb{N}, \mathbb{N}_0, \mathbb{Z}, \mathbb{Q}, \mathbb{R}: \emptyset$$

$$\mathbb{C}: -3i; 2i$$

- $z^2 - 3iz + 10 = 0$



$$z^2 - 3iz + 10 = 0$$

$$\Delta = 9i^2 - 40 = -49 = (7i)^2$$

$$z = \frac{3i \pm 7i}{2}$$

$$z = \frac{3i - 7i}{2} \quad \vee \quad z = \frac{3i + 7i}{2}$$

$$z = -2i \quad \quad \quad z = 5i$$

$$\mathbb{N}, \mathbb{N}_0, \mathbb{Z}, \mathbb{Q}, \mathbb{R}: \emptyset$$

$$\mathbb{C}: -2i; 5i$$

Przykład Rozłożyć wielomiany zespolone na czynniki liniowe

- $w(z) = z^2 - 1$
- $w(z) = z^2 + 1$
- $w(z) = z^2 - 2z + 10$
- $w(z) = z^4 - 16$
- $w(z) = z^2 - 2z + 4$
- $w(z) = (z^3 - 2z)(2z - 1)(z^2 + 9)$
- $w(z) = z^2 + iz + 6$
- $w(z) = z^2 - 3iz + 10$

Rozwiązanie

- $w(z) = z^2 - 1$

Miejsca zerowe wielomianu: 1

$$\Rightarrow w(z) = (z - 1)(z + 1)$$

- $w(z) = z^2 + 1$

Miejsca zerowe wielomianu: $\pm i$

$$\Rightarrow w(z) = (z - i)(z + i)$$

- $w(z) = z^2 - 2z + 10$

Miejsca zerowe wielomianu: $1 \pm 3i$

$$\Rightarrow w(z) = (z - 1 - 3i)(z - 1 + 3i)$$

- $w(z) = z^4 - 16$

Miejsca zerowe wielomianu: $\pm 2; \pm 2i$

$$\Rightarrow w(x) = (z - 2)(z + 2)(z - 2i)(z + 2i)$$

- $w(z) = z^2 - 2z + 4$

Miejsca zerowe wielomianu: $1 \pm i\sqrt{3}$

$$\Rightarrow w(z) = (z - 1 - i\sqrt{3})(z - 1 + i\sqrt{3})$$



$$\bullet \quad w(z) = (z^3 - 2z)(2z - 1)(z^2 + 9)$$

$$w(z) = z(z^2 - 2)(2z - 1)(z^2 + 9)$$

Miejsca zerowe wielomianu: $0; \pm\sqrt{2}; \frac{1}{2}; \pm 3i$

$$\begin{aligned} w_n(z) = a_n z^n + \dots + a_1 z + a_0 &\Rightarrow a_6 = 2 \Rightarrow w_n(z) = a_n(z \dots) \dots \\ w(z) = 2z^6 + \dots &\Rightarrow w(z) = 2(z \dots) \dots \end{aligned}$$

$$\Rightarrow w(z) = 2z(z - \sqrt{2})(z + \sqrt{2})\left(z - \frac{1}{2}\right)(z - 3i)(z + 3i)$$

$$\bullet \quad w(z) = z^2 + iz + 6$$

Miejsca zerowe wielomianu: $-3i; 2i$

$$\Rightarrow w(z) = (z + 3i)(z - 2i)$$

$$\bullet \quad w(z) = z^2 - 3iz + 10$$

Miejsca zerowe wielomianu: $-2i; 5i$

$$\Rightarrow w(z) = (z + 2i)(z - 5i)$$

1.4. RACHUNEK RÓŻNICZKOWY FUNKCJI JEDNEJ ZMIENNEJ I JEGO ZASTOSOWANIE TECHNICZNE

Podstawowe pojęcia i twierdzenia

Definicja Pochodną funkcji $f(x)$ w punkcie x_0 nazywa się granica ilorazu różnicowego

$\lim_{\Delta x \rightarrow 0} \frac{\Delta f}{\Delta x} = \lim_{x \rightarrow x_0} \frac{f(x) - f(x_0)}{x - x_0}$ (o ile istnieje i skończona), gdzie $\Delta x = x - x_0$ – to jest przyrost argumentu, wtedy jak $\Delta f = f(x) - f(x_0)$ – przyrost funkcji. Innymi słowy, pochodną nazywa się granica ilorazu przyrostu funkcji do przyrostu argumentu przy tym, że przyrost argumentu dąży do zera o ile taka granica istnieje i skończona. Pochodną funkcji $f(x)$ w punkcie x_0 oznaczamy przez $f'(x_0)$ lub $\frac{df}{dx}|_{x=x_0}$.

Definicja Funkcja, która ma pochodną w punkcie x_0 nazywa się *różniczkowalna w punkcie x_0* . Funkcja, która różniczkowalna w każdym punkcie ze swojej dziedziny nazywa się *różniczkowalna*.

Dalej w danym rozdziale wszystkie funkcje będą różniczkowalne.

Dla obliczania pochodnej korzystamy z wymienionych niżej wzorów i reguł różniczkowania.

Wzory różniczkowania:

- $(C)' = 0, C \in \mathbb{R}$
- $(x^n)' = n \cdot x^{n-1}, n \neq 0$
 - $(x)' = 1$
 - $(x^2)' = 2x$



- $(\sqrt{x})' = \frac{1}{2\sqrt{x}}$
- $\left(\frac{1}{x}\right)' = -\frac{1}{x^2}$
- $(a^x)' = a^x \ln a, a > 0, a \neq 1$
 - $(e^x)' = e^x$
- $(\log_a x)' = \frac{1}{x \ln a}, a > 0, a \neq 1$
 - $(\ln x)' = \frac{1}{x}$
- $(\sin x)' = \cos x$
- $(\cos x)' = -\sin x$
- $(\operatorname{tg} x)' = \frac{1}{\cos^2 x}$
- $(\operatorname{ctg} x)' = \frac{-1}{\sin^2 x}$
- $(\arcsin x)' = \frac{1}{\sqrt{1-x^2}}$
- $(\arccos x)' = \frac{-1}{\sqrt{1-x^2}}$
- $(\operatorname{arctg} x)' = \frac{1}{1+x^2}$
- $(\operatorname{arcctg} x)' = \frac{-1}{1+x^2}$
- $(\sinh x)' = \cosh x$
- $(\cosh x)' = \sinh x$
- $(\operatorname{th} x)' = \frac{1}{\cosh^2 x}$
- $(\operatorname{ctgh} x)' = -\frac{1}{\sinh^2 x}$

Reguły różniczkowania: jeżeli funkcje u, v różniczkowalne, to

- $(C \cdot u)' = C \cdot u', C \in \mathbb{R}$ (stały współczynnik)
- $(u \pm v)' = u' \pm v'$ (pochodna sumy / różnicy)
- $(u \cdot v)' = u' \cdot v + u \cdot v'$ (pochodna iloczynu)
- $\left(\frac{u}{v}\right)' = \frac{u' \cdot v - u \cdot v'}{v^2}, v \neq 0$ (pochodna ilorazu)
- $u(v) = u' \cdot v'$ (pochodna funkcji złożonej)

Przykład Obliczyć pochodną funkcji $f(x) = 5x + 2$ w punktach $x_0 = 0$ i $x_0 = 3$ według definicji.

Rozwiązanie

- $x_0 = 0$:



$$\Delta x = x - x_0 = x - 0 = x$$

$$f(x) = 5x + 2$$

$$f(x_0) = f(0) = 2$$

$$\Delta f = f(x) - f(x_0) = 5x + 2 - 2 = 5x$$

$$\lim_{\Delta x \rightarrow 0} \frac{\Delta f}{\Delta x} = \lim_{x \rightarrow 0} \frac{5x}{x} = \lim_{x \rightarrow 0} (5) = 5$$

- $x_0 = 3$:

$$\Delta x = x - x_0 = x - 3$$

$$f(x) = 5x + 2$$

$$f(x_0) = f(3) = 17$$

$$\Delta f = f(x) - f(x_0) = 5x + 2 - 17 = 5x - 15$$

$$\lim_{\Delta x \rightarrow 0} \frac{\Delta f}{\Delta x} = \lim_{x \rightarrow 3} \frac{5x - 15}{x - 3} = \lim_{x \rightarrow 3} \frac{5(x - 3)}{x - 3} = \lim_{x \rightarrow 3} (5) = 5$$

Przykład Znaleźć pochodne następujących funkcji, wykorzystując wzory różniczkowania

- $y_1 = x^5$
- $y_2 = 5^x$
- $y_3 = 5^5$
- $y_4 = \log_5 x$
- $y_5 = \log_5 2$
- $y_6 = e^5$
- $y_7 = \sqrt[5]{x}$

Rozwiązanie

- $y_1 = x^5$
 $y'_1 = 5x^4$
- $y_2 = 5^x$
 $y'_2 = 5^x \ln 5$
- $y_3 = 5^5$
 $y'_3 = 0$
- $y_4 = \log_5 x$
 $y'_4 = \frac{1}{x \ln 5}$
- $y_5 = \log_5 2$
 $y'_5 = 0$
- $y_6 = e^5$
 $y'_6 = 0$
- $y_7 = \sqrt[5]{x} = x^{1/5}$
 $y'_7 = \frac{1}{5} x^{-4/5} = \frac{1}{5 \sqrt[5]{x^4}}$

Przykład Znaleźć pochodne funkcji w punktach (o ile zaznaczone), wykorzystując wzory i reguły różniczkowania

- $y_1 = x^3$
- $y_2 = \cos x$
- $y_3 = x^3 + \cos x$
- $y_4 = x^3 \cos x, x_0 = 0$
- $y_5 = \frac{x^3}{\cos x}$



- $y_6 = \frac{x^2}{2x-3}, x_0 = 2$

- $y_8 = 3^x(2x - \cosh x), x_0 = 0$

- $y_7 = x^5(e^x + 2 \ln x), x_0 = 1$

Rozwiązanie

- $y_1 = x^3$

$$y'_1 = 3x^2$$

- $y_2 = \cos x$

$$y'_2 = -\sin x$$

- $y_3 = x^3 + \cos x$

$$y'_3 = 3x^2 - \sin x$$

- $y_4 = x^3 \cos x, x_0 = 0$

$$y'_4 = 3x^2 \cos x - x^2 \sin x$$

$$y'(0) = 0$$

- $y_5 = \frac{x^3}{\cos x}$

$$y'_5 = \frac{3x^2 \cos x + x^3 \sin x}{\cos^2 x}$$

- $y_6 = \frac{x^2}{2x-3}, x_0 = 2$

$$y'_6 = \frac{2x(2x-3) - x^2(2)}{(2x-3)^2} = \frac{2x^2 - 6x}{(2x-5)^2}$$

$$y'_6(2) = -4$$

- $y_7 = x^5(e^x + 2 \ln x), x_0 = 1$

$$y'_7 = 5x^4(e^x + 2 \ln x) + x^5 \left(e^x + \frac{2}{x} \right) = (5x^4 + x^5)e^x + 10x^4 \ln x + 2x^4$$

$$y'_7(1) = 6e + 2$$

- $y_8 = 3^x(2x - \cosh x), x_0 = 0$

$$y'_8 = 3^x \ln 3 (2x - \cosh x) + 3^x(2 - \sinh x)$$

$$y'_8(0) = -\ln 3 + 2$$

Przykład Znaleźć pochodne funkcji złożonych

- $y_1 = \cos x^2$

- $y_5 = e^{x^2+2x}$

- $y_2 = \cos^2 x$

- $y_6 = (x^2 + 2x)^e$

- $y_3 = \sinh e^x$

- $y_7 = \ln(\sinh x)$

- $y_4 = e^{\sinh x}$

- $y_8 = \ln(\sinh(\ln x))$

*Rozwiązanie*

$$\bullet \quad y_1 = \cos x^2 = \cos(x^2)$$

$$y_1' = -\sin x^2 \cdot (x^2)' = -2x \sin x^2$$

$$\bullet \quad y_2 = \cos^2 x = (\cos x)^2$$

$$y_2' = 2 \cos x \cdot (\cos x)' = -2 \cos x \sin x = -\sin 2x$$

$$\bullet \quad y_3 = \sinh(e^x)$$

$$y_3' = \cosh e^x \cdot (e^x)' = e^x \cosh e^x$$

$$\bullet \quad y_4 = e^{\sinh x}$$

$$y_4' = e^{\sinh x} \cdot (\sinh x)' = e^{\sinh x} \cosh x$$

$$\bullet \quad y_5 = e^{x^2+2x}$$

$$y_5' = e^{x^2+2x} \cdot (x^2 + 2x)' = (2x + 2)e^{x^2+2x}$$

$$\bullet \quad y_6 = (x^2 + 2x)^e$$

$$y_6' = e(x^2 + 2x)^{e-1} \cdot (x^2 + 2x)' = e(2x + 2)(x^2 + 2x)^{e-1}$$

$$\bullet \quad y_7 = \ln(\sinh x)$$

$$y_7' = \frac{1}{\sinh x} \cdot (\sinh x)' = \frac{\cosh x}{\sinh x} = \operatorname{ctgh} x$$

$$\bullet \quad y_8 = \ln(\sinh(\ln x))$$

$$\begin{aligned} y_8' &= \frac{1}{\sinh(\ln x)} \cdot (\sinh(\ln x))' = \frac{1}{\sinh(\ln x)} \cdot \cosh(\ln x) \cdot (\ln x)' \\ &= \frac{\cosh(\ln x)}{x \sinh(\ln x)} = \frac{1}{x} \operatorname{ctgh} x \end{aligned}$$

Definicja Drugą pochodną funkcji $f(x)$ nazywa się pochodną pochodnej tej funkcji $(f')'$ (o ile ona istnieje); trzecią pochodną funkcji $f(x)$ jest pochodną drugiej pochodnej $(f'')'$; itd.; n -tą pochodną funkcji $f(x)$ nazywa się pochodną $(n-1)$ -ej pochodnej funkcji $f(x)$.

Pochodne wyższych rzędów oznaczamy przez $f'' \equiv \frac{d^2 f}{dx^2}$, $f''' \equiv \frac{d^3 f}{dx^3}$, $f^{(n)} \equiv \frac{d^n f}{dx^n}$.

Przykład Znaleźć pochodne wskazanego rzędu w punktach (o ile zaznaczone)

$$y = e^{5x}, y^{(n)}, y'(2), y''(1), y^{(3)}(0), y^{(n)}(0)$$

Rozwiązanie

$$y = e^{5x}$$

$$y' = 5e^{5x}$$

$$y''' = 5^3 e^{5x}$$

$$y^{(n)} = 5^n e^{5x},$$

$$y'(2) = 5e^{10}$$

$$y'''(0) = 125$$

$$n \in \mathbb{N}$$

$$y'' = 5^2 e^{5x}$$

$$\dots$$

$$y^{(n)}(0) = 5^n$$

$$y''(1) = 25e^5$$



Twierdzenie (Reguła de L'Hospitala) Jeżeli pod granicą jest jeden z symboli nieoznaczonych $\left[\frac{0}{0}, \frac{\infty}{\infty}\right]$, to granica nie ulegnie zmian przy osobnym różniczkowaniu licznika i mianownika, mianowicie

$$\lim_{x \rightarrow \dots} \frac{f(x)}{g(x)} = \left[\frac{0}{0} \vee \frac{\infty}{\infty} \right] = \lim_{x \rightarrow \dots} \frac{f'(x)}{g'(x)}$$

Przykład Obliczyć granicy funkcji

$$\bullet \lim_{x \rightarrow \infty} \frac{x^2 + 2x + 1}{5x^2 - 4}$$

$$\bullet \lim_{x \rightarrow 0} \frac{\operatorname{tg} 7x}{x}$$

$$\bullet \lim_{x \rightarrow 1} \frac{x^3 + 2x^2 + x}{x^2 - 3x + 2}$$

$$\bullet \lim_{x \rightarrow 0} \frac{\operatorname{arctg} x}{2x}$$

Rozwiązanie

$$\bullet \lim_{x \rightarrow \infty} \frac{x^2 + 2x + 1}{5x^2 - 4} = \left[\frac{\infty}{\infty} \right] = \lim_{x \rightarrow \infty} \frac{(x^2 + 2x + 1)'}{(5x^2 - 4)'} = \lim_{x \rightarrow \infty} \frac{2x + 2}{10x} = \left[\frac{\infty}{\infty} \right] = \lim_{x \rightarrow \infty} \frac{(2x + 2)'}{(10x)'} = \lim_{x \rightarrow \infty} \frac{2}{10} = 0,2$$

$$\bullet \lim_{x \rightarrow 1} \frac{x^3 - 2x^2 + x}{x^2 - 3x + 2} = \left[\frac{0}{0} \right] = \lim_{x \rightarrow 1} \frac{(x^3 - 2x^2 + x)'}{(x^2 - 3x + 2)'} = \lim_{x \rightarrow 1} \frac{3x^2 - 4x + 1}{2x - 3} = \frac{0}{-1} = 0$$

$$\bullet \lim_{x \rightarrow 0} \frac{\operatorname{tg} 7x}{x} = \left[\frac{0}{0} \right] = \lim_{x \rightarrow 0} \frac{(\operatorname{tg} 7x)'}{(x)'} = \lim_{x \rightarrow 0} \frac{\frac{7}{\cos^2 7x}}{1} = \lim_{x \rightarrow 0} \frac{7}{\cos^2 7x} = 7$$

$$\bullet \lim_{x \rightarrow 0} \frac{\operatorname{arctg} x}{2x} = \left[\frac{0}{0} \right] = \lim_{x \rightarrow 0} \frac{(\operatorname{arctg} x)'}{(2x)'} = \lim_{x \rightarrow 0} \frac{\frac{1}{1+x^2}}{2} = \lim_{x \rightarrow 0} \frac{1}{2(1+x^2)} = \frac{1}{2}$$

Równanie stycznej do wykresu funkcji

Styczną do wykresu funkcji ciągłej $y = f(x)$ w punkcie $(x_0; y_0)$, $y_0 = f(x_0)$ można znaleźć ze wzoru

$$y - y_0 = f'(x_0)(x - x_0)$$

Monotoniczność funkcji, optymalizacja

Monotoniczność funkcji oznacza, że funkcja jest rosnąca ($\nearrow$), niemalejąca ($\nearrow$), malejąca ($\searrow$), nierosnąca ($\searrow$) lub stała.

Ekstrema funkcji – to maksymalna i minimalna wartości funkcji. Są różne rodzaje ekstrema – ekstrema lokalne, ekstrema absolutne (globalne), w domkniętym przedziale, właściwe lokalne lub absolutne. Zwykle badamy ekstrema lokalne w całej dziedzinie funkcji oraz ekstrema absolutne w domkniętym przedziale. Dalej, gdy nie będzie zaznaczono innego, mówimy o ekstremach lokalnych.

Ekstremum (lokalny) funkcji: (twierdzenie Fermata, warunek konieczny) jeśli funkcja ma ekstremum w punkcie x_0 oraz jest różniczkowalna w tym punkcie, to $f'(x_0) = 0$. Punkty, w których $f'(x_0) = 0$ zwane są *punktami stacjonarnymi* funkcji $f(x)$. Inaczej mówiąc, punkty



podejrzane o ekstremum funkcji szukamy wśród punktów stacjonarnych tej funkcji oraz punktów (z dziedziny funkcji), w których pochodna nie istnieje.

Monotoniczność funkcji (warunek dostateczny): jeśli funkcja $f(x)$ jest różniczkowalna w punkcie x_0 oraz

- $f'(x_0) > 0$, to funkcja $f(x)$ jest rosnąca w punkcie x_0 ;
- $f'(x_0) < 0$, to funkcja $f(x)$ jest malejąca w punkcie x_0 .

Monotoniczność w punkcie oznacza monotoniczność w jego dość małym otoczeniu.

Algorytm 1 szukania ekstremów (oraz przedziałów monotoniczności) funkcji $f(x)$:

- wypisujemy dziedzinę D_f funkcji $f(x)$;
- obliczamy pochodną funkcji $f'(x)$;
- szukamy miejsca zerowe pochodnej $f'(x) = 0$ oraz punkty, w których $\nexists f'(x)$, z czego otrzymujemy kandydatów na ekstremum;
- badamy znak pochodnej w dość małym otoczeniu każdego z otrzymanych kandydatów, tak więc:
 - jeśli pochodna $f'(x)$ w punkcie x_0 zmienia znak z „+” na „-”, to x_0 jest punktem maksimum;
 - jeśli pochodna $f'(x)$ w punkcie x_0 zmienia znak z „-” na „+”, to x_0 jest punktem minimum.

Algorytm 2 szukania ekstremów funkcji $f(x)$:

- wypisujemy dziedzinę D_f funkcji $f(x)$;
- obliczamy pochodną funkcji $f'(x)$;
- szukamy miejsca zerowe pochodnej $f'(x) = 0$ oraz punkty, w których $\nexists f'(x)$, z czego otrzymujemy kandydatów na ekstremum;
- obliczamy drugą pochodną w znalezionych kandydatach na ekstremum, tak więc:
 - jeśli druga pochodna $f''(x) > 0$, to x_0 jest punktem minimum;
 - jeśli druga pochodna $f''(x) < 0$, to x_0 jest punktem maksimum.

Twierdzenie (Weierstrassa o kresach) Ciągła funkcja określona w domkniętym przedziale $f: [a, b] \rightarrow \mathbb{R}$ osiąga swoje maksymalną i minimalną wartości w tym przedziale.

Innymi słowy, ciągła funkcja w domkniętym przedziale zawsze ma dokładnie jeden minimum i jeden maksimum (tak zwane ekstrema absolutne w domkniętym przedziale).

Przykład Znaleźć przedziały monotoniczności oraz ekstrema lokalne funkcji

- $f_1(x) = \frac{e^x}{x^3}$
- $f_2(x) = \frac{3x+3}{x^2}$



Rozwiązanie

- $f_1(x) = \frac{e^x}{x^3}$

$$\text{Dziedzina } D_{f_1} = \mathbb{R} \setminus \{0\}$$

$$f_1'(x) = \frac{e^x x^3 - e^x 3x^2}{x^6} = \frac{e^x(x-3)}{x^4}$$

$$f_1'(x) = 0 \Leftrightarrow \frac{e^x(x-3)}{x^4} = 0$$

$$e^x(x-3) = 0 \wedge x^4 \neq 0$$

$$x = 3 \wedge x \neq 0$$

Pochodna $f_1'(x)$ nie istnieje w punkcie $x = 0$ ale $0 \notin D_{f_1}$, co oznacza, że ona istnieje we wszystkich punktach z dziedziny funkcji

x	$(-\infty; 0)$	$(0; 3)$	3	$(3; \infty)$
$f_1'(x)$	< 0	< 0	0	> 0
$f_1(x)$	$\searrow\searrow$	$\searrow\searrow$	loc min $f_{1\min} = \frac{e^3}{27}$	$\nearrow\nearrow$

$$f_1(3) = \frac{e^3}{3^3} = \frac{e^3}{27} \approx 0,744$$

- $f_2(x) = \frac{3x+3}{x^2}$

$$\text{Dziedzina } D_{f_2} = \mathbb{R} \setminus \{0\}$$

$$f_2'(x) = \left(\frac{3x+3}{x^2}\right)' = \frac{3x^2 - 2x(3x+3)}{x^4} = \frac{-3x-6}{x^3}$$

$$f_2'(x) = 0 \Leftrightarrow \frac{-3x-6}{x^3} = 0$$

$$-3x-6 = 0 \wedge x^3 \neq 0$$

$$x = -2 \wedge x \neq 0$$

$$x = -2$$

$$f_2'(x) = \frac{-3(x+2)}{x^3}$$

Pochodna $f_2'(x)$ nie istnieje w punkcie $x = 0$ ale $0 \notin D_{f_2}$, co oznacza, że ona istnieje we wszystkich punktach z dziedziny funkcji

x	$(-\infty; -2)$	-2	$(-2; 0)$	$(0; \infty)$
$f_2'(x)$	< 0	0	> 0	< 0
$f_2(x)$	$\searrow\searrow$	loc min $f_{2\min} = -0,75$	$\nearrow\nearrow$	$\searrow\searrow$

$$f_2(-2) = -0,75$$



Przykład Wyznaczyć ekstrema absolutne w domkniętym przedziale

- $f_1(x) = x^3 - 9x^2 + 24x, x \in \langle 0, 3 \rangle$
- $f_2(x) = \frac{e^x}{x^3}, x \in \langle 1, 4 \rangle$

Rozwiązanie

- $f_1(x) = x^3 - 9x^2 + 24x, x \in \langle 0, 3 \rangle$

Dziedzina $D_{f_1} = \langle 0, 3 \rangle$

$$f_1'(x) = 3x^2 - 18x + 24 = 3(x^2 - 6x + 8) = 0 \Leftrightarrow x = 2 \vee x = 4$$

$$f_1'(x) = 0 \vee \nexists f_1'(x) \Leftrightarrow x = 2 \vee x = 4$$

$$2 \in \langle 0, 3 \rangle, 4 \notin \langle 0, 3 \rangle$$

$$f_1(0) = 0$$

$$f_1(2) = 20$$

$$f_1(3) = 18$$

$$f_{1\min} = f_1(0) = 0$$

$$f_{1\max} = f_1(2) = 20$$

- $f_2(x) = \frac{e^x}{x^3}, x \in \langle 1, 4 \rangle$

Dziedzina $D_{f_2} = \langle 1, 4 \rangle \subset \mathbb{R} \setminus \{0\}$

$$f_2'(x) = 0 \vee \nexists f_2'(x) \Leftrightarrow x = 3$$

(punkt $x = 0$ nie rozstrzygamy, bo $0 \notin D_{f_2}$)

$$3 \in \langle 1, 4 \rangle$$

$$f_2(1) = e$$

$$f_2(3) = \frac{e^3}{27} \approx 0,744$$

$$f_2(4) = \frac{e^4}{64} \approx 0,853$$

$$f_{2\min} = f_2(3) = \frac{e^3}{27}$$

$$f_{2\max} = f_2(1) = e$$

Wklęsłość, wypukłość funkcji, punkty przegięcia

Wklęsłość i wypukłość funkcji to są własności o wykresu funkcji a stycznej do niego.

Definicja Funkcja nazywa się *wypukła (wklęsła)* w pewnym przedziale, jeżeli styczna w każdym punkcie z tego przedziału leży pod (nad) wykresem samej funkcji.

Również są inne określenia, tak więc wypukła $\equiv$ wypukła w dół, wklęsła $\equiv$ wypukła w górę.



Twierdzenie (warunek dostateczny) Jeśli funkcja $f(x)$ dwukrotnie różniczkowalna w przedziale (a, b) jest

- wypukła w (a, b) , to dla każdego $x_0 \in (a, b)$ zachodzi $f''(x_0) > 0$;
- wklęsła w (a, b) , to dla każdego $x_0 \in (a, b)$ zachodzi $f''(x_0) < 0$.

Definicja *Punktami przegięcia* zwane są punkty, w których funkcja z jednej strony jest wypukła lecz z drugiej strony – wklęsła.

Punkty przegięcia funkcji (warunek konieczny): jeśli x_0 jest punktem przegięcia dwukrotnie różniczkowalnej funkcji $f(x)$, to $f''(x_0) = 0$. Inaczej mówiąc, kandydatów na punkty przegięcia funkcji szukamy wśród miejsc zerowych pochodnej drugiego rzędu tej funkcji oraz punktów (z dziedziny funkcji), w których druga pochodna nie istnieje.

Algorytm 1 szukania punktów przegięcia (oraz przedziałów wklęsłości i wypukłości) funkcji $f(x)$:

- wypisujemy dziedzinę D_f funkcji $f(x)$;
- obliczamy drugą pochodną funkcji $f''(x)$;
- szukamy miejsca zerowe drugiej pochodnej $f''(x) = 0$ oraz punkty, w których $\nexists f''(x)$, z czego otrzymujemy kandydatów na punkty przegięcia;
- badamy znak drugiej pochodnej w dość małym otoczeniu każdego z otrzymanych kandydatów, tak więc:
 - jeśli $f''(x)$ w punkcie x_0 zmienia znak z „+” na „-”, to x_0 jest punktem maksimum;
 - jeśli $f''(x)$ w punkcie x_0 zmienia znak z „-” na „+”, to x_0 jest punktem minimum.

Przykład Znaleźć przedziały wklęsłości, wypukłości oraz punkty przegięcia funkcji

- $f_1(x) = x^3 - 9x^2 + 24x$
- $f_2(x) = \frac{e^x}{x^3}$

Rozwiązanie

- $f_1(x) = x^3 - 9x^2 + 24x$

Dziedzina $D_{f_1} = \mathbb{R}$

$$f_1'(x) = 3x^2 - 18x + 24$$

$$f_1''(x) = 6x - 18$$

Druga pochodna $f_1''(x)$ istnieje we wszystkich punktach z dziedziny funkcji

$$f_1''(x) = 0 \Leftrightarrow x = 3$$

x	$(-\infty; 3)$	3	$(3; +\infty)$
-----	----------------	---	----------------



		punkt przegięcia	
$f_1''(x)$	< 0	0	> 0
$f_1(x)$	wklęsła	18	wypukła

- $f_2(x) = \frac{e^x}{x^3}$

Dziedzina $D_{f_2} = \mathbb{R} \setminus \{0\}$

$$f_2'(x) = \frac{e^x(x-3)}{x^4} = \frac{xe^x - 3e^x}{x^4}$$

$$f_2''(x) = \left(\frac{xe^x - 3e^x}{x^4}\right)' = \frac{(e^x + xe^x - 3e^x)x^4 - 4x^3(xe^x - 3e^x)}{x^8}$$

$$= \frac{e^x(x^2 - 6x + 12)}{x^5}$$

Druga pochodna $f_2''(x)$ istnieje we wszystkich punktach z dziedziny funkcji

$$f_2''(x) = 0 \Leftrightarrow \frac{e^x(x^2 - 6x + 12)}{x^5} = 0$$

$$x^2 - 6x + 12 = 0 \wedge x \neq 0$$

$$\forall x: x^2 - 6x + 12 > 0$$

x	$(-\infty; 0)$	$0 \notin D_{f_2}$ punkt przegięcia	$(0; +\infty)$
$f_2''(x)$	< 0	$\nexists$	> 0
$f_2(x)$	wklęsła	$\nexists$	wypukła

Asymptoty wykresu funkcji

Definicja *Asymptotą* wykresu funkcji nazywa się prosta, jeżeli odległość od punktu wykresu a prostą dąży do zera przy oddalaniu się tego punktu w sposób nieograniczony.

Są dwa rodzaje asymptot wykresu funkcji:

- pionowa: prosta $y = a$ jest asymptotą pionową wykresu funkcji $y = f(x)$, jeśli
 - $\lim_{x \rightarrow a^-} f(x) = \begin{bmatrix} +\infty \\ -\infty \end{bmatrix} \Rightarrow$ asymptota lewostronna
 - $\lim_{x \rightarrow a^+} f(x) = \begin{bmatrix} +\infty \\ -\infty \end{bmatrix} \Rightarrow$ asymptota prawostronna
 - $\lim_{x \rightarrow a^-} f(x) = \begin{bmatrix} +\infty \\ -\infty \end{bmatrix}$ oraz $\lim_{x \rightarrow a^+} f(x) = \begin{bmatrix} +\infty \\ -\infty \end{bmatrix} \Rightarrow$ asymptota obustronna
- ukośna (w tym pozioma): prosta $y = kx + b$ jest asymptotą ukośną wykresu funkcji $y = f(x)$, jeśli istnieją skończone granice
 - $k = \lim_{x \rightarrow -\infty} \frac{f(x)}{x} \Rightarrow$ asymptota lewostronna
 - $b = \lim_{x \rightarrow -\infty} (f(x) - kx)$



$$\bullet \quad \begin{aligned} k &= \lim_{x \rightarrow +\infty} \frac{f(x)}{x} \\ b &= \lim_{x \rightarrow +\infty} (f(x) - kx) \end{aligned} \Rightarrow \text{asymptota prawostronna}$$

Niektóre zastosowania rachunku różniczkowego funkcji jednej zmiennej

Rachunek różniczkowy funkcji jednej zmiennej gra istotną rolę w wielu innych naukach, umożliwiając analizę zmian. Jeśli mamy jakąkolwiek wartość w postaci funkcji, to za pomocą pochodnej możemy rozwiązać zadanie optymalizacji tej funkcji, znaleźć przedziały, w których funkcją jest rosnąca bądź malejąca. Pochodna jest instrumentem mierzącym zmiany funkcji, mianowicie jak zmienia się funkcja przy małych zmianach argumentu. Z innej strony elastyczność E_f funkcji $f = f(x)$ – to jest stosunek względnej (lub procentowej) zmianę samej funkcji $\frac{\Delta f(x)}{f(x)}$ do względnej (procentowej) zmiany argumentu $\frac{\Delta x}{x}$, mianowicie

$$E_f = \frac{\frac{\Delta f(x)}{f(x)}}{\frac{\Delta x}{x}} = \frac{\Delta f(x)}{\Delta x} \cdot \frac{x}{f(x)} \approx \frac{x}{f(x)} \cdot f'(x)$$

Innymi słowy, z tego, że elastyczność funkcji $f(x)$ w punkcie x_0 wynosi np. 5, dowiemy się, że wartość funkcji w punkcie x_0 wzrośnie o 5% przy zwiększeniu argumentu x o 1%.

W przypadku ruchu jednowymiarowego (w kinematyce), prędkość chwilowa $v(t) = s'(t)$ jest pochodną funkcji drogi $s(t)$ względem czasu, przyspieszenie chwilowe z kolei $a(t) = v'(t) = s''(t)$ jest pochodną prędkości lub drugą pochodną drogi. To jest związane z tym, że prędkością jest szybkość zmiany położenia, wtedy jak przyspieszeniem występuje szybkość zmiany prędkości. Również ważna jest sytuacja, kiedy ciało porusza się jednostajnie zmiennie, co oznacza stałość przyspieszenia $a = \text{const.}$ Wtedy prędkość będzie określona wzorem $v(t) = at + v_0$, a droga $s(t) = v_0 t + \frac{at^2}{2}$, gdzie v_0 jest prędkością początkową.

W sterowaniu maszynami pochodna również znajduje swoje zastosowanie. W ten sposób, pozycja robota oraz jej zmiana może być opisana przez pochodną względem czasu (drugie prawo Newtona)

$$F = m \frac{d^2 s}{dt^2}$$

gdzie F to siła, m – masa, $s = s(t)$ – przemieszczenie, t – czas, $a(t) = s''(t)$ – przyspieszenie.

Innym ważnym zadaniem jest płynność ruchu robota na skrzyżowaniach, które jest wprost związane z minimalizacją prędkości lub przyspieszenia.

Oprócz tego rachunek różniczkowy jest podstawa rachunku całkowego oraz równań różniczkowych jak zwyczajnych tak i w pochodnych cząstkowych. Równania różniczkowe w swoją kolej „nieskończenie wiele” zastosowań w fizyce (równanie Newtona, Maxwella,



Schrödingera), inżynierii i mechanice (opisują ruch mechanizmów), automatyce i robotyce (sterowanie automatyczne, grafice komputerowej) i wiele innych.

Przykład Wyznaczyć elastyczność funkcji $f(x) = \frac{e^x}{x^3}$ w punktach $x_1 = 4$ i $x_2 = 2$.

Rozwiązanie

$$f(4) = \frac{e^4}{64}$$

$$f'(x) = \frac{e^x(x-3)}{x^4}$$

$$f'(x_1) = f'(4) = \frac{e^4}{256}$$

$$E_f(x_1) = \frac{x_1}{f(x_1)} \cdot f'(x_1)$$

$$E_f(4) = \frac{4}{\frac{e^4}{64}} \cdot \frac{e^4}{256} = 1$$

Ostatnie oznacza, że zwiększenie argumentu o 1% powoduje zwiększenie wartości funkcji o 1%.

$$f(2) = \frac{e^2}{8}$$

$$f'(x) = \frac{e^x(x-3)}{x^4}$$

$$f'(x_2) = f'(2) = \frac{-e^2}{16}$$

$$E_f(x_2) = \frac{x_2}{f(x_2)} \cdot f'(x_2)$$

$$E_f(2) = \frac{2}{\frac{e^2}{8}} \cdot \frac{-e^2}{16} = -1$$

Ostatnie oznacza, że zwiększenie argumentu o 1% powoduje zmniejszenie wartości funkcji o 1%.

Przykład Prostoliniowy ruch samochodu określony równaniem $s(t) = 0,2t^3 + 0,01t$ (czas jest mierzony w sekundach, droga – w metrach). Znaleźć

- o ile przemieści się w czasie 30 sekund;
- jaką będzie miał prędkość chwilową po 30 sekundach ruchu;
- jakie będzie przyspieszenie po 10 sekundach ruchu?

Rozwiązanie

- o ile przemieści się w czasie 30 sekund



$$s(30) = 180,3 \text{ m}$$

- jaką będzie miał prędkość chwilową po 30 sek. ruchu

$$v(t) = s'(t) = (0,2t^2 + 0,01t)' = 0,4t + 0,01$$

$$v(30) = 12,01 \text{ m/s}$$

- jakie będzie przyspieszenie po 10 sekundach ruchu

$$a(t) = v'(t) = (0,4t + 0,01)' \equiv 0,4 \text{ m/s}^2$$

Przyspieszenie jest stałe i wynosi $0,4 \text{ m/s}^2$.

Przykład Trajektoria ruchu punktu materialnego określona równaniem $s(t) = 0,25t^2 + 10t$ (czas jest mierzony w sekundach, droga – w metrach). Znaleźć

- o ile przemieści się obiekt po 5 sekundach ruchu;
- ile wynosi prędkość początkowa;
- ile wynosi prędkość chwilowa w momencie $t = 10 \text{ s}$;
- czy porusza się obiekt jednostajnie zmiennie i jakie jest w tym przypadku przyspieszenie?

Rozwiązanie

- o ile przemieści się obiekt po 5 sekundach ruchu

$$s(5) = 56,25 \text{ m}$$

- ile wynosi prędkość początkowa

$$v(t) = s'(t) = (0,25t^2 + 10t)' = 0,5t + 10$$

$$v(0) = 10 \text{ m/s} = 36 \text{ km/h}$$

- ile wynosi prędkość chwilowa w momencie $t = 10 \text{ s}$

$$v(t) = 0,5t + 10$$

$$v(10) = 15 \text{ m/s} = 54 \text{ km/h}$$

- czy porusza się obiekt jednostajnie zmiennie i jakie jest w tym przypadku przyspieszenie

$$a(t) = v'(t) = (0,5t + 10)' = 0,5 = \text{const}$$

Tak, obiekt porusza się ze stałym przyspieszeniem, które wynosi $a = 0,5 \text{ m/s}^2$.

Przykład Droga hamowania samochodu na suchym asfalcie opisana wzorem

$$s_1(t) = 50t - t^2$$

Znaleźć prędkość początkowa oraz kiedy samochód skończy jazdę i o ile przemieści się od początku hamowania?



Rozwiązanie

Niech droga będzie mierzona w tzw. jednostkach drogowych (jed.d.), wtedy jak czas – jednostkach czasowych (jed.cz.).

Samochód skończy jazdę wtedy i tylko wtedy, kiedy go prędkość będzie wynosiła 0.

$$s_1(t) = 50t - t^2$$

$$v_1(t) = s'_1(t) = 50 - 2t$$

$$v_1(0) = 50$$

Prędkość początkowa to jest 50 jed.d./jed.cz.

$$v'_1(t) = 0 \Leftrightarrow 50 - 2t = 0 \Leftrightarrow t = 25$$

Czas hamowania na suchym asfalcie wynosi 25 jed.cz.

$$s_1(25) = 625$$

Samochód przemieści się o 625 jed.d. od początku hamowania.

Przykład Droga hamowania samochodu na mokrym asfalcie opisana wzorem

$$s_2(t) = 50t - 0,1t^2$$

Znaleźć prędkość początkowa oraz kiedy samochód skończy jazdę i o ile przemieści się od początku hamowania?

Rozwiązanie

Niech droga będzie mierzona w tzw. jednostkach drogowych (jed.d.), wtedy jak czas – jednostkach czasowych (jed.cz.).

Samochód skończy jazdę wtedy i tylko wtedy, kiedy go prędkość będzie wynosiła 0.

Obejrzymy przypadki hamowanie w jasną pogodę i po deszczu.

Hamowanie na mokrym asfalcie:

$$s_2(t) = 50t - 0,1t^2$$

$$v_2(t) = s'_2(t) = 50 - 0,2t$$

$$v_2(0) = 50$$

Prędkość początkowa to jest 50 jed.d./jed.cz.

$$v'_2(t) = 0 \Leftrightarrow 50 - 0,2t = 0 \Leftrightarrow t = 250$$

Czas hamowania na suchym asfalcie wynosi 250 jed.cz.

$$s_2(250) = 6250$$

Samochód przemieści się o 625 jed.d. jednostek drogowych od początku hamowania.

Przykład Obiekt porusza się prostoliniowo zgodnie z trajektorią $s(t) = e^t - 1$ (czas jest mierzony w sekundach, droga – w metrach). Znaleźć

- o ile przemieści się obiekt po 5 sekundach ruchu;



- ile wynosi prędkość chwilowa w momencie $t = 5$ s;
- czy porusza się obiekt jednostajnie zmiennie.

Rozwiązanie

- o ile przemieści się obiekt po 5 sekundach ruchu

$$s(5) = e^5 - 1 \approx 147,41 \text{ m}$$

- ile wynosi prędkość chwilowa w momencie $t = 5$ s

$$v(t) = s'(t) = (e^t - 1)' = e^t$$

$$v(5) = e^5 \approx 148,41 \text{ m/s}$$

- czy porusza się obiekt jednostajnie zmiennie

$$a(t) = v'(t) = (e^t)' = e^t \neq \text{const}$$

Nie, ponieważ przyspieszenie nie jest stałą, a jest funkcja wykładniczą.

Różniczka funkcji jednej zmiennej

Różniczką funkcji jednej zmiennej nazywa się zmiana funkcji względem zmiany argumentu. Innymi słowy, różniczka funkcji $f(x)$ nazywa się

$$df(x) = f'(x)dx$$

Zastosowania różniczki powstają w obliczeniach przybliżonych oraz w rachunku całkowym funkcji jednej zmiennej, które będą dokładnie opisane w kolejnych rozdziałach.

1.5. RACHUNEK CAŁKOWY FUNKCJI JEDNEJ ZMIENNEJ I JEGO

ZASTOSOWANIE TECHNICZNE

Rachunek całkowy wraz z rachunkiem różniczkowym tworzą podstawowy oraz jeden z najważniejszych działów analizy matematycznej, zwany „Rachunkiem”. Całkowanie funkcji jest działaniem odwrotnym do różniczkowania, ale w przeciwieństwie do pochodnej, całka nie jest taka jednoznaczna, mianowicie jest określona z dokładnością do stałej.

Całka nieoznaczona

Definicja Funkcją pierwotną funkcji $f(x)$ nazywa się taka funkcja $F(x)$ (o ile ona istnieje), pochodna której jest równa danej funkcji $f(x)$, tzn. $F'(x) = f(x)$.

Przykład Sprawdzić, czy są wymienione niżej funkcji $F_i(x)$ pierwotnymi odpowiednich funkcji $f_i(x)$

- $F_1(x) = -\sin x, f_1(x) = \cos x$
- $F_2(x) = x^2, f_2(x) = x$
- $F_3(x) = 2^x, f_3(x) = \frac{1}{2^x}$
- $F_4(x) = \arccos x, f_4(x) = \arcsin x$
- $F_5(x) = e^x, F_6(x) = e^x + 1, F_7(x) = e^x - 1, F_8(x) = e^x + \text{const},$



$$f_5(x) = e^x$$

Rozwiązanie

- $F_1'(x) = (-\sin x)' = -(-\cos x) = \cos x = f_1(x)$

jest

- $F_2'(x) = (x^2)' = 2x \neq f_2(x) = x$

nie jest

- $F_3'(x) = (2^x)' = 2^x \ln 2 \neq f_3(x) = \frac{1}{2^x}$

nie jest

- $F_4'(x) = (\arccos x)' = -\frac{1}{\sqrt{1-x^2}} \neq f_4(x) = \arcsin x$

nie jest

- $F_5(x) = e^x, F_6(x) = e^x + 1, F_7(x) = e^x - 1, F_8(x) = e^x + \text{const},$

$$f_5(x) = e^x$$

$$F_5'(x) = (e^x)' = e^x = f_5(x)$$

$$F_6'(x) = (e^x + 1)' = e^x = f_5(x)$$

$$F_7'(x) = (e^x - 1)' = e^x = f_5(x)$$

$$F_8'(x) = (e^x + \text{const})' = e^x = f_5(x)$$

tak, $F_{5-8}(x)$ są pierwotnymi funkcjami $f_5(x)$.

Łatwo zauważyć, że funkcja pierwotna określona z dokładnością do stałej $C \in \mathbb{R}$, mianowicie jeżeli $F(x)$ jest pierwotną funkcji $f(x)$, to dla jakiejkolwiek liczby rzeczywistej C funkcja $F(x) + C$ również będzie pierwotną funkcji $f(x)$.

Definicja Zbiór wszystkich funkcji pierwotnych $\{F(x) + C | C \in \mathbb{R}\}$ nazywa się *całką nieoznaczoną* funkcji $f(x)$. Oznaczamy przez $\int f(x) dx = F(x) + C, C \in \mathbb{R}$. Funkcja $f(x)$ nazywa się *funkcją podcałkową*, wtedy jak $f(x)dx$ – *wyrazem podcałkowym*. Funkcja, dla której istnieje funkcja pierwotna, nazywa się *całkowalna*.

Twierdzenie Funkcja ciągła (w przedziale) jest całkowalna (w tym przedziale).

Dalej, jeżeli nie będzie zaznaczono przeciwnego, niech każda funkcja będzie ciągła, z czego będzie wynikała jej całkowalność.

Tak samo jak w rachunku różniczkowym, rzadko korzystamy z definicji aby znaleźć całkę. Natomiast wykorzystujemy wzory, reguły (własności) oraz metody całkowania.

Wzory całkowania: poniższe wzory prawdziwe w przedziałach, w których odpowiednie funkcje ciągłe

- $\int x^n dx = \frac{x^{n+1}}{n+1} + C, n \neq -1, C \in \mathbb{R}$



- $\int 1 dx = x + C, C \in \mathbb{R}$
- $\int A dx = Ax + C, A = \text{const}, C \in \mathbb{R}$
- $\int \frac{1}{x} dx = \ln |x| + C, C \in \mathbb{R}$
- $\int a^x dx = \frac{a^x}{\ln a} + C, a > 0, a \neq 1, C \in \mathbb{R}$
 - $\int e^x dx = e^x + C, C \in \mathbb{R}$
- $\int \sin x dx = -\cos x + C, C \in \mathbb{R}$
- $\int \cos x dx = \sin x + C, C \in \mathbb{R}$
- $\int \frac{1}{\sin^2 x} dx = -\text{ctg } x + C, C \in \mathbb{R}$
- $\int \frac{1}{\cos^2 x} dx = \text{tg } x + C, C \in \mathbb{R}$
- $\int \sinh x dx = \cosh x + C, C \in \mathbb{R}$
- $\int \cosh x dx = \sinh x + C, C \in \mathbb{R}$
- $\int \frac{1}{\sinh^2 x} dx = -\text{ctgh } x + C, C \in \mathbb{R}$
- $\int \frac{1}{\cosh^2 x} dx = \text{tgh } x + C, C \in \mathbb{R}$
- $\int \frac{1}{x^2 + a^2} dx = \frac{1}{a} \arctg \frac{x}{a} + C, a \neq 0, C \in \mathbb{R}$
- $\int \frac{1}{x^2 - a^2} dx = \frac{1}{2a} \ln \left| \frac{x-a}{x+a} \right| + C, a \neq 0, C \in \mathbb{R}$
(tzw. logarytm wysoki)
- $\int \frac{1}{\sqrt{a^2 - x^2}} dx = \arcsin \frac{x}{a} + C, a \neq 0, C \in \mathbb{R}$
- $\int \frac{1}{\sqrt{x^2 + b}} dx = \ln |x + \sqrt{x^2 + b}| + C, b \neq 0, C \in \mathbb{R}$
(tzw. logarytm długi)

Własności całkowania (liniowość): $u = u(x), v = v(x)$

- $\int Au dx = A \cdot \int u dx, A \in \mathbb{R}$
(stały współczynnik)
- $\int (u \pm v) dx = \int u dx \pm \int v dx$
(całka sumy / różnicy)

Metody całkowania:

- całkowanie przez podstawienie
 - całkowanie bezpośrednie
- całkowanie przez części
- całkowanie wymierne
- całkowanie pierwiastkowe
- całkowanie trygonometryczne

Przykład Znaleźć całki nieoznaczone



- | | | |
|--|--|-------------------------------------|
| • $\int 5x^4 dx$ | • $\int \left(\frac{1}{x} - x\right) dx$ | • $\int \frac{1}{x^2+2} dx$ |
| • $\int (x-3) dx$ | • $\int \left(\frac{x^3}{2} + \frac{2}{x^3} - \frac{2}{x}\right) dx$ | • $\int \frac{1}{\sqrt{x^2+3}} dx$ |
| • $\int (5x^4 + e^x) dx$ | • $\int \frac{1}{x^2-16} dx$ | • $\int \frac{1}{\sqrt{x^2-12}} dx$ |
| • $\int 2 \cos x dx$ | • $\int \frac{1}{x^2+16} dx$ | • $\int \frac{1}{\sqrt{25-x^2}} dx$ |
| • $\int \frac{2}{\sin^2 x} dx$ | • $\int \frac{1}{x^2-2} dx$ | |
| • $\int \left(x^2 + \frac{1}{x^2}\right) dx$ | | |

Rozwiązanie

- $\int 5x^4 dx = 5 \int x^4 dx = 5 \cdot \frac{x^5}{5} + C = x^5 + C, C \in \mathbb{R}$
- $\int (x-3) dx = \int x dx - 3 \int 1 dx = \frac{x^2}{2} - 3x + C, C \in \mathbb{R}$
- $\int (5x^4 + e^x) dx = 5 \int x^4 dx + \int e^x dx = x^5 + e^x + C, C \in \mathbb{R}$
- $\int 2 \cos x dx = 2 \int \cos x dx = 2 \sin x + C, C \in \mathbb{R}$
- $\int \frac{2}{\sin^2 x} dx = 2 \int \frac{1}{\sin^2 x} dx = -2 \operatorname{ctg} x + C, C \in \mathbb{R}$
- $\int \left(x^2 + \frac{1}{x^2}\right) dx = \int x^2 dx + \int x^{-2} dx = \frac{x^3}{3} + \frac{x^{-1}}{-1} + C = \frac{x^3}{3} - \frac{1}{x} + C, C \in \mathbb{R}$
- $\int \left(\frac{1}{x} - x\right) dx = \int \frac{1}{x} dx - \int x dx = \ln|x| - \frac{x^2}{2} + C, C \in \mathbb{R}$
- $\int \left(\frac{x^3}{2} + \frac{2}{x^3} - \frac{2}{x}\right) dx = \frac{1}{2} \int x^3 dx + 2 \int x^{-3} dx - 2 \int \frac{1}{x} dx = \frac{1}{2} \cdot \frac{x^4}{4} + 2 \cdot \frac{x^{-2}}{-2} - 2 \ln x + C = \frac{x^4}{8} - \frac{1}{x^2} - 2 \ln x + C, C \in \mathbb{R}$
- $\int \frac{1}{x^2-16} dx = \int \frac{1}{x^2-4^2} dx = \frac{1}{2 \cdot 4} \ln \left| \frac{x-4}{x+4} \right| + C = \frac{1}{8} \ln \left| \frac{x-4}{x+4} \right| + C, C \in \mathbb{R}$
- $\int \frac{1}{x^2+16} dx = \int \frac{1}{x^2+4^2} dx = \frac{1}{4} \operatorname{arctg} \frac{x}{4} + C, C \in \mathbb{R}$
- $\int \frac{1}{x^2-2} dx = \int \frac{1}{x^2-(\sqrt{2})^2} dx = \frac{1}{2\sqrt{2}} \ln \left| \frac{x-\sqrt{2}}{x+\sqrt{2}} \right| + C, C \in \mathbb{R}$
- $\int \frac{1}{x^2+2} dx = \int \frac{1}{x^2+(\sqrt{2})^2} dx = \frac{1}{\sqrt{2}} \operatorname{arctg} \frac{x}{\sqrt{2}} + C, C \in \mathbb{R}$
- $\int \frac{1}{\sqrt{x^2+3}} dx = \ln|x + \sqrt{x^2+3}| + C, C \in \mathbb{R}$
- $\int \frac{1}{\sqrt{x^2-12}} dx = \ln|x + \sqrt{x^2-12}| + C, C \in \mathbb{R}$
- $\int \frac{1}{\sqrt{25-x^2}} dx = \int \frac{1}{\sqrt{5^2-x^2}} dx = \arcsin \frac{x}{5} + C, C \in \mathbb{R}$

Całkowanie bezpośrednie:

$$\int f(x) dx = F(x) + C, C \in \mathbb{R} \Rightarrow \int f(kx+b) dx = \frac{1}{k} F(kx+b) + C, C \in \mathbb{R}$$



Słownie: jeśli we wzorach całkowania argument funkcji podcałkowej x zastąpimy przez funkcję liniową $kx + b$, to po całkowaniu otrzymamy współczynnik $\frac{1}{k}$.

Przykład Znaleźć całki nieoznaczone

- $\int (5 - 7x)^n dx,$
 $n \neq -1$
- $\int e^{3x+1} dx$
- $\int \frac{1}{\sinh^2 2x} dx$
- $\int \frac{1}{5-7x} dx$
- $\int \sinh 5x dx$
- $\int \frac{1}{x^2+2x+5} dx$
- $\int \frac{1}{7x+5} dx$
- $\int \cos(5 - 7x) dx$
- $\int \frac{1}{x^2+2x-3} dx$
- $\int \frac{1}{7x+5} dx$
- $\int \frac{1}{\cos^2 3x} dx$
- $\int \frac{1}{\sqrt{4x^2+4x+10}} dx$
- $\int 5^{3-x} dx$
- $\int \cosh(1 - x) dx$

Rozwiązanie

- $\int (5 - 7x)^n dx = \left[\begin{matrix} kx + b = -7x + 5 \\ k = -7 \end{matrix} \right] = \frac{1}{-7} \cdot \frac{(5-7x)^{n+1}}{n+1} + C = \frac{(5-7x)^{n+1}}{-7(n+1)} + C, C \in \mathbb{R}$
- $\int \frac{1}{5-7x} dx = \left[\begin{matrix} kx + b = -7x + 5 \\ k = -7 \end{matrix} \right] = \frac{1}{-7} \ln|5 - 7x| + C, C \in \mathbb{R}$
- $\int \frac{1}{7x+5} dx = \left[\begin{matrix} kx + b = 7x + 5 \\ k = 7 \end{matrix} \right] = \frac{1}{7} \ln|7x + 5| + C, C \in \mathbb{R}$
- $\int 5^{3-x} dx = \left[\begin{matrix} kx + b = -x + 3 \\ k = -1 \end{matrix} \right] = \frac{1}{-1} \cdot \frac{5^{3-x}}{\ln 5} + C = -\frac{5^{3-x}}{\ln 5} + C, C \in \mathbb{R}$
- $\int e^{3x+1} dx = \left[\begin{matrix} kx + b = 3x + 1 \\ k = 3 \end{matrix} \right] = \frac{1}{3} e^{3x+1} + C, C \in \mathbb{R}$
- $\int \sinh 5x dx = \left[\begin{matrix} kx + b = 5x \\ k = 5 \end{matrix} \right] = \frac{1}{5} \cosh 5x + C, C \in \mathbb{R}$
- $\int \cos(5 - 7x) dx = \left[\begin{matrix} kx + b = -7x + 5 \\ k = -7 \end{matrix} \right] = \frac{1}{-7} \sin(5 - 7x) + C = -\frac{1}{7} \sin(5 - 7x) + C, C \in \mathbb{R}$
- $\int \frac{1}{\cos^2 3x} dx = \left[\begin{matrix} kx + b = 3x \\ k = 3 \end{matrix} \right] = \frac{1}{3} \operatorname{tg} 3x + C, C \in \mathbb{R}$
- $\int \cosh(1 - x) dx = \left[\begin{matrix} kx + b = -x + 1 \\ k = -1 \end{matrix} \right] = \frac{1}{-1} \sinh(1 - x) + C = -\sinh(1 - x) + C, C \in \mathbb{R}$
- $\int \frac{1}{\sinh^2 2x} dx = \left[\begin{matrix} kx + b = 2x \\ k = 2 \end{matrix} \right] = \frac{1}{2} (-\operatorname{ctg} 2x) + C = -\frac{1}{2} \operatorname{ctg} 2x + C, C \in \mathbb{R}$
- $\int \frac{1}{x^2+2x+5} dx = \left[\begin{matrix} x^2 + 2x + 5 \\ x^2 + 2x + 1 + 4 \\ (x+1)^2 + 2^2 \end{matrix} \right] = \int \frac{1}{(x+1)^2+2^2} dx = \left[\begin{matrix} kx + b = x + 1 \\ k = 1 \end{matrix} \right] = \arctg(x + 1) + C, C \in \mathbb{R}$



$$\begin{aligned}
 \bullet \quad \int \frac{1}{x^2+2x-3} dx &= \left[\begin{array}{l} x^2 + 2x - 3 \\ x^2 + 2x + 1 - 4 \\ (x+1)^2 - 2^2 \end{array} \right] = \int \frac{1}{(x+1)^2 - 2^2} dx = \left[\begin{array}{l} kx + b = x + 1 \\ k = 1 \end{array} \right] = \\
 \frac{1}{2} \ln \left| \frac{x+1-2}{x+1+2} \right| + C &= \frac{1}{2} \ln \left| \frac{x-1}{x+3} \right| + C, C \in \mathbb{R} \\
 \bullet \quad \int \frac{1}{\sqrt{4x^2+4x+10}} dx &= \left[\begin{array}{l} 4x^2 + 4x + 10 = \\ (2x+1)^2 + 9 \end{array} \right] = \int \frac{1}{\sqrt{(2x+1)^2+9}} dx = \\
 \left[\begin{array}{l} kx + b = 2x + 1 \\ k = 2 \end{array} \right] &= \frac{1}{2} \ln |2x+1 + \sqrt{(2x+1)^2+9}| + C = \frac{1}{2} \ln |2x+1 + \\
 \sqrt{4x^2+4x+10}| &+ C, C \in \mathbb{R}
 \end{aligned}$$

Całkowanie przez podstawienie (metoda zamiany zmiennej):

$$\begin{aligned}
 \int f(g(x))g'(x) dx &= \left[\begin{array}{l} g(x) = t \\ d(g(x)) = dt \\ g'(x)dx = dt \end{array} \right] = \int f(t) dt = F(t) + C = \left[t = g(x) \right] \\
 &= F(g(x)) + C, C \in \mathbb{R}
 \end{aligned}$$

Schemat: krok 1: wybieramy funkcję, którą chcemy zamienić; krok 2: zamieniamy; krok 3: różniczkujemy zamianę; krok 4: otrzymujemy wyraz, na który zamieni się różniczka; krok 5: przepisujemy całkę przez nowy argument; krok 6: całkujemy; krok 7: wracamy do początkowego argumentu; krok 8: otrzymujemy odpowiedź. Uwaga do kroków 1 i 4: zamianę wybieramy sami tak aby w wyniku całka była przepisana w całości w nowej zmiennej (nie może być w jednej całce jednocześnie dwa argumenty) oraz funkcja podcałkowa była taką, która nadaje się do całkowania.

Całkowanie bezpośrednie jest jednym z przypadków całkowania przez podstawienie, mianowicie jeśli przy całkowaniu funkcji zamiast całkowania bezpośredniego można skorzystać z zamiany $kx + b = t$, co daje ten sam wynik.

Przykład Znaleźć całki nieoznaczone

$$\begin{aligned}
 \bullet \quad \int 2x \cos(x^2 - 1) dx & \qquad \bullet \quad \int \operatorname{tg} x dx \\
 \bullet \quad \int (2x + 3)e^{x^2+3x-1} dx & \qquad \bullet \quad \int \sin^3 x \cos x dx
 \end{aligned}$$

Rozwiązanie

$$\begin{aligned}
 \bullet \quad \int 2x \cos(x^2 - 1) dx &= \left[\begin{array}{l} x^2 - 1 = t \\ d(x^2 - 1) = dt \\ (x^2 - 1)'dx = dt \\ 2x dx = dt \end{array} \right] = \int \cos t dt = \sin t + C = \left[t = \right. \\
 x^2 - 1 &] = \sin(x^2 - 1) + C, C \in \mathbb{R}
 \end{aligned}$$



$$\bullet \int (2x+3)e^{x^2+3x-1} dx = \left[\begin{array}{l} x^2+3x-1=t \\ d(x^2+3x-1)=dt \\ (x^2+3x-1)'dx=dt \\ (2x+3)dx=dt \end{array} \right] = \int e^t dt = e^t + C = \llbracket t =$$

$$x^2+3x-1 \rrbracket = e^{x^2+3x-1} + C, C \in \mathbb{R}$$

$$\bullet \int \operatorname{tg} x dx = \int \frac{\sin x}{\cos x} dx = \left[\begin{array}{l} \cos x = t \\ d(\cos x) = dt \\ (\cos x)'dx = dt \\ -\sin x dx = dt \\ \sin x dx = -dt \end{array} \right] = -\int \frac{1}{t} dt = -\ln|t| + C = \llbracket t =$$

$$\cos x \rrbracket = -\ln|\cos x| + C, C \in \mathbb{R}$$

$$\bullet \int \sin^3 x \cos x dx = \left[\begin{array}{l} \sin x = t \\ d(\sin x) = dt \\ (\sin x)'dx = dt \\ \cos x dx = dt \end{array} \right] = \int t^3 dt = \frac{t^4}{4} + C = \llbracket t = \sin x \rrbracket =$$

$$\frac{\sin^4 x}{4} + C, C \in \mathbb{R}$$

Wyłączanie z różniczki: przy wyłączaniu funkcji z różniczki, różniczkujemy ją, tzn. szukamy pochodną

$$f(x) = f'(x)dx$$

Włączanie do różniczki: przy włączaniu funkcji do różniczki, całkujemy ją, tzn. szukamy funkcję pierwotną

$$f'(x)dx = df(x) \text{ lub } g(x)dx = dG(x)$$

gdzie $G(x)$ jest funkcją pierwotną funkcji $g(x)$.

Wskazówka: często są wykorzystywane wymienione niżej wzory na włączanie do różniczki (niech $a, b \in \mathbb{R}, a \neq 0$):

- $e^x dx = d(e^x)$
- $e^{ax+b} dx = d\left(\frac{1}{a}e^{ax+b}\right) = \frac{1}{a}d(e^{ax+b})$
- $x^n dx = d\left(\frac{x^{n+1}}{n+1}\right) = \frac{1}{n+1}d(x^{n+1})$
- $\sin x dx = d(-\cos x) = -d(\cos x)$
- $\sin(ax+b) dx = d\left(-\frac{1}{a}\cos(ax+b)\right) = -\frac{1}{a}d(\cos(ax+b))$
- $\cos x dx = d(\sin x)$
- $\cos(ax+b) dx = d\left(\frac{1}{a}\sin(ax+b)\right) = \frac{1}{a}d(\sin(ax+b))$
- $\frac{1}{x} dx = d(\ln x)$

Całkowanie przez części: $u = u(x), v = v(x)$



$$\int u dv = uv - \int v du$$

Przykład Znaleźć całki nieoznaczone

- $\int 2x \cos x dx$
- $\int \ln x dx$
- $\int x^3 e^x dx$
- $\int \sqrt{1-x^2} dx$

Rozwiązanie

$$\begin{aligned} \bullet \int 2x \cos x dx &= [\cos x dx = d(\sin x)] = \int 2x d(\sin x) = \\ &= \left[\begin{array}{l} u = 2x \\ v = \sin x \end{array} \right] = 2x \sin x - \int \sin x d(2x) = 2x \sin x - 2 \int \sin x dx = \\ &= 2x \sin x + 2 \cos x + C, C \in \mathbb{R} \end{aligned}$$

$$\begin{aligned} \bullet \int \ln x dx &= \left[\begin{array}{l} u = \ln x \\ v = x \end{array} \right] = x \ln x - \int x d(\ln x) = x \ln x - \\ &= \int x \frac{1}{x} dx + C = x \ln x - \int 1 dx + C = x \ln x - x + C = x(\ln x - 1) + C, C \in \mathbb{R} \end{aligned}$$

$$\begin{aligned} \bullet \int x^3 e^x dx &= [e^x dx = d(e^x)] = \int x^3 d(e^x) = \left[\begin{array}{l} u = x^3 \\ v = e^x \end{array} \right] = \\ &= x^3 e^x - \int e^x d(x^3) = x^3 e^x - \int 3x^2 e^x dx = x^3 e^x - 3 \int x^2 e^x dx \quad [\equiv], \\ \int x^2 e^x dx &= \int x^2 de^x = \left[\begin{array}{l} u = x^2 \\ v = e^x \end{array} \right] = x^2 e^x - \int e^x d(x^2) = \\ &= x^2 e^x - 2 \int x e^x dx = x^2 e^x - 2(xe^x - e^x) = x^2 e^x - 2xe^x + 2e^x, \text{ bo} \\ \int x e^x dx &= \int x d(e^x) = \left[\begin{array}{l} u = x \\ v = e^x \end{array} \right] = x e^x - \int e^x d(x) = x e^x - \\ &= \int e^x dx = x e^x - e^x, \\ [\equiv] x^3 e^x - 3(x^2 e^x - 2x e^x + 2e^x) + C &= e^x(x^3 - 3x^2 + 6x - 6), C \in \mathbb{R} \end{aligned}$$

$$\begin{aligned} \bullet I = \int \sqrt{1-x^2} dx &= \left[\begin{array}{l} u = \sqrt{1-x^2} \\ v = x \end{array} \right] = x\sqrt{1-x^2} - \int x d(\sqrt{1-x^2}) = \\ &= x\sqrt{1-x^2} - \int x \frac{1}{2\sqrt{1-x^2}} (-2x) dx = x\sqrt{1-x^2} - \int \frac{-x^2}{\sqrt{1-x^2}} dx = x\sqrt{1-x^2} - \\ &= \int \frac{1-x^2-1}{\sqrt{1-x^2}} dx = x\sqrt{1-x^2} - \int \left(\frac{1-x^2}{\sqrt{1-x^2}} - \frac{1}{\sqrt{1-x^2}} \right) dx = x\sqrt{1-x^2} - \int \left(\sqrt{1-x^2} - \right. \\ &= \left. \frac{1}{\sqrt{1-x^2}} \right) dx = x\sqrt{1-x^2} - I + \arcsin x \end{aligned}$$

$$I = x\sqrt{1-x^2} - I + \arcsin x$$



$$2I = x\sqrt{1-x^2} + \arcsin x$$

$$I = \frac{1}{2} \left(x\sqrt{1-x^2} + \arcsin x \right) + C$$

$$\int \sqrt{1-x^2} dx = \frac{1}{2} \left(x\sqrt{1-x^2} + \arcsin x \right) + C, C \in \mathbb{R}$$

Całka oznaczona

Definicja Niech $x_1, x_2 \in \mathbb{R}$, wtedy pod *całką oznaczoną* funkcji ciągłej $f(x)$ w przedziale $\langle x_1; x_2 \rangle$ będziemy rozumieć różnicę wartości $F(x_2)$ i $F(x_1)$ funkcji pierwotnej, tzn.

$$\int_{x_1}^{x_2} f(x) dx = [F(x)]_{x_1}^{x_2} = F(x_2) - F(x_1)$$

Ostania równość nazywa się *wzorem Newtona-Leibniza*.

Funkcja $f(x)$ jest *całkowalna* w przedziale $\langle x_1; x_2 \rangle$, gdy ona ma ograniczoną funkcję pierwotną w tym przedziale. Warunkiem koniecznym całkowalności funkcji jest jej ograniczoność, wtedy jak warunkiem dostatecznym jest ciągłość funkcji.

Przykład Obliczyć całki oznaczone

- $\int_0^1 5x^4 dx$
- $\int_{-1}^2 (x-3) dx$
- $\int_{-1}^1 (5x^4 + e^x) dx$
- $\int_0^{\pi/2} 2 \cos x dx$
- $\int_{\pi/4}^{3\pi/4} \frac{2}{\sin^2 x} dx$
- $\int_1^2 \left(\frac{1}{x} - x \right) dx$

Rozwiązanie

Poniższe rozwiązania są tylko częścią od całkowitego rozwiązania, początek którego można znaleźć w poprzednich przykładach (całka nieoznaczona)

- $\int_0^1 5x^4 dx = [x^5]_0^1 = (1) - (0) = 1$
- $\int_{-1}^2 (x-3) dx = \left[\frac{x^2}{2} - 3x \right]_{-1}^2 = \left(\frac{4}{2} - 6 \right) - \left(\frac{1}{2} - 3 \right) = \frac{15}{2}$
- $\int_{-1}^1 (5x^4 + e^x) dx = [x^5 + e^x]_{-1}^1 = (1 + e) - \left(-1 + \frac{1}{e} \right) = e - \frac{1}{e} + 2 = \frac{e^2 + 2e + 1}{e}$
- $\int_0^{\pi/2} 2 \cos x dx = [2 \sin x]_0^{\pi/2} = \left(2 \sin \frac{\pi}{2} \right) - (2 \sin 0) = 2$
- $\int_{\pi/4}^{3\pi/4} \frac{2}{\sin^2 x} dx = [-2 \operatorname{ctg} x]_{\pi/4}^{3\pi/4} = \left(-2 \operatorname{ctg} \frac{3\pi}{4} \right) - \left(-2 \operatorname{ctg} \frac{\pi}{4} \right) = 2 + 2 = 4$
- $\int_1^2 \left(\frac{1}{x} - x \right) dx = \left[\ln|x| - \frac{x^2}{2} \right]_1^2 = \left(\ln 2 - \frac{4}{2} \right) - \left(\ln 1 - \frac{1}{2} \right) = \ln 2 - \frac{3}{2}$

Własności:

- $\int_{x_1}^{x_2} A f(x) dx = A \cdot \int_{x_1}^{x_2} f(x) dx, A \in \mathbb{R}$



(stały współczynnik)

$$\bullet \int_{x_1}^{x_2} (f(x) \pm g(x)) dx = \int_{x_1}^{x_2} f(x) dx \pm \int_{x_1}^{x_2} g(x) dx$$

(całka sumy / różnicy lub addytywność względem funkcji podcałkowej)

$$\bullet \int_{x_1}^{x_2} f(x) dx = - \int_{x_2}^{x_1} f(x) dx$$

$$\bullet \int_{x_1}^{x_2} f(x) dx = \int_{x_1}^a f(x) dx + \int_a^{x_2} f(x) dx, a \in \langle x_1; x_2 \rangle$$

(addytywność względem przedziału całkowania)

$$\bullet \int_a^a f(x) dx = 0$$

$$\bullet \int_{-a}^a f(x) dx = 0 \text{ jeżeli funkcja } f(x) \text{ nieparzysta}$$

Przykład Obliczyć całki oznaczone

$$\bullet \int_0^1 (5 - 7x)^8 dx$$

$$\bullet \int_0^2 e^{3x+1} dx$$

$$\bullet \int_{-1}^0 \sinh 5x dx$$

Rozwiązanie

Poniższe rozwiązania są tylko częścią od całkowitego rozwiązania, początek którego można znaleźć w poprzednich przykładach (całka nieoznaczona)

$$\bullet \int_0^1 (5 - 7x)^8 dx = \left[-\frac{(5-7x)^9}{63} \right]_0^1 = \left[-\frac{(-2)^9}{63} \right] - \left[-\frac{5^9}{63} \right] = \frac{279\,091}{9}$$

$$\bullet \int_0^2 e^{3x+1} dx = \left[\frac{1}{3} e^{3x+1} \right]_0^2 = \left[\frac{1}{3} e^7 \right] - \left[\frac{1}{3} e \right] = \frac{e^7 - e}{3}$$

$$\bullet \int_{-1}^0 \sinh 5x dx = \left[\frac{1}{5} \cosh 5x \right]_{-1}^0 = \left[\frac{1}{5} \cosh 0 \right] - \left[\frac{1}{5} \cosh(-1) \right] = \llbracket \cosh(-x) = \cosh x \rrbracket = \frac{1 - \cosh 1}{5}$$

Całkowanie przez podstawienie (metoda zamiany zmiennej):

$$\begin{aligned} \int_{x_1}^{x_2} f(g(x))g'(x) dx &= \left[\begin{array}{ll} g(x) = t & t = g(x) \\ d(g(x)) = dt & t_1 = g(x_1) \\ g'(x)dx = dt & t_2 = g(x_2) \end{array} \right] = \int_{t_1}^{t_2} f(t) dt = [F(t)]_{t_1}^{t_2} \\ &= F(t_2) - F(t_1) \end{aligned}$$

Wskazówka: możemy obliczyć całkę oznaczoną za pomocą metody zamiany zmiennej na dwa sposoby, mianowicie wykorzystujemy schemat dla całkowania przez podstawienia jak dla całki nieoznaczonej, w którym dodatkowo

- przy zamianie zmiennej wyznaczamy nowe granice całkowania, tzn.

$$\begin{aligned} \int_{x_1}^{x_2} f(g(x))g'(x) dx &= \left[\begin{array}{ll} g(x) = t & t = g(x) \\ d(g(x)) = dt & t_1 = g(x_1) \\ g'(x)dx = dt & t_2 = g(x_2) \end{array} \right] = \int_{t_1}^{t_2} f(t) dt = [F(t)]_{t_1}^{t_2} \\ &= F(t_2) - F(t_1) \end{aligned}$$

lub



- nie obliczamy nowych granic całkowania, natomiast, zanim przejdziemy do podstawiania granic całkowania do funkcji pierwotnej, musimy wrócić do początkowego argumentu x , tzn.

$$\begin{aligned} \int_{x_1}^{x_2} f(g(x))g'(x) dx &= \left[\begin{array}{l} g(x) = t \\ d(g(x)) = dt \\ g'(x)dx = dt \end{array} \right] = \int f(t) dt = F(t) = [t = g(x)] \\ &= [F(g(x))]_{x_1}^{x_2} = F(g(x_2)) - F(g(x_1)) \end{aligned}$$

Przykład Obliczyć całki oznaczone

$$\bullet \int_0^1 (2x+3)e^{x^2+3x-1} dx \qquad \bullet \int_{-\pi/2}^0 \sin^3 x \cos x dx$$

Rozwiązanie

Poniższe rozwiązania są tylko częścią od całkowitego rozwiązania, początek którego można znaleźć w poprzednich przykładach (całka nieoznaczona)

$$\bullet \int_0^1 (2x+3)e^{x^2+3x-1} dx = \left[\begin{array}{l} t = x^2 + 3x - 1 \\ t_1 = -1 \\ t_2 = 3 \end{array} \right] = \int_{-1}^3 e^t dt = [e^t]_{-1}^3 = e^3 - e^{-1} = \frac{e^4 - 1}{e}$$

$$\bullet \int_{-\pi/2}^0 \sin^3 x \cos x dx = \left[\begin{array}{l} t = \sin x \\ t_1 = -1 \\ t_2 = 0 \end{array} \right] = \int_{-1}^0 t^3 dt = \left[\frac{t^4}{4} \right]_{-1}^0 = 0 - \frac{1}{4} = -\frac{1}{4}$$

Całkowanie przez części: $u = u(x), v = v(x), x \in \langle x_1; x_2 \rangle$

$$\int_{x_1}^{x_2} u dv = [uv]_{x_1}^{x_2} - \int_{x_1}^{x_2} v du$$

lub

$$\int_{x_1}^{x_2} u dv = \left[uv - \int v du \right]_{x_1}^{x_2}$$

Różnica między powyższymi wzorami polega na tym, że granicy całkowania podstawiamy w trakcie obliczeń (pierwszy wzór) lub na samym końcu (drugi wzór).

Przykład Znaleźć całki oznaczone

$$\bullet \int_1^e \ln x dx \qquad \bullet \int_1^3 x^3 e^x dx$$

Rozwiązanie

Poniższe rozwiązania są tylko częścią od całkowitego rozwiązania, początek którego można znaleźć w poprzednich przykładach (całka nieoznaczona)



- $$\int_1^e \ln x \, dx = [x \ln x]_1^e - \int_1^e x \, d(\ln x) = [e \ln e - \ln 1] - \int_1^e \left(x \cdot \frac{1}{x}\right) dx = e - \int_1^e 1 \, dx = e - [x]_1^e = e - (e - 1) = 1$$
- $$\int_1^3 x^3 e^x \, dx = [e^x(x^3 - 3x^2 + 6x - 6)]_1^3 = [e^3(27 - 27 + 18 - 6)] - [e(1 - 3 + 6 - 6)] = 12e^3 + 2e$$

Całka niewłaściwa

Całka niewłaściwa jest rozszerzeniem całki oznaczonej na przedziały nieograniczone (tzn. $\langle x_1; +\infty \rangle$, $(-\infty; x_2]$ lub $(-\infty; +\infty)$, $x_1, x_2 \in \mathbb{R}$) bądź na przedziały ograniczone $\langle x_1; x_2 \rangle$, ale takie, w których funkcja $f(x)$ podcałkowa jest nieograniczona. Nieograniczoność funkcji $f(x)$ w punkcie x_0 oznacza, że $\lim_{x \rightarrow x_0} f(x) = +\infty$ lub $\lim_{x \rightarrow x_0} f(x) = -\infty$. Tak więc są dwa rodzaje całek niewłaściwych, które są określone w wymienione niżej sposoby (niech $D_f = \mathbb{R}$)

- nieograniczoność przedziału całkowania:

$$\begin{aligned} \int_{x_1}^{+\infty} f(x) \, dx &= \lim_{x_2 \rightarrow +\infty} \left(\int_{x_1}^{x_2} f(x) \, dx \right), \\ \int_{-\infty}^{x_2} f(x) \, dx &= \lim_{x_1 \rightarrow -\infty} \left(\int_{x_1}^{x_2} f(x) \, dx \right), \\ \int_{-\infty}^{+\infty} f(x) \, dx &= \int_{-\infty}^{x_1} f(x) \, dx + \int_{x_1}^{+\infty} f(x) \, dx; \end{aligned}$$

- nieograniczoność funkcji podcałkowej:

dla funkcji $f(x)$, $x \in \langle x_1; x_2 \rangle$ nieograniczonej w punkcie x_1 :

$$\int_{x_1}^{x_2} f(x) \, dx = \lim_{\varepsilon \rightarrow 0^+} \left(\int_{x_1 + \varepsilon}^{x_2} f(x) \, dx \right)$$

dla funkcji $f(x)$, $x \in \langle x_1; x_2 \rangle$ nieograniczonej w punkcie x_2 :

$$\int_{x_1}^{x_2} f(x) \, dx = \lim_{\varepsilon \rightarrow 0^+} \left(\int_{x_1}^{x_2 - \varepsilon} f(x) \, dx \right)$$

dla funkcji $f(x)$, $x \in \langle x_1; x_2 \rangle$ nieograniczonej w punkcie wewnętrznym $x_0 \in (x_1; x_2)$:

$$\int_{x_1}^{x_2} f(x) \, dx = \lim_{\varepsilon \rightarrow 0} \left(\int_{x_1}^{x_0 - \varepsilon} f(x) \, dx \right) + \lim_{\delta \rightarrow 0} \left(\int_{x_0 + \delta}^{x_2} f(x) \, dx \right)$$

Możliwe również połączenie dwóch rodzajów całek niewłaściwych.

Definicja Całka niewłaściwa nazywa się *zbieżna*, jeżeli odpowiednia granica bądź granicy istnieją i skończone. Całka niewłaściwa, która nie jest zbieżna, nazywa się *rozbieżna*.



Przykład Sprawdzić zbieżność całek niewłaściwych (niektóre przykłady wzięte z [2]) i w przypadku zbieżności obliczyć

$$\bullet \int_0^{+\infty} e^x dx \quad \bullet \int_{-\infty}^0 e^x dx \quad \bullet \int_{-\infty}^{+\infty} e^x dx$$

Rozwiązanie

$$\bullet \int_0^{+\infty} e^x dx = \lim_{x_2 \rightarrow +\infty} \left(\int_0^{x_2} e^x dx \right) = \lim_{x_2 \rightarrow +\infty} (e^{x_2} - 1) = [\infty - 1] = \infty$$

rozbieżna

$$\bullet \int_{-\infty}^0 e^x dx = \lim_{x_1 \rightarrow -\infty} \left(\int_{x_1}^0 e^x dx \right) = \lim_{x_1 \rightarrow -\infty} (1 - e^{x_1}) = \llbracket 1 - e^{-\infty} = 1 - \frac{1}{e^{\infty}} = 1 - 0 \rrbracket = 1$$

zbieżna do 1

$$\bullet \int_{-\infty}^{+\infty} e^x dx = \int_{-\infty}^0 e^x dx + \int_0^{+\infty} e^x dx$$

rozbieżna, bo $\int_0^{+\infty} e^x dx$ jest rozbieżna

Niektóre zastosowania rachunku całkowego funkcji jednej zmiennej

Niech dalej każda funkcja będzie ciągła w odpowiednim przedziale. Wtedy z ciągłości funkcji wynika jej całkowalność w tym przedziale.

Pole figur płaskich: pole S figury płaskiej ograniczonej krzywą $y = f(x)$ a osią Ox w przedziale $\langle a; b \rangle$ (co równoważne dodatkowym ograniczeniu prostymi $y = a$ i $y = b$) wynosi

- $S = \int_a^b f(x) dx$, jeżeli $f(x) \geq 0, x \in \langle a; b \rangle$ (tzn. wykres funkcji znajduje się nad osią Ox w przedziale $\langle a; b \rangle$);
- $S = -\int_a^b f(x) dx$, jeżeli $f(x) \leq 0, x \in \langle a; b \rangle$ (tzn. wykres funkcji znajduje się pod osią Ox w przedziale $\langle a; b \rangle$);
- $S = \int_a^c f(x) dx + (-\int_c^b f(x) dx)$, jeżeli $f(x) \geq 0, x \in \langle a; c \rangle$ i $f(x) \leq 0, x \in \langle c; b \rangle$ (tzn. wykres funkcji znajduje się nad osią Ox w przedziale $\langle a; c \rangle$ i pod osią Ox w przedziale $\langle c; b \rangle$).

Pole S figury płaskiej ograniczonej krzywymi $y = f_1(x)$ i $y = f_2(x)$ w przedziale $\langle a; b \rangle$, takimi że dla wszystkich $x \in \langle a; b \rangle$ zachodzi $f_2(x) \geq f_1(x)$, wynosi

$$\bullet S = \int_a^b (f_2(x) - f_1(x)) dx$$

Przykład Obliczyć pola wymienionych niżej figur płaskich

- $F_1 = \{(x; y) | -1 \leq x \leq 1, x^2 - 1 \leq y \leq 0\}$
- $F_2 = \{(x; y) | 0 \leq x \leq \pi, x - \pi \leq y \leq \sin x\}$

*Rozwiązanie*

$$\bullet F_1 = \{(x; y) | -1 \leq x \leq 1, x^2 - 1 \leq y \leq 0\}$$

$$f_1(x) = x^2 - 1, x \in \langle 0; 1 \rangle$$

$$S_1 = - \int_0^1 f_2(x) dx = - \int_0^1 (x^2 - 1) dx = \left[-\frac{x^2}{3} + x \right]_0^1 = \left(-\frac{1}{3} + 1 \right) + 0 = \frac{2}{3} \\ \approx 0,67$$

$$\bullet F_2 = \{(x; y) | 0 \leq x \leq \pi, x - \pi \leq y \leq \sin x\}$$

$$f_{2,1}(x) = x - \pi, f_{2,2}(x) = \sin x, x \in \langle 0; \pi \rangle$$

$$S_2 = \int_0^\pi (f_{2,2}(x) - f_{2,1}(x)) dx = \int_0^\pi (\sin x - (x - \pi)) dx \\ = \left[-\cos x - \frac{x^2}{2} + \pi x \right]_0^\pi = \left(-\cos \pi - \frac{\pi^2}{2} + \pi^2 \right) - (-\cos 0) \\ = \frac{\pi^2}{2} + 2 \approx 6,93$$

Długość łuku krzywej płaskiej: długość L łuku krzywej funkcji $y = f(x)$, $x \in \langle a; b \rangle$ wynosi

$$L = \int_a^b \sqrt{1 + (f'(x))^2} dx$$

Przykład Obliczyć długość krzywej opisanej wzorem $f(x) = \cosh x$, $x \in \langle 0; 1 \rangle$.

Rozwiązanie

$$f(x) = \cosh x, x \in \langle 0; 1 \rangle$$

$$L = \int_0^1 \sqrt{1 + (f'(x))^2} dx = \int_0^1 \sqrt{1 + (\sinh x)^2} dx = \int_0^1 \sqrt{\cosh^2 x} dx \\ = \int_0^1 \cosh x dx = [\sinh x]_0^1 = \sinh 1 - \sinh 0 = \sinh 1 \\ = \frac{e - e^{-1}}{2} = \frac{e^2 - 1}{2e} \approx 8,68$$

Współrzędne środka ciężkości (tzw. momenty statyczne) **jednorodnej figury płaskiej:**
współrzędne środka ciężkości $(x_s; y_s)$ figury płaskiej F o polu S

$$\bullet F = \{(x; y) | a \leq x \leq b, 0 \leq y \leq f(x)\} \text{ można znaleźć ze wzorów}$$



$$x_s = \frac{1}{S} \int_a^b x f(x) dx$$

$$y_s = \frac{1}{2S} \int_a^b f^2(x) dx$$

(figura ograniczona wykresem funkcji $y = f(x)$, $x \in \langle a; b \rangle$, osią Ox oraz prostymi $x = a$ i $x = b$)

- $F = \{(x; y) | a \leq x \leq b, f_1(x) \leq y \leq f_2(x)\}$ można znaleźć ze wzorów

$$x_s = \frac{1}{S} \int_a^b x(f_2(x) - f_1(x)) dx$$

$$y_s = \frac{1}{2S} \int_a^b (f_2^2(x) - f_1^2(x)) dx$$

(figura ograniczona wykresami funkcji $y = f_1(x)$, $y = f_2(x)$, $f_2(x) \geq f_1(x)$, $x \in \langle a; b \rangle$ oraz prostymi $x = a$ i $x = b$)

Przykład Znaleźć momenty statyczne $(x_s; y_s)$ figury płaskiej

$$F = \{(x; y) | 0 \leq x \leq 3, 0 \leq y \leq e^x - 1\}$$

Rozwiązanie

$$F = \{(x; y) | 0 \leq x \leq 3, 0 \leq y \leq e^x - 1\}$$

$$f(x) = e^x - 1$$

$$S = \int_0^3 f(x) dx = \int_0^3 (e^x - 1) dx = [e^x - x]_0^3 = (e^3 - 3) - 1 = e^3 - 4$$

$$x_s = \frac{1}{S} \int_0^3 x f(x) dx = \frac{1}{e^3 - 4} \int_0^3 x (e^x - 1) dx \quad \square,$$

$$\int x e^x dx = \int [e^x dx = d(e^x)] = \int x d(e^x) = \left[\begin{array}{l} u = x \\ v = e^x \end{array} \int u dv = vu - \int v du \right] = x e^x -$$

$$\int e^x d(x) = x e^x - e^x + C = (x - 1) e^x + C, C \in \mathbb{R},$$

$$\square \frac{1}{e^3 - 4} \left[(x - 1) e^x - \frac{x^2}{2} \right]_0^3 = \frac{\left(2e^3 - \frac{9}{2} \right) - (-1)}{e^3 - 4} = \frac{4e^3 - 7}{2e^3 - 8} \approx 2,28$$



$$\begin{aligned}
 y_s &= \frac{1}{2S} \int_0^3 f^2(x) dx = \frac{1}{2(e^3 - 4)} \int_0^3 (e^x - 1)^2 dx \\
 &= \frac{1}{2(e^3 - 4)} \int_0^3 (e^{2x} - 2e^x + 1) dx \\
 &= \frac{1}{2(e^3 - 4)} \left[\frac{1}{2} e^{2x} - 2e^x + x \right]_0^3 = \frac{\left(\frac{1}{2} e^6 - 2e^3 + 3\right) - \left(\frac{1}{2} - 2\right)}{2(e^3 - 4)} \\
 &= (e^6 - 4e^3 + 9)/4(e^3 - 4) \approx 5,16
 \end{aligned}$$

Współrzędne środka ciężkości $\left(\frac{4e^3-7}{2e^3-8}; \frac{e^6-4e^3+9}{4(e^3-4)}\right)$

Pole powierzchni obrotowej: pole S powierzchni powstałej przez obrót wykresu funkcji $y = f(x)$, $x \in \langle a; b \rangle$ wokół

- osi Ox wynosi:

- $S = 2\pi \int_a^b f(x) \sqrt{1 + (f'(x))^2} dx$, jeżeli $f(x) \geq 0$, $x \in \langle a; b \rangle$
- $S = 2\pi \int_a^b |f(x)| \sqrt{1 + (f'(x))^2} dx$

- osi Oy wynosi (przy warunku dodatkowym $a \geq 0$):

$$S = 2\pi \int_a^b x \sqrt{1 + (f'(x))^2} dx$$

Przykład Obliczyć pola powierzchni obrotowej powstałej przez obrót wykresu funkcji $f(x) = \cosh x$, $0 \leq x \leq 1$ wokół osi Oy .

Rozwiązanie

$$f(x) = \cosh, 0 \leq x \leq 1 \text{ wokół osi } Oy$$



$$\begin{aligned}
S &= 2\pi \int_0^1 x \sqrt{1 + (f_2'(x))^2} dx = 2\pi \int_0^1 x \sqrt{1 + (\sinh x)^2} dx = 2\pi \int_0^1 x \cosh x dx \\
&= 2\pi \int_0^1 x d \sinh x = \left[\int u dv = uv - \int v du \right] \\
&\quad \begin{matrix} u = x \\ v = \sinh x \end{matrix} \\
&= 2\pi \left[x \sinh x - \int \sinh x dx \right]_0^1 = 2\pi [x \sinh x - \cosh x]_0^1 \\
&= 2\pi ((\sinh 1 - \cosh 1) - (0 - \cosh 0)) \\
&= 2\pi \left(\frac{e - e^{-1}}{2} - \frac{e + e^{-1}}{2} + 1 \right) = \pi(2 - 2e^{-1}) = \frac{2\pi(e - 1)}{e} \\
&\approx 3,97
\end{aligned}$$

Objętość bryły obrotowej: objętość V bryły obrotowej powstałej przez obrót figury płaskiej $F = \{(x; y) | a \leq x \leq b, 0 \leq y \leq f(x)\}$ wokół

- osi Ox wynosi:

$$V = \pi \int_a^b f^2(x) dx$$

- osi Oy wynosi (przy warunku dodatkowym $a \geq 0$):

$$V = 2\pi \int_a^b x f(x) dx$$

Jeśli figura ograniczona wykresami funkcji $f_{1,2}(x), x \in \langle a; b \rangle$, takimi że $f_1(x) \leq f_2(x), x \in \langle a; b \rangle$, mianowicie może być zapisana w postaci

$$F = \{(x; y) | a \leq x \leq b, f_1(x) \leq y \leq f_2(x)\}$$

to objętość bryły obrotowej powstałej przez obrót figury F wokół osi Ox wynosi

$$V = \pi \int_a^b (f_1^2(x) - f_2^2(x)) dx$$

Przykład Obliczyć objętości brył obrotowych powstałych przez obroty figur płaskich wokół zaznaczonych osi

- $F_1 = \{(x; y) | 0 \leq x \leq 1, 0 \leq y \leq x\}$ wokół osi Oy
- $F_2 = \{(x; y) | 0 \leq x \leq 1, 0 \leq y \leq 5x^2\}$ wokół osi Ox

Rozwiązanie

- $F_1 = \{(x; y) | 0 \leq x \leq 1, 0 \leq y \leq x\}$ wokół osi Oy

$$f_1(x) = x, x \in \langle 0; 1 \rangle$$



$$V_1 = 2\pi \int_0^1 x f_1(x) dx = 2\pi \int_0^1 x^2 dx = 2\pi \left[\frac{x^3}{3} \right]_0^1 = \frac{2\pi}{3} \approx 2,09$$

- $F_2 = \{(x; y) | 0 \leq x \leq 1, 0 \leq y \leq 5x^2\}$ wokół osi Ox

$$f_2(x) = 5x^2, x \in \langle 0; 1 \rangle$$

$$V_2 = \pi \int_0^1 f_2^2(x) dx = \pi \int_0^1 25x^4 dx = \pi [5x^5]_0^1 = 5\pi \approx 15,71$$

Współrzędne środka ciężkości jednorodnej bryły obrotowej: współrzędne środków ciężkości $(x_s; y_s)$ jednorodnej bryły obrotowej o objętości V , która powstała przez obrót figury

$$F = \{(x; y) | a \leq x \leq b, 0 \leq y \leq f(x)\}$$

wokół osi Ox można znaleźć ze wzorów

$$x_s = \frac{\pi}{V} \int_a^b x f(x) dx$$

$$y_s = 0$$

(współrzędna $y_s = 0$ zerowa, bo bryła jest symetryczna względem osi Ox).

Przykład Znaleźć współrzędne środka ciężkości bryły obrotowej, która powstała przez obrót figury płaskiej $F = \{(x; y) | 0 \leq x \leq 1, 0 \leq y \leq 5x^2\}$ wokół osi Ox .

Rozwiązanie

$$F = \{(x; y) | 0 \leq x \leq 1, 0 \leq y \leq 5x^2\}$$

$$V = 5\pi \text{ (z poprzednich przykładów)}$$

$$x_s = \frac{\pi}{V} \int_0^1 x f(x) dx = \frac{\pi}{5\pi} \int_0^1 x \cdot 5x^2 dx = \frac{1}{5} \int_0^1 5x^3 dx = \left[\frac{x^4}{4} \right]_0^1 = \frac{1}{4} - 0 = \frac{1}{4}$$

Współrzędne środka ciężkości $\left(\frac{1}{4}; 0\right)$.

Długość drogi w ruchu zmiennym: ponieważ całkowanie jest działaniem odwrotnym do różniczkowania, mając funkcję (w tym funkcję stałą) przyspieszenia ruchu pewnego ciała, można znaleźć prędkość, z czego w swoją kolej – drogę, mianowicie

$$v(t) = \int a(t) dt + v_0,$$

gdzie v_0 – to jest prędkość początkowa

$$s(t) = \int v(t) dt$$

W przypadku, gdy cząstka porusza się jednostajnie zmiennie (co oznacza, że przyspieszenie stałe) oraz zakładając, że początkowa droga wynosi zero, otrzymamy wzory



$$v(t) = \int a \, dt + v_0 = at + v_0$$

$$s(t) = \int v(t) \, dt = \int (at + v_0) \, dt = \frac{at^2}{2} + v_0 t$$

Jeżeli cząstka porusza się ze stałą prędkością $v = \text{const}$ a prędkość początkowa jest zerowa otrzymamy wzór

$$s(t) = \int v(t) \, dt = \int v \, dt = vt \Rightarrow S = v \cdot t$$

Przykład Obiekt porusza się ze stałym przyspieszeniem $a = 0,5 \, \text{m/s}^2$. Zakładając zerowość prędkości początkowej, znaleźć o ile przemieści się obiekt przez $t = 10 \, \text{s}$.

Rozwiązanie

$$s(t) = \frac{at^2}{2} + v_0 t$$

$$a = 0,5, v_0 = 0$$

$$s(t) = \frac{t^2}{4}$$

$$s(12) = \frac{10^2}{4} = 25 \, \text{m}$$

Przykład Obiekt porusza się ze zmiennym przyspieszeniem $a(t) = 0,02t \, \text{m/s}^2$. Zakładając zerowość prędkości początkowej, znaleźć o ile przemieści się obiekt przez $t = 5 \, \text{s}$.

Rozwiązanie

$$a(t) = t, v_0 = 0$$

$$v(t) = \int a(t) \, dt + v_0$$

$$v(t) = \int 0,02t \, dt = \frac{0,02t^2}{2} = 0,01t^2$$

$$s(t) = \int v(t) \, dt$$

$$s(t) = \int 0,01t^2 \, dt = \frac{0,01t^3}{3} = \frac{t^3}{300}$$

$$s(12) = \frac{10^3}{300} = \frac{1}{3} \approx 0,33 \, \text{m}$$

Przykład Obiekt porusza się ze zmiennym przyspieszeniem $a(t) = 0,24(t^3 - t) \, \text{m/s}^2$. Zakładając zerowość prędkości początkowej, znaleźć o ile przemieści się obiekt przez $t = 10 \, \text{s}$.

Rozwiązanie

$$a(t) = 0,24(t^3 - t), v_0 = 0$$



$$v(t) = \int a(t) dt + v_0$$

$$v(t) = \int (0,24t^3 - 0,24t) dt = 0,06t^4 - 0,12t^2$$

$$s(t) = \int v(t) dt$$

$$s(t) = \int (0,06t^4 - 0,12t^2) dt = 0,012t^5 - 0,04t^3$$

$$s(10) = 0,012 \cdot 10^5 - 0,04 \cdot 10^3 = 1\,160 \text{ m} = 1,16 \text{ km}$$

1.6. FUNKCJE DWU LUB WIĘCEJ ZMIENNYCH

Podstawowe pojęcia

Niech $n \in \mathbb{N}, n \geq 2$.

Definicja Funkcją n zmiennych nazywa się odwzorowanie $z = f(\bar{x})$ zbioru $D \subseteq \mathbb{R}^n$ (zwane dziedziną)

$$f : \begin{matrix} D \rightarrow \mathbb{R} \\ (x_1, x_2, \dots, x_n) \mapsto f(x_1, x_2, \dots, x_n) \end{matrix}$$

które każdemu elementowi $\bar{x} = (x_1, x_2, \dots, x_n)$ przyporządkuje dokładnie jedną wartość $f(x_1, x_2, \dots, x_n) \in \mathbb{R}$. $\bar{x} = (x_1, x_2, \dots, x_n)$ nazywa się *argumentem* funkcji, $x_1, x_2, \dots, x_n$ – *zmiennymi niezależnymi*, $f(x_1, x_2, \dots, x_n)$ – *wartością* funkcji, z – *zmienną zależną*.

W przestrzeni $\mathbb{R}^n$ odległość pomiędzy punktami $P_1(x_1^{(1)}, x_2^{(1)}, \dots, x_n^{(1)})$, $P_2(x_1^{(2)}, x_2^{(2)}, \dots, x_n^{(2)})$ określona wzorem

$$d(P_1; P_2) = \sqrt{(x_1^{(1)} - x_1^{(2)})^2 + (x_2^{(1)} - x_2^{(2)})^2 + \dots + (x_n^{(1)} - x_n^{(2)})^2}$$

Tak więc w przypadku przestrzeni

- $\mathbb{R}^1 = \mathbb{R}$: $d(P_1; P_2) = \sqrt{(x_1^{(1)} - x_1^{(2)})^2} = |x_1^{(1)} - x_1^{(2)}|$
- $\mathbb{R}^2$: $d(P_1; P_2) = \sqrt{(x_1^{(1)} - x_1^{(2)})^2 + (x_2^{(1)} - x_2^{(2)})^2}$

Dalej obejrzymy przypadek funkcji dwu zmiennych $z = f(x, y)$ (kiedy liczba zmiennych większa od dwóch, teoria przedstawia się w analogiczny sposób).

Definicja (Heinego) *Granicą funkcji* dwu zmiennych $z = f(x, y)$ w punkcie $P_0 = (x_0, y_0)$ nazywa się g , jeżeli dla każdego ciągu liczbowego $\{P_n = (x_n, y_n)\}$, dążącego do P_0 , zachodzi $\lim_{n \rightarrow \infty} f(x_n, y_n) = g$. Oznaczamy granicę funkcji dwu zmiennych przez $\lim_{\substack{x \rightarrow x_0 \\ y \rightarrow y_0}} f(x, y)$.



$$\text{Innymi słowy, } \lim_{\substack{x \rightarrow x_0 \\ y \rightarrow y_0}} f(x, y) = g \Leftrightarrow \left[\begin{matrix} P_n(x_n, y_n) \rightarrow P_o(x_0, y_0) \\ n \rightarrow \infty \end{matrix} \Rightarrow f(x_n, y_n) \rightarrow g \right]$$

Definicja Funkcja dwu zmiennych $y = f(x, y)$ nazywa się *ciągłą* w punkcie (x_0, y_0) , jeśli $\lim_{\substack{x \rightarrow x_0 \\ y \rightarrow y_0}} f(x, y) = f(x_0, y_0)$.

Inaczej mówiąc, funkcja ciągła w punkcie, jeśli granica w punkcie równa się wartość funkcji w tym punkcie.

Definicja *Pochodną cząstkową funkcji* dwu zmiennych $f(x, y)$ względem zmiennej x w punkcie (x_0, y_0) nazywa się granica (o ile ona istnieje)

$$\lim_{\Delta x \rightarrow 0} \frac{f(x_0 + \Delta x, y_0) - f(x_0, y_0)}{\Delta x}$$

Oznaczamy: $\frac{\partial f}{\partial x} \big|_{(x_0, y_0)}, f'_x(x_0, y_0), f_x(x_0, y_0)$

Definicja *Pochodną cząstkową funkcji* dwu zmiennych $f(x, y)$ względem zmiennej y w punkcie (x_0, y_0) nazywa się granica (o ile ona istnieje)

$$\lim_{\Delta y \rightarrow 0} \frac{f(x_0, y_0 + \Delta y) - f(x_0, y_0)}{\Delta y}$$

Oznaczamy: $\frac{\partial f}{\partial y} \big|_{(x_0, y_0)}, f'_y(x_0, y_0), f_y(x_0, y_0)$

Podobnie do pochodnej funkcji jednej zmiennej, dla obliczania pochodnych cząstkowych, nie koniecznym jest wykorzystanie definicji. Natomiast, dla pochodnej cząstkowej względem jednej zmiennej różniczkujemy funkcje jako funkcje jednej zmiennej, zakładając że druga zmienna jest stałą. Tak więc dla różniczkowania funkcji $f(x, y)$ względem zmiennej x różniczkujemy funkcje $f(x, \text{const})$, a dla różniczkowania funkcji $f(x, y)$ względem zmiennej y różniczkujemy funkcje $f(\text{const}, y)$.

Dla funkcji dwu zmiennych również określone pojęcia pochodnych wyższych rzędów, tzn. pochodnych od pochodnych. W porównaniu do funkcji jednej zmiennej, pochodnych wyższych rzędów jest więcej. Tak na przykład, pochodne funkcji $f(x, y)$ rzędu drugiego:

$$\bullet \frac{\partial^2 f}{\partial x^2} = \frac{\partial}{\partial x} \left(\frac{\partial f}{\partial x} \right) \quad \bullet \frac{\partial^2 f}{\partial y^2} = \frac{\partial}{\partial y} \left(\frac{\partial f}{\partial y} \right) \quad \bullet \frac{\partial^2 f}{\partial x \partial y} = \frac{\partial}{\partial x} \left(\frac{\partial f}{\partial y} \right) \quad \bullet \frac{\partial^2 f}{\partial y \partial x} = \frac{\partial}{\partial y} \left(\frac{\partial f}{\partial x} \right)$$

dla funkcji ciągłej ostatnie dwie pochodne cząstkowe (zwane *pochodnymi mieszanymi*) są równe sobie (twierdzenie Schwarz'a).

Pochodne drugiego rzędu oznaczamy przez

$$\bullet \frac{\partial^2 f}{\partial x^2} \equiv f''_{xx} \equiv f_{xx} \quad \bullet \frac{\partial^2 f}{\partial x \partial y} \equiv f''_{xy} \equiv f_{xy}$$

$$\bullet \frac{\partial^2 f}{\partial y^2} \equiv f''_{yy} \equiv f_{yy} \quad \bullet \frac{\partial^2 f}{\partial y \partial x} \equiv f''_{yx} \equiv f_{yx}$$

Pochodne drugiego rzędu w punkcie $(x_0; y_0)$ oznaczamy przez

$$\bullet \frac{\partial^2 f}{\partial x^2} \bigg|_{(x_0; y_0)} \equiv f''_{xx}(x_0; y_0) \equiv f_{xx}(x_0; y_0)$$



- $\frac{\partial^2 f}{\partial y^2} \Big|_{(x_0; y_0)} \equiv f''_{yy}(x_0; y_0) \equiv f_{yy}(x_0; y_0)$
- $\frac{\partial^2 f}{\partial x \partial y} \Big|_{(x_0; y_0)} \equiv f''_{xy}(x_0; y_0) \equiv f_{xy}(x_0; y_0)$
- $\frac{\partial^2 f}{\partial y \partial x} \Big|_{(x_0; y_0)} \equiv f''_{yx}(x_0; y_0) \equiv f_{yx}(x_0; y_0)$

Niech dalej dla każdej funkcji istnieją wszystkie pochodne cząstkowe potrzebnego rzędu (co oznacza, że funkcja będzie *różniczkowalna* odpowiedniego rzędu).

Definicja *Różniczką zupełną rzędu pierwszego* funkcji n zmiennych $f(x_1, x_2, \dots, x_n)$ w punkcie $P_0(x_1^{(0)}, x_2^{(0)}, \dots, x_n^{(0)})$ nazywa się $df(P_0) = \frac{\partial f}{\partial x_1} dx_1 + \frac{\partial f}{\partial x_2} dx_2 + \dots + \frac{\partial f}{\partial x_n} dx_n$. Różniczką zupełną rzędu pierwszego funkcji dwu zmiennych $f(x_1, x_2)$ w punkcie $P_0(x_1^{(0)}, x_2^{(0)})$ nazywa się $df(P_0) = \frac{\partial f}{\partial x_1} dx_1 + \frac{\partial f}{\partial x_2} dx_2$. Różniczką zupełną rzędu drugiego funkcji dwu zmiennych $f(x, y)$ w punkcie $P_0(x_0, y_0)$ nazywa się $df(P_0) = \frac{\partial^2 f}{\partial x^2} \Big|_{P_0} dx^2 + 2 \frac{\partial^2 f}{\partial x \partial y} \Big|_{P_0} dx dy + \frac{\partial^2 f}{\partial y^2} \Big|_{P_0} dy^2$.

Definicja *Gradientem* funkcji n zmiennych $f(x_1, x_2, \dots, x_n)$ nazywa się wektor $\left[\frac{\partial f}{\partial x_1}; \frac{\partial f}{\partial x_2}; \dots; \frac{\partial f}{\partial x_n} \right]$. Gradientem funkcji dwu zmiennych $f(x_1, x_2)$ nazywa się wektor $\left[\frac{\partial f}{\partial x_1}; \frac{\partial f}{\partial x_2} \right]$. Oznaczamy gradient funkcji f przez ∇f lub $\text{grad } f$.

Gradient (w punkcie) wskazuje kierunek największego wzrostu funkcji (w tym punkcie). Długość wektora gradient wskazuje wielkość wzrostu (tzn. o ile wzrośnie funkcja przy jednostkowym wzroście argumentu). Zbiór wektorów $\text{grad } f$ tworzy tzw. *pole gradientu*.

Jeśli na przykład, funkcja $f(x, y, z)$ jest funkcją temperatury, to $\nabla f(P_0) = \left[\frac{\partial f}{\partial x} \Big|_{P_0}, \frac{\partial f}{\partial y} \Big|_{P_0}, \frac{\partial f}{\partial z} \Big|_{P_0} \right]$ jest wskaźnikiem największego wzrostu temperatury w punkcie $P_0(x_0, y_0, z_0)$, a moduł tego wektora $|\nabla f|_{P_0} = \sqrt{\left(\frac{\partial f}{\partial x} \Big|_{P_0} \right)^2 + \left(\frac{\partial f}{\partial y} \Big|_{P_0} \right)^2 + \left(\frac{\partial f}{\partial z} \Big|_{P_0} \right)^2}$ odpowie na pytanie o ile wzrośnie temperatura gdy przemieścimy się na jednostkę długości w kierunku wektora $\nabla f(P_0)$.

Przykład Znaleźć wszystkie pochodne cząstkowe rzędu pierwszego wymienionych niżej funkcji wielu zmiennych (w zaznaczonych punktach)

- $f(x, y) = 3x^2 + 2y^3 - 5$
- $f(x, y) = x^2y + xy^2$
- $f(x, y) = \sin(2xy^5)$
- $f(x, y) = e^{xy}, P_0(1; 0)$
- $f(x, y, z) = x + y^2 + z^3, P_0(3; 2; 1)$
- $f(x, y, z) = xy + y^3z^2 - 2x^2z$



- $f(x, y, z) = \ln(xyz), P_0(1; 1; 1)$

- $f(x_1, x_2, x_3, x_4) = e^{x_1+2x_2} - x_3^2 + x_3x_4$

Rozwiązanie

- $f(x, y) = 3x^2 + 2y^3 - 5$

$$f'_x = [y = C = \text{const}] = [3x^2 + 2C^3 - 5]'_x = 6x = [C = y] = 6x$$

$$f'_y = [x = C = \text{const}] = [3C^2 + 2y^3 - 5]'_y = 6y^2 = [C = x] = 6y^2$$

- $f(x, y) = x^2y + xy^2$

$$f'_x = [y = C = \text{const}] = [x^2C + xC^2]'_x = 2xC + C^2 = [C = y] = 2xy + y^2$$

$$f'_y = [x = C = \text{const}] = [C^2y + Cy^2]'_y = C^2 + 2Cy = [C = x] = x^2 + 2xy$$

- $f(x, y) = \sin(2xy^5)$

$$f'_x = [\sin(2xy^5)]'_x = \cos(2xy^5) \cdot (2xy^5)'_x = 2y^5 \cos(2xy^5)$$

$$f'_y = [\sin(2xy^5)]'_y = \cos(2xy^5) (2xy^5)'_y = 10xy^4 \cos(2xy^5)$$

- $f(x, y) = e^{xy}, P_0(1; 0)$

$$f'_x = [e^{xy}]'_x = e^{xy} [xy]'_x = ye^{xy}$$

$$f'_y = [e^{xy}]'_y = e^{xy} [xy]'_y = xe^{xy}$$

$$f'_x|_{P_0} = (ye^{xy})|_{P_0} = \left[\begin{matrix} x_0 = 1 \\ y_0 = 0 \end{matrix} \right] = 0$$

$$f'_y|_{P_0} = (xe^{xy})|_{P_0} = \left[\begin{matrix} x_0 = 1 \\ y_0 = 0 \end{matrix} \right] = 1$$

- $f(x, y, z) = x + y^2 + z^3, P_0(3; 2; 1)$

$$f'_x = \left[\begin{matrix} y = C_1 = \text{const} \\ z = C_2 = \text{const} \end{matrix} \right] = [x + C_1^2 + C_2^3]'_x = 1 = \left[\begin{matrix} C_1 = y \\ C_2 = z \end{matrix} \right] = 1$$

$$f'_y = \left[\begin{matrix} x = C_1 = \text{const} \\ z = C_2 = \text{const} \end{matrix} \right] = [C_1 + y^2 + C_2^3]'_y = 2y = \left[\begin{matrix} C_1 = x \\ C_2 = z \end{matrix} \right] = 2y$$

$$f'_z = \left[\begin{matrix} x = C_1 = \text{const} \\ y = C_2 = \text{const} \end{matrix} \right] = [C_1 + C_2^2 + z^3]'_z = 3z^2 = \left[\begin{matrix} C_1 = x \\ C_2 = y \end{matrix} \right] = 3z^2$$

$$f'_x|_{P_0} = (1)|_{P_0} = \left[\begin{matrix} x_0 = 3 \\ y_0 = 2 \\ z_0 = 1 \end{matrix} \right] = 1$$

$$f'_y|_{P_0} = (2y)|_{P_0} = \left[\begin{matrix} x_0 = 3 \\ y_0 = 2 \\ z_0 = 1 \end{matrix} \right] = 4$$

$$f'_z|_{P_0} = (3z^2)|_{P_0} = \left[\begin{matrix} x_0 = 3 \\ y_0 = 2 \\ z_0 = 1 \end{matrix} \right] = 3$$

- $f(x, y, z) = xy + y^3z^2 - 2x^2z$

$$f'_x = [xy + y^3z^2 - 2x^2z]'_x = y + 0 - 4xz = y - 4xz$$

$$f'_y = [xy + y^3z^2 - 2x^2z]'_y = x + 3y^2z - 0 = x + 3y^2z$$

$$f'_z = [xy + y^3z^2 - 2x^2z]'_z = 0 + 2y^3z - 2x^2 = 2y^3z - 2x^2$$



- $f(x, y, z) = \ln(xyz), P_0(1; 1; 1)$

$$f'_x = [\ln(xyz)]'_x = \frac{1}{xyz} \cdot [xyz]'_x = \frac{yz}{xyz} = \frac{1}{x}$$

$$f'_y = [\ln(xyz)]'_y = \frac{1}{xyz} \cdot [xyz]'_y = \frac{xz}{xyz} = \frac{1}{y}$$

$$f'_z = [\ln(xyz)]'_z = \frac{1}{xyz} \cdot [xyz]'_z = \frac{xy}{xyz} = \frac{1}{z}$$

$$f'_x|_{P_0} = \left(\frac{1}{x}\right)|_{P_0} = \left[\begin{matrix} x_0 = 1 \\ y_0 = 1 \\ z_0 = 1 \end{matrix}\right] = 1$$

$$f'_y|_{P_0} = \left(\frac{1}{y}\right)|_{P_0} = \left[\begin{matrix} x_0 = 1 \\ y_0 = 1 \\ z_0 = 1 \end{matrix}\right] = 1$$

$$f'_z|_{P_0} = \left(\frac{1}{z}\right)|_{P_0} = \left[\begin{matrix} x_0 = 1 \\ y_0 = 1 \\ z_0 = 1 \end{matrix}\right] = 1$$

- $f(x_1, x_2, x_3, x_4) = e^{x_1+2x_2} - x_3^2 + x_3x_4$

$$f'_{x_1} = [e^{x_1+2x_2} - x_3^2 + x_3x_4]'_{x_1} = e^{x_1+2x_2}[x_1 + 2x_2]'_{x_1} - 0 + 0 = e^{x_1+2x_2}$$

$$f'_{x_2} = [e^{x_1+2x_2} - x_3^2 + x_3x_4]'_{x_2} = e^{x_1+2x_2}[x_1 + 2x_2]'_{x_2} - 0 + 0 = 2e^{x_1+2x_2}$$

$$f'_{x_3} = [e^{x_1+2x_2} - x_3^2 + x_3x_4]'_{x_3} = 0 - 2x_3 + x_4 = x_4 - 2x_3$$

$$f'_{x_4} = [e^{x_1+2x_2} - x_3^2 + x_3x_4]'_{x_4} = 0 - 0 + x_3 = x_3$$

Przykład Znaleźć wszystkie pochodne cząstkowe rzędu drugiego wymienionych niżej funkcji dwóch i trzech zmiennych

- $f(x, y) = x^2y + xy^2$

- $f(x, y) = e^{xy}$

- $f(x, y) = \sin(2xy^5)$

- $f(x, y, z) = xy + y^3z^2 - 2x^2z$

Rozwiązanie

- $f(x, y) = x^2y + xy^2$

$$f''_{xx} = (f'_x)'_x = (2xy + y^2)'_x = 2y$$

$$f''_{xy} = (f'_x)'_y = (2xy + y^2)'_y = 2x + 2y$$

$$f''_{yx} = (f'_y)'_x = (x^2 + 2xy)'_x = 2x + 2y$$

$$f''_{yy} = (f'_y)'_y = (x^2 + 2xy)'_y = 2x$$

$$f''_{xy} \equiv f''_{yz} \equiv 2x + 2y$$

- $f(x, y) = \sin(2xy^5)$

$$f''_{xx} = (f'_x)'_x = (2y^5 \cos(2xy^5))'_x = 2y^5(-\sin(2xy^5)) \cdot 2y^5 = -4y^{10} \sin(2xy^5)$$

$$f''_{xy} = (f'_x)'_y = (2y^5 \cos(2xy^5))'_y = 10y^4 \cos(2xy^5) - 20xy^9 \sin(2xy^5)$$



$$f''_{yx} = (f'_y)'_x = (10xy^4 \cos(2xy^5))'_x = 10y^4 \cos(2xy^5) - 20xy^9 \sin(2xy^5)$$

$$f''_{yy} = (f'_y)'_y = (10xy^4 \cos(2xy^5))'_y = 40xy^3 \cos(2xy^5) - 100x^2y^8 \sin(2xy^5)$$

$$f''_{xy} \equiv f''_{yz} \equiv 10y^4 \cos(2xy^5) - 20xy^9 \sin(2xy^5)$$

- $f(x, y) = e^{xy}$

$$f''_{xx} = (f'_x)'_x = (ye^{xy})'_x = y^2 e^{xy}$$

$$f''_{xy} = (f'_x)'_y = (ye^{xy})'_y = e^{xy} + xye^{xy}$$

$$f''_{yx} = (f'_y)'_x = (xe^{xy})'_x = e^{xy} + xye^{xy}$$

$$f''_{yy} = (f'_y)'_y = (xe^{xy})'_y = x^2 e^{xy}$$

$$f''_{xy} \equiv f''_{yz} \equiv e^{xy} + xye^{xy}$$

- $f(x, y, z) = xy + y^3z^2 - 2x^2z$

$$f''_{xx} = (f'_x)'_x = (y - 4xz)'_x = -4z$$

$$f''_{xy} = (f'_x)'_y = (y - 4xz)'_y = 1$$

$$f''_{xz} = (f'_x)'_z = (y - 4xz)'_z = -4x$$

$$f''_{yx} = (f'_y)'_x = (x + 3y^2z^2)'_x = 1$$

$$f''_{yy} = (f'_y)'_y = (x + 3y^2z^2)'_y = 6yz^2$$

$$f''_{yz} = (f'_y)'_z = (x + 3y^2z^2)'_z = 6y^2z$$

$$f''_{zx} = (f'_z)'_x = (2y^3z - 2x^2)'_x = -4x$$

$$f''_{zy} = (f'_z)'_y = (2y^3z - 2x^2)'_y = 6y^2z$$

$$f''_{zz} = (f'_z)'_z = (2y^3z - 2x^2)'_z = 2y^3$$

$$f''_{xy} \equiv f''_{yx} \equiv 1$$

$$f''_{xz} \equiv f''_{zx} \equiv -4x$$

$$f''_{yz} \equiv f''_{zy} \equiv 6y^2z$$

Przykład Znaleźć gradient funkcji dwu zmiennych w zaznaczonych punktach

- $f(x, y) = x^2 + y^2, P_0(0; 0)$
- $f(x, y) = x^2 - y^2, P_0(1; -1)$

Rozwiązanie

- $f(x, y) = x^2 + y^2, P_0(0; 0)$

$$f'_x = 2x$$

$$f'_x(0; 0) = 0$$

$$f'_y = 2y$$

$$f'_y(0; 0) = 0$$

$$\text{grad } f|_{P_0} = (0; 0)$$

- $f(x, y) = x^2 - y^2, P_0(1; -1)$



$$f'_x = 2x$$

$$f'_x(1; -1) = 2$$

$$f'_y = -2y$$

$$f'_y(1; -1) = 2$$

$$\text{grad } f|_{P_0} = (2; 2)$$

Optimalizacja funkcji dwu zmiennych

Podobnie do funkcji jednej zmiennej, mówiąc o optymalizacji funkcji wielu zmiennych, myślimy o poszukiwaniu ekstremów lokalnych, największy, najmniejszych wartości (w domkniętym obszarze) bądź o znalezieniu ekstremów warunkowych (ostatniego nie ma dla funkcji jednej zmiennej). W poniższym materiale wykorzystujemy niektóre pojęcia algebry liniowej, dokładną informację o których można znaleźć w następnym rozdziale (Działanie na macierzach. Obliczanie wyznacznika i macierzy odwrotnej. Operacje elementarne. Rozwiązywanie układów równań liniowych).

Definicja *Macierzą Hessego* (lub hessianem) funkcji dwu zmiennych $f(x, y)$ w punkcie $P_0(x_0, y_0)$ nazywa się macierz pochodnych cząstkowych rzędu drugiego obliczonych w punkcie P_0 , mianowicie

$$H_f(P_0) = \left[\begin{array}{cc} \frac{\partial^2 f}{\partial x^2} & \frac{\partial^2 f}{\partial x \partial y} \\ \frac{\partial^2 f}{\partial y \partial x} & \frac{\partial^2 f}{\partial y^2} \end{array} \right]_{P_0}$$

Oznaczając $A = \frac{\partial^2 f}{\partial x^2}$, $B = \frac{\partial^2 f}{\partial x \partial y} \equiv \frac{\partial^2 f}{\partial y \partial x}$, $C = \frac{\partial^2 f}{\partial y^2}$, dla funkcji która ma ciągle pochodne cząstkowe otrzymamy $H_f(P_0) = \left[\begin{array}{cc} A & B \\ B & C \end{array} \right]_{P_0}$.

$$H_f = \left[\begin{array}{cc} \frac{\partial^2 f}{\partial x^2} & \frac{\partial^2 f}{\partial x \partial y} \\ \frac{\partial^2 f}{\partial y \partial x} & \frac{\partial^2 f}{\partial y^2} \end{array} \right] \equiv \left[\begin{array}{cc} f''_{xx} & f''_{xy} \\ f''_{yx} & f''_{yy} \end{array} \right]$$

Definicja *Punktami stacjonarnymi* funkcji dwu zmiennych $f(x, y)$ nazywają się punkty z dziedziny funkcji, w których

$$\frac{\partial f}{\partial x} = \frac{\partial f}{\partial y} = 0$$

Ekstremum lokalne:



- warunek konieczny: jeśli funkcja dwu zmiennych posiada w punkcie $P_0(x_0, y_0)$

ekstremum lokalne, to P_0 jest punktem stacjonarnym, tzn. $\begin{cases} \frac{\partial f}{\partial x}|_{P_0} = 0 \\ \frac{\partial f}{\partial y}|_{P_0} = 0 \end{cases}$, lub

pochodne cząstkowe w punkcie P_0 nie istnieją

- warunek dostateczny (dla punktów P_0 spełniających warunek konieczny):

$$\det H_f(P_0) > 0 \quad \wedge \quad \begin{cases} A|_{P_0} = \frac{\partial^2 f}{\partial x^2}|_{P_0} > 0 \Rightarrow \min \\ A|_{P_0} = \frac{\partial^2 f}{\partial x^2}|_{P_0} < 0 \Rightarrow \max \end{cases}$$

$\det H_f(P_0) < 0 \Rightarrow$ funkcja nie posiada extr w punkcie P_0

$$\begin{cases} \det H_f(P_0) > 0 \quad \wedge \quad A|_{P_0} = \frac{\partial^2 f}{\partial x^2}|_{P_0} = 0 \\ \det H_f(P_0) = 0 \end{cases} \Rightarrow \text{brak odpowiedzi}$$

Brak odpowiedzi oznacza, że powyższy warunek dostateczny nie rozstrzyga istnienia ekstremów w odpowiednim przypadku.

Przykład Zbadać ekstrema lokalne poniższych funkcji dwu zmiennych [2]

- $f(x, y) = x^3 + y^3 + 3xy$
- $f(x, y) = 3 \ln \frac{x}{6} + 2 \ln y + \ln(12 - x - y)$

Rozwiązanie

- $f(x, y) = x^3 + y^3 + 3xy$

$$\begin{cases} f'_x = 3x^2 + 3y = 0 \\ f'_y = 3y^2 + 3x = 0 \end{cases} \Rightarrow (0; 0), (-1; -1)$$

$$f''_{xx} = 6x$$

$$f''_{xx}(0; 0) = 0$$

$$f''_{xx}(-1; -1) = -6$$

$$f''_{xy} \equiv f''_{yx} = 3$$

$$f''_{xy}(0; 0) = f''_{yx}(0; 0) = 3$$

$$f''_{xy}(-1; -1) = f''_{yx}(-1; -1) = 3$$

$$f''_{yy} = 6y$$

$$f''_{yy}(0; 0) = 0$$

$$f''_{yy}(-1; -1) = -6$$

$$H_f(0; 0) = \begin{bmatrix} 0 & 3 \\ 3 & 0 \end{bmatrix}$$

$$\det H_f(0; 0) = \begin{vmatrix} 0 & 3 \\ 3 & 0 \end{vmatrix} = -9$$



$$\det H_f(0; 0) = \begin{vmatrix} 0 & 3 \\ 3 & 0 \end{vmatrix} = -9 < 0 \Rightarrow \text{funkcja } f(x; y) \text{ nie osiąga extr w } (0; 0)$$

$$H_f(-1; -1) = \begin{bmatrix} -6 & 3 \\ 3 & -6 \end{bmatrix}$$

$$\det H_f(-1; -1) = \begin{vmatrix} -6 & 3 \\ 3 & -6 \end{vmatrix} = 27$$

$$\det H_f(-1; -1) = 27 > 0$$

$$f''_{xx}(-1; -1) = -6 < 0 \Rightarrow \text{funkcja } f(x; y) \text{ osiąga max w } (-1; -1)$$

$$f_{\max} = f(-1; -1) = 1$$

- $f(x, y) = 3 \ln \frac{x}{6} + 2 \ln y + \ln(12 - x - y)$

$$f(x, y) = 3 \ln x - 3 \ln 6 + 2 \ln y + \ln(12 - x - y)$$

$$\begin{cases} f'_x = \frac{3}{x} - \frac{1}{12 - x - y} = 0 \\ f'_y = \frac{2}{y} - \frac{1}{12 - x - y} = 0 \end{cases} \Rightarrow (6; 4)$$

$$f''_{xx} = -\frac{3}{x^2} - \frac{1}{(12 - x - y)^2}$$

$$f''_{xx}(6; 4) = -\frac{1}{3}$$

$$f''_{xy} \equiv f''_{yx} = -\frac{1}{(12 - x - y)^2}$$

$$f''_{xy}(6; 4) = f''_{yx}(6; 4) = -\frac{1}{4}$$

$$f''_{yy} = -\frac{2}{y^2} - \frac{1}{(12 - x - y)^2}$$

$$f''_{yy}(6; 4) = -\frac{3}{8}$$

$$H_f(6; 4) = \begin{bmatrix} -\frac{1}{3} & -\frac{1}{4} \\ -\frac{1}{4} & -\frac{3}{8} \end{bmatrix}$$

$$\det H_f(6; 4) = \begin{vmatrix} -\frac{1}{3} & -\frac{1}{4} \\ -\frac{1}{4} & -\frac{3}{8} \end{vmatrix} = \frac{1}{16}$$

$$\det H_f(6; 4) = \frac{1}{16} > 0$$

$$f''_{xx}(6; 4) = -\frac{1}{3} < 0 \Rightarrow \text{funkcja } f(x; y) \text{ osiąga max w } (6; 4)$$

$$f_{\max} = f(6; 4) = 5 \ln 2$$

Największa, najmniejsza wartości w domkniętym obszarze: podobnie do funkcji jednej zmiennej, aby znaleźć największą i najmniejszą wartości funkcji w domkniętym



obszarze, działamy według schematu: (1) szukamy punkty stacjonarne funkcji, (2) spośród znalezionych punktów wybieramy takie, które leżą wewnątrz obszaru, (3) obliczamy wartości funkcji w otrzymanych punktach oraz na brzegach obszaru, (4) ze wszystkich otrzymanych wartości wybieramy największą i najmniejszą wartości funkcji.

Badanie funkcji na brzegach obszaru potrzebuje dużo obliczeń, ponieważ nawet w przypadku obszaru prostokątnego, otrzymujemy cztery brzegi i na każdym z nich jak wynik zadanie największej, najmniejszej wartości funkcji jednej zmiennej w domkniętym przedziale.

Przykład Znaleźć największą i najmniejszą wartości poniższych funkcji w zaznaczonych domkniętych obszarach [2]

$$\bullet f(x, y) = x^2 + y^2 - xy - x, D = \{(x, y) \mid 0 \leq x \leq 1, 0 \leq y \leq 1\}$$

$$\bullet f(x, y) = x^2 - y^2, D = \{(x, y) \mid x^2 + y^2 \leq 1\}$$

Rozwiązanie

$$\bullet f(x, y) = x^2 + y^2 - xy - x, D = \{(x, y) \mid 0 \leq x \leq 1, 0 \leq y \leq 1\}$$

$$\begin{aligned} f'_x &= 2x - y - 1 = 0 \\ f'_y &= 2y - x = 0 \end{aligned} \Rightarrow \left(\frac{2}{3}; \frac{1}{3}\right) \in D$$

$$f\left(\frac{2}{3}; \frac{1}{3}\right) = -\frac{1}{3}$$

$$f(0; y) = y^2, y \in \langle 0; 1 \rangle$$

$$g_1(y) = f(0; y) = y^2, y \in \langle 0; 1 \rangle$$

$$\begin{aligned} g_1(0) &= 0 \\ g_1(1) &= 1 \end{aligned} \Rightarrow \begin{aligned} g_{1\min} &= g_1(0) = f(0; 0) = 0 \\ g_{1\max} &= g_1(1) = f(0; 1) = 1 \end{aligned} \quad y \in \langle 0; 1 \rangle$$

$$f(1; y) = y^2 - y, y \in \langle 0; 1 \rangle$$

$$g_2(y) = y^2 - y = \left(y - \frac{1}{2}\right)^2 - \frac{1}{4}, y \in \langle 0; 1 \rangle$$

$$\begin{aligned} g_2(0) &= 0 \\ g_2\left(\frac{1}{2}\right) &= -\frac{1}{4} \\ g_2(1) &= 0 \end{aligned} \Rightarrow \begin{aligned} g_{2\max} &= g_2(0) = g_2(1) = f(1; 0) = f(1; 1) = 0 \\ g_{2\min} &= g_2\left(\frac{1}{2}\right) = f\left(1; \frac{1}{2}\right) = -\frac{1}{4} \end{aligned} \quad y \in \langle 0; 1 \rangle$$

$$f(x, 0) = x^2 - x, x \in \langle 0; 1 \rangle$$

$$g_3(x) = x^2 - x = \left(x - \frac{1}{2}\right)^2 - \frac{1}{4}$$

$$\begin{aligned} g_3(0) &= 0 \\ g_3\left(\frac{1}{2}\right) &= -\frac{1}{4} \\ g_3(1) &= 0 \end{aligned} \Rightarrow \begin{aligned} g_{3\min} &= g_3\left(\frac{1}{2}\right) = f\left(\frac{1}{2}; 0\right) = -\frac{1}{4} \\ g_{3\max} &= g_3(0) = g_3(1) = f(0; 0) = f(0; 1) = 0 \end{aligned} \quad x \in \langle 0; 1 \rangle$$



$$f(x; 1) = x^2 - 2x + 1$$

$$g_4(x) = x^2 - 2x + 1 = (x - 1)^2$$

$$\begin{aligned} g_4(0) = 1 \\ g_4(1) = 0 \end{aligned} \Rightarrow \begin{aligned} g_{4\min} = g_4(1) = f(1; 1) = 0 \\ g_{4\max} = g_4(0) = f(0; 1) = 1 \end{aligned} \quad x \in \langle 0; 1 \rangle$$

$$g_{4\min} = g_4(1) = f(1; 1) = 0$$

Do porównani są poniższe wartości:

$$f\left(\frac{2}{3}; \frac{1}{3}\right) = -\frac{1}{3};$$

$$f\left(\frac{1}{2}; 0\right) = -\frac{1}{4};$$

$$f(0; 0) = 0;$$

$$f(0; 0) = f(0; 1) = 0;$$

$$f(0; 1) = 1;$$

$$f(1; 1) = 0;$$

$$f(1; 0) = f(1; 1) = 0;$$

$$f(0; 1) = 1$$

$$f\left(1; \frac{1}{2}\right) = -\frac{1}{4};$$

$$f_{\max} = f(0; 1) = 1$$

$$f_{\min} = f\left(\frac{2}{3}; \frac{1}{3}\right) = -\frac{1}{3} \quad x \in D$$

- $f(x, y) = x^2 - y^2, D = \{(x, y) | x^2 + y^2 \leq 1\}$

$$\begin{aligned} f'_x = 2x = 0 \\ f'_y = 2y = 0 \end{aligned} \Rightarrow (0; 0) \in D$$

$$f(0; 0) = 0$$

Brzeg obszaru $D: x^2 + y^2 = 1 \Leftrightarrow y^2 = 1 - x^2 \Leftrightarrow y = \pm\sqrt{1 - x^2}, x \in \langle -1; 1 \rangle$

$$\Rightarrow g(x) = f\left(x, \pm\sqrt{1 - x^2}\right) = x^2 - (1 - x^2) = 2x^2 - 1, x \in \langle -1; 1 \rangle$$

$$g(x) = 2x^2 - 1, x \in \langle -1; 1 \rangle$$

$$g'(x) = 4x = 0 \Rightarrow x = 0 \in \langle -1; 1 \rangle$$

$$g(-1) = 1$$

$$g(0) = -1 \Rightarrow \begin{aligned} g_{\min} = g(0) = f(0; \pm 1) = -1 \\ g_{\max} = g(\pm 1) = f(\pm 1; 0) = 1 \end{aligned} \quad x \in \langle -1; 1 \rangle$$

$$g(1) = 1$$

Do porównani są poniższe wartości:

$$f(0; 0) = 0; f(0; 1) = -1; f(\pm 1; 0) = 1$$

$$\begin{aligned} f_{\max} = f(\pm 1; 0) = 1 \\ f_{\min} = f(0; \pm 1) = -1 \end{aligned} \quad x \in D$$

Ekstremum warunkowe: *ekstremami warunkowymi* nazywają się ekstrema funkcji (max, min) jeśli argumenty są związane dodatkowymi warunkami. Tak więc, dla funkcji dwu zmiennych zagadnienie poszukiwania ekstremów warunkowych wygląda następująco

$$\begin{aligned} f(x, y) &\rightarrow \text{extr} \\ \psi(x, y) &= 0 \end{aligned}$$



Ekstremum warunkowe jest uogólnieniem ekstremum lokalnego.

Przy zwiększaniu liczby zmiennych, liczba warunków dodatkowych również może (ale nie musi) ulec zwiększeniu, ale dla funkcji n zmiennych maksymalna liczba warunków dodatkowych wynosi $(n - 1)$.

Schemat badania ekstremum warunkowego: (0) sprawdzamy czy nie można przejść do funkcji jednej zmiennej, podstawiając jedną ze zmiennych z warunku dodatkowego; jeśli nie ma takiej możliwości, to rozwiązujemy według schematu; (1) formujemy funkcję pomocniczą, zwaną później *funkcją Lagrange'a* (lub *Lagranżjanem*)

$$L(x, y, \lambda) = f(x, y) + \lambda \psi(x, y)$$

(2) badamy ekstrema funkcji $L(x, y, \lambda)$ podobnie do funkcji trzech zmiennych, co oznacza

- poszukiwanie punktów stacjonarnych $\frac{\partial L}{\partial x} = \frac{\partial L}{\partial y} = \frac{\partial L}{\partial \lambda} = 0$, z czego otrzymujemy punkty podejrzone o ekstremum
- obliczamy wszystkie pochodne cząstkowe rzędu drugiego Lagranżjana (łącznie 9 szt.)
- układamy macierz Hessego (jako dla funkcji trzech zmiennych x, y i λ)

$$H_L = \begin{bmatrix} L''_{xx} & L''_{xy} & L''_{x\lambda} \\ L''_{yx} & L''_{yy} & L''_{y\lambda} \\ L''_{\lambda x} & L''_{\lambda y} & L''_{\lambda\lambda} \end{bmatrix} \Rightarrow H_L = \begin{bmatrix} L''_{xx} & L''_{xy} & \psi'_x \\ L''_{yx} & L''_{yy} & \psi'_y \\ \psi'_x & \psi'_y & 0 \end{bmatrix}$$

i obliczamy jej wyznacznik we wszystkich punktach podejrzanym o ekstremum, wtedy

$$\det H_L(P_0) > 0 \Rightarrow \max$$

$$\det H_L(P_0) < 0 \Rightarrow \min$$

$$\det H_L(P_0) = 0 \Rightarrow \text{brak odpowiedzi}$$

Przykład Znaleźć ekstrema warunkowe [2]

- $f(x, y) = \frac{1}{x} + \frac{1}{y}$, przy warunku $xy = 1$
- $f(x, y) = x^3 + y^3$, przy warunku $x + y = 2$

Rozwiązanie

- $f(x, y) = \frac{1}{x} + \frac{1}{y}$, przy warunku $xy = 1$

$$xy = 1 \Rightarrow \begin{cases} x \neq 0 \\ y \neq 0 \end{cases} \Rightarrow \left[xy = 1 \Leftrightarrow y = \frac{1}{x} \right]$$

$y = \frac{1}{x}$ podstawiamy do funkcji $f(x, y)$ i otrzymujemy funkcję jednej zmiennej x ,

którą oznaczmy przez $g(x) \equiv f\left(x, \frac{1}{x}\right)$

$$g(x) = \frac{1}{x} + x = \frac{x^2 + 1}{x}$$



$$g'(x) = \frac{x^2 - 1}{x^2}$$

$$g'(x) = 0 \Leftrightarrow \begin{cases} x = \pm 1 \\ x \neq 0 \end{cases}$$

x	$(-\infty; -1)$	-1	$(-1; 0)$	$(0; 1)$	1	$(1; +\infty)$
$g'(x)$	> 0	0	< 0	< 0	0	> 0
$g(x)$	$\nearrow \nearrow$	loc min $g_{\min} =$	$\searrow \searrow$	$\searrow \searrow$	loc max $g_{\max} = 4$	$\nearrow \nearrow$

$$g_{\min} = g(-1) = -2 \quad \Rightarrow \quad f_{\min} = f(-1; -1) = -2$$

$$g_{\max} = g(1) = 2 \quad \Rightarrow \quad f_{\max} = f(1; 1) = 2$$

- $f(x, y) = x^3 + y^3$, przy warunku $x + y = 2$

W tym przykładzie również można przejść do badania ekstremum funkcji jednej zmiennej, ale zrobmy według całego schematu ((1),(2)).

$$f(x, y) = x^3 + y^3$$

$$\psi(x, y) = x + y - 2$$

$$L(x, y, \lambda) = f(x, y) + \lambda \psi(x, y)$$

$$L(x, y, \lambda) = x^3 + y^3 + \lambda(x + y - 2)$$

$$\begin{cases} L'_x = 3x^2 + \lambda = 0 \\ L'_y = 3y^2 + \lambda = 0 \\ L'_\lambda = x + y - 2 = 0 \end{cases} \Rightarrow \begin{cases} x = 1 \\ y = 1 \\ \lambda = -3 \end{cases} \Rightarrow \begin{matrix} (1; 1; -3) \\ P_0(1; 1), \lambda = -3 \end{matrix}$$

$$L''_{xx} = 6x \Rightarrow L''_{xx}(1; 1; -3) = 6$$

$$L''_{yx} = L''_{xy} = 0 \Rightarrow L''_{yx}(1; 1; -3) = L''_{xy}(1; 1; -3) = 0$$

$$L''_{yy} = 6y \Rightarrow L''_{yy}(1; 1; -3) = 6$$

$$\psi'_x = 1 \Rightarrow \psi'_x(1; 1; -3) = 1$$

$$\psi'_y = 1 \Rightarrow \psi'_y(1; 1; -3) = 1$$

$$H_L(1; 1; -3) = \begin{bmatrix} 6 & 0 & 1 \\ 0 & 6 & 1 \\ 1 & 1 & 0 \end{bmatrix}$$

$$\det H_L(1; 1; -3) = \begin{vmatrix} 6 & 0 & 1 \\ 0 & 6 & 1 \\ 1 & 1 & 0 \end{vmatrix} = -12$$

$$\det H_L(1; 1; -3) = -12 < 0 \Rightarrow f_{\min} = f(1; 1) = 2$$



1.7. DZIAŁANIE NA MACIERZACH. OBLICZANIE WYZNACZNIKA I MACIERZY ODWROTNEJ. OPERACJE ELEMENTARNE. ROZWIĄZYWANIE UKŁADÓW RÓWNAŃ LINIOWYCH

Macierze

Definicja *Macierz*ą nazywa się prostokątna tablica liczb. Zapisujemy macierze w postaci $A = [a_{ij}]_{n \times m}$ (lub bez zaznaczania tzw. wymierności $A = [a_{ij}]$), gdzie $a_{ij}, i = \overline{1, n}, j = \overline{1, m}$ zwane elementami macierzy A , tzn.

$$A = \begin{bmatrix} a_{11} & a_{12} & \dots & a_{1m} \\ a_{21} & a_{22} & \dots & a_{2m} \\ \dots & \dots & \dots & \dots \\ a_{n1} & a_{n2} & \dots & a_{nm} \end{bmatrix}$$

Rzędy poziome $[a_{11} \ a_{12} \ \dots \ a_{1m}], [a_{21} \ a_{22} \ \dots \ a_{2m}], \dots, [a_{n1} \ a_{n2} \ \dots \ a_{nm}]$ są

zwane *wierszami* macierzy, rzędy pionowe $\begin{bmatrix} a_{11} \\ a_{21} \\ \dots \\ a_{n1} \end{bmatrix}, \begin{bmatrix} a_{12} \\ a_{22} \\ \dots \\ a_{n2} \end{bmatrix}, \dots, \begin{bmatrix} a_{1m} \\ a_{2m} \\ \dots \\ a_{nm} \end{bmatrix}$ – *kolumnami*. Wskaźnik ij

elementu a_{ij} macierzy A pokazuje jego przynależność do i -ego wiersza a j -ej kolumny.

Definicja *Wymiernością* macierzy (lub macierz jest o wymiarach) nazywa się $n \times m$, gdzie n tj. liczba wierszy, m tj. liczba kolumn; oznaczamy przez $\dim A$.

Definicja Macierze $A = [a_{ij}]$ i $B = [b_{ij}]$ zwane *równymi*, gdy $\dim A = \dim B$ oraz elementy na odpowiednich miejscach równe.

Definicja Gdy liczba wierszy i liczba kolumn są równe sobie, macierz A nazywa się *kwadratowa* i w tym przypadku mówimy, że A jest *macierzą n -go stopnia*. Dla macierzy $A = [a_{ij}]$ n -go stopnia elementy $a_{ii}, i = \overline{1, n}$ tworzą przekątną główną, wtedy jak elementy $a_{i, n-i+1}, i = \overline{1, n}$ – przekątną drugą, innymi słowy

przekątna główna

$$\begin{bmatrix} a_{11} & \cdot & \cdot & \cdot & \cdot \\ \cdot & a_{22} & \cdot & \cdot & \cdot \\ \cdot & \cdot & a_{33} & \cdot & \cdot \\ \cdot & \cdot & \cdot & \dots & \cdot \\ \cdot & \cdot & \cdot & \cdot & a_{nn} \end{bmatrix}$$

przekątna druga

$$\begin{bmatrix} \cdot & \cdot & \cdot & \cdot & a_{1n} \\ \cdot & \cdot & \cdot & a_{2, n-1} & \cdot \\ \cdot & \cdot & \dots & \cdot & \cdot \\ \cdot & a_{n-1, 2} & \cdot & \cdot & \cdot \\ a_{n1} & \cdot & \cdot & \cdot & \cdot \end{bmatrix}$$

Definicja *Macierz*ą *zerową* nazywa się macierz, elementami której występują liczby zerowe. Oznaczamy macierz zerową przez $O_{n \times m}$ lub bez zaznaczania wymierności przez O .

Definicja *Macierz*ą *jednostkową* (tzw. *jedynką macierzową*) nazywa się macierz kwadratowa, w której na przekątnej głównej są 1, a pozostałe elementy zerowe. Oznaczamy macierz jednostkową n -go stopnia przez I_n lub bez zaznaczania wymierności przez I .

Przykład Macierze jednostkowe i zerowe o zaznaczonych wymiarach wyglądają tak



$$O_{21} = \begin{bmatrix} 0 \\ 0 \end{bmatrix}$$

$$O_{22} = \begin{bmatrix} 0 & 0 \\ 0 & 0 \end{bmatrix}$$

$$O_{23} = \begin{bmatrix} 0 & 0 & 0 \\ 0 & 0 & 0 \end{bmatrix}$$

$$I_1 = [1]$$

$$I_2 = \begin{bmatrix} 1 & 0 \\ 0 & 1 \end{bmatrix}$$

$$I_3 = \begin{bmatrix} 1 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & 1 \end{bmatrix}$$

Działanie na macierzach:

- Transponowanie macierzy: macierz $A^T = [a_{ji}]_{m \times n}$ nazywa się *macierzą transponowaną* macierzy $A = [a_{ij}]_{n \times m}$. Inaczej mówiąc macierz transponowaną otrzymamy z macierzy A przestawiając wiersze z kolumnami, lub przepisując 1-szy wiersz jako 1-szą kolumnę, 2-gi wiersz jako 2-gą kolumnę itd., ostatni wiersz jako ostatnią kolumnę. Takie przekształcenie macierzy można sobie wyobrazić jako symetrię elementów względem przekątnej głównej, dlatego macierzy które nie zmieniają się po transponowaniu zwane *symetryczne*.
Niektóre własności transponowania: (1) $(A^T)^T = A$; (2) $\dim A = n \times m \Rightarrow \dim A^T = m \times n$.
- *Mnożeniem macierzy* $A = [a_{ij}]_{n \times m}$ przez liczbę $c \in \mathbb{R}$ nazywa się macierz $cA = [ca_{ij}]_{n \times m}$. Wskazówka: mnożymy każdy element macierzy przez liczbę.
- *Sumą, różnicą macierzy* $A = [a_{ij}]_{n \times m}$ i $B = [b_{ij}]_{n \times m}$ o jednakowych wymiarach nazywa się macierz $C = [a_{ij} \pm b_{ij}]_{n \times m}$. Wskazówka: dodajemy, odejmujemy elementy na odpowiadających sobie miejscach. Macierz zerowa jest elementem neutralnym dodawania bądź odejmowania macierzy. Elementem przeciwnym dodawania do macierzy A występuje macierz $(-1)A$.
- *Mnożeniem macierzy* $A = [a_{ij}]_{n \times k}$ przez macierz $B = [b_{ij}]_{k \times m}$ (liczba kolumn macierzy A równa się liczbie wierszy macierzy B) nazywa się macierz $C = [c_{ij}]_{n \times m}$, w której element c_{ij} jest iloczynem i -go wiersza macierzy A a j -ej kolumny macierzy B (w sensie iloczynu skalarnego wektorów).

Innymi słowy,

$$A \cdot B = C$$

$$[a_{ij}]_{n \times k} \cdot [b_{ij}]_{k \times m} = [c_{ij}]_{n \times m}, \text{ gdzie}$$

$$c_{ij} = [a_{i1} \quad a_{i2} \quad \dots \quad a_{ik}] \cdot \begin{bmatrix} b_{1j} \\ b_{2j} \\ \dots \\ b_{kj} \end{bmatrix} = a_{i1}b_{1j} + a_{i2}b_{2j} + \dots + a_{ik}b_{kj}$$



Macierz jednostkowa jest elementem neutralnym mnożenia dwóch macierzy.

Mnożenie macierzy przez macierz na ogół nie jest przemienne, tzn. $A \cdot B \neq B \cdot A$ (często jest sytuacja, kiedy iloczyn AB istnieje, natomiast w innej kolejności BA nie istnieje lub na odwrót).

Dla znalezienia iloczynu AB można wykorzystać poniższy schemat

$$\begin{array}{c|c} & B \\ \hline A & AB \end{array}$$

- *Potęgowanie macierzy*: $A^2 = A \cdot A$, $A^3 = A^2 \cdot A$, $A^4 = A^3 \cdot A$, itd.
- *Operacje elementarne* na wierszach macierzy to są wymienione niżej działania:
 - pomnożenie wiersza macierzy przez liczbę niezerową
 - zamiana dwóch wierszy miejscami
 - dodawanie do jednego wiersza innego pomnożonego przez liczbę

W ten sam sposób można określić operacje elementarne na kolumnach macierzy.

Przykład Dla poniższych macierzy A_i , $i = 1; 2$, znaleźć: (1) wymiarność: $\dim A_i$; (2) elementy $a_{13}, a_{31}, a_{22}, a_{24}, a_{15}$ (gdy takie istnieją); (3) transponowanie: A_i^T

$$\bullet \quad A_1 = \begin{bmatrix} 1 & \pi & 0 \\ 5 & -1 & 2 \\ 3 & 0 & -3 \end{bmatrix} \quad \bullet \quad A_2 = \begin{bmatrix} 0 & 1 & 3 & 0 & -2 \\ 1 & 0 & -5 & e & 7 \end{bmatrix}$$

Rozwiązanie

$$\bullet \quad A_1 = \begin{bmatrix} 1 & \pi & 0 \\ 5 & -1 & 2 \\ 3 & 0 & -3 \end{bmatrix} \Rightarrow \dim A_1 = 3 \times 3 \quad \begin{cases} a_{13} = 0 \\ a_{31} = 3 \\ a_{22} = -1 \\ \nexists a_{24} \\ \nexists a_{15} \end{cases}$$

$$A_1^T = \begin{bmatrix} 1 & \pi & 0 \\ 5 & -1 & 2 \\ 3 & 0 & -3 \end{bmatrix}^T = \begin{bmatrix} 1 & 5 & 3 \\ \pi & -1 & 0 \\ 0 & 2 & -3 \end{bmatrix}$$

$$\bullet \quad A_2 = \begin{bmatrix} 0 & 1 & 3 & 0 & -2 \\ 1 & 0 & -5 & e & 7 \end{bmatrix} \Rightarrow \dim A_2 = 2 \times 5 \quad \begin{cases} a_{13} = 3 \\ \nexists a_{31} \\ a_{22} = 0 \\ a_{24} = e \\ a_{15} = -2 \end{cases}$$

$$A_2^T = \begin{bmatrix} 0 & 1 & 3 & 0 & -2 \\ 1 & 0 & -5 & e & 7 \end{bmatrix}^T = \begin{bmatrix} 0 & 1 \\ 1 & 0 \\ 3 & -5 \\ 0 & e \\ -2 & 7 \end{bmatrix}$$

Przykład Dla macierzy kwadratowych A i B znaleźć macierz $3A - 2B^T + I_2$, gdzie

$$A = \begin{bmatrix} 1 & 2 \\ 3 & 4 \end{bmatrix}, B = \begin{bmatrix} -1 & 1 \\ 2 & 0 \end{bmatrix}$$



Rozwiązanie

$$A = \begin{bmatrix} 1 & 2 \\ 3 & 4 \end{bmatrix} \Rightarrow 3A = \begin{bmatrix} 3 & 6 \\ 9 & 12 \end{bmatrix}$$

$$B = \begin{bmatrix} -1 & 1 \\ 2 & 0 \end{bmatrix} \Rightarrow B^T = \begin{bmatrix} -1 & 2 \\ 1 & 0 \end{bmatrix} \Rightarrow 2B^T = \begin{bmatrix} -2 & 4 \\ 2 & 0 \end{bmatrix}$$

$$3A - 2B^T + I_2 = \begin{bmatrix} 3 & 6 \\ 9 & 12 \end{bmatrix} - \begin{bmatrix} -2 & 4 \\ 2 & 0 \end{bmatrix} + \begin{bmatrix} 1 & 0 \\ 0 & 1 \end{bmatrix} = \begin{bmatrix} 6 & 2 \\ 7 & 13 \end{bmatrix}$$

Przykład Dla macierzy $A = \begin{bmatrix} 2 & -1 \\ 0 & 3 \end{bmatrix}, B = \begin{bmatrix} 0 & -1 \\ 2 & -1 \end{bmatrix}$ znaleźć $AB, BA, A^2, B^2, A^2 - B^2, A^T, B^T, A^T + B^T$.

Rozwiązanie

$$\bullet AB = \begin{bmatrix} 2 & -1 \\ 0 & 3 \end{bmatrix} \cdot \begin{bmatrix} 0 & -1 \\ 2 & -1 \end{bmatrix} = \begin{bmatrix} 2 \cdot 0 + (-1) \cdot 2 & 2 \cdot (-1) + (-1) \cdot (-1) \\ 0 \cdot 0 + 3 \cdot 2 & 0 \cdot (-1) + 3 \cdot (-1) \end{bmatrix} = \begin{bmatrix} -2 & -1 \\ 6 & -3 \end{bmatrix}$$

$$\bullet BA = \begin{bmatrix} 0 & -1 \\ 2 & -1 \end{bmatrix} \cdot \begin{bmatrix} 2 & -1 \\ 0 & 3 \end{bmatrix} = \begin{bmatrix} 0 \cdot 2 + (-1) \cdot 0 & 0 \cdot (-1) + (-1) \cdot 3 \\ 2 \cdot 2 + (-1) \cdot 0 & 2 \cdot (-1) + (-1) \cdot 3 \end{bmatrix} = \begin{bmatrix} 0 & -3 \\ 4 & -5 \end{bmatrix}$$

$$\bullet A^2 = \begin{bmatrix} 2 & -1 \\ 0 & 3 \end{bmatrix} \cdot \begin{bmatrix} 2 & -1 \\ 0 & 3 \end{bmatrix} = \begin{bmatrix} 2 \cdot 2 + (-1) \cdot 0 & 2 \cdot (-1) + (-1) \cdot 3 \\ 0 \cdot 2 + 3 \cdot 0 & 0 \cdot (-1) + 3 \cdot 3 \end{bmatrix} = \begin{bmatrix} 4 & -5 \\ 0 & 9 \end{bmatrix}$$

$$\bullet B^2 = \begin{bmatrix} 0 & -1 \\ 2 & -1 \end{bmatrix} \cdot \begin{bmatrix} 0 & -1 \\ 2 & -1 \end{bmatrix} = \begin{bmatrix} 0 \cdot 0 + (-1) \cdot 2 & 0 \cdot (-1) + (-1) \cdot (-1) \\ 2 \cdot 0 + (-1) \cdot 2 & 2 \cdot (-1) + (-1) \cdot (-1) \end{bmatrix} = \begin{bmatrix} -2 & 1 \\ -2 & -1 \end{bmatrix}$$

$$\bullet A^2 - B^2 = \begin{bmatrix} 4 & -5 \\ 0 & 9 \end{bmatrix} - \begin{bmatrix} -2 & 1 \\ -2 & -1 \end{bmatrix} = \begin{bmatrix} 6 & -6 \\ 2 & 10 \end{bmatrix}$$

$$\bullet A^T = \begin{bmatrix} 2 & 0 \\ -1 & 3 \end{bmatrix}$$

$$\bullet B^T = \begin{bmatrix} 0 & 2 \\ -1 & -1 \end{bmatrix}$$

$$\bullet A^T + B^T = \begin{bmatrix} 2 & 0 \\ -1 & 3 \end{bmatrix} + \begin{bmatrix} 0 & 2 \\ -1 & -1 \end{bmatrix} = \begin{bmatrix} 2 & 2 \\ -2 & 2 \end{bmatrix}$$

Wyznaczniki stopnia 1, 2 i 3

Definicja (rekurencyjna) *Wyznacznikiem* $\det A$ lub $|A|$ macierzy kwadratowej A n -go stopnia ($n = 1, 2, 3$) nazywa się liczba określona w wymieniony niżej sposób

- $n = 1: \det[a_{11}] = a_{11}$
- $n = 2: \begin{vmatrix} a_{11} & a_{12} \\ a_{21} & a_{22} \end{vmatrix} = a_{11} \cdot a_{22} - a_{12} \cdot a_{21}$



$$\bullet \quad n = 3: \begin{vmatrix} a_{11} & a_{12} & a_{13} \\ a_{21} & a_{22} & a_{23} \\ a_{31} & a_{32} & a_{33} \end{vmatrix} = a_{11} \cdot a_{22} \cdot a_{33} + a_{12} \cdot a_{23} \cdot a_{31} + a_{21} \cdot a_{32} \cdot a_{13} -$$

$$a_{31} \cdot a_{22} \cdot a_{13} - a_{21} \cdot a_{12} \cdot a_{33} - a_{32} \cdot a_{23} \cdot a_{11}$$

Wskazówka: przy obliczeniu wyznaczników 3-go stopnia również można wykorzystać metodę Sarrusa (są dwie wersje) bądź metodę trójkątów

- metoda Sarrusa

$$\begin{vmatrix} a_{11} & a_{12} & a_{13} \\ a_{21} & a_{22} & a_{23} \\ a_{31} & a_{32} & a_{33} \end{vmatrix} = \begin{vmatrix} a_{11} & a_{12} \\ a_{21} & a_{22} \\ a_{31} & a_{32} \end{vmatrix} =$$

$$= \begin{vmatrix} a_{11} & & & & \\ & a_{22} & & & \\ & & a_{33} & & \end{vmatrix} + \begin{vmatrix} & a_{12} & & & \\ & & a_{23} & & \\ & & & a_{31} & \end{vmatrix} + \begin{vmatrix} & & a_{13} & & \\ & a_{21} & & & \\ & & a_{32} & & \end{vmatrix} -$$

$$- \begin{vmatrix} & & & a_{13} & \\ & a_{22} & & & \\ & & a_{32} & & \end{vmatrix} - \begin{vmatrix} & & & & a_{11} \\ & & a_{23} & & \\ & a_{31} & & & \end{vmatrix} - \begin{vmatrix} & & & & \\ & & a_{23} & & \\ & a_{31} & & & \end{vmatrix} =$$

$$= a_{11}a_{22}a_{33} + a_{12}a_{23}a_{31} + a_{13}a_{21}a_{32} - a_{31}a_{22}a_{13} - a_{32}a_{23}a_{11} - a_{33}a_{21}a_{12}$$

lub

$$\begin{vmatrix} a_{11} & a_{12} & a_{13} \\ a_{21} & a_{22} & a_{23} \\ a_{31} & a_{32} & a_{33} \end{vmatrix} =$$

$$\begin{vmatrix} a_{11} & a_{12} & a_{13} \\ a_{21} & a_{22} & a_{23} \end{vmatrix} =$$

$$= \begin{vmatrix} a_{11} & & & & \\ & a_{22} & & & \\ & & a_{33} & & \end{vmatrix} + \begin{vmatrix} & a_{21} & & & \\ & & a_{32} & & \\ & & & a_{13} & \end{vmatrix} + \begin{vmatrix} & & a_{31} & & \\ & a_{12} & & & \\ & & a_{23} & & \end{vmatrix} -$$

$$- \begin{vmatrix} & & & a_{13} & \\ & a_{22} & & & \\ & & a_{32} & & \end{vmatrix} - \begin{vmatrix} & & & & a_{11} \\ & & a_{23} & & \\ & a_{31} & & & \end{vmatrix} - \begin{vmatrix} & & & & \\ & & a_{33} & & \\ & a_{21} & & & \end{vmatrix} =$$

$$= a_{11}a_{22}a_{33} + a_{21}a_{32}a_{13} + a_{31}a_{12}a_{23} - a_{31}a_{22}a_{13} - a_{11}a_{32}a_{23} - a_{21}a_{12}a_{33}$$

- metoda trójkątów

$$\begin{vmatrix} a_{11} & a_{12} & a_{13} \\ a_{21} & a_{22} & a_{23} \\ a_{31} & a_{32} & a_{33} \end{vmatrix} =$$

$$= \begin{vmatrix} a_{11} & & & & \\ & a_{22} & & & \\ & & a_{33} & & \end{vmatrix} + \begin{vmatrix} & a_{12} & & & \\ & & a_{23} & & \\ & & & a_{31} & \end{vmatrix} + \begin{vmatrix} & & a_{13} & & \\ & a_{21} & & & \\ & & a_{32} & & \end{vmatrix} -$$



$$- \begin{vmatrix} a_{13} & a_{12} & a_{11} \\ a_{22} & a_{21} & a_{23} \\ a_{31} & a_{33} & a_{32} \end{vmatrix} =$$

$$= a_{11}a_{22}a_{33} + a_{12}a_{23}a_{31} + a_{21}a_{32}a_{13} - a_{31}a_{22}a_{13} - a_{21}a_{12}a_{33} - a_{32}a_{23}a_{11}$$

Przykład Obliczyć wyznaczniki

$$\bullet \begin{vmatrix} 5 & -1 \\ 1 & 2 \end{vmatrix}$$

$$\bullet \begin{vmatrix} \sin 1 & -\cos 1 \\ \cos 1 & \sin 1 \end{vmatrix}$$

$$\bullet \begin{vmatrix} 1 & 1 & 2 \\ 3 & 5 & 8 \\ 13 & 21 & 34 \end{vmatrix}$$

Rozwiązanie

$$\bullet \begin{vmatrix} 5 & -1 \\ 1 & 2 \end{vmatrix} = 5 \cdot 2 - (-1) \cdot 1 = 8$$

$$\bullet \begin{vmatrix} \sin 1 & -\cos 1 \\ \cos 1 & \sin 1 \end{vmatrix} = \sin 1 \cdot \sin 1 - (-\cos 1) \cdot \cos 1 = \sin^2 1 + \cos^2 1 = 1$$

$$\bullet \begin{vmatrix} 1 & 1 & 2 \\ 3 & 5 & 8 \\ 13 & 21 & 34 \end{vmatrix} = 1 \cdot 5 \cdot 34 + 3 \cdot 21 \cdot 2 + 1 \cdot 8 \cdot 13 - 13 \cdot 5 \cdot 2 - 21 \cdot 8 \cdot 1 - 3 \cdot 1 \cdot 34 = 0$$

Przykład Rozwiązać równanie

$$\begin{vmatrix} 1 & 0 & 0 \\ a & 5-x & 0 \\ b & c & x+3 \end{vmatrix} = 0$$

Rozwiązanie

$$\begin{vmatrix} 1 & 0 & 0 \\ a & 5-x & 0 \\ b & c & x+3 \end{vmatrix} = 0$$

$$\begin{vmatrix} 1 & 0 & 0 \\ a & 5-x & 0 \\ b & c & x+3 \end{vmatrix} = (5-x)(x+3) + 0 + 0 - 0 - 0 - 0 = 0$$

$$(5-x)(x+3) = 0$$

$$x = -3 \vee x = 5$$

Dopełnienia algebraiczne, wyznaczniki ≥ 4 stopni, macierz odwrotna

Definicja *Minorem* (lub podwyznacznikiem) M_{ij} macierzy kwadratowej nazywa się wyznacznik macierzy po wykreśleniu i -go wiersza a j -ej kolumny. *Dopełnieniem algebraicznym* D_{ij} nazywa się liczba $D_{ij} = (-1)^{i+j} \cdot M_{ij}$.

Dopełnienie algebraiczne D_{ij} różni się od minoru M_{ij} tylko znakiem i to nie zawsze.

Macierz kwadratowa ma tyle minorów (dopełnień algebraicznych) ile ma elementów. Inaczej mówiąc, macierz n -go stopnia ma dokładnie n^2 minorów i n^2 dopełnień algebraicznych.



Definicja Niech A jest macierzą n -go stopnia. Wtedy macierz elementami której występują dopełnienia algebraiczne

$$[D_{ij}] = \begin{bmatrix} D_{11} & D_{12} & \dots & D_{1n} \\ D_{21} & D_{22} & \dots & D_{2n} \\ \dots & \dots & \dots & \dots \\ D_{n1} & D_{n2} & \dots & D_{nn} \end{bmatrix}$$

nazywa się *macierzą dopełnień algebraicznych*. Transponowanie macierzy dopełnień algebraicznych nazywa się *macierzą dołączoną* macierzy A

$$A^D = [D_{ij}]^T$$

Macierz dopełnień algebraicznych oraz macierz dołączona macierzy kwadratowej zawsze istnieją.

Definicja Macierz A^{-1} nazywa się *odwrotną macierzą* macierzy n -go stopnia A (o ile istnieje), jeśli

$$AA^{-1} = A^{-1}A = I_n$$

Twierdzenie Macierz odwrotna macierzy A istnieje wtedy i tylko wtedy, gdy wyznacznik macierzy A różni się od zera.

Innymi słowy, $\exists A^{-1} \Leftrightarrow \det A \neq 0$.

Metody na znalezienie macierzy odwrotnej (pod warunkiem, że ona istnieje):

- metoda macierzy dołączonej (lub dopełnień algebraicznych): macierz odwrotną można znaleźć ze wzoru

$$A^{-1} = \frac{1}{\det A} \cdot A^D = \frac{1}{\det A} [D_{ij}]^T$$

- metoda macierzy klatkowej (lub eliminacji Gaussa): (1) układamy macierz klatkową $[A \mid I_n]$, dopisując obok macierzy n -go stopnia A macierz jednostkową I_n ; (2) za pomocą operacji elementarnych na wierszach macierzy klatkowej (tzw. operacji Gaussa), w lewej części robimy macierz jednostkową I_n (jeżeli macierz odwrotna A^{-1} istnieje, to odpowiednia lista operacji elementarnych zawsze istnieje) $[A \mid I_n] \sim \dots \sim [I_n \mid \dots]$; (3) wtedy po prawej stronie otrzymanej macierzy klatkowej będzie macierz odwrotna, tzn.

$$[A \mid I_n] \sim \text{operacje elementarne na wierszach} \sim [I_n \mid A^{-1}]$$

(przy wykorzystaniu metody macierzy klatkowej pamiętamy, że operacje elementarne wykonujemy tylko i wyłącznie na wierszach oraz na obu częściach macierzy klatkowej).



Przykład Dla macierzy $A = \begin{bmatrix} 2 & 3 \\ -3 & 4 \end{bmatrix}$ znaleźć dopełnienia algebraiczne $D_{ij}, i, j = \overline{1,2}$ oraz macierz odwrotną A^{-1} (gdy ona istnieje) oraz sprawdzić czy znaleziona macierz będzie odwrotną do macierzy A .

Rozwiązanie

$$A = \begin{bmatrix} 2 & 3 \\ -3 & 4 \end{bmatrix}$$

$$M_{11} = 4$$

$$D_{11} = (-1)^{1+1}M_{11} = 4$$

$$M_{12} = -3$$

$$D_{12} = (-1)^{1+2}M_{12} = 3$$

$$M_{21} = 3$$

$$D_{21} = (-1)^{2+1}M_{21} = -3$$

$$M_{22} = 2$$

$$D_{22} = (-1)^{2+2}M_{22} = 2$$

$$\det A = \begin{vmatrix} 2 & 3 \\ -3 & 4 \end{vmatrix} = 17 \neq 0 \Rightarrow \exists A^{-1}$$

$$[D_{ij}] = \begin{bmatrix} 4 & 3 \\ -3 & 2 \end{bmatrix}$$

$$A^D = \begin{bmatrix} 4 & 3 \\ -3 & 2 \end{bmatrix}^T = \begin{bmatrix} 4 & -3 \\ 3 & 2 \end{bmatrix}$$

$$A^{-1} = \frac{1}{17} \begin{bmatrix} 4 & -3 \\ 3 & 2 \end{bmatrix} = \begin{bmatrix} \frac{4}{17} & -\frac{3}{17} \\ \frac{3}{17} & \frac{2}{17} \end{bmatrix}$$

Sprawdźmy, czy macierz $A^{-1} = \frac{1}{17} \begin{bmatrix} 4 & -3 \\ 3 & 2 \end{bmatrix}$ naprawdę będzie odwrotną $A = \begin{bmatrix} 2 & 3 \\ -3 & 4 \end{bmatrix}$:

$$\begin{aligned} A \cdot A^{-1} &= \begin{bmatrix} 2 & 3 \\ -3 & 4 \end{bmatrix} \cdot \frac{1}{17} \begin{bmatrix} 4 & -3 \\ 3 & 2 \end{bmatrix} = \frac{1}{17} \begin{bmatrix} 2 \cdot 4 + 3 \cdot 3 & 2 \cdot (-3) + 3 \cdot 2 \\ (-3) \cdot 4 + 4 \cdot 3 & (-3) \cdot (-3) + 4 \cdot 2 \end{bmatrix} \\ &= \frac{1}{17} \begin{bmatrix} 17 & 0 \\ 0 & 17 \end{bmatrix} = \begin{bmatrix} 1 & 0 \\ 0 & 1 \end{bmatrix} \end{aligned}$$

$$\begin{aligned} A^{-1} \cdot A &= \frac{1}{17} \begin{bmatrix} 4 & -3 \\ 3 & 2 \end{bmatrix} \cdot \begin{bmatrix} 2 & 3 \\ -3 & 4 \end{bmatrix} = \frac{1}{17} \begin{bmatrix} 4 \cdot 2 + (-3) \cdot (-3) & 4 \cdot 3 + (-3) \cdot 4 \\ 3 \cdot 2 + 2 \cdot (-3) & 3 \cdot 3 + 2 \cdot 4 \end{bmatrix} \\ &= \frac{1}{17} \begin{bmatrix} 17 & 0 \\ 0 & 17 \end{bmatrix} = \begin{bmatrix} 1 & 0 \\ 0 & 1 \end{bmatrix} \end{aligned}$$

Tak, macierz A^{-1} jest odwrotną macierzy A .

Przykład Dla macierzy $A = \begin{bmatrix} 3 & 4 & 4 \\ 1 & -1 & 0 \\ 0 & 2 & 1 \end{bmatrix}$ znaleźć dopełnienia algebraiczne $D_{ij}, i, j = \overline{1,3}$ oraz macierz odwrotną A^{-1} (o ile ona istnieje).

Rozwiązanie

$$M_{11} = \begin{vmatrix} -1 & 0 \\ 2 & 1 \end{vmatrix} = -1$$

$$M_{12} = \begin{vmatrix} 1 & 0 \\ 0 & 1 \end{vmatrix} = 1$$

$$M_{13} = \begin{vmatrix} 1 & -1 \\ 0 & 2 \end{vmatrix} = 2$$



$$M_{21} = \begin{vmatrix} 4 & 4 \\ 2 & 1 \end{vmatrix} = -4$$

$$M_{23} = \begin{vmatrix} 3 & 4 \\ 0 & 2 \end{vmatrix} = 6$$

$$M_{32} = \begin{vmatrix} 3 & 4 \\ 1 & 0 \end{vmatrix} = -4$$

$$M_{22} = \begin{vmatrix} 3 & 4 \\ 0 & 1 \end{vmatrix} = 3$$

$$M_{31} = \begin{vmatrix} 4 & 4 \\ -1 & 0 \end{vmatrix} = 4$$

$$M_{33} = \begin{vmatrix} 3 & 4 \\ 1 & -1 \end{vmatrix} = -7$$

$$D_{11} = (-1)^{1+1}M_{11} = -1$$

$$D_{21} = (-1)^{2+1}M_{21} = 4$$

$$D_{31} = (-1)^{3+1}M_{31} = 4$$

$$D_{12} = (-1)^{1+2}M_{12} = -1$$

$$D_{22} = (-1)^{2+2}M_{22} = 3$$

$$D_{32} = (-1)^{3+2}M_{32} = 4$$

$$D_{13} = (-1)^{1+3}M_{13} = 2$$

$$D_{23} = (-1)^{2+3}M_{23} = -6$$

$$D_{33} = (-1)^{3+3}M_{33} = -7$$

$$\det A = \begin{vmatrix} 3 & 4 & 4 \\ 1 & -1 & 0 \\ 0 & 2 & 1 \end{vmatrix} = 1 \neq 0 \Rightarrow \exists A^{-1}$$

$$[D_{ij}] = \begin{bmatrix} -1 & -1 & 2 \\ 4 & 3 & -6 \\ 4 & 4 & -7 \end{bmatrix}$$

$$A^D = \begin{bmatrix} -1 & -1 & 2 \\ 4 & 3 & -6 \\ 4 & 4 & -7 \end{bmatrix}^T = \begin{bmatrix} -1 & 4 & 4 \\ -1 & 3 & 4 \\ 2 & -6 & -7 \end{bmatrix}$$

$$A^{-1} = \frac{1}{1} \begin{bmatrix} -1 & 4 & 4 \\ -1 & 3 & 4 \\ 2 & -6 & -7 \end{bmatrix} = \begin{bmatrix} -1 & 4 & 4 \\ -1 & 3 & 4 \\ 2 & -6 & -7 \end{bmatrix}$$

Rozwinięcie Laplace'a: obliczyć wyznaczniki 4-go i wzwyż stopni można za pomocą rozwinięcia Laplace'a, w którym: (1) wybieramy wiersz bądź kolumnę według której będziemy rozwijać wyznacznik (wiersz lub kolumna, gdzie znajduje się największa ilość liczb zerowych); (2) rozwijamy według wybranego wiersza lub kolumny jak niżej

- rozwinięcie według i -go wiersza

$$\begin{vmatrix} a_{11} & a_{12} & a_{13} & a_{14} & \dots & a_{1n} \\ a_{21} & a_{22} & a_{23} & a_{24} & \dots & a_{2n} \\ \dots & \dots & \dots & \dots & \dots & \dots \\ \mathbf{a_{i1}} & \mathbf{a_{i2}} & \mathbf{a_{i3}} & \mathbf{a_{i4}} & \dots & \mathbf{a_{in}} \\ \dots & \dots & \dots & \dots & \dots & \dots \\ a_{n1} & a_{n2} & a_{n3} & a_{n4} & \dots & a_{nn} \end{vmatrix}$$

$$= a_{i1}D_{i1} + a_{i2}D_{i2} + a_{i3}D_{i3} + a_{i4}D_{i4} + \dots + a_{in}D_{in}$$

- rozwinięcie według j -ej kolumny

$$\begin{vmatrix} a_{11} & a_{12} & \dots & \mathbf{a_{1j}} & \dots & a_{1n} \\ a_{21} & a_{22} & \dots & \mathbf{a_{2j}} & \dots & a_{2n} \\ a_{31} & a_{32} & \dots & \mathbf{a_{3j}} & \dots & a_{3n} \\ a_{41} & a_{42} & \dots & \mathbf{a_{4j}} & \dots & a_{4n} \\ \dots & \dots & \dots & \dots & \dots & \dots \\ a_{n1} & a_{n2} & \dots & \mathbf{a_{nj}} & \dots & a_{nn} \end{vmatrix}$$

$$= a_{1j}D_{1j} + a_{2j}D_{2j} + a_{3j}D_{3j} + a_{4j}D_{4j} + \dots + a_{nj}D_{nj}$$

Rozwinięcie Laplace'a również ma miejsce dla wyznaczników 2-go i 3-go stopni.

Przykład Obliczyć wyznacznik 4-go stopnia



$$\begin{vmatrix} 1 & 0 & -2 & 0 \\ 0 & 2 & 0 & 1 \\ 1 & 0 & 0 & 1 \\ -3 & 1 & 3 & 5 \end{vmatrix}$$

Rozwiązanie

$$\begin{vmatrix} 1 & 0 & -2 & 0 \\ 0 & 2 & 0 & 1 \\ 1 & 0 & 0 & 1 \\ -3 & 1 & 3 & 5 \end{vmatrix} = [\text{rozwijamy według 2-go wiersza}] = 0 \cdot D_{21} + 2 \cdot D_{22} + 0 \cdot$$

$$D_{23} + 1 \cdot D_{24} = 2D_{22} + D_{24} = \llbracket D_{ij} = (-1)^{i+j} M_{ij} \rrbracket = 2(-1)^{2+2} M_{22} + (-1)^{2+4} M_{24} =$$

$$2M_{22} + M_{24} \equiv,$$

$$M_{22} = \begin{vmatrix} 1 & -2 & 0 \\ 1 & 0 & 1 \\ -3 & 3 & 5 \end{vmatrix} = 13$$

$$\equiv 2 \cdot 13 + (-2) = 24$$

$$M_{24} = \begin{vmatrix} 1 & 0 & -2 \\ 1 & 0 & 0 \\ -3 & 1 & 3 \end{vmatrix} = -2$$

Równania macierzowe**Przykład** Rozwiązać równanie macierzowe $XA = 3B + C^T$, gdzie

$$A = \begin{bmatrix} 3 & 4 & 4 \\ 1 & -1 & 0 \\ 0 & 2 & 1 \end{bmatrix}, B = [1 \quad -1 \quad 3], C = \begin{bmatrix} 1 \\ 0 \\ 2 \end{bmatrix}$$

Rozwiązanie

$$XA = 3B + C^T \mid \cdot A^{-1} \text{ po lewej stronie}$$

$$XAA^{-1} = (3B + C^T)A^{-1}$$

$$XI_3 = (3B + C^T)A^{-1}$$

$$X = (3B + C^T)A^{-1}$$

$$B = [1 \quad -1 \quad 3] \Rightarrow 3B = [3 \quad -3 \quad 9]$$

$$C = \begin{bmatrix} 1 \\ 0 \\ 2 \end{bmatrix} \Rightarrow C^T = [1 \quad 0 \quad 2]$$

$$3B + C^T = [3 \quad -3 \quad 9] + [1 \quad 0 \quad 2] = [4 \quad -3 \quad 11]$$

$$A = \begin{bmatrix} 3 & 4 & 4 \\ 1 & -1 & 0 \\ 0 & 2 & 1 \end{bmatrix} \Rightarrow A^{-1} = \begin{bmatrix} -1 & 4 & 4 \\ -1 & 3 & 4 \\ 2 & -6 & -7 \end{bmatrix}$$

$$X = (3B + C^T)A^{-1} = [4 \quad -3 \quad 11] \cdot \begin{bmatrix} -1 & 4 & 4 \\ -1 & 3 & 4 \\ 2 & -6 & -7 \end{bmatrix} = [21 \quad -59 \quad -73]$$

Układy równań liniowych

Definicja Układem n równań z n niewiadomymi $(x_1; x_2; \dots; x_n)$ gdy $n = \overline{1,3}$ oraz w ogólnej postaci nazywają się poniższe układy równań

- $n = 1$: $a_{11}x = b_1$



- $n = 2$:
$$\begin{cases} a_{11}x_1 + a_{12}x_2 = b_1 \\ a_{21}x_1 + a_{22}x_2 = b_2 \end{cases}$$
- $n = 3$:
$$\begin{cases} a_{11}x_1 + a_{12}x_2 + a_{13}x_3 = b_1 \\ a_{21}x_1 + a_{22}x_2 + a_{23}x_3 = b_2 \\ a_{31}x_1 + a_{32}x_2 + a_{33}x_3 = b_3 \end{cases}$$
- n :
$$\begin{cases} a_{11}x_1 + a_{12}x_2 + \dots + a_{1n}x_n = b_1 \\ a_{21}x_1 + a_{22}x_2 + \dots + a_{2n}x_n = b_2 \\ \dots \\ a_{n1}x_1 + a_{n2}x_2 + \dots + a_{nn}x_n = b_n \end{cases}$$

Liczby $a_{ij} \in \mathbb{R}$ zwane są *współczynnikami* układu równań, $b_i \in \mathbb{R}$ – *wyrazami wolnymi*.

Wektor liczb $(x_1^{(0)}; x_2^{(0)}; \dots; x_n^{(0)})$ nazywa się *rozwiązaniem układu równań* (z n niewiadomymi), jeżeli podstawiając wartości $x_i^{(0)}, i = \overline{1, n}$ zamiast odpowiednich zmiennych $x_i, i = \overline{1, n}$ otrzymamy prawidłowe równości. Układ równań z zależności od ilości rozwiązań nazywa się

- *oznaczony*, gdy ma dokładnie jedno rozwiązanie;
- *nieoznaczony*, gdy ma nieskończenie wiele rozwiązań;
- *sprzeczny*, gdy nie ma rozwiązań.

Różne postaci układów równań liniowych (z trzema i z n niewiadomymi):

$$\begin{cases} a_{11}x_1 + a_{12}x_2 + a_{13}x_3 = b_1 \\ a_{21}x_1 + a_{22}x_2 + a_{23}x_3 = b_2 \\ a_{31}x_1 + a_{32}x_2 + a_{33}x_3 = b_3 \end{cases}$$

$$\left[\begin{array}{ccc|c} a_{11} & a_{12} & a_{13} & b_1 \\ a_{21} & a_{22} & a_{23} & b_2 \\ a_{31} & a_{32} & a_{33} & b_3 \end{array} \right]$$

$$\begin{bmatrix} a_{11} & a_{12} & a_{13} \\ a_{21} & a_{22} & a_{23} \\ a_{31} & a_{32} & a_{33} \end{bmatrix} \begin{bmatrix} x_1 \\ x_2 \\ x_3 \end{bmatrix} = \begin{bmatrix} b_1 \\ b_2 \\ b_3 \end{bmatrix}$$

$$\begin{cases} a_{11}x_1 + a_{12}x_2 + \dots + a_{1n}x_n = b_1 \\ a_{21}x_1 + a_{22}x_2 + \dots + a_{2n}x_n = b_2 \\ \dots \\ a_{n1}x_1 + a_{n2}x_2 + \dots + a_{nn}x_n = b_n \end{cases}$$

$$\left[\begin{array}{cccc|c} a_{11} & a_{12} & \dots & a_{1n} & b_1 \\ a_{21} & a_{22} & \dots & a_{2n} & b_2 \\ \dots & \dots & \dots & \dots & \dots \\ a_{n1} & a_{n2} & \dots & a_{nn} & b_n \end{array} \right]$$

$$\begin{bmatrix} a_{11} & a_{12} & \dots & a_{1n} \\ a_{21} & a_{22} & \dots & a_{2n} \\ \dots & \dots & \dots & \dots \\ a_{n1} & a_{n2} & \dots & a_{nn} \end{bmatrix} \begin{bmatrix} x_1 \\ x_2 \\ \dots \\ x_n \end{bmatrix} = \begin{bmatrix} b_1 \\ b_2 \\ \dots \\ b_n \end{bmatrix}$$

gdzie macierz $\begin{bmatrix} a_{11} & a_{12} & \dots & a_{1n} \\ a_{21} & a_{22} & \dots & a_{2n} \\ \dots & \dots & \dots & \dots \\ a_{n1} & a_{n2} & \dots & a_{nn} \end{bmatrix}$ nazywa się *macierzą współczynników*,

$\begin{bmatrix} a_{11} & a_{12} & \dots & a_{1n} & b_1 \\ a_{21} & a_{22} & \dots & a_{2n} & b_2 \\ \dots & \dots & \dots & \dots & \dots \\ a_{n1} & a_{n2} & \dots & a_{nn} & b_n \end{bmatrix}$ – *macierzą rozszerzoną* układu równań liniowych, wektor $\begin{bmatrix} x_1 \\ x_2 \\ \dots \\ x_n \end{bmatrix}$



nazywa się *wektorem zmiennych*, wektor $\begin{bmatrix} b_1 \\ b_2 \\ \dots \\ b_n \end{bmatrix}$ nazywa się *wektorem wyrazów wolnych*, a

kolumna $\begin{bmatrix} \dots & \begin{bmatrix} b_1 \\ b_2 \\ \dots \\ b_n \end{bmatrix} \end{bmatrix}$ – *kolumną wyrazów wolnych*.

Dla rozwiązywania układu równań liniowych zwykle wykorzystujemy jedną z trzech poniższych metod:

- metoda Cramera (dla układów oznaczonych bądź sprzecznych oraz tylko i wyłącznie gdy liczba równań i liczba niewiadomych są równe sobie);
- metoda macierzy odwrotnej (dla układów równań oznaczonych oraz tylko i wyłącznie gdy liczba równań i liczba niewiadomych są równe sobie);
- metoda eliminacji Gaussa (dla wszystkich układów równań, w tym dla takich, w których liczba równań i liczba niewiadomych różnią się).

Metoda Cramera: głównie polega na obliczaniu $(n + 1)$ wyznaczników n -go stopnia w przypadku układu równań z n niewiadomymi. Dla układu równań z trzema niewiadomymi metoda Cramera wygląda następująco:

$$\begin{cases} a_{11}x_1 + a_{12}x_2 + a_{13}x_3 = b_1 \\ a_{21}x_1 + a_{22}x_2 + a_{23}x_3 = b_2 \\ a_{31}x_1 + a_{32}x_2 + a_{33}x_3 = b_3 \end{cases}$$

- przepisujemy układ za pomocą macierzy rozszerzonej:

$$\left[\begin{array}{ccc|c} a_{11} & a_{12} & a_{13} & b_1 \\ a_{21} & a_{22} & a_{23} & b_2 \\ a_{31} & a_{32} & a_{33} & b_3 \end{array} \right]$$

- obliczamy wyznaczniki:

$$\Delta = \begin{vmatrix} a_{11} & a_{12} & a_{13} \\ a_{21} & a_{22} & a_{23} \\ a_{31} & a_{32} & a_{33} \end{vmatrix}$$

$$\Delta_1 = \begin{vmatrix} b_1 & a_{12} & a_{13} \\ b_2 & a_{22} & a_{23} \\ b_3 & a_{32} & a_{33} \end{vmatrix}$$

$$\Delta_2 = \begin{vmatrix} a_{11} & b_1 & a_{13} \\ a_{21} & b_2 & a_{23} \\ a_{31} & b_3 & a_{33} \end{vmatrix}$$

$$\Delta_3 = \begin{vmatrix} a_{11} & a_{12} & b_1 \\ a_{21} & a_{22} & b_2 \\ a_{31} & a_{32} & b_3 \end{vmatrix}$$

- w zależności od wartości wyznaczników szukamy rozwiązanie układu równań (o ile ono istnieje):
 - $\Delta \neq 0 \Rightarrow$ układ równań oznaczony, a rozwiązaniem występują liczby



$$x_1 = \frac{\Delta_1}{\Delta}, x_2 = \frac{\Delta_2}{\Delta}, x_3 = \frac{\Delta_3}{\Delta}$$

- $\Delta = 0$ oraz choćby jeden wyznacznik $\Delta_i \neq 0$ ($\exists i \in \{1; 2; 3\}$) $\Rightarrow$ układ równań sprzeczny;
- $\Delta = 0$ oraz $\Delta_1 = \Delta_2 = \Delta_3 = 0 \Rightarrow$ układ równań nieoznaczony, tzn. ma nieskończenie wiele rozwiązań, ale odpowiedzi na pytania jakie są te rozwiązania metoda Cramera nie daje.

W przypadku zwiększenia bądź zmniejszenia liczby niewiadomych wraz z odpowiednią zmianą liczby równań sama metoda nie ulegnie zmian, a różnica polega w tym, że zwiększy się tylko liczba wyznaczników do obliczania i ich stopni.

Metoda macierzy odwrotnej: głównie polega na poszukiwaniu macierzy odwrotnej (o ile ona istnieje). Dla układu równań z trzema niewiadomymi metoda macierzy odwrotnej wygląda następująco:

$$\begin{cases} a_{11}x_1 + a_{12}x_2 + a_{13}x_3 = b_1 \\ a_{21}x_1 + a_{22}x_2 + a_{23}x_3 = b_2 \\ a_{31}x_1 + a_{32}x_2 + a_{33}x_3 = b_3 \end{cases}$$

- przepisujemy układ w postaci macierzowej:

$$\begin{bmatrix} a_{11} & a_{12} & a_{13} \\ a_{21} & a_{22} & a_{23} \\ a_{31} & a_{32} & a_{33} \end{bmatrix} \begin{bmatrix} x_1 \\ x_2 \\ x_3 \end{bmatrix} = \begin{bmatrix} b_1 \\ b_2 \\ b_3 \end{bmatrix}$$

- układamy równanie macierzowe:

$$A = \begin{bmatrix} a_{11} & a_{12} & a_{13} \\ a_{21} & a_{22} & a_{23} \\ a_{31} & a_{32} & a_{33} \end{bmatrix}, X = \begin{bmatrix} x_1 \\ x_2 \\ x_3 \end{bmatrix}, B = \begin{bmatrix} b_1 \\ b_2 \\ b_3 \end{bmatrix}$$

$$AX = B$$

- rozwiązujemy równanie macierzowe, z czego otrzymujemy rozwiązanie układu równań liniowych:

$$AX = B \Rightarrow X = A^{-1}B \Rightarrow \begin{bmatrix} x_1 \\ x_2 \\ x_3 \end{bmatrix} = X = A^{-1}B$$

Definicja Układy równań nazywają się *równoważnymi*, jeśli zbiory ich rozwiązań są równe.

Twierdzenie Przy wykonaniu poniższych operacji na równaniach, ciągle otrzymujemy równoważne układy równań:

- tzw. operacje elementarne:
 - pomnożenie równania przez liczbę niezerową,



- zamiana dwóch równań miejscami,
- dodawanie do jednego równania innego pomnożonego przez liczbę;
- wykreślenie równania $0 = 0$.

Wymienione wyżej operacje elementarne ekwiwalentne operacjom elementarnym na wierszach macierzy rozszerzonej układu równań.

Metoda eliminacji Gaussa: głównie polega na przekształcaniu macierzy rozszerzonej układu równań za pomocą operacji elementarnych. Metoda Gaussa działa na wszystkich (względem liczby rozwiązań) układach równań, w tym na takich, w których liczba równań i liczba niewiadomych różnią się. Schemat metody eliminacji Gaussa (krótsza wersja):

- przepisujemy układ za pomocą macierzy rozszerzonej;
- wykorzystując operacje elementarne tylko i wyłącznie na wierszach macierzy rozszerzonej układu równań oraz wykreślając zerowe wiersze, na miejscu macierzy współczynników robimy macierz schodkową (tzn. macierz, w której pod przekątną główną mamy zera);
- z otrzymanej macierzy schodkowej wypisujemy układ równań, który rozwiązujemy od dołu (tzn. zaczynając od ostatniego równania).

Pod minorem (bez zaznaczenia indeksów) będziemy również rozumieć wyznacznik macierzy po wykreśleniu jednego lub więcej wierszy lub/i kolumn. Taki minor jest określony zarówno dla macierzy kwadratowej, jak i takiej, która nie jest kwadratowa.

Definicja *Rzędem* macierzy A nazywa się liczba $r = \text{rank}(A)$, taka że z macierzy A można wybrać (wykreślając wiersze i kolumny) choćby jeden niezerowy minor stopnia r , wtedy jak wszystkie minory stopni $> r$ będą zerowe.

Twierdzenie Operacje elementarne na wierszach lub na kolumnach macierzy oraz wykreślenie zerowego wiersza lub kolumny nie zmieniają rzędu macierzy.

Rząd macierzy można szukać na dwa sposoby:

- najpierw obliczamy minory największego możliwego stopnia póki aż otrzymamy liczbę niezerową; jeśli taki minor istnieje, to rząd macierzy równa się stopniowi tego minoru; jeżeli taki minor nie istnieje, powtarzamy procedurę dla minorów stopni o 1 mniejszych dopóki nie znajdziemy niezerowy minor, stopień którego będzie równać się rzędowi macierzy;
- za pomocą operacji elementarnych na wierszach bądź na kolumnach macierzy przekształcić ją w macierz schodkową, wtedy liczba niezerowych wierszy będzie równać się rzędowi macierzy.



Twierdzenie (Kroneckera-Capellego) Układ równań ma rozwiązania wtedy i tylko wtedy, gdy rzędy macierzy współczynników i macierzy rozszerzonej układu równań są równe sobie.

W praktyce powyższe twierdzenie oznacza, że jeśli po wykonaniu operacji elementarnych na wierszach macierzy rozszerzonej otrzymamy wiersz, w którym na miejscu współczynników będą liczby zerowe i odpowiedni wyraz wolny również wynosi 0, to układ równań ma rozwiązania, a jeśli w wierszu współczynnika będą zerowe, a wyraz wolny – nie, to układ równań będzie sprzeczny.

Przykład Znaleźć ilości rozwiązań poniższych równań liniowych

$$\bullet \begin{cases} 5x - 3y = 2 \\ x + y = 2 \end{cases}$$

$$\bullet \begin{cases} 3x + 3y = 6 \\ x + y = 2 \end{cases}$$

$$\bullet \begin{cases} 3x + 3y = 3 \\ x + y = 2 \end{cases}$$

Rozwiązanie

Dla rozwiązania wykorzystamy metodę Cramera

$$\bullet \begin{cases} 5x - 3y = 2 \\ x + y = 2 \end{cases}$$

$$\begin{bmatrix} 5 & -3 & 2 \\ 1 & 1 & 2 \end{bmatrix}$$

$$\Delta = \begin{vmatrix} 5 & -3 \\ 1 & 1 \end{vmatrix} = 5 - (-3) = 8 \neq 0 \Rightarrow \text{układ równań ma dokładnie jedno rozwiązanie}$$

$$\bullet \begin{cases} 3x + 3y = 6 \\ x + y = 2 \end{cases}$$

$$\begin{bmatrix} 3 & 3 & 6 \\ 1 & 1 & 2 \end{bmatrix}$$

$$\Delta = \begin{vmatrix} 3 & 3 \\ 1 & 1 \end{vmatrix} = 3 - 3 = 0 \Rightarrow \text{układ równań ma nieskończenie wiele rozwiązań}$$

$$\Delta_x = \begin{vmatrix} 6 & 3 \\ 2 & 1 \end{vmatrix} = 6 - 6 = 0$$

$$\Delta_y = \begin{vmatrix} 3 & 6 \\ 1 & 2 \end{vmatrix} = 6 - 6 = 0$$

$\Rightarrow$ układ równań nie ma rozwiązań

$$\bullet \begin{cases} 3x + 3y = 3 \\ x + y = 2 \end{cases}$$

$$\begin{bmatrix} 3 & 3 & 3 \\ 1 & 1 & 2 \end{bmatrix}$$

$$\Delta = \begin{vmatrix} 3 & 3 \\ 1 & 1 \end{vmatrix} = 6 - 6 = 0 \Rightarrow \text{układ równań nie jest oznaczony}$$

$$\Delta_x = \begin{vmatrix} 3 & 3 \\ 2 & 1 \end{vmatrix} = 3 - 6 = -3 \neq 0 \Rightarrow \text{układ równań nieoznaczony}$$

Przykład Rozwiązać oznaczony układ równań liniowych



$$\begin{cases} x_1 - 2x_2 + x_3 = 2 \\ 2x_2 - x_3 = -1 \\ x_1 + x_2 + x_3 = 2 \end{cases}$$

Rozwiązanie

Dla rozwiązania wykorzystamy metodę Cramera

$$\begin{cases} x_1 - 2x_2 + x_3 = 2 \\ 2x_2 - x_3 = -1 \\ x_1 + x_2 + x_3 = 2 \end{cases}$$

$$\begin{cases} x_1 - 2x_2 + x_3 = 2 \\ 2x_2 + 0x_3 - x_3 = -1 \\ x_1 + x_2 + x_3 = 2 \end{cases}$$

$$\left[\begin{array}{ccc|c} 1 & -2 & 1 & 2 \\ 0 & 2 & -1 & -1 \\ 1 & 1 & 1 & 2 \end{array} \right]$$

$$\Delta = \begin{vmatrix} 1 & -2 & 1 \\ 0 & 2 & -1 \\ 1 & 1 & 1 \end{vmatrix} = 3 \neq 0$$

$$\Delta_1 = \begin{vmatrix} 2 & -2 & 1 \\ -1 & 2 & -1 \\ 2 & 1 & 1 \end{vmatrix} = 3$$

$$x_1 = \frac{\Delta_1}{\Delta} = \frac{3}{3} = 1$$

$$\Delta_2 = \begin{vmatrix} 1 & 2 & 1 \\ 0 & -1 & -1 \\ 1 & 2 & 1 \end{vmatrix} = 0$$

$$x_2 = \frac{\Delta_2}{\Delta} = \frac{0}{3} = 0$$

$$\Delta_3 = \begin{vmatrix} 1 & -2 & 2 \\ 0 & 2 & -1 \\ 1 & 1 & 2 \end{vmatrix} = 3$$

$$x_3 = \frac{\Delta_3}{\Delta} = \frac{3}{3} = 1$$

Odpowiedź (1; 0; 1).

Przykład Sprawdzić czy jest poniższy układ równań oznaczony. Jeśli tak, to rozwiązać

$$\begin{cases} 3x_1 + 4x_2 + 4x_3 = 0 \\ x_1 - x_2 = -1 \\ 2x_2 + x_3 = 1 \end{cases}$$

Rozwiązanie

Dla rozwiązania wykorzystamy metodę macierzy odwrotnej

$$\begin{cases} 3x_1 + 4x_2 + 4x_3 = 0 \\ x_1 - x_2 = -1 \\ 2x_2 + x_3 = 1 \end{cases} \Leftrightarrow \begin{bmatrix} 3 & 4 & 4 \\ 1 & -1 & 0 \\ 0 & 2 & 1 \end{bmatrix} \begin{bmatrix} x_1 \\ x_2 \\ x_3 \end{bmatrix} = \begin{bmatrix} 0 \\ -1 \\ 1 \end{bmatrix} \Leftrightarrow AX = B, \text{ gdzie}$$



$$A = \begin{bmatrix} 3 & 4 & 4 \\ 1 & -1 & 0 \\ 0 & 2 & 1 \end{bmatrix}, X = \begin{bmatrix} x_1 \\ x_2 \\ x_3 \end{bmatrix}, B = \begin{bmatrix} 0 \\ -1 \\ 1 \end{bmatrix}$$

$$AX = B \Rightarrow X = A^{-1}B$$

$$A = \begin{bmatrix} 3 & 4 & 4 \\ 1 & -1 & 0 \\ 0 & 2 & 1 \end{bmatrix} \Rightarrow A^{-1} = \begin{bmatrix} -1 & 4 & 4 \\ -1 & 3 & 4 \\ 2 & -6 & -7 \end{bmatrix} \Rightarrow$$

$$X = A^{-1}B = \begin{bmatrix} -1 & 4 & 4 \\ -1 & 3 & 4 \\ 2 & -6 & -7 \end{bmatrix} \cdot \begin{bmatrix} 0 \\ -1 \\ 1 \end{bmatrix} = \begin{bmatrix} 0 \\ 1 \\ -1 \end{bmatrix} \Rightarrow$$

$$\begin{bmatrix} x_1 \\ x_2 \\ x_3 \end{bmatrix} = \begin{bmatrix} 0 \\ 1 \\ -1 \end{bmatrix}$$

Odpowiedź (0; 1; -1).

Przykład Rozwiązać oznaczony układ równań liniowych

$$\begin{cases} x_1 + 2x_2 + 3x_3 + 4x_4 + 5x_5 = 2 \\ 2x_1 + 3x_2 + 7x_3 + 10x_4 + 13x_5 = 12 \\ 3x_1 + 5x_2 + 11x_3 + 16x_4 + 21x_5 = 17 \\ 2x_1 - 7x_2 + 7x_3 + 7x_4 + 2x_5 = 57 \\ x_1 + 4x_2 + 5x_3 + 3x_4 + 10x_5 = 7 \end{cases}$$

Rozwiązanie

Ponieważ liczba zmiennych i równań wynosi 5, to aby nie obliczać sześciu wyznaczników 5-go stopnia wykorzystamy metodę eliminacji Gaussa

$$\begin{cases} x_1 + 2x_2 + 3x_3 + 4x_4 + 5x_5 = 2 \\ 2x_1 + 3x_2 + 7x_3 + 10x_4 + 13x_5 = 12 \\ 3x_1 + 5x_2 + 11x_3 + 16x_4 + 21x_5 = 17 \\ 2x_1 - 7x_2 + 7x_3 + 7x_4 + 2x_5 = 57 \\ x_1 + 4x_2 + 5x_3 + 3x_4 + 10x_5 = 7 \end{cases}$$

$$\begin{aligned} & \left[\begin{array}{ccccc|c} 1 & 2 & 3 & 4 & 5 & 2 \\ 2 & 3 & 7 & 10 & 13 & 12 \\ 3 & 5 & 11 & 16 & 21 & 17 \\ 2 & -7 & 7 & 7 & 2 & 57 \\ 1 & 4 & 5 & 3 & 10 & 7 \end{array} \right] \begin{array}{l} w_1(-2) + w_2 \\ w_1(-3) + w_3 \\ \sim \\ w_1(-2) + w_4 \\ w_1(-1) + w_5 \end{array} \left[\begin{array}{ccccc|c} 1 & 2 & 3 & 4 & 5 & 2 \\ 0 & -1 & 1 & 2 & 3 & 8 \\ 0 & -1 & 2 & 4 & 6 & 11 \\ 0 & -11 & 1 & -1 & -8 & 53 \\ 0 & 2 & 2 & -1 & 5 & 5 \end{array} \right] \begin{array}{l} w_2(-1) \\ \sim \\ \sim \end{array} \\ & \left[\begin{array}{ccccc|c} 1 & 2 & 3 & 4 & 5 & 2 \\ 0 & 1 & -1 & -2 & -3 & -8 \\ 0 & -1 & 2 & 4 & 6 & 11 \\ 0 & -11 & 1 & -1 & -8 & 53 \\ 0 & 2 & 2 & -1 & 5 & 5 \end{array} \right] \begin{array}{l} w_2 \cdot 1 + w_3 \\ w_2 \cdot 11 + w_4 \\ \sim \\ w_2(-2) + w_5 \end{array} \left[\begin{array}{ccccc|c} 1 & 2 & 3 & 4 & 5 & 2 \\ 0 & 1 & -1 & -2 & -3 & -8 \\ 0 & 0 & 1 & 2 & 3 & 3 \\ 0 & 0 & -10 & -23 & -41 & -35 \\ 0 & 0 & 4 & 3 & 11 & 21 \end{array} \right] \begin{array}{l} w_3 \cdot 10 + w_4 \\ w_3(-4) + w_5 \\ \sim \end{array} \\ & \left[\begin{array}{ccccc|c} 1 & 2 & 3 & 4 & 5 & 2 \\ 0 & 1 & -1 & -2 & -3 & -8 \\ 0 & 0 & 1 & 2 & 3 & 3 \\ 0 & 0 & 0 & -3 & -11 & -5 \\ 0 & 0 & 0 & -5 & -1 & 9 \end{array} \right] \begin{array}{l} w_4(-2) + w_5 \\ \sim \end{array} \left[\begin{array}{ccccc|c} 1 & 2 & 3 & 4 & 5 & 2 \\ 0 & 1 & -1 & -2 & -3 & -8 \\ 0 & 0 & 1 & 2 & 3 & 3 \\ 0 & 0 & 0 & -3 & -11 & -5 \\ 0 & 0 & 0 & 1 & 21 & 19 \end{array} \right] \begin{array}{l} w_4 \leftrightarrow w_5 \\ \sim \end{array}$$



$$\left[\begin{array}{ccccc|c} 1 & 2 & 3 & 4 & 5 & 2 \\ 0 & 1 & -1 & -2 & -3 & -8 \\ 0 & 0 & 1 & 2 & 3 & 3 \\ 0 & 0 & 0 & 1 & 21 & 19 \\ 0 & 0 & 0 & -3 & -11 & -5 \end{array} \right] \xrightarrow{w_4 \cdot 3 + w_5} \left[\begin{array}{ccccc|c} 1 & 2 & 3 & 4 & 5 & 2 \\ 0 & 1 & -1 & -2 & -3 & -8 \\ 0 & 0 & 1 & 2 & 3 & 3 \\ 0 & 0 & 0 & 1 & 21 & 19 \\ 0 & 0 & 0 & 0 & 52 & 52 \end{array} \right] \xrightarrow{w_5 \cdot \frac{1}{52}}$$

$$\left[\begin{array}{ccccc|c} 1 & 2 & 3 & 4 & 5 & 2 \\ 0 & 1 & -1 & -2 & -3 & -8 \\ 0 & 0 & 1 & 2 & 3 & 3 \\ 0 & 0 & 0 & 1 & 21 & 19 \\ 0 & 0 & 0 & 0 & 1 & 1 \end{array} \right] \Rightarrow \begin{cases} x_1 + 2x_2 + 3x_3 + 4x_4 + 5x_5 = 2 \\ x_2 - x_3 - 2x_4 - 3x_5 = -8 \\ x_3 + 2x_4 + 3x_5 = 3 \\ x_4 + 21x_5 = 19 \\ x_5 = 1 \end{cases}$$

$$x_5 = 1 \Rightarrow \begin{pmatrix} x_4 + 21x_5 = 19 \\ x_4 = 19 - 21x_5 \\ x_4 = -2 \end{pmatrix} \Rightarrow \begin{pmatrix} x_3 + 2x_4 + 3x_5 = 3 \\ x_3 = 3 - 2x_4 - 3x_5 \\ x_3 = 4 \end{pmatrix} \Rightarrow$$

$$\begin{pmatrix} x_2 - x_3 - 2x_4 - 3x_5 = -8 \\ x_2 = -8 + x_3 + 2x_4 + 3x_5 \\ x_2 = -5 \end{pmatrix} \Rightarrow \begin{pmatrix} x_1 + 2x_2 + 3x_3 + 4x_4 + 5x_5 = 2 \\ x_1 = 2 - 2x_2 - 3x_3 - 4x_4 - 5x_5 \\ x_1 = 3 \end{pmatrix}$$

Odpowiedź (3; -5; 4; -2; 1).

Przykład Udowodnić, że układ równań sprzeczny

$$\begin{cases} x_1 + 2x_2 - 3x_3 = -1 \\ 5x_1 + 8x_2 + x_3 = 5 \\ 2x_1 + 2x_2 + 10x_3 = 2 \end{cases}$$

Rozwiązanie

Wykorzystamy metodę eliminacji Gaussa

$$\begin{cases} x_1 + 2x_2 - 3x_3 = -1 \\ 5x_1 + 8x_2 + x_3 = 5 \\ 2x_1 + 2x_2 + 10x_3 = 2 \end{cases} \Leftrightarrow \left[\begin{array}{ccc|c} 1 & 2 & -3 & -1 \\ 5 & 8 & 1 & 5 \\ 2 & 2 & 10 & 2 \end{array} \right]$$

$$\left[\begin{array}{ccc|c} 1 & 2 & -3 & -1 \\ 5 & 8 & 1 & 5 \\ 2 & 2 & 10 & 2 \end{array} \right] \xrightarrow{w_1 \cdot (-5) + w_2} \left[\begin{array}{ccc|c} 1 & 2 & -3 & -1 \\ 0 & -2 & 16 & 10 \\ 2 & 2 & 10 & 2 \end{array} \right] \xrightarrow{w_1 \cdot (-2) + w_3} \left[\begin{array}{ccc|c} 1 & 2 & -3 & -1 \\ 0 & -2 & 16 & 10 \\ 0 & -2 & 16 & 4 \end{array} \right] \xrightarrow{w_2 \cdot \left(-\frac{1}{2}\right)} \left[\begin{array}{ccc|c} 1 & 2 & -3 & -1 \\ 0 & 1 & -8 & -5 \\ 0 & -2 & 16 & 6 \end{array} \right] \xrightarrow{w_2 \cdot 2 + w_3}$$

$$\left[\begin{array}{ccc|c} 1 & 2 & -3 & 1 \\ 0 & 1 & -8 & -5 \\ 0 & 0 & 0 & -4 \end{array} \right] \Rightarrow \begin{cases} x_1 + 2x_2 - 3x_3 = -1 \\ x_2 - 8x_3 = -5 \\ 0 = -4 \end{cases}$$

Otrzymaliśmy sprzeczne równanie $0 = -4$, co oznacza że układ równań również jest sprzeczny.

Z innej strony rząd macierzy współczynników wynosi 2, bo

$$\text{rz} \begin{bmatrix} 1 & 2 & -3 \\ 5 & 8 & 1 \\ 2 & 2 & 10 \end{bmatrix} = \text{rz} \begin{bmatrix} 1 & 2 & -3 \\ 0 & 1 & -8 \\ 0 & 0 & 0 \end{bmatrix} = 2$$

a rząd macierzy rozszerzonej układu równań wynosi 3, bo

$$\text{rz} \left[\begin{array}{ccc|c} 1 & 2 & -3 & -1 \\ 5 & 8 & 1 & 5 \\ 2 & 2 & 10 & 2 \end{array} \right] = \text{rz} \left[\begin{array}{ccc|c} 1 & 2 & -3 & 1 \\ 0 & 1 & -8 & -5 \\ 0 & 0 & 0 & -4 \end{array} \right]$$



co według twierdzenia Kroneckera-Capellego oznacza, że układ równań nie ma rozwiązań.

Odpowiedź $\emptyset$.

Przykład Rozwiązać nieoznaczony układ równań liniowych

$$\begin{cases} x_1 + 2x_2 - 3x_3 = -1 \\ 5x_1 + 8x_2 + x_3 = 5 \\ 2x_1 + 2x_2 + 10x_3 = 8 \end{cases}$$

Rozwiązanie

Ponieważ układ równań jest nieoznaczony, musimy wykorzystać metodę eliminacji Gaussa

$$\left[\begin{array}{ccc|c} 1 & 2 & -3 & -1 \\ 5 & 8 & 1 & 5 \\ 2 & 2 & 10 & 8 \end{array} \right] \xrightarrow{w_1 \cdot (-5) + w_2} \left[\begin{array}{ccc|c} 1 & 2 & -3 & -1 \\ 0 & -2 & 16 & 10 \\ 2 & 2 & 10 & 8 \end{array} \right] \xrightarrow{w_1 \cdot (-2) + w_3} \left[\begin{array}{ccc|c} 1 & 2 & -3 & -1 \\ 0 & -2 & 16 & 10 \\ 0 & 0 & 0 & 0 \end{array} \right] \sim$$

$$\left[\begin{array}{ccc|c} 1 & 2 & -3 & -1 \\ 0 & 1 & -8 & -5 \end{array} \right] \Rightarrow \begin{cases} x_1 + 2x_2 + 3x_3 = -1 \\ x_2 - 8x_3 = -5 \end{cases} \Rightarrow \begin{cases} x_1 = -13t + 9 \\ x_2 = 8t - 5 \\ x_3 = t \end{cases}$$

Odpowiedź $\{(-13t + 9; 8t - 5; t) | t \in \mathbb{R}\}$.

1.8. RÓWNANIA RÓŻNICZKOWE ZWYCZAJNE

Podstawowe pojęcia i twierdzenia

Definicja Równanie postaci $F(x, y, y', y'', \dots, y^{(n)}) = 0$, zawierające niezależną zmienną x , niewiadomą funkcję $y = y(x)$ wraz z jej pochodnymi $y', y'', \dots, y^{(n)}$, nazywa się *równaniem różniczkowym zwyczajnym*. Rząd najwyższej pochodnej, zawierającej się w równaniu, nazywa się *rzędem równania różniczkowego*.

Zwyczajność równania różniczkowego oznacza, że niewiadoma funkcja jest funkcją jednej zmiennej (dla funkcji dwu i więcej zmiennych mamy natomiast równania różniczkowe cząstkowe).

Równanie różniczkowe zwyczajne niech będzie dalej zwane RR.

W trakcie rozwiązywania RR pamiętamy różne notacji pochodnej, tzn. $y^{(n)} = \frac{d^n y}{dx^n}$

$$y' = \frac{dy}{dx}$$

$$y'' = \frac{d^2 y}{dx^2}$$

$$y''' = \frac{d^3 y}{dx^3}$$

Definicja Zbiór funkcji n -krotnie różniczkowalnych, z których każda spełnia RR, nazywa się *rozwiązaniem ogólnym*, a poszczególna funkcja nazywa się *rozwiązaniem szczególnym* równania różniczkowego.

Zwykle rozwiązanie ogólne równania różniczkowego rzędu n zależy od n stałych (tzw. stałych całkowania) $C_1, C_2, \dots, C_n$, ponieważ całka nieoznaczona określa się z dokładnością do



stałej. A rozwiązanie szczególne można otrzymać z rozwiązania ogólnego podstawiając odpowiednie wartości parametrów.

Rozwiązanie ogólne rozwiązania różniczkowego $y' = f(x, y)$ nazywa się jeszcze *całką ogólną RR*, a rozwiązanie szczególne – *całką szczególną RR*.

Definicja Zagadnieniem Cauchy'ego (zagadnieniem początkowym) nazywa się zadanie polegające na znalezieniu szczególnego rozwiązania równania różniczkowego rzędu n , które również spełnia n warunków początkowych.

Tak na przykład, zagadnienie Cauchy'ego rzędu $k \in \{1; 2; 3\}$ wygląda następująco:

$$\bullet \begin{cases} F(x, y, y') = 0 \\ y(x_0) = y_0 \end{cases} \quad \bullet \begin{cases} F(x, y, y', y'') = 0 \\ y(x_0) = y_0 \\ y'(x_0) = y_1 \end{cases} \quad \bullet \begin{cases} F(x, y, y', y'', y''') = 0 \\ y(x_0) = y_0 \\ y'(x_0) = y_1 \\ y''(x_0) = y_2 \end{cases}$$

Przykład Czy jest funkcja y rozwiązaniem RR?

$$\bullet \begin{cases} y' = 5x^4 + 3e^x \\ y = x^4 + e^x \end{cases}$$

$$2\sqrt{x}y' = 1$$

$$\bullet \begin{cases} y = \sqrt{x} + C \\ C \in \mathbb{R} \end{cases}$$

$$\bullet \begin{cases} y'' - 2y' + y = 0 \\ y = 3e^x + xe^x + 1 \end{cases}$$

$$y'' = 4e^{2x}$$

$$\bullet \begin{cases} y = e^{2x} + C_1x + C_2 \\ C_1, C_2 \in \mathbb{R} \end{cases}$$

Rozwiązanie

$$\bullet \begin{cases} y' = 5x^4 + 3e^x \\ y = x^4 + e^x \end{cases}$$

$$y = x^4 + e^x \Rightarrow y' = (x^4 + e^x)' = 4x^3 + 3e^x \neq 5x^4 + 3e^x$$

nie jest rozwiązaniem RR

$$2\sqrt{x}y' = 1$$

$$\bullet \begin{cases} y = \sqrt{x} + C \\ C \in \mathbb{R} \end{cases}$$

$$y = \sqrt{x} + C \Rightarrow y' = (\sqrt{x} + C)' = \frac{1}{2\sqrt{x}} \Rightarrow 2\sqrt{x}y' = 2\sqrt{x} \cdot \frac{1}{2\sqrt{x}} = 1$$

jest rozwiązaniem RR

$$\bullet \begin{cases} y'' - 2y' + y = 0 \\ y = 3e^x + xe^x + 1 \end{cases}$$

$$y = 3e^x + xe^x + 1 \Rightarrow \begin{aligned} y' &= 3e^x + e^x + xe^x = 4e^x + xe^x \\ y'' &= 4e^x + e^x + xe^x = 5e^x + xe^x \end{aligned} \Rightarrow$$

$$y'' - 2y' + y = 5e^x + xe^x - 2(4e^x + xe^x) + 3e^x + xe^x + 1 = 1 \neq 0$$

nie jest rozwiązaniem RR



$$y'' = 4e^{2x}$$

- $y = e^{2x} + C_1x + C_2$
 $C_1, C_2 \in \mathbb{R}$

$$y = e^{2x} + C_1x + C_2 \Rightarrow \begin{aligned} y' &= (e^{2x} + C_1x + C_2)' = 2e^{2x} + C_1 \\ y'' &= (2e^{2x} + C_1)' = 4e^{2x} \end{aligned}$$

jest rozwiązaniem RR

Przykład Czy spełnia rozwiązanie RR zaznaczony warunek początkowy?

$$2\sqrt{x}y' = 1$$

- $y = \sqrt{x}$
 $y(0) = 0$

$$y'' = 4e^{2x}$$

- $y = e^{2x} + x$
 $y(1) = 2$

Rozwiązanie

$$2\sqrt{x}y' = 1$$

- $y = \sqrt{x}$
 $y(0) = 0$

$$y = \sqrt{x} \Rightarrow y(0) = \sqrt{0} = 0$$

spełnia warunek początkowy

$$y'' = 4e^{2x}$$

- $y = e^{2x} + x$
 $y(1) = 2$

$$y = e^{2x} + x \Rightarrow y(1) = e^2 + 1 \neq 2$$

nie spełnia warunku początkowego

Rozwiązywanie wybranych równań różniczkowych nieliniowych rzędu pierwszego

Definicja RR o zmiennych rozdzielonych nazywa się równanie postaci $f(y)y' = g(x)$.

Rozwiązanie RR o zmiennych rozdzielonych otrzymamy po całkowaniu obustronnym, które polega głównie na całkowaniu równości i wyłączaniu funkcji y z otrzymanej równości (o ile to jest możliwe). Inaczej mówiąc,

$$f(y)y' = g(x)$$

$$f(y)dy = g(x)dx$$

$$\int f(y)dy = \int g(x)dx$$

$$F(y) = G(x)$$

$$y = F^{-1}(G(x))$$

Przykład Znaleźć rozwiązanie ogólne RR o zmiennych rozdzielonych

- $y' = 6x^2 - \sin x + 1$

- $dy = 3x^2dx$



- $e^y dy = (e^x + 4x) dx$

- $y' = e^{\cos x} \sin x$

Rozwiązanie

- $y' = 6x^2 - \sin x + 1$

$$\frac{dy}{dx} = 6x^2 - \sin x + 1$$

$$dy = (6x^2 - \sin x + 1) dx$$

$$\int 1 dy = \int (6x^2 - \sin x + 1) dx$$

$$y + C_1 = 2x^3 + \cos x + x + C_2, C_{1,2} \in \mathbb{R}$$

niech $C = C_2 - C_1$

$$y = 2x^3 + \cos x + x + C, C \in \mathbb{R}$$

- $dy = 3x^2 dx$

$$dy = 3x^2 dx$$

$$\int 1 dy = \int 3x^2 dx$$

$$y = x^3 + C, C \in \mathbb{R}$$

- $e^y dy = (e^x + 4x) dx$

$$e^y dy = (e^x + 4x) dx$$

$$\int e^y dy = \int (e^x + 4x) dx$$

$$e^y = e^x + 2x^2 + C, C \in \mathbb{R}$$

$$\ln e^y = \ln(e^x + 2x^2 + C)$$

$$y = \ln(e^x + 2x^2 + C), C \in \mathbb{R}$$

- $y' = e^{\cos x} \sin x$

$$\frac{dy}{dx} = e^{\cos x} \sin x$$

$$dy = e^{\cos x} \sin x dx$$

$$\int 1 dy = \int e^{\cos x} \sin x dx$$

$$\begin{aligned} \int e^{\cos x} \sin x dx &= \left[\begin{array}{l} \cos x = t \\ d(\cos x) = dt \\ -\sin x dx = dt \\ \sin x dx = -dt \end{array} \right] = - \int e^t dt = -e^t + C = \llbracket t = \cos x \rrbracket \\ &= -e^{\cos x} + C \end{aligned}$$

$$y = -e^{\cos x} + C, C \in \mathbb{R}$$

Przykład Znaleźć rozwiązanie szczególne RR



- $y' = 6x^2 - \sin x + 1$

- $dy = 3x^2 dx$

- $e^y dy = (e^x + 4x) dx$

- $y' = e^{\cos x} \sin x$

Rozwiązanie

Rozwiązania szczególne otrzymujemy z rozwiązań ogólnych, podstawiając jakiekolwiek dopuszczalne wartości stałej C . Tak więc,

- $y' = 6x^2 - \sin x + 1$

Rozwiązanie ogólne RR: $y = 2x^3 + \cos x + x + C, C \in \mathbb{R}$

Rozwiązanie szczególne: $C = 0 \Rightarrow y = 2x^3 + \cos x + x$

- $dy = 3x^2 dx$

Rozwiązanie ogólne RR: $y = x^3 + C, C \in \mathbb{R}$

Rozwiązanie szczególne: $C = 1 \Rightarrow y = x^3 + 1$

- $e^y dy = (e^x + 4x) dx$

Rozwiązanie ogólne RR: $y = \ln(e^x + 2x^2 + C), C \in \mathbb{R}$

Rozwiązanie szczególne: $C = -2 \Rightarrow y = \ln(e^x + 2x^2 - 2)$

- $y' = e^{\cos x} \sin x$

Rozwiązanie ogólne RR: $y = -e^{\cos x} + C, C \in \mathbb{R}$

Rozwiązanie szczególne: $C = e \Rightarrow y = e - e^{\cos x}$

Przykład Rozwiązać zagadnienie Cauchy'ego

- $y' = 6x^2 - \sin x + 1$
 $y(0) = 7$

- $dy = 3x^2 dx$
 $y(1) = 0$

- $e^y dy = (e^x + 4x) dx$
 $y(0) = 1$

- $y' = e^{\cos x} \sin x$
 $y\left(\frac{\pi}{2}\right) = 3$

Rozwiązanie

Rozwiązania ogólne RR już mamy otrzymane w poprzednich przykładach. Na ogół, poniższe rozwiązania są tylko częścią od całego rozwiązania zagadnienia Cauchy'ego.

- $y' = 6x^2 - \sin x + 1$
 $y(0) = 7$

Rozwiązanie ogólne RR: $y = 2x^3 + \cos x + x + C, C \in \mathbb{R}$

$$y(0) = 0 + \cos 0 + C = 1 + C = 7$$

$$1 + C = 7 \Rightarrow C = 6$$

Rozwiązanie zagadnienia Cauchy'ego: $y = 2x^3 + \cos x + x + 6$



- $dy = 3x^2 dx$
 $y(1) = 0$

Rozwiązanie ogólne RR: $y = x^3 + C, C \in \mathbb{R}$

$$y(1) = 1 + C = 0$$

$$1 + C = 0 \Rightarrow C = -1$$

Rozwiązanie zagadnienia Cauchy'ego: $y = x^3 - 1$

- $e^y dy = (e^x + 4x) dx$
 $y(0) = 1$

Rozwiązanie ogólne RR: $y = \ln(e^x + 2x^2 + C), C \in \mathbb{R}$

$$y(0) = \ln(1 + 0 + C) = 1$$

$$\ln(1 + C) = 1 \Rightarrow \ln(1 + C) = \ln e \Rightarrow C = e - 1$$

Rozwiązanie zagadnienia Cauchy'ego: $y = \ln(e^x + 2x^2 + e - 1)$

- $y' = e^{\cos x} \sin x$
 $y\left(\frac{\pi}{2}\right) = 3$

Rozwiązanie ogólne RR: $y = -e^{\cos x} + C, C \in \mathbb{R}$

$$y\left(\frac{\pi}{2}\right) = -e^{\cos \frac{\pi}{2}} + C = -e^0 + C = C - 1 = 3$$

$$C - 1 = 3 \Rightarrow C = 4$$

Rozwiązanie zagadnienia Cauchy'ego: $y = 4 - e^{\cos x}$

Definicja RR liniowym rzędu pierwszego nazywa się równanie postaci

$$y' + p(x)y = q(x)$$

Jeżeli $q(x) \equiv 0$, to otrzymane równanie $y' + p(x)y = 0$ nazywa się RR *jednorodnym*. Jeżeli $q(x) \neq 0$, to otrzymane równanie nazywa się RR *niejednorodnym*.

Równanie $y' + p(x)y = 0$ jest również RR o zmiennych rozdzielonych.

Twierdzenie Rozwiązanie ogólne równania różniczkowego liniowego rzędu pierwszego można znaleźć ze wzoru

$$y = e^{-P(x)} \left(\int q(x) e^{P(x)} dx + C \right), C \in \mathbb{R}$$

gdzie $P(x) = \int p(x) dx$.

W ostatnim wzorze, obliczając całki nieoznaczone, nie wpisujemy stałych, ponieważ stała C już jest w rozwiązaniu RR.

Jest również inne metody dla rozwiązywania RR liniowego rzędu pierwszego.

Przykład Rozwiązać RR liniowe rzędu pierwszego

- $y' + y = 1$

- $y' + 3x^2 y = 3x^2$



- $y' + y = 2e^{3x}$

- $y' + 3y = \cos x$

Rozwiązanie

- $y' + y = 1$

$$\begin{aligned} y' + p(x)y &= q(x) \\ y' + y &= 1 \end{aligned} \Rightarrow \begin{aligned} p(x) &= 1 \\ q(x) &= 1 \end{aligned}$$

$$P(x) = \int 1 \, dx = x$$

$$\int q(x)e^{P(x)} \, dx = \int e^x \, dx = e^x$$

$$y = e^{-x}(e^x + C) = Ce^{-x} + 1, C \in \mathbb{R}$$

- $y' + 3x^2y = 3x^2$

$$\begin{aligned} y' + p(x)y &= q(x) \\ y' + 3x^2y &= 3x^2 \end{aligned} \Rightarrow \begin{aligned} p(x) &= 3x^2 \\ q(x) &= 3x^2 \end{aligned}$$

$$P(x) = \int 3x^2 \, dx = x^3$$

$$\begin{aligned} \int q(x)e^{P(x)} \, dx &= \int 3x^2 e^{x^3} \, dx = \left[\begin{array}{l} x^3 = t \\ d(x^3) = dt \\ 3x^2 dx = dt \end{array} \right] = \int e^t \, dt = e^t = \llbracket t = x^3 \rrbracket \\ &= e^{x^3} \end{aligned}$$

$$y = e^{-x^3}(e^{x^3} + C) = Ce^{-x^3} + 1, C \in \mathbb{R}$$

- $y' + y = 2e^{3x}$

$$\begin{aligned} y' + p(x)y &= q(x) \\ y' + y &= 2e^{3x} \end{aligned} \Rightarrow \begin{aligned} p(x) &= 1 \\ q(x) &= 2e^{3x} \end{aligned}$$

$$P(x) = \int 1 \, dx = x$$

$$\int q(x)e^{P(x)} \, dx = \int 2e^{3x} \cdot e^x \, dx = \int 2e^{4x} \, dx = \frac{1}{2}e^{4x}$$

$$y = e^{-x}\left(\frac{1}{2}e^{4x} + C\right) = \frac{1}{2}e^{3x} + Ce^{-x}, C \in \mathbb{R}$$

- $y' + 3y = \cos x$

$$\begin{aligned} y' + p(x)y &= q(x) \\ y' + 3y &= \cos x \end{aligned} \Rightarrow \begin{aligned} p(x) &= 3 \\ q(x) &= \cos x \end{aligned}$$

$$P(x) = \int 3 \, dx = 3x$$

$$\int q(x)e^{P(x)} \, dx = \int \cos x \cdot e^{3x} \, dx \equiv;$$



$$\begin{aligned}
I &= \int \cos x \cdot e^{3x} dx = \int e^{3x} d(\sin x) = e^{3x} \sin x - \int \sin x de^{3x} \\
&= e^{3x} \sin x - \int 3e^{3x} \sin x dx = e^{3x} \sin x - \int 3e^{3x} d(-\cos x) \\
&= e^{3x} \sin x + 3e^{3x} \cos x - \int \cos x d3e^{3x} \\
&= e^{3x} \sin x + 3e^{3x} \cos x - \int 9e^{3x} \cos x dx \\
&= e^{3x} \sin x + 3e^{3x} \cos x - 9I \\
I &= e^{3x} \sin x + 3e^{3x} \cos x - 9I \\
10I &= e^{3x} \sin x + 3e^{3x} \cos x \\
I &= \frac{1}{10} (e^{3x} \sin x + 3e^{3x} \cos x);
\end{aligned}$$

$$\begin{aligned}
&\equiv \frac{1}{10} (e^{3x} \sin x + 3e^{3x} \cos x) = \frac{e^{3x}(\sin x + 3 \cos x)}{10} \\
y &= e^{-3x} \left(\frac{e^{3x}(\sin x + 3 \cos x)}{10} + C \right) = \frac{\sin x + 3 \cos x}{10} + Ce^{-3x}, C \in \mathbb{R}
\end{aligned}$$

Rozwiązanie jednorodnych równań różniczkowych liniowych wyższych rzędów o stałych współczynnikach

Definicja RR postaci $ay'' + by' + cy = f(x)$, $a, b, c \in \mathbb{R}$ nazywa się *RR liniowym rzędu drugiego o stałych współczynnikach* (a, b, c) . Jeżeli dodatkowo $f(x) \equiv 0$, otrzymane równanie $ay'' + by' + cy = 0$ nazywa się *jednorodne*. Jeżeli $f(x) \neq 0$, to otrzymane nazywa się *niejednorodne*.

Definicja RR postaci $a_n y^{(n)} + a_{n-1} y^{(n-1)} + \dots + a_1 y' + a_0 y = f(x)$, $a_i \in \mathbb{R}, i = \overline{0, n}$ nazywa się *RR liniowym rzędu n o stałych współczynnikach* $(a_0, a_1, \dots, a_n)$. Jeżeli dodatkowo $f(x) \equiv 0$, otrzymane równanie nazywa się *jednorodne*, a jeżeli $f(x) \neq 0$, to – *niejednorodne*.

Definicja Równaniem charakterystycznym RR liniowego rzędu drugiego

$$ay'' + by' + cy = f(x), a, b, c \in \mathbb{R}$$

nazywa się równanie

$$a\lambda^2 + b\lambda + c = 0$$

które otrzymujemy z RR po podstawieniu $y'' = \lambda^2$, $y' = \lambda$, $y = 1$.

Definicja Równaniem charakterystycznym RR liniowego rzędu n

$$a_n y^{(n)} + a_{n-1} y^{(n-1)} + \dots + a_1 y' + a_0 y = 0, a_i \in \mathbb{R}, i = \overline{0, n}$$

nazywa się równanie

$$a_n \lambda^n + a_{n-1} \lambda^{n-1} + \dots + a_1 \lambda + a_0 = 0$$



które otrzymujemy z RR po podstawieniu $y^{(n)} = \lambda^n$, $y^{(n-1)} = \lambda^{n-1}$, ..., $y' = \lambda$, $y = 1$.

Równanie charakterystyczne rozwiązujemy w zbiorze liczb zespolonych, co oznacza, że równanie charakterystyczne rzędu n zawsze musi mieć dokładnie n rozwiązań (rozwiązanie krotności np. 2 jest to samo, co dwa jednakowe rozwiązania).

Rozwiązania bazowe RR: W zależności od znalezionych rozwiązań równania charakterystycznego, otrzymujemy takie tzw. rozwiązania bazowe równania różniczkowego:

- jeżeli pierwiastki rzeczywiste, to
 - λ jest pierwiastkiem krotności 1 $\Rightarrow$ jedno rozwiązanie bazowe $y = e^{\lambda x}$
 - λ jest pierwiastkiem krotności 2 $\Rightarrow$ dwa rozwiązania bazowe
$$\begin{cases} y_1 = e^{\lambda x} \\ y_2 = x e^{\lambda x} \end{cases}$$
 - λ jest pierwiastkiem krotności 3 $\Rightarrow$ trzy rozwiązania bazowe
$$\begin{cases} y_1 = e^{\lambda x} \\ y_2 = x e^{\lambda x} \\ y_3 = x^2 e^{\lambda x} \end{cases}$$
 - λ jest pierwiastkiem krotności k $\Rightarrow k$ rozwiązań bazowych
$$\begin{cases} y_1 = e^{\lambda x} \\ y_2 = x e^{\lambda x} \\ y_3 = x^2 e^{\lambda x} \\ \dots \\ y_k = x^{k-1} e^{\lambda x} \end{cases}$$
- jeżeli pierwiastki zespolone, to (i tj. jednostka urojona, $i^2 = -1$)
 - $\lambda = \alpha \pm i\beta$ są rozwiązaniami krotności 1 $\Rightarrow$ dwa rozwiązania bazowe
$$\begin{cases} y_1 = e^{\alpha x} \cos \beta x \\ y_2 = e^{\alpha x} \sin \beta x \end{cases}$$
 - $\lambda = \alpha \pm i\beta$ są rozwiązaniami krotności 2 $\Rightarrow$ cztery rozwiązania bazowe
$$\begin{cases} y_1 = e^{\alpha x} \cos \beta x \\ y_2 = e^{\alpha x} \sin \beta x \\ y_3 = x e^{\alpha x} \cos \beta x \\ y_4 = x e^{\alpha x} \sin \beta x \end{cases}$$
 - $\lambda = \alpha \pm i\beta$ są rozwiązaniami krotności 3 $\Rightarrow$ sześć rozwiązań bazowych
$$\begin{cases} y_1 = e^{\alpha x} \cos \beta x \\ y_2 = e^{\alpha x} \sin \beta x \\ y_3 = x e^{\alpha x} \cos \beta x \\ y_4 = x e^{\alpha x} \sin \beta x \\ y_5 = x^2 e^{\alpha x} \cos \beta x \\ y_6 = x^2 e^{\alpha x} \sin \beta x \end{cases}$$
 - $\lambda = \alpha \pm i\beta$ są rozwiązaniami krotności k $\Rightarrow 2k$ rozwiązań bazowych



$$\left[\begin{array}{l} y_1 = e^{\alpha x} \cos \beta x \\ y_2 = e^{\alpha x} \sin \beta x \\ y_3 = x e^{\alpha x} \cos \beta x \\ y_4 = x e^{\alpha x} \sin \beta x \\ y_5 = x^2 e^{\alpha x} \cos \beta x \\ y_6 = x^2 e^{\alpha x} \sin \beta x \\ \vdots \\ y_{2k-1} = x^{k-1} e^{\alpha x} \cos \beta x \\ y_{2k} = x^{k-1} e^{\alpha x} \sin \beta x \end{array} \right.$$

Twierdzenie Schemat rozwiązania jednorodnych RR liniowych wyższych rzędów:

$$a_n y^{(n)} + a_{n-1} y^{(n-1)} + \dots + a_1 y' + a_0 y = 0, a_i \in \mathbb{R}, i = \overline{0, n}$$

- układamy równanie charakterystyczne

$$a_n \lambda^n + a_{n-1} \lambda^{n-1} + \dots + a_1 \lambda + a_0 = 0$$

- rozwiązujemy równanie charakterystyczne $\Rightarrow n$ rozwiązań $\lambda_1, \lambda_2, \dots, \lambda_n$
- generujemy funkcje bazowe $y_1, y_2, \dots, y_n$ w zależności od otrzymanych wyżej rozwiązań równania charakterystycznego
- wypisujemy kombinację liniową rozwiązań bazowych

$$y = C_1 y_1 + C_2 y_2 + \dots + C_n y_n$$

która zależy od n stałych $C_i \in \mathbb{R}, i = \overline{0, n}$ oraz jest rozwiązaniem ogólnym jednorodnego RR rzędu n .

Przykład Rozwiązać RR liniowe wyższych rzędów

- $y'' + 5y' + 6y = 0$
- $y'' - 6y' + 9y = 0$
- $y''' - 5y'' + 4y' = 0$
- $y'' + y = 0$
- $y'' - 4y' + 13y = 0$
- $y^{IV} + 9y'' = 0$

Rozwiązanie

- $y'' + 5y' + 6y = 0$

$$\lambda^2 + 5\lambda + 6 = 0$$

$$\begin{array}{l} \lambda_1 = -3 \\ \lambda_2 = -2 \end{array} \Rightarrow \begin{array}{l} y_1 = e^{-3x} \\ y_2 = e^{-2x} \end{array}$$

$$y = C_1 e^{-3x} + C_2 e^{-2x}, C_{1,2} \in \mathbb{R}$$

- $y'' - 6y' + 9y = 0$

$$\lambda^2 - 6\lambda + 9 = 0$$

$$\lambda_{1,2} = 3 \Rightarrow \begin{array}{l} y_1 = e^{3x} \\ y_2 = x e^{3x} \end{array}$$

$$y = C_1 e^{3x} + C_2 x e^{3x}, C_{1,2} \in \mathbb{R}$$

- $y''' - 5y'' + 4y' = 0$



$$\lambda^3 - 5\lambda^2 + 4\lambda = 0$$

$$\lambda(\lambda^2 - 5\lambda + 4) = 0$$

$$\begin{array}{ll} \lambda_1 = 0 & y_1 = e^{0x} = 1 \\ \lambda_2 = 1 & \Rightarrow y_2 = e^x \\ \lambda_3 = 4 & y_3 = e^{4x} \end{array}$$

$$y = C_1 + C_2 e^x + C_3 e^{4x}, C_1, C_2, C_3 \in \mathbb{R}$$

- $y'' + y = 0$

$$\lambda^2 + 1 = 0$$

$$\lambda^1 = -1$$

$$\begin{array}{ll} \lambda_{1,2} = \pm i & y_1 = e^{0x} \cos(1 \cdot x) = \cos x \\ \lambda_{1,2} = 0 \pm 1i & \Rightarrow y_2 = e^{0x} \sin(1 \cdot x) = \sin x \end{array}$$

$$y = C_1 \cos x + C_2 \sin x, C_{1,2} \in \mathbb{R}$$

- $y'' - 4y' + 13y = 0$

$$\lambda^2 - 4\lambda + 13 = 0$$

$$\Delta = -36 = (6i)^2 \Rightarrow \sqrt{\Delta} = \pm 6i$$

$$\lambda_{1,2} = \frac{4 \pm 6i}{2} = 2 \pm 3i \Rightarrow \begin{array}{l} y_1 = e^{2x} \cos 3x \\ y_2 = e^{2x} \sin 3x \end{array}$$

$$y = C_1 e^{2x} \cos 3x + C_2 e^{2x} \sin 3x, C_{1,2} \in \mathbb{R}$$

- $y^{IV} + 9y'' = 0$

$$\lambda^4 + 9\lambda^2 = 0$$

$$\lambda^2(\lambda^2 + 9) = 0$$

$$\begin{array}{ll} \lambda^2 = 0 & \lambda^2 + 9 = 0 \\ \lambda_{1,2} = 0 & \vee \begin{array}{l} \lambda^2 = -9 \\ \lambda_{3,4} = \pm 3i \end{array} \end{array}$$

$$\begin{array}{ll} \lambda_{1,2} = 0 & \Rightarrow \begin{array}{l} y_1 = 1 \\ y_2 = x \\ y_3 = \cos 3x \\ y_4 = \sin 3x \end{array} \\ \lambda_{3,4} = \pm 3i & \end{array}$$

$$y = C_1 + C_2 x + C_3 \cos 3x + C_4 \sin 3x, C_1, C_2, C_3, C_4 \in \mathbb{R}$$

Przykład Rozwiązać zagadnienia Cauchy'ego

- $\begin{array}{l} y'' + 5y' + 6y = 0 \\ y(0) = 5 \\ y'(0) = -12 \end{array}$

- $\begin{array}{l} y''' - 5y'' + 4y' = 0 \\ y(0) = 2 \\ y'(0) = 4 \\ y''(0) = 16 \end{array}$

Rozwiązanie



Rozwiązania ogólne RR już mamy otrzymane w poprzednich przykładach. Na ogół, poniższe rozwiązania są tylko częścią od całego rozwiązania zagadnienia Cauchy'ego.

$$\begin{aligned} y'' + 5y' + 6y &= 0 \\ y(0) &= 5 \\ y'(0) &= -12 \end{aligned}$$

Rozwiązanie ogólne RR: $y = C_1 e^{-3x} + C_2 e^{-2x}, C_{1,2} \in \mathbb{R}$

$$y' = -3C_1 e^{-3x} - 2C_2 e^{-2x}$$

$$\begin{aligned} y(0) &= C_1 + C_2 \\ y'(0) &= -3C_1 - 2C_2 \end{aligned} \Rightarrow \begin{cases} C_1 + C_2 = 5 \\ -3C_1 - 2C_2 = -12 \end{cases} \Rightarrow \begin{aligned} C_1 &= 2 \\ C_2 &= 3 \end{aligned} \Rightarrow$$

rozwiązanie zagadnienia Cauchy'ego: $y = 2e^{-3x} + 3e^{-2x}$

$$\begin{aligned} y''' - 5y'' + 4y' &= 0 \\ y(0) &= 2 \\ y'(0) &= 4 \\ y''(0) &= 16 \end{aligned}$$

Rozwiązanie ogólne RR: $y = C_1 + C_2 e^x + C_3 e^{4x}, C_1, C_2, C_3 \in \mathbb{R}$

$$y' = C_2 e^x + 4C_3 e^{4x}$$

$$y'' = C_2 e^x + 16C_3 e^{4x}$$

$$\begin{aligned} y(0) &= C_1 + C_2 + C_3 = 2 \\ y'(0) &= C_2 + 4C_3 = 4 \\ y''(0) &= C_2 + 16C_3 = 16 \end{aligned} \Rightarrow \begin{cases} C_1 + C_2 + C_3 = 2 \\ C_2 + 4C_3 = 4 \\ C_2 + 16C_3 = 16 \end{cases} \Rightarrow \begin{aligned} C_1 &= 1 \\ C_2 &= 0 \\ C_3 &= 1 \end{aligned} \Rightarrow$$

rozwiązanie zagadnienia Cauchy'ego: $y = 1 + e^{4x}$

Interpretacja geometryczna RR rzędu pierwszego

Niech dane jest RR rzędu pierwszego $x' = f(t, x)$, rozwiązaniem którego niech będzie funkcja $x = x(t)$, przechodząca przez punkt $(t_0; x_0), x_0 = x(t_0)$.

Tak na przykład, jeśli jednowymiarowy ruch punktu materialnego określony równaniem różniczkowym $x''(t) = \frac{F}{m}$, to, zakładając warunki początkowe $\begin{matrix} x(t_0) = x_0 \\ x'(t_0) = x_1 \end{matrix}$, rozwiązanie zagadnienia Cauchy'ego określa trajektorię ruchu. Bez warunków początkowych, otrzymamy cały zbiór trajektorii (w zależności od stałych całkowania).

Z innej strony z rachunku różniczkowego funkcji jednej zmiennej wiadomo, że $y'(x_0) = \tan \alpha$, gdzie α – tj. kąt nachylenia stycznej do wykresu funkcji $y = y(x)$ w punkcie $(x_0; y_0)$ a dodatniego kierunku osi Ox . Wykorzystując to, można zdefiniować pole kierunków RR i krzywe całkowe RR [12].

Niektóre zastosowania równań różniczkowych

Równania różniczkowe mają swoje zastosowanie niemal wszędzie, gdzie jest zmiana, która w swoją kolej określa się różniczką (lub pochodną). Jak już było zaznaczono w



rozdziałach różniczkowego i całkowego rachunków, dynamika ruchu jednowymiarowego punktu materialnego obiektu według drugiego prawa Newtona [13]:

$$F = ma$$

Ostatnia równość, po podstawieniu zamiast przyspieszenia a pochodną prędkości $v'(t)$ oraz drugą pochodną drogi $s''(t)$, przepisuje się w postaciach

$$ma = F \Rightarrow mv'(t) = F \Rightarrow ms''(t) = F$$

Zakładając, że siła nie zmienia się, tzn. $F = \text{const}$, otrzymamy równanie takie różniczkowe rzędu drugiego

$$s''(t) = \frac{F}{m}$$

Rozwiązaniem którego jest funkcja $s(t) = \frac{F}{2m}t^2 + C_1t + C_2, C_{1,2} \in \mathbb{R}$

Przykład Znaleźć trajektorię ruchu jednowymiarowego punktu materialnego masy $m = 3$ na które działa stała siła $F = 12$.

Rozwiązanie

Zgodnie z drugim prawem Newtona $F = ma \Leftrightarrow a = \frac{F}{m}$. Podstawiając $a = x''(t)$ otrzymamy równanie różniczkowe drugiego rzędu $x''(t) = \frac{F}{m}$.

$$\begin{matrix} m = 3 \\ F = 12 \end{matrix} \Rightarrow x''(t) = \frac{12}{3} = 4$$

$$x''(t) = 4$$

$$\int x''(t) dt = \int 4 dt$$

$$x'(t) = 4t + C_1$$

$$\int x'(t) dt = \int (4t + C_1) dt$$

$$x(t) = 2t^2 + C_1t + C_2, C_{1,2} \in \mathbb{R}$$

Trajektoria ruchu jednowymiarowego punktu materialnego jest określona funkcją $x(t) = 2t^2 + C_1t + C_2$, gdzie $C_{1,2}$ – to są tzw. stałe całkowania, tzn. dowolne stałe rzeczywiste.

Innym przykładem jest równanie związane z drganiem harmonicznym (tzn. drganiem ciała, na które działa siła $\vec{F}$) [13], które w jednowymiarowym przypadku ma postać

$$F = -kx(t)$$

gdzie $x(t)$ – tj. wartość wychylenia z położenia równowagi (zależy od czasu), $k > 0$ – stała, a znam minus związany z tym, że siła F skierowana przeciwnie do wychylenia.

Przykład Rozwiązać poniższe równanie jednowymiarowego drgania harmonicznego ciała z masą $m = 1$



$$F = -9x(t)$$

Rozwiązanie

Zgodnie z drugim prawem Newtona

$$F = ma = \llbracket a = x''(t) \rrbracket = mx''(t) = -9x(t)$$

$$mx''(t) = -9x(t)$$

$$x''(t) = -9x(t)$$

$$x''(t) + 9x(t) = 0$$

$$\lambda^2 + 9 = 0$$

$$\lambda_{1,2} = \pm 3i$$

$$\lambda_{1,2} = \pm 3i \Rightarrow \begin{aligned} x_1(t) &= \cos 3t \\ x_2(t) &= \sin 3t \end{aligned}$$

$$x(t) = C_1 x_1(t) + C_2 x_2(t), C_{1,2} \in \mathbb{R}$$

$$x(t) = C_1 \cos 3t + C_2 \sin 3t, C_{1,2} \in \mathbb{R}$$

Istnieje φ , takie że $\cos \varphi = \frac{C_1}{\sqrt{C_1^2 + C_2^2}}$, $\sin \varphi = \frac{C_2}{\sqrt{C_1^2 + C_2^2}}$, skąd wynikają poniższe równości

$$\begin{aligned} x(t) &= \sqrt{C_1^2 + C_2^2} \left(\frac{C_1}{\sqrt{C_1^2 + C_2^2}} \cos 3t + \frac{C_2}{\sqrt{C_1^2 + C_2^2}} \sin 3t \right) \\ &= \sqrt{C_1^2 + C_2^2} (\cos \varphi \cos 3t + \sin \varphi \sin 3t) \\ &= \llbracket \cos \alpha \cos \beta + \sin \alpha \sin \beta = \cos(\alpha - \beta) \rrbracket = \sqrt{C_1^2 + C_2^2} \cdot \cos(\varphi - 3t) \\ &= \llbracket \cos(-\gamma) = \cos \gamma \rrbracket = \sqrt{C_1^2 + C_2^2} \cdot \cos(3t - \varphi) \end{aligned}$$

Tak więc, funkcja drgania harmonicznego ciała z masą $m = 1$ ma postać

$$x(t) = \sqrt{C_1^2 + C_2^2} \cdot \cos(3t - \varphi)$$

gdzie $\sqrt{C_1^2 + C_2^2}$ określa amplitudę drgań, a φ – ich fazę.

1.9. OBLICZANIE CAŁEK PODWÓJNYCH, POTRÓJNYCH, KRZYWOLINIOWYCH I POWIERZCHNIOWYCH

Definicja Całką podwójną funkcji $f(x, y)$ w domkniętym obszarze $D^2 \subset \mathbb{R}^2$ nazywa się całka

$$\iint_{D^2} f(x, y) dx dy$$

Całkę podwójną można obliczyć z jednej z wn. równości (w zależności od obszaru D^2):



$$\bullet \quad D^2 = \left\{ (x; y) \in \mathbb{R}^2 \mid \begin{array}{l} x_1 \leq x \leq x_2 \\ g_1(x) \leq y \leq g_2(x) \end{array} \right\}:$$

$$\iint_{D^2} f(x, y) \, dx \, dy = \int_{x_1}^{x_2} \int_{g_1(x)}^{g_2(x)} f(x; y) \, dy \, dx$$

$$\bullet \quad D^2 = \left\{ (x; y) \in \mathbb{R}^2 \mid \begin{array}{l} y_1 \leq y \leq y_2 \\ h_1(y) \leq x \leq h_2(y) \end{array} \right\}:$$

$$\iint_{D^2} f(x, y) \, dx \, dy = \int_{y_1}^{y_2} \int_{h_1(y)}^{h_2(y)} f(x; y) \, dx \, dy$$

Długość odcinka $\langle x_1, x_2 \rangle \subset \mathbb{R}$ można znaleźć za pomocą całki $\int_{x_1}^{x_2} 1 \, dx = x_2 - x_1$.

Twierdzenie Pole płaskiego obszaru $D^2 \subset \mathbb{R}^2$ można znaleźć za pomocą całki podwójnej, mianowicie

$$\iint_{D^2} 1 \, dx \, dy$$

Definicja Całką potrójną funkcji $f(x, y, z)$ w domkniętym obszarze $D^3 \subset \mathbb{R}^3$ nazywa się całka

$$\iiint_{D^3} f(x, y, z) \, dx \, dy \, dz$$

$$\text{Niech } D^3 = \left\{ (x; y; z) \in \mathbb{R}^3 \mid \begin{array}{l} x_1 \leq x \leq x_2 \\ g_1(x) \leq y \leq g_2(x) \\ h_1(x, y) \leq z \leq h_2(x, y) \end{array} \right\}, \text{ wtedy całkę potrójną można}$$

obliczyć z równości

$$\iiint_{D^3} f(x, y, z) \, dx \, dy \, dz = \int_{x_1}^{x_2} \int_{g_1(x)}^{g_2(x)} \int_{h_1(x, y)}^{h_2(x, y)} f(x; y) \, dz \, dy \, dx$$

Twierdzenie Objętość bryły $D^3 \subset \mathbb{R}^3$ można znaleźć za pomocą całki potrójnej, mianowicie

$$\iiint_{D^3} 1 \, dx \, dy \, dz$$

Definicja Całką n -krotną funkcji $f(x_1, x_2, \dots, x_n)$ nazywa się całka

$$\int \int \dots \int_{D^n} f(x_1, x_2, \dots, x_n) \, dx_1 \, dx_2 \, \dots \, dx_n$$



Całką krzywoliniową nazywa się całka, w której funkcja podcałkowa przyjmuje swoje wartości wzdłuż krzywej. W przypadku krzywej zamkniętej, otrzymamy tzw. *całkę okrężną*.
Całką powierzchniową nazywa się całka z powierzchnią w jakości obszaru całkowania.

Przykład Obliczyć całki podwójne:

- $\iint_{D^2} (3x^2 + 8xy) dx dy, D^2 = \left\{ (x; y) \left| \begin{array}{l} 0 \leq x \leq 1 \\ x^2 \leq y \leq x \end{array} \right. \right\}$
- $\iint_{D^2} 4xy dx dy, D^2 = \left\{ (x; y) \left| \begin{array}{l} 1 \leq y \leq 2 \\ y^2 \leq x \leq y^3 \end{array} \right. \right\}$

Rozwiązanie

- $\iint_{D^2} (3x^2 + 8xy) dx dy = \int_0^1 \left(\int_{x^2}^x (3x^2 + 8xy) dy \right) dx \quad \equiv;$
 $\int_{x^2}^x (3x^2 + 8xy) dy = [3x^2 y + 4xy^2]_{y=x^2}^{y=x} = (3x^3 + 4x^3) - (3x^4 + 4x^5) =$
 $7x^3 - 3x^4 - 4x^5;$
 $\equiv \int_0^1 (7x^3 - 3x^4 - 4x^5) dx = \left[\frac{7}{4}x^4 - \frac{3}{5}x^5 - \frac{2}{3}x^6 \right]_0^1 = \left(\frac{7}{4} - \frac{3}{5} - \frac{2}{3} \right) - 0 = \frac{29}{60}$
- $\iint_{D^2} 4xy dx dy = \int_1^2 \left(\int_{y^2}^{y^3} 4xy dx \right) dy \quad \equiv;$
 $\int_{y^2}^{y^3} 4xy dx = [2x^2 y]_{x=y^2}^{x=y^3} = 2y^7 - 2y^5;$
 $\equiv \int_1^2 (2y^7 - 2y^5) dy = \left[\frac{1}{4}y^8 - \frac{1}{3}y^6 \right]_1^2 = 42,75$

Przykład Wykorzystując całkę podwójną, znaleźć pole koła o promieniu $r = 1$.

Rozwiązanie

$$\begin{aligned} x^2 + y^2 &= r^2 \\ x^2 + y^2 &= 1 \\ y &= \pm \sqrt{1 - x^2}, -1 \leq x \leq 1 \end{aligned} \Rightarrow D = \left\{ (x, y) \left| \begin{array}{l} -1 \leq x \leq 1 \\ -\sqrt{1 - x^2} \leq y \leq \sqrt{1 - x^2} \end{array} \right. \right\}$$

Wtedy pole koła równa się

$$\iint_D 1 dx dy = \int_{-1}^1 \left(\int_{-\sqrt{1-x^2}}^{\sqrt{1-x^2}} 1 dy \right) dx = \int_{-1}^1 \left([y]_{y=-\sqrt{1-x^2}}^{y=\sqrt{1-x^2}} \right) dx = \int_{-1}^1 (2\sqrt{1-x^2}) dx \quad \equiv;$$

z jednego z przykładów w rozdziale Rachunek całkowy wynika, że

$$\int \sqrt{1-x^2} dx = \frac{1}{2} (x\sqrt{1-x^2} + \arcsin x) + C$$

a więc

$$\equiv 2 \cdot \frac{1}{2} \left[x\sqrt{1-x^2} + \arcsin x \right]_{-1}^1 = \arcsin 1 - \arcsin(-1) = \frac{\pi}{2} - \left(-\frac{\pi}{2} \right) = \pi$$

Przykład Wykorzystując całkę potrójną, znaleźć wzór na obliczanie objętości walca o wysokości h i średnicy r .



Rozwiązanie

$$\frac{r}{h} \Rightarrow D = \left\{ (x, y) \left| \begin{array}{l} -r \leq x \leq r \\ -\sqrt{r^2 - x^2} \leq y \leq \sqrt{r^2 - x^2} \\ 0 \leq z \leq h \end{array} \right. \right\}$$

Wtedy objętość walca równa się

$$\begin{aligned} \iiint_D 1 \, dx \, dy \, dz &= \int_0^h \int_{-r}^r \int_{-\sqrt{r^2 - x^2}}^{\sqrt{r^2 - x^2}} 1 \, dy \, dx \, dz \quad \boxed{\Xi_1}; \\ \int_{-\sqrt{r^2 - x^2}}^{\sqrt{r^2 - x^2}} 1 \, dy &= [y]_{y=-\sqrt{r^2 - x^2}}^{y=\sqrt{r^2 - x^2}} = \sqrt{r^2 - x^2} - (-\sqrt{r^2 - x^2}) = 2\sqrt{r^2 - x^2} \\ \boxed{\Xi_1} \int_0^h \int_{-r}^r 2\sqrt{r^2 - x^2} \, dx \, dz &\quad \boxed{\Xi_2} \end{aligned}$$

podobnie do jednego z przykładów rozdziału Rachunek całkowy, można znaleźć całkę

$$\begin{aligned} \int \sqrt{r^2 - x^2} \, dx &= \frac{1}{2} \left(x\sqrt{r^2 - x^2} + r^2 \arcsin \frac{x}{r} \right) + C \Rightarrow \\ \int_{-r}^r 2\sqrt{r^2 - x^2} \, dx &= \left[x\sqrt{r^2 - x^2} + r^2 \arcsin \frac{x}{r} \right]_{x=-r}^{x=r} = r^2 \arcsin 1 - r^2 \arcsin(-1) = \pi r^2; \\ \boxed{\Xi_2} \int_0^h \pi r^2 \, dz &= [\pi r^2 z]_0^h = \pi r^2 h \end{aligned}$$

1.10. OBLICZANIE CAŁEK ORAZ ROZWIĄZYWANIE RÓWNAŃ NIELINIOWYCH PRZY POMOCY METOD NUMERYCZNYCH

Obliczanie całek oznaczonych

Metody numeryczne najczęściej wykorzystujemy do aproksymacji funkcji, obliczania całek oznaczonych oraz rozwiązywania równań różniczkowych bądź ich układów, w przypadku gdy to nie jest możliwe (lub bardzo skomplikowane) aby znaleźć dokładne rozwiązanie analityczne. Natomiast za pomocą przybliżonych metod obliczeniowych (w tym ww. metod numerycznych), otrzymamy rozwiązanie na tyle bliskie do dokładnego na ile jest to potrzebne (od poziomu aproksymacji zależy ilość odpowiednich obliczeń).

Tak, dla obliczania całki oznaczonej można wykorzystać metodę prostokątów (lewych, prawych bądź średnich), metodę trapezów, metodę Simpsona lub metodę Gaussa. Wszystkie metody numeryczne obliczanie całek oznaczonych głównie polegają na obliczaniu pola figury płaskiej ograniczonej wykresem funkcji a osią Ox w odpowiednim odcinku.



Jedną z najłatwiejszych z ww. metod jest metoda prostokątów, która również nie wykazuje większych trudności przy napisaniu programu komputerowego (w np. C++, Python).

Metoda prostokątów średnich. Krótki schemat metody na przykładzie obliczania całki oznaczonej $\int_a^b f(x) dx$ funkcji $f(x)$ całkowalnej w przedziale $\langle a; b \rangle$, takiej że $f(x) \geq 0, x \in \langle a; b \rangle$ wygląda następująco:

$$f(x) \geq 0, x \in \langle a; b \rangle$$

$$\int_a^b f(x) dx \approx ?$$

- przedział całkowania $\langle a; b \rangle$ dzielimy na n odcinków jednakowej długości $l_n = \frac{b-a}{n}$ (liczbę odcinków n wybieramy sami)

$$\langle a; b \rangle = \langle a; a + l_n \rangle \cup \langle a + l_n; a + 2l_n \rangle \cup \langle a + 2l_n; a + 3l_n \rangle \cup \dots \cup \langle b - l_n; b \rangle$$

$$\langle a; b \rangle = \bigcup_{k=1}^n \langle a + (k-1)l_n; a + kl_n \rangle$$

- w każdym odcinku szukamy wartość funkcji $f(x_k)$ dokładnie w jego środku, tzn. w punktach $x_k = \frac{a+(k-1)l_n+a+kl_n}{2} = a + (k-0,5)l_n, k = \overline{1, n}$, oraz szkicujemy prostokąty, w których jedną stroną występuje odcinek $\langle a + kl_n; a + (k+1)l_n \rangle$, a długość drugiej strony równa się ww. środkowej wartości funkcji w tym odcinku
- obliczamy pola S_k wyżej zaprojektowanych prostokątów i sumujemy otrzymane wartości

$$S_k = l_n f(x_k) = l_n f(a + (k-0,5)l_n), k = \overline{1, n}$$

$$\sum_{k=1}^n S_k = \sum_{k=1}^n l_n f(a + (k-0,5)l_n) = l_n \sum_{k=1}^n f(a + (k-0,5)l_n)$$

- wtedy z zastosowania całki oznaczonej (rozdział Rachunek całkowy funkcji jednej zmiennej) wynika przybliżona równość

$$\int_a^b f(x) dx \approx \sum_{k=1}^n S_k$$

z czego otrzymujemy wzór na obliczenie przybliżone całki oznaczonej funkcji $f(x)$ o dodatnich wartościach w przedziale $\langle a; b \rangle$ (liczbę n wybieramy sami)

$$\int_a^b f(x) dx \approx l_n \sum_{k=1}^n f(a + (k-0,5)l_n)$$

$$l_n = \frac{b-a}{n}$$



Przy wybieraniu liczby odcinków n , wykorzystujemy zasadę, że wartości funkcji na końcach każdego z otrzymanych odcinków nie zbyt wiele różnią się. Ogólnie im większa liczba n , tym lepszą aproksymację otrzymamy.

Podobnie do metody prostokątów średnich, za pomocą metody prostokątów lewych / prawych otrzymamy przybliżoną wartość całki oznaczonej

$$\int_a^b f(x) dx \approx l_n \sum_{k=1}^n f(a + (k-1)l_n) \quad \left(\int_a^b f(x) dx \approx l_n \sum_{k=1}^n f(a + kl_n) \right)$$

$$l_n = \frac{b-a}{n} \quad l_n = \frac{b-a}{n}$$

Przykład Obliczyć całkę oznaczoną za pomocą metody prostokątów $\int_0^1 3x^2 dx$, wykorzystując przy tym różne liczby odcinków. Porównać otrzymane odpowiedzi z dokładnym rozwiązaniem $\int_0^1 3x^2 dx = 1$.

Rozwiązanie

Do obliczania całki wykorzystamy środkową wersję metody prostokątów.

$$\int_a^b f(x) dx \Rightarrow \begin{matrix} a = 0 \\ b = 1 \\ f(x) = 3x^2 \end{matrix}$$

$$\int_0^1 3x^2 dx$$

- przedział całkowania dzielimy na 5 odcinków:

wtedy długość odcinków wynosi $l_5 = \frac{1-0}{5} = 0,2$

$$\langle 0; 1 \rangle = \langle 0; 0,2 \rangle \cup \langle 0,2; 0,4 \rangle \cup \langle 0,4; 0,6 \rangle \cup \langle 0,6; 0,8 \rangle \cup \langle 0,8; 1 \rangle$$

x_k – środek odcinka $k = \overline{1,5}$	$f(x_k) = 3x_k^2$ $k = \overline{1,5}$	$S_k = l_5 f(x_k) = 0,6x_k^2$ $k = \overline{1,5}$
0,1	0,03	0,006
0,3	0,27	0,054
0,5	0,75	0,15
0,7	1,47	0,294
0,9	2,43	0,486

$$S = \sum_{k=1}^5 S_k = 0,99$$

Różnica pomiędzy dokładną a przybliżoną wartościami całki oznaczonej równa się

$$|1 - 0,99| = 0,01$$

- przedział całkowania dzielimy na 10 odcinków:



wtedy długość odcinków wynosi $l_{10} = \frac{1-0}{10} = 0,1$

$$\begin{aligned} \langle 0; 1 \rangle &= \langle 0; 0,1 \rangle \cup \langle 0,1; 0,2 \rangle \cup \langle 0,2; 0,3 \rangle \cup \langle 0,3; 0,4 \rangle \cup \langle 0,4; 0,5 \rangle \cup \langle 0,5; 0,6 \rangle \\ &\cup \langle 0,6; 0,7 \rangle \cup \langle 0,7; 0,8 \rangle \cup \langle 0,8; 0,9 \rangle \cup \langle 0,9; 1 \rangle \end{aligned}$$

x_k – środek odcinka $k = \overline{1,10}$	$f(x_k) = 3x_k^2$ $k = \overline{1,10}$	$S_k = l_{10}f(x_k) = 0,3x_k^2$ $k = \overline{1,10}$
0,05	0,0075	0,00075
0,15	0,0675	0,00675
0,25	0,1875	0,01875
0,35	0,3675	0,03675
0,45	0,6075	0,06075
0,55	0,9075	0,09075
0,65	1,2675	0,12675
0,75	1,6875	0,16875
0,85	2,1675	0,21675
0,95	2,7075	0,27075

$$S = \sum_{k=1}^{10} S_k = 0,9975$$

Różnica pomiędzy dokładną a przybliżoną wartościami całki oznaczonej równa się

$$|1 - 0,9975| = 0,0025$$

Rozwiązywanie równań różniczkowych

(w tym nieliniowych)

Do metod numerycznych rozwiązywania równań różniczkowych należą takie, jak metoda Eulera, metoda Rungego-Kuty, metoda Adamsa ta inne. Obejrzymy metodę Eulera (zwaną jeszcze metodą Eulera-Taylora) [16] dla równania różniczkowego rzędu pierwszego $y' = f(x, y)$ z warunkiem początkowym $y(x_0) = y_0$

$$y' = f(x, y)$$

$$\frac{dy}{dx} = f(x, y)$$

$$dy = f(x, y)dx$$

$$\int_{y_i}^{y_{i+1}} 1 dy = \int_{x_i}^{x_{i+1}} f(x, y) dx$$

$$[y]_{y_n}^{y_{n+1}} = \int_{x_i}^{x_{i+1}} f(x, y) dx$$



$$y_{n+1} - y_n = \int_{x_i}^{x_{i+1}} f(x, y) dx$$

$$y_{n+1} = y_n + \int_{x_i}^{x_{i+1}} f(x, y) dx$$

Po wykorzystaniu metody prostokątów lewych do obliczania całki $\int_{x_i}^{x_{i+1}} f(x, y) dx$ z liczbą odcinków $n = 1$, otrzymamy równość $\int_{x_i}^{x_{i+1}} f(x, y) dx \approx (x_{i+1} - x_i)f(x_i, y_i)$, a więc

$$y_{n+1} = y_n + (x_{i+1} - x_i)f(x_i, y_i)$$

Przykład Znaleźć rozwiązanie przybliżone równania różniczkowego wykorzystując metodę Eulera

$$e^y dy = (e^x + 4x) dx$$

przy warunku początkowym $y(0) = 1$. Porównać otrzymaną odpowiedź z rozwiązaniem analitycznym $y = \ln(e^x + 2x^2 + e - 1)$.

Rozwiązanie

$$\frac{dy}{dx} = e^x + 4x$$

$$\begin{aligned} \frac{dy}{dx} &= f(x, y) \\ \frac{dy}{dx} &= (e^x + 4x)e^{-y} \end{aligned} \Rightarrow f(x, y) = (e^x + 4x)e^{-y}$$

Wybierając krok $h = 0,2$ otrzymamy

$$\begin{aligned} x_0 &= 0 \\ y_0 &= y(x_0) = 1 \\ x_{i+1} &= x_i + h \\ y_{i+1} &= y_i + hf(x_i, y_i) \quad i \geq 1 \end{aligned}$$

$x_0 = 0$	$y_0 = y(x_0) = 1$	$f(x_0, y_0) = e^{-1} \approx 0,368$
$x_{i+1} = x_i + h$	$y_{i+1} = y_i + hf(x_i, y_i)$	$f(x, y) = (e^x + 4x)e^{-y}$
$x_1 = 0,2$	$y_1 = y_0 + 0,2f(x_0, y_0) \approx 1,074$	$f(x_1, y_1) \approx 0,691$
$x_2 = 0,4$	$y_2 = y_1 + 0,2f(x_1, y_1) \approx 1,212$	$f(x_2, y_2) \approx 0,92$
$x_3 = 0,6$	$y_3 = y_2 + 0,2f(x_2, y_2) \approx 1,396$	$f(x_3, y_3) \approx 1,046$
$x_4 = 0,8$	$y_4 = y_3 + 0,2f(x_3, y_3) \approx 1,605$	$f(x_4, y_4) \approx 1,09$
$x_5 = 1$	$y_5 = y_4 + 0,2f(x_4, y_4) \approx 1,823$	$f(x_5, y_5) \approx 1,085$
$x_6 = 1,2$	$y_6 = y_5 + 0,2f(x_5, y_5) \approx 2,04$	$f(x_6, y_6) \approx 1,056$
$x_7 = 1,4$	$y_7 = y_6 + 0,2f(x_6, y_6) \approx 2,251$	$f(x_7, y_7) \approx 1,016$
$x_8 = 1,6$	$y_8 = y_7 + 0,2f(x_7, y_7) \approx 2,454$	$f(x_8, y_8) \approx 0,975$
$x_9 = 1,8$	$y_9 = y_8 + 0,2f(x_8, y_8) \approx 2,65$	$f(x_9, y_9) \approx 0,937$



$x_{10} = 2$	$y_{10} = y_9 + 0,2f(x_9, y_9) \approx 2,837$	$f(x_{10}, y_{10}) \approx 0,901$
...	...	...

W powyższej tabeli dane dotyczące przybliżenia rozwiązania równania różniczkowego

$$\begin{aligned} e^y dy &= (e^x + 4x) dx \\ y(0) &= 1 \end{aligned}$$

rozwiązaniem analitycznym którego jest funkcja

$$y = \ln(e^x + 2x^2 + e - 1).$$

W następna tabela porówna wartości rozwiązania przybliżonego y_i a odpowiedni wartości dokładnego rozwiązania analitycznego:

argument x_i	aproksymacja y_{i+1}	przybliżone wartości dokładnego rozwiązania $y(x_i)$, gdzie $y = \ln(e^x + 2x^2 + e - 1)$	dokładność aproksymacji $ y_{i+1} - y(x_i) $
$x_0 = 0$	$y_0 = 1$	$y(x_0) = 1$	0
$x_1 = 0,2$	$y_1 \approx 1,074$	$y(x_1) = \ln\left(e + \sqrt[5]{e} - \frac{23}{25}\right) \approx 1,105$	0,031
$x_2 = 0,4$	$y_2 \approx 1,212$	$y(x_2) = \ln\left(e + \sqrt[5]{e^2} - \frac{17}{25}\right) \approx 1,261$	0,049
$x_3 = 0,6$	$y_3 \approx 1,396$	$y(x_3) = \ln\left(e + \sqrt[5]{e^3} - \frac{7}{25}\right) \approx 1,45$	0,054
$x_4 = 0,8$	$y_4 \approx 1,605$	$y(x_4) = \ln\left(e + \sqrt[5]{e^4} + \frac{7}{25}\right) \approx 1,653$	0,045
$x_5 = 1$	$y_5 \approx 1,823$	$y(x_5) = \ln(2e + 1) \approx 1,862$	0,039
$x_6 = 1,2$	$y_6 \approx 2,04$	$y(x_6) = \ln\left(e + \sqrt[5]{e^6} + \frac{47}{25}\right) \approx 2,07$	0,03
$x_7 = 1,4$	$y_7 \approx 2,251$	$y(x_7) = \ln\left(e + \sqrt[5]{e^7} + \frac{73}{25}\right) \approx 2,271$	0,216
$x_8 = 1,6$	$y_8 \approx 2,454$	$y(x_8) = \ln\left(e + \sqrt[5]{e^8} + \frac{103}{25}\right) \approx 2,467$	0,013
$x_9 = 1,8$	$y_9 \approx 2,65$	$y(x_9) = \ln\left(e + \sqrt[5]{e^9} + \frac{137}{25}\right) \approx 2,657$	0,007
$x_{10} = 2$	$y_{10} \approx 2,837$	$y(x_{10}) = \ln(e + e^2 + 7) \approx 2,84$	0,003
...	...	...	...

Zwiększając krok h pogorsza się sytuacja z dokładnością aproksymacji. Natomiast przy zmniejszeniu kroku h otrzymujemy wartości bliższe do odpowiednich wartości dokładnego rozwiązania równania różniczkowego.

1.11. Zakończenie

W powyższych materiałach przedstawiono podstawowe rozdziały zarówno z zakresu algebry liniowej, tak i analizy matematycznej, które są niezbędne dla różnego rodzaju badań w



fizyce, technice i innych naukach potrzebnych do studiów na kierunku Mechanika i Budowa Maszyn.

Większość rozdziałów napisano bez oparcia na konkretne źródło, ale są również takie, do napisania których było wykorzystano odpowiednie książki (cytowania zaznaczone) lub strony internetowe Wikipedia Wolna encyklopedia, wyznacznik.pl, Khan Academy, chatGPT.

1.12. Bibliografia

2. Crilly, T. (2023). Matematyka. 50 idei, które powinieneś znać. Warszawa: Wydawnictwo Naukowe PWN SA.
3. Wrociński, I. (2015). Matematyka dla logistyków. Poznań: Wyższa Szkoła Logistyki.
4. Kaczyński, A. (2021). Ćwiczenia z podstaw matematyki wyższej. Algebra liniowa. Geometria analityczna. Optymalizacja liniowa. Warszawa: Oficyna Wydawnicza Politechniki Warszawskiej.
5. Kalińska, K. (2011). Matematyka: przykłady i zadania. Włocławek: PWSZ.
6. Hącia, L. (2006). Matematyka: kurs inżynierski. Piła: PWSZ.
7. Stankiewicz, W. (2012). Zadania z matematyki dla wyższych uczelni technicznych część A/B. Warszawa: Wydawnictwo Naukowe PWN.
8. Decewicz, G., Żakowski, W. (2017). Matematyka: analiza matematyczna. Część 1. Warszawa: Wydawnictwo Naukowe PWN SA.
9. Kołodziej, W., Żakowski, W. (2017). Matematyka: analiza matematyczna. Część 2. Warszawa: Wydawnictwo Naukowe PWN SA.
10. Krywicki, W., Włodarski, L. (2019/2018). Analiza matematyczna w zadaniach, część I/II. Warszawa: Wydawnictwo Naukowe PWN.
11. Grązewicz, W. (2015) Równania różniczkowe. Materiały pomocnicze do wykładu z równań różniczkowych dla studentów Automatyki i Robotyki. Gdańsk: Politechnika Gdańska, Centrum Nauczania Matematyki i Kształcenia na Odległość.
12. Janus, J., Vladimirov, V. (2024). Równania różniczkowe zwyczajne. Pobrano z: <https://epodreczniki.open.agh.edu.pl/handbook/25>. Dostęp: 24.09.2024.
13. Radożycki, T. (2023). Rozwiązujemy zadania z analizy matematycznej. Część 2. Rzeszów: wydawnictwo Oświatowe FORSZE.
14. Frontczak, P. (2020). Metody numeryczne. Wykład nr 8: Rozwiązywanie równań różniczkowych (zwyczajnych). Pobrano z: <https://www.youtube.com/watch?v=GozT57U5IP0>. Dostęp: 29.09.2024.



Paweł Sobczak

2. FIZYKA W MECHANICE I BUDOWIE MASZYN

2.1. Wstęp

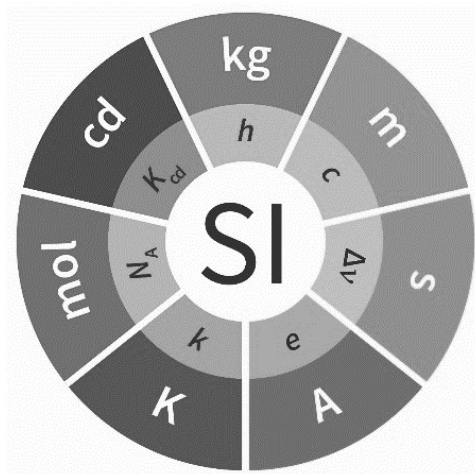
Fizyka to nauka przyrodnicza zajmująca się badaniem fundamentalnych i uniwersalnych cech świata materialnego oraz zachodzących w nim zjawisk. Jej głównym celem jest odkrywanie praw rządzących przyrodą, które mają wpływ na wszystkie procesy fizyczne. Gdy zrozumiemy prawa fizyki i zjawiska zachodzące we Wszechświecie, będziemy mogli zatasować tę wiedzę w technice i inżynierii. Niniejszy rozdział jest jedynie pewnym wstępem do fizyki którą poznacie Państwo na zajęciach. To przypomnienie pewnych wielkości fizycznych, praw i definicji. Ze względu na ograniczenia objętościowe poruszone zostały tylko wybrane aspekty, a wszystkie poznacie Państwo na zajęciach. Zajęcia z przedmiotu fizyka będą się składać z wykładów, ćwiczeń na których będziemy rozwiązywać zadania i problemy inżynierskie oraz laboratoriów na których będziecie Państwo dokonywać pomiarów wielkości fizycznych, opracowywania wyników pomiarów, analizy danych, czy oszacowania niepewności pomiarowych. Dlatego rozdział zawiera zarówno wybrane wiadomości teoretyczne, przykładowy protokół z laboratoriów i rozwiązanie zadania.

2.2. Wybrane zagadania fizyczne

2.2.1. Wielkości fizyczne, jednostki

Prawa fizyki wyrażają związki między różnymi wielkościami fizycznymi poprzez różnego rodzaju wzory, czy równania. Warto zaznaczyć, że wyróżniamy wielkości podstawowe, takie jak np. masę, czas, długość (m , t , d) oraz wielkości pochodne, np. prędkość, siłę (v , F). Wielkości pochodne powstają z pewnych działań na wielkościach podstawowych, np. $v = s/t$. Wielkości fizyczne są wyrażone w pewnych jednostkach np. kilogramach, sekundach, metrach. Wyróżniamy jednostki podstawowe (np. kg, m, s) i pochodne (np. m/s, kg·m/s²) oraz wtórne (np. J, N). W Polsce obowiązuje Międzynarodowy Układ Jednostek Miar SI¹, który stosowany jest na całym Świecie, jako podstawowy język nauki, technologii, przemysłu i ogólnoswiatowego handlu. Jest to spójny układ siedmiu jednostek podstawowych (Rys. 1.): kilograma, metra, sekundy, ampera, kelwina, mola, kandeli.

¹ <https://www.gum.gov.pl/pl/redefinicja-si/tablice-z-ukladem-si/3269,Nowe-tablice-z-ukladem-Jednostek-Miar-SI-po-redefinicji-juz-dostepne.html> Dostęp: 20.09.2024.



Rys. 1. Jednostki podstawowe układu SI.

<https://www.gum.gov.pl/pl/redefinicja-si/tablice-z-ukladem-si/3269,Nowe-tablice-z-ukladem-Jednostek-Miar-SI-po-redefinicji-juz-dostepne.html> Dostęp: 20.09.2024.

Do układu SI zaliczamy także jednostki pochodne oraz pochodne o nazwach specjalnych. Na Rys. 1. przedstawione zostały również stałe (np. h , c , e) na których podstawie są definiowane właśnie jednostki podstawowe. Fizyka opisuje obiekty różnej wielkości, dlatego jest konieczność stosowania różnych przedrostków jednostek miar (np. mili, mikro, nano), co zostało pokazane na rysunku Rys. 2. Przedrostki określające wielokrotne i podwielokrotne jednostki miar.

Czynnik	Nazwa	Symbol	Czynnik	Nazwa	Symbol
10^1	deka	da	10^{-1}	decy	d
10^2	hekto	h	10^{-2}	centy	c
10^3	kilo	k	10^{-3}	mili	m
10^6	mega	M	10^{-6}	mikro	μ
10^9	giga	G	10^{-9}	nano	n
10^{12}	tera	T	10^{-12}	piko	p

Rys. 2. Przedrostki jednostki miar.

<https://www.gum.gov.pl/pl/redefinicja-si/przedrostki-si/5474,Przedrostki-SI.html> Dostęp: 20.09.2024.

Wielkości fizyczne dzielimy na wielkości wektorowe i skalarnie. Wielkości skalarnie (np. masa, objętość, czas, ładunek, temperatura, praca), by je wyrazić wystarczy podać wartość i jednostkę. Z kolei wielkości wektorowe (np. prędkość, przyspieszenie, siła, pęd, natężenie pola), by je opisać nie wystarczy podać wartość i jednostkę. Do pełnego opisu potrzebują wartości, kierunku, zwrotu i punktu przyłożenia. Wielkości skalarnie (skalary), to porostu liczby. Zupełnie



inne podejście stosuje się do wielkości wektorowych (wektorów), które można dodawać, odejmować, rozkładać na składowe, czyli te operacje które poznaliście Państwo w szkole średniej. Wielkości wektorowe zaznacza się w tekście poprzez pogrubienie lub strzałkę nad daną wielkością. W tym miejscu jednak chciałby zwrócić uwagę na mnożenie wektorów. Wyróżniamy dwa sposoby mnożenia wektorów. Zakładamy, że mamy dwa wektory $\mathbf{a}$ i $\mathbf{b}$.

- Pierwszy sposób, tak zwany **iloczyn skalarny**. Wyniku mnożenia dwóch wektorów w sposób skalarny powstaje zwykła liczba (skalar). Iloczyn skalarny zaznacza się we wzorze poprzez kropkę ($\cdot$)

$$\mathbf{a} \cdot \mathbf{b} = |\mathbf{a}| \cdot |\mathbf{b}| \cos \alpha.$$

Przykładem może być praca (W)

$$W = \vec{F} \cdot \vec{s} \cos \alpha.$$

- Drugi sposób, to tak zwany **iloczyn wektorowy**. Wyniku mnożenia dwóch wektorów otrzymujemy nowy wektor $\mathbf{c}$, którego kierunek i zwrot określa reguła prawej dłoni lub śruby prawoskrętnej. Iloczyn wektorowy zaznacza się we wzorze poprzez znak ($\times$)

$$\mathbf{c} = \mathbf{a} \times \mathbf{b} = |\mathbf{a}| \cdot |\mathbf{b}| \sin \alpha.$$

Wektor $\mathbf{c}$ jest prostopadły do płaszczyzny wyznaczonej przez wektory $\mathbf{a}$ i $\mathbf{b}$.

2.2.2. Praca, moc energia

Pracę (W) wykonaną przez **stałą siłę** $\mathbf{F}$ definiuje się jako iloczyn skalarny tej siły $\mathbf{F}$ i wektora przesunięcia $\mathbf{s}$.

$$W = \mathbf{F} \cdot \mathbf{s} = F s \cos \alpha \quad [J]$$

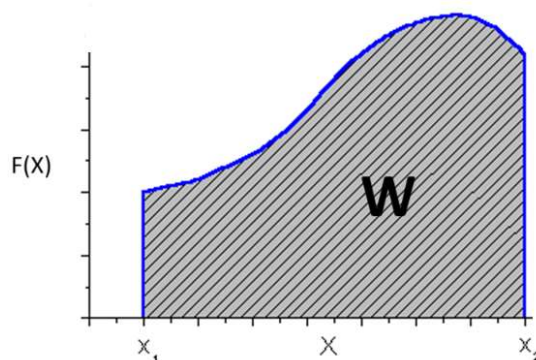
Praca jest dodatnia gdy kąt $\alpha < 90^\circ$, gdy $\alpha = 90^\circ$ praca wynosi 0 (wektor siły i przesunięcia jest prostopadły, $\cos(90^\circ)$ wynosi 0). Maksymalna wartość pracy jest gdy wektory siły i przesunięcia są równoległe i mają ten sam zwrot ($\cos(0^\circ)$ wynosi 1). Gdy jeszcze uwzględnimy siłę tarcia $\mathbf{T}$, która ma przeciwny zwrot, co siłą $\mathbf{F}$, wówczas praca ma wartość ujemną ($\cos(180^\circ)$ wynosi -1) i wtedy mówimy o pracy wykonanej przez siłę tarcia $\mathbf{T}$.

$$W = \mathbf{T} \cdot \mathbf{s} = T s \cos \alpha \quad [J]$$

Pracę (W) wykonaną przez **siłę zmienną** $\mathbf{F}$ można wyrazić poprzez całkę. Pole powierzchni pod krzywą na wykresie $F(X)$ równa się liczbowo pracy (W) wykonanej przez siłę $\mathbf{F}$ na



odcinku $x_1 - x_2$. Przykładowa praca wykonana przez siłę zmienną F została przedstawiona na Rys. 3.



Rys. 3. Pole powierzchni pod krzywą na wykresie $F(X)$ równa się pracy (W).

Opracowanie własne.

$$W = \int_{x_1}^{x_2} F(x) dx.$$

Moc (P) definiujemy jako ilość wykonanej pracy W lub przekazanej energii do czasu t w jakim została ona wykonana

$$P = \frac{W}{t} \left[\frac{J}{1s} = 1W \right].$$

Moc 1W jest bardzo mała, dlatego bardzo często stosuje się jednostkę kilowata 1kW. W przypadku urządzeń mechanicznych moc też podawana jest w koniach mechanicznych (KM). Wartość 1KM, to w przybliżeniu 0,74 kW.

Gdy ciało jest w ruchu to posiada energię kinetyczną (E_k). Rozważamy energię kinetyczną ciała poruszającego się ze stałą prędkością v . Połowę iloczynu masy ciała i kwadratu prędkości nazywamy energią kinetyczną (E_k) ciała o masie m

$$E_k = \frac{mv^2}{2} [J].$$

Gdy ciało jest na pewnej wysokości h względem podłoża, to wówczas ma nagromadzoną pracę W w postaci energii potencjalnej grawitacji (E_p). Rozważamy, że ciało znajduje się blisko powierzchni Ziemi, wówczas energię potencjalną grawitacji możemy zapisać w postaci iloczynu masy m , przyspieszenia ziemskiego g i wysokości h na jakiej znajduje się ciało

$$E_p = mgh [J].$$



Można wyróżnić jeszcze inne formy energii (np. energię potencjalną sprężystości), jednak nie będziemy ich tutaj rozważać. Warto jednak zaznaczyć, że jedną formę energii można zamienić w inną, każda przemiana jednak związana jest ze stratą. Suma energii potencjalnej grawitacji (E_p) i energii kinetycznej (E_k) nazywa się energią mechaniczną (E_m) lub całkowitą, której wartość jest stała w pewnym izolowanym układzie

$$E_m = E_p + E_k.$$

2.2.3. Wybrane zagadnienia z kinematyki

Kinematyka jest to dział fizyki zajmujący się opisem ruchu ciał (bez badania przyczyn ruchu). Z kinematyką są związane następujące pojęcia:

- **ruch** jest to zmiana położenia jednego ciała względem drugich wraz z upływem czasu,
- układ tych drugich ciał nazywamy **układem odniesienia**,
- **ruch** jest pojęciem **względny**, zależy od wyboru układu odniesienia,
- **punkt materialny** to obiekt obdarzony masą, którego rozmiary (objętość) możemy zaniedbać. Taki opis jest prawidłowy dla ruchu postępowego gdy ciała nie obracają się, ani nie wykonują drgań.

Prędkość (v) definiujemy jako zmianę położenia ciała w jednostce czasu

$$v = \frac{x - x_0}{t - t_0} \left[\frac{m}{s} \right].$$

Gdy wykreślimy na układzie współrzędnym zależność zmiany położenia w czasie, to od kąta nachylenia prostej $x(t)$ zależy wartość prędkości.

Prędkość chwilowa jest pochodną drogi względem czasu. Aby określić wartość prędkości chwilowej w danym punkcie na wykresie $x(t)$, należy narysować styczną do krzywej w tym punkcie. Współczynnik nachylenia stycznej (tangens kąta nachylenia) jest równy prędkości chwilowej w wybranym momencie

$$v = \frac{dx}{dt}.$$

Prędkość średnią oblicza się gdy ciało porusza się ruchem zmiennym, tzn. ciało ciągle zmienia wartość swojej prędkości w czasie. Można ją obliczyć jako iloraz całkowitej zmiany położenia (przesunięcia) do całkowitego czasu, w którym ta zmiana zaszła:

$$v_{sr} = \frac{\Delta x}{\Delta t} = \frac{s_1 + s_2 + \dots}{t_1 + t_2 + \dots}.$$



Prędkość średnia opisuje ogólne tempo zmiany położenia w czasie, nie uwzględnia ona chwilowych zmian prędkości w trakcie ruchu – odzwierciedla jedynie wynik końcowy, jak szybko i w jakim kierunku przemieścił się obiekt w danym czasie.

Przyspieszenie (a) określa tempo zmiany prędkości. Będziemy rozważać przyspieszenie bądź opóźnienie jednostajne. Jeżeli ciało przyspiesza lub hamuje i jego prędkość zmienia się jednostajnie z czasem, to przyspieszenie a tego ciała jest stałe. Gdy:

- ($a > 0$) to ruch jednostajnie przyspieszony (prędkość wzrasta),
- ($a < 0$) to ruch jednostajnie opóźniony (prędkość maleje).

$$a = \frac{\Delta v}{\Delta t} \left[\frac{m}{s^2} \right] \quad \text{lub} \quad a = \frac{dv}{dt}$$

By opisać prędkość (v) ciała i położenie (s) ciała w ruchu jednostajnym zmiennym wykorzystuje się wzory:

$$v = v_0 + at \left[\frac{m}{s} \right],$$

$$s = s_0 + v_0 t + \frac{at^2}{2} [m].$$

Zakładamy, że przyspieszanie ma wartość stałą $a = \text{const.}$

Swobodny spadek ciał w pobliżu powierzchni Ziemi, to przykład ruchu jednostajnie zmiennego. Przyspieszenie którego doznaje swobodnie spadające ciało nazywane jest **przyspieszeniem ziemskim** (grawitacyjnym) $g = 9,81 \left[\frac{m}{s^2} \right]$. Swobodny spadek możemy analizować na dwa sposoby: 1) korzystając z przemiany energii potencjalnej grawitacji w energię kinetyczną, 2) korzystając ze wzorów na ruch jednostajnie zmienny, gdzie $a = g$. Poniżej są wyprowadzone wzory na czas (t) spadku ciała, wysokość (h) na której znajduje się ciało oraz prędkość (v) jaką uzyska ciało przed uderzeniem w ziemię (tzw. prędkość końcową)

$$t = \sqrt{\frac{2h}{g}} [s], \quad h = \frac{gt^2}{2} [m], \quad v_k = \sqrt{2gh} \left[\frac{m}{s} \right].$$

2.2.4. Wybrane zagadnienia z dynamiki

Dynamika jest to dział fizyki zajmuje się opisem ruchu uwzględniając przyczyny ruchu.

Wyróżniamy:

- mechanikę klasyczną – stosujemy gdy ciała poruszają się z małymi prędkościami (w porównaniu z prędkością światła c),



- mechanikę kwantową – stosujemy gdy ciała poruszają się z prędkościami o porównywalnej wartości z prędkością światła c .

Pęd ciała ($\mathbf{p}$) definiujemy jako iloczyn jego masy i prędkości (wektorowej)

$$\mathbf{p} = m \cdot \mathbf{v} \left[kg \frac{m}{s} \right].$$

Definicja siły ($\mathbf{F}$) – jeżeli na ciało o masie m działa siła $\mathbf{F}$, to definiujemy ją jako zmianę pędu w czasie

$$\mathbf{F} = \frac{d\mathbf{p}}{dt} [N].$$

Zakładamy, że m ma wartość stałą. Jeżeli wprowadzimy do wzoru na siłę ($\mathbf{F}$) wyżej zdefiniowany pęd ($\mathbf{p}$) otrzymamy znany już wzór ze szkoły podstawowej, że $\mathbf{F} = m \cdot \mathbf{a}$,

$$\mathbf{F} = \frac{d(m\mathbf{v})}{dt} = \frac{dm}{dt} \mathbf{v} + m \frac{d\mathbf{v}}{dt} \qquad \mathbf{F} = m \frac{d\mathbf{v}}{dt} = m\mathbf{a} .$$

Zasady dynamiki Newtona:

- **I zasada dynamiki Newtona**

Ciało, na które nie działa żadna siła (lub gdy siła wypadkowa jest równa zero) pozostaje w spoczynku lub porusza się ze stałą prędkością po linii prostej.

- **II zasada dynamiki Newtona**

Tempo zmian pędu ciała jest równe sile wypadkowej działającej na to ciało. Dla ciała o stałej masie sprowadza się to do iloczynu masy i przyspieszenia ciała.

- **III zasada dynamiki Newtona**

Gdy dwa ciała oddziałują ze sobą, siła, którą ciało pierwsze wywiera na ciało drugie, jest równa co do wartości i przeciwnie skierowana do siły, którą ciało drugie wywiera na ciało pierwsze.

Układ inercjalny i nieinercjalny:

- I zasada dynamiki stwierdza, że jeżeli na ciało nie działa żadna siła (lub gdy $\mathbf{F}_w = 0$), to istnieje taki układ odniesienia, w którym to ciało spoczywa lub porusza się ruchem jednostajnym prostoliniowym. Taki układ nazywamy **układem inercjalnym**,
- Zasady dynamiki są spełnione w układzie inercjalnym,
- Układ odniesienia, w którym na ciało działają **siły bezwładności** nosi nazwę **układu nieinercyjnego**,
- W układach inercjalnych rządzą takie same prawa.



Siły bezwładności:

- Iloczyn masy i przyspieszenia (ze znakiem minus) nazywamy siłą bezwładności F_b

$$F_b = -m \cdot a [N],$$

- minus** oznacza, że **zwrot siły bezwładności jest przeciwny do zwrotu przyspieszenia układu**,
- Siła pojawiająca się w nieinercyjnym układzie odniesienia, będąca wynikiem przyspieszenia tego układu,
- Siła bezwładności jest **siłą pozorną**, która pojawia się tylko w układach nieinercyjnych, czyli takich, które poruszają się z przyspieszeniem. W odróżnieniu od rzeczywistych sił, siły bezwładności nie wynikają z bezpośredniego oddziaływania między ciałami, lecz są efektem przyspieszenia układu odniesienia, w którym dokonujemy obserwacji,
- Przykładem siły bezwładności jest **Siła Coriolisa**.

Pojęcia tarcia:

- Maksymalna **siła tarcia statycznego** T_s jest równa tej krytycznej sile, którą musimy przyłożyć, żeby ruszyć ciało z miejsca,
- Dla suchej powierzchni T_s spełnia dwa prawa:
 - T_s jest w przybliżeniu niezależna od wielkości pola powierzchni styku ciał;
 - T_s jest proporcjonalna do siły z jaką jedna powierzchnia naciska na drugą;
- Współczynnik tarcia statycznego** μ_s

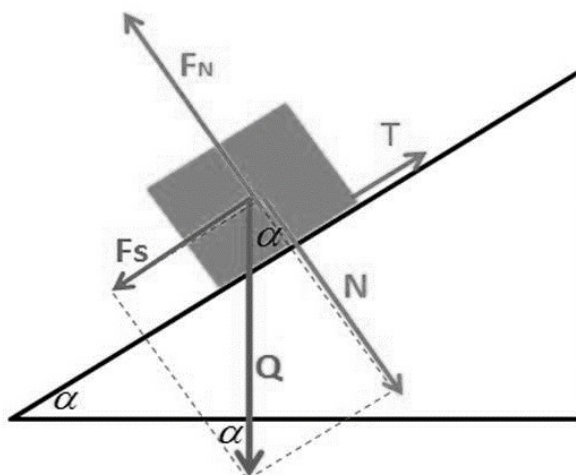
$$\mu_s = \frac{T_s}{F_N},$$

- Gdy działająca siła F jest większa od T_s to ciało zostanie wprowadzone w ruch – wówczas występuje **tarcie kinetyczne**,
- Tarcie kinetyczne oprócz dwóch praw dla tarcia statycznego spełnia jeszcze jedno prawo:
 - T_k nie zależy od prędkości względnej poruszania się powierzchni;
- Współczynnik tarcia kinetycznego** μ_k

$$\mu_k = \frac{T_k}{F_N},$$

- $\mu_k < \mu_s$ – współczynnik tarcia kinetycznego jest mniejszy od współczynnik tarcia statycznego.

Równia pochyła – analiza ruchu ciała na równi pochyłej, wyznaczenie siły wypadkowej i przyspieszenia ciała z uwzględnieniem tarcia.



Rys. 4. Rozkład sił na równi pochyłej.

<https://fizyka-kursy.pl/blog/rownia-pochyla> Dostęp: 20.09.2024.

Ciało o masie m znajduje się na równi pochyłej pod kątem α do poziomu. Na ciało działa ciężar Q (siła ciężkości) pionowo w dół. Q rozkłada się na dwie składowe, siłę nacisku N i siłę suwającą ciało z równi pochyłej F_s . Zgodnie z III zasadą dynamiki Newtona reakcją na siłę N jest siła F_N o tym samym kierunku jednak przeciwnym zwrocie. Siły N i F_N równoważą się wzajemnie. Na ciało jeszcze działa siła tarcia T . Ciężar możemy wyrazić wzorem

$$Q = m \cdot g \text{ [N]}.$$

Po rozłożeniu ciężaru (Q) na składowe mamy (N i F_s). N jest równoważone przez F_N . Pozostaje jedynie siła suwająca ciało z równi F_s , która jest pomniejszona o siłę tarcia T . Siłę wypadkową możemy zapisać

$$F_w = F_s - T.$$

$$\sin \alpha = \frac{F_s}{Q} \quad F_s = Q \sin \alpha$$

$$\cos \alpha = \frac{N}{Q} \quad N = Q \cos \alpha$$

Siła tarcia wynosi $T = f \cdot N$, gdzie f to współczynnik tarcia. Podstawiamy teraz do wzoru na siłę wypadkową wartość siły suwającej i tarcia

$$F_w = Q \sin \alpha - f \cdot N,$$

$$F_w = Q \sin \alpha - f \cdot Q \cos \alpha,$$



$$F_w = mg \cdot \sin \alpha - fmg \cdot \cos \alpha,$$

$$F_w = mg(\sin \alpha - f \cdot \cos \alpha) [N].$$

Gdy obliczyliśmy wartość siły wypadkowej dla ciała znajdującego się na równi pochyłej możemy teraz obliczyć wartość przyspieszenia jakiego dozna ciało. Przyspieszenie definiuje się jako wartość siły wypadkowej przez masę ciała

$$a = \frac{F_w}{m},$$

$$a = \frac{mg(\sin \alpha - f \cdot \cos \alpha) [N]}{m},$$

$$a = g(\sin \alpha - f \cdot \cos \alpha) \left[\frac{m}{s^2} \right].$$

Prawo powszechnego ciążenia: każde dwa ciała o masach m_1 i m_2 przyciągają się wzajemnie siłą grawitacji wprost proporcjonalną do iloczynu mas, a odwrotnie proporcjonalną do kwadratu odległości między nimi

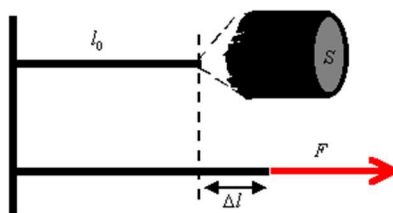
$$F_1 = F_2 = G \frac{m_1 m_2}{r^2},$$

gdzie G to stała grawitacji, $G = 6,67 \cdot 10^{-11} \text{ Nm}^2/\text{kg}^2$. Jak już zostało to wyżej wspomniane, ciężar ciała (siła grawitacji) wyraża się wzorem

$$Q = m \cdot g [N].$$

Elementy wytrzymałościowe ciał stałych

Prawo Hooke'a: Przyrost długości Δl , jakiego doznaje ciało sprężyste rozciągane siłą osiową F , jest wprost proporcjonalny do wartości tej siły i do długości początkowej l_0 ciała oraz odwrotnie proporcjonalny do jego pola przekroju poprzecznego S , a ponadto jest zależny od rodzaju materiału. Zachowanie ciała sprężystego na które działa osiowa siła rozciągająca zostało pokazane na Rys. 5.



Rys. 5. Rozciągane siłą osiową próbki, prawo Hooke'a.



Rodzaj materiału określa współczynnik proporcjonalności k lub moduł Younga E

$$\Delta l = k \frac{Fl_0}{S}$$

k – współczynnikiem proporcjonalności

E – moduł sprężystości podłużnej, moduł Younga [Pa]

Związek pomiędzy współczynnikiem proporcjonalności, a modulem sprężystości podłużnej, modulem Younga określa wzór

$$\frac{1}{k} = E.$$

Pojęcie naprężenia (σ) – określa się jako wartość liczbowa siły przypadającej na jednostkę pola przekroju poprzecznego odkształcanego ciała

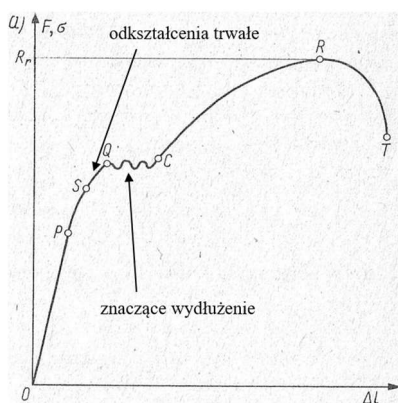
$$\sigma = \frac{F}{S} \text{ [Pa]} .$$

Prawo Hooke'a można wyrazić przy użyciu naprężenia i modułu sprężystości podłużnej, modułu Younga w następujący sposób

$$\Delta l = \frac{\sigma l_0}{E}.$$

Różne materiały są badane w maszynach wytrzymałościowych, które pozwalają zbadać własności poszczególnych materiałów, zachowanie materiału (próbki) w zależności od siły i naprężenia oraz wyznaczyć wartość naprężenia, po przekroczeniu którego następuje rozerwanie materiału. Maszyny podczas próby generują wykresy podobne do tego pokazanego na Rys. 6. Na wykresie można dostrzec charakterystyczne punkty:

- P – granica proporcjonalności,
- S – granica sprężystości,
- Q – granica płynności ,
- R – punkt, w którym jest osiągnięte naprężenie R_r zwanego **wytrzymałością na rozciąganie**,
- T – rozerwanie.



Rys. 6. Badanie wytrzymałościowe próbki materiału.

Kamiński, W., Kamiński, Z. (2020). Fizyka dla kandydatów na wyższe uczelnie techniczne tom I, PWN.

Prawo Hooke'a ma tylko zastosowanie na odcinku **OP**. R_r to wartość naprężenia, po przekroczeniu którego następuje rozerwanie materiału, która np. dla stali wynosi ok. $R_r = 3 \cdot 10^8$ Pa. Wytrzymałość materiału na rozciąganie zależy od:

- jednorodności materiału,
- obecności rys, pęcherzyków gazu,
- temperatury,
- czasu działania siły, jej zmienności, itp.

Ze względu bezpieczeństwa w obliczeniach konstrukcyjnych wprowadza się współczynnik bezpieczeństwa n . Określa on ile razy naprężenie dopuszczalne k_r , które może wystąpić w materiale, powinno być mniejsze od jego wytrzymałości R_r . Dla przeciętnej konstrukcji n przyjmuje wartość od 5 do 10.

$$k_r = \frac{R_r}{n}$$

2.3. Wprowadzenie do pracowni fizycznej

2.3.1. Niepewności pomiarowe

Każde ciało fizyczne lub zjawisko fizyczne ma wiele cech mierzalnych, które będziemy mierzyć na pracowni fizycznej i wyrażać w jednostkach układu SI. Pomiarów będziemy dokonywać za pomocą narzędzi pomiarowych, takich jak: suwmiarka, waga, termometr itp. Pomiary dzielimy na:

- **bezpośrednie** – wynik pomiaru odczytujemy bezpośrednio z narzędzia pomiarowego (np. pomiar czasu stoperem, pomiar długości suwmiarką itp.),



- **pośrednie** – wielkość którą chcemy wyznaczyć nie uzyskujemy bezpośrednio z narzędzia pomiarowego lecz obliczeń ze wzoru po uprzednim pomiarze kilku wielkości w sposób bezpośredni (np. pomiar prędkości ciała, na początku mierzymy bezpośrednio przemieszczenie s i czas trwania ruchu t , a następnie za pomocą obliczeń wyznaczamy prędkość $v = s/t$).

Każda mierzona wielkość fizyczna ma swoją wartość rzeczywistą (x_o) – oczekiwaną, z pomiaru uzyskujemy wartości pomiarowe (x). Różnica wartości rzeczywistej i zmierzonej nazywa się błędem rzeczywistym lub rzeczywistą niepewnością pomiarową

$$x_o - x = \Delta x_{rz}.$$

Warto zaznaczyć, że **wartości rzeczywiste wielkości fizycznych są nieznane**. Szacuje się przedział $x \pm \Delta x$ w którym z dużym założonym prawdopodobieństwem mieścić się wartość rzeczywista x_o . Połowę szerokości tego przedziału będziemy nazywać **niepewnością pomiarową** Δx pomiaru x . Niepewność pomiarową dzielimy na:

- **systematyczną** związaną z aparaturą pomiarową i zastosowaną metodę pomiarową,
- **przypadkową** związaną tylko i wyłącznie z przedmiotem pomiaru (przedmiotem badań), a nie z przyrządem pomiarowym.

Całkowita niepewność systematyczna wynosi:

$$\Delta x = \Delta x_d + \Delta x_k + \Delta x_o ,$$

gdzie:

Δx_d – to wartość najmniejszej podziałki na skali urządzenia (np. w przypadku linijki to 1mm),

Δx_k – to dokładność wzorowania zależna od zakresu i klasy narzędzia pomiarowego,

Δx_o – to wartość niepewności pochodzącej od obserwatora (połowa szerokości wahań wskazówki, połowa najmniejszej podziałki na skali).

Z kolei miara niepewności przypadkowej jest odchylenie standardowe pojedynczego pomiaru S_x lub odchylenie standardowe wartości średniej. Maksymalna niepewność przypadkowa wynosi

$$S_{\bar{x}} = 3S_x ,$$

gdzie:



$$S_{\bar{x}} = \sqrt{\frac{1}{n(n-1)} \sum_{i=1}^n (x_i - \bar{x})^2}$$

gdzie: $\bar{x} = \frac{1}{n} \sum_{i=1}^n x_i$ – średnia arytmetyczna (z pomiarów bezpośrednich)

x_i – wartości kolejnych pomiarów,

n – liczba pomiarów.

Zatem całkowitą niepewność pomiarową obliczamy

$$\Delta x = \Delta x_d + \Delta x_k + \Delta x_o + 3S_x.$$

Po obliczeniu całkowitej niepewności pomiarowej dokonujemy stosownego (zgodnie z zasadami zaokrąglania pomiarów) zaokrąglenia wartości niepewności oraz wyniku pomiaru (wartości średniej) do takiego miejsca znaczącego jakim jest ostatnie miejsce znaczące niepewności pomiarowej i zestawiamy

$$x = (\bar{x} - \Delta x_c) [\text{jednostka}].$$

W przypadku pomiarów pośrednich niepewność pomiarową oblicza się ze wzoru

$$\Delta z = \frac{z_{\max} - z_{\min}}{2}$$

gdzie: $z_{\max}$ i $z_{\min}$ oznaczają odpowiednio maksymalną i minimalną wartość pomiaru pośredniego.

W przypadku gdy funkcję Z da się przedstawić w postaci iloczynowej, jak poniżej

$$z = C * x_1^{a_1} * x_2^{a_2} * \dots * x_n^{a_n}$$

do obliczenia niepewności całkowitej pomiaru pośredniego można zastosować wzór uproszczony:

$$\Delta z = z \left(\left[a_1 \frac{\Delta x_1}{x_1} \right] + \left[a_2 \frac{\Delta x_2}{x_2} \right] + \dots + \left[a_n \frac{\Delta x_n}{x_n} \right] \right).$$

2.3.2. Przykładowy protokół z pracowni fizycznej

Imię i nazwisko:	Kierunek:	Podpis:
Rok: Semestr:		
Data opracowania:	Ćwiczenie prowadził:	Ocena:



I. **Temat:** Wyznaczenie współczynnika załamania światła metodą kąta najmniejszego odchylenia.

II. Wstęp:

1. Cele doświadczenia:

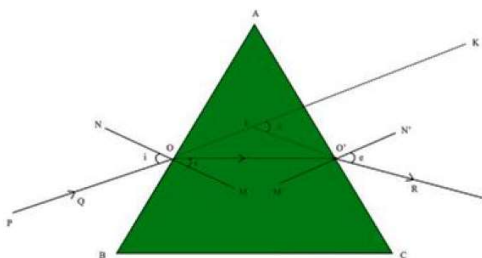
- ✓ Wyznaczenie współczynnika załamania światła metodą kąta najmniejszego odchylenia za pomocą spektrometru i pryzmatu, przy użyciu lampy rtęciowej.
 - Pomiar kąta najmniejszego odchylenia
 - Pomiar kąta łamiącego w pryzmacie

Załamanie światła następuje na granicy dwóch ośrodków przezroczystych o różnej gęstości. Miarą załamania światła jest współczynnik załamania n , który wyznacza się z zależności:

$$n = \frac{\sin \alpha}{\sin \beta}$$

α – kąt padania światła

β – kąt załamania światła



Rys. 1. Załamanie światła w pryzmacie

- Kąt padania światła (angle of incident i) – kąt QON = α
- Kąt załamania światła (angle of refraction r) – kąt MOO' = β
- Kąt najmniejszego odchylenia (angle of deviation δ) – kąt SLK = δ
- Kąt łamiący pryzmatu (=angle of refractive A) – kąt BAC = φ

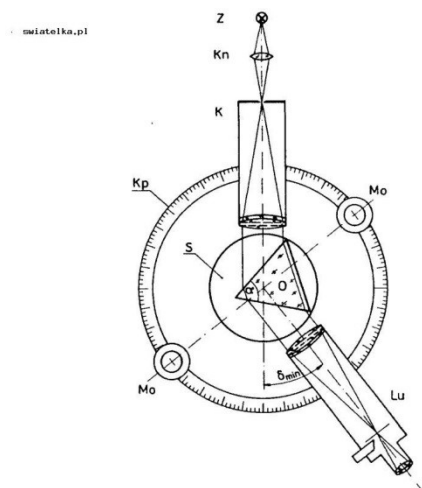
Z powyższego rysunku wynika kilka zależności:

$$\beta = \frac{1}{2}\varphi \quad \alpha = \frac{1}{2}(\delta + \varphi) \quad \delta = 2(\alpha - \beta)$$

$$n = \frac{\sin \frac{1}{2}(\delta + \varphi)}{\sin \frac{1}{2}\varphi} \quad \varphi = \frac{1}{2}(\gamma_l + \gamma_p) \quad \delta = |\varepsilon_l - \varepsilon_o| \text{ lub } \delta = |\varepsilon_p - \varepsilon_o|$$

2. Przebieg doświadczeń:

Do pomiaru kąta łamiącego i najmniejszego odchylenia używamy spektrometru, który składa się z następujących części: kolimatora, stolika, lunety i kątomierza.



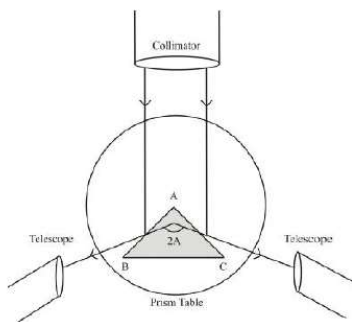
Rys. 6.5. Spektrometr: Z – lampa spektralna, Kn – kondensor, K – kolimator, Kp – koło podziałowe, S – stolik, Lu – luneta, Mo – mikroskopy odczytowe skali, O – badany pryzmat

Rys. 2. Budowa spektrometru

Doświadczenie 1.

Wyznaczenie kąta łamiącego φ w pryzmacie

Uruchom lampę rtęciową. Kolimator spektrometru umieść blisko źródła światła. Usuń szczelinę z lamy, jeżeli intensywność światła jest niska. Ustaw odpowiednią szerokość szczeliny kolimatora (szczelina musi być wąska, tak by uzyskać intensywną wiązkę światła). Następnie obróć lunetę do bezpośredniego widzenia szczeliny (na wprost kolimatora) i ustaw okular lunety, aby uzyskać wyraźny obraz szczeliny. Ustaw linię pionową na środku szczeliny. Umieść pryzmat na stole jak to zostało pokazane na Rys. 3. (chropowatą ścianką do lunety). Pryzmat powinien być ustawiony na stoliku w taki sposób, aby światło z kolimatora biegło w przybliżeniu równoległe do dwusiecznej kąta łamiącego. Wiązka światła rozdzieli się na dwie wiązki tworzące ze sobą kąt $\gamma_l + \gamma_p$. Mierzmy kąty γ_l i γ_p .



Rys. 4. Pomiar kąta łamiącego w pryzmacie

Mierzmy kąty γ_l i γ_p naprowadzając lunetę na kierunek badanego promienia i doprowadzamy do położenia, w którym obraz szczeliny pokryje się ze środkiem pola widzenia, oznaczonym pionową kreską (z lewej i prawej strony pryzmatu).

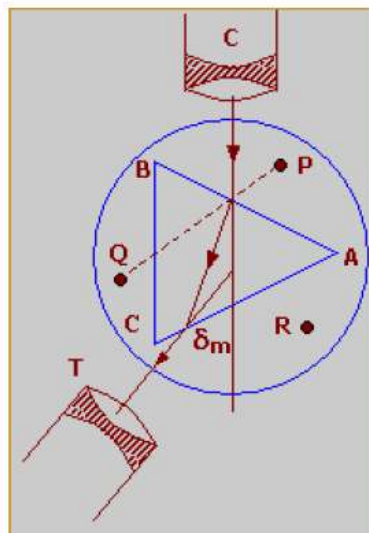
$$\varphi = \frac{1}{2}(\gamma_l + \gamma_p)$$

Doświadczenie 2.

Wyznaczanie kąta najmniejszego odchylenia δ w pryzmacie

Uruchom lampę rtęciową. Kolimator spektrometru umieść blisko źródła światła. Usuń szczelinę z lampy, jeżeli intensywność światła jest niska. Ustaw odpowiednią szerokość szczeliny kolimatora (szczelina musi być wąska, tak by uzyskać intensywną wiązkę światła). Następnie obróć lunetę do bezpośredniego widzenia szczeliny (na wprost kolimatora) i ustaw okular lunety, aby uzyskać wyraźny obraz szczeliny. Ustaw linię pionową na środku szczeliny, a następnie dokonaj pomiaru kąta ε_0 .

W celu wyznaczenia kąta najmniejszego odchylenia ustawiamy pryzmat na stoliku spektrometru tak, aby wiązka z kolimatora biegła prostopadle do dwusiecznej kąta łamiącego (zgodnie z Rys. 5). Następnie obracaj lunetę w lewą stronę, tak by poszukać czystego widma światła białego. Wybierz sobie jedną linię widmową (np. czerwoną). Obserwując widmo w lunecie obracaj stolikiem optycznym w lewo. Zaobserwuj taką sytuację, że poszczególne linie widmowe zatrzymają się, a następnie zaczną zmieniać swój kierunek. Wówczas naprowadzamy lunetę na wybraną linię widmową i doprowadzamy do położenia, w którym środek wybranej linii widmowej pokrywa się z pionową kreską, odczytujemy kąt ε_l .



Rys. 5. Pomiar kąta najmniejszego odchylenia w pryzmacie

III. Pomiary:

Doświadczenie 1.

Wyznaczenie kąta łamiącego φ w pryzmacie

Mierzymy kąty γ_l i γ_p .

$$\Delta\gamma_{l_d} = 0,5^\circ \quad \Delta\gamma_{l_o} = \frac{1}{2} \cdot 0,5 = 0,25^\circ$$

$$\Delta\gamma_{p_d} = 0,5^\circ \quad \Delta\gamma_p = \frac{1}{2} \cdot 0,5 = 0,25^\circ$$

Lp.	γ_l	γ_p
1.	57,5°	62,5°
2.	57,0°	63,0°
3.	57,5°	62,5°
4.	57,0°	63,0°
5.	57,0°	63,5°

Doświadczenie 2.

Wyznaczanie kąta najmniejszego odchylenia δ w pryzmacie

Mierzymy kąty ε_o i ε_l .

$$\Delta\varepsilon_{o_d} = 0,5^\circ \quad \Delta\varepsilon_{o_o} = \frac{1}{2} \cdot 0,5 = 0,25^\circ$$

$$\Delta\varepsilon_{l_d} = 0,5^\circ \quad \Delta\varepsilon_{l_o} = \frac{1}{2} \cdot 0,5 = 0,25^\circ$$



Lp.	ε_l	ε_o
1.	0,0°	38,0°
2.	0,0°	38,0°
3.	0,0°	38,0°
4.	0,0°	38,0°
5.	0,0°	38,5°

IV. Opracowanie wyników:**Doświadczenie 1.**

$$\bar{\gamma}_l = \frac{\gamma_1 + \gamma_2 + \gamma_3 + \gamma_4 + \gamma_5}{5} = \frac{57,5^\circ + 57,0^\circ + 57,5^\circ + 57,0^\circ + 57,0^\circ}{5} = 57,2^\circ$$

$$s_{\bar{\gamma}_l} = \frac{0,273861279}{\sqrt{5}} = 0,122474487$$

$$\Delta\gamma_l = \Delta\gamma_{l_d} + \Delta\gamma_{l_o} + 3 * s_{\bar{\gamma}_l} = 0,5 + 0,25 + 3 * 0,122474487 = 1,117423461 \approx 1,1^\circ$$

$$\bar{\gamma}_p = \frac{\gamma_1 + \gamma_2 + \gamma_3 + \gamma_4 + \gamma_5}{5} = \frac{62,5^\circ + 63,0^\circ + 62,5^\circ + 63,0^\circ + 63,5^\circ}{5} = 62,9^\circ$$

$$s_{\bar{\gamma}_p} = \frac{0,418330013}{\sqrt{5}} = 0,187082869$$

$$\Delta\gamma_p = \Delta\gamma_{p_d} + \Delta\gamma_{p_o} + 3 * s_{\bar{\gamma}_p} = 0,5 + 0,25 + 3 * 0,187082869 = 1,311248607 \approx 1,3^\circ$$

Wynik doświadczenia:

$$\bar{\gamma}_l = (57,2^\circ \mp 1,1^\circ)$$

$$\bar{\gamma}_p = (62,9^\circ \mp 1,3^\circ)$$

Doświadczenie 2.

$$\bar{\varepsilon}_l = \frac{\varepsilon_1 + \varepsilon_2 + \varepsilon_3 + \varepsilon_4 + \varepsilon_5}{5} = 0^\circ$$

$$s_{\bar{\varepsilon}_l} = \frac{0}{\sqrt{5}} = 0$$

$$\Delta\varepsilon_l = \Delta\varepsilon_{l_d} + \Delta\varepsilon_{l_o} + 3 * s_{\bar{\varepsilon}_l} = 0,5 + 0,25 + 3 * 0 = 0,750^\circ \approx 0,8^\circ$$



$$\bar{\varepsilon}_o = \frac{\varepsilon_1 + \varepsilon_2 + \varepsilon_3 + \varepsilon_4 + \varepsilon_5}{5} = \frac{38,0^\circ + 38,0^\circ + 38,0^\circ + 38,0^\circ + 38,5^\circ}{5} = 38,1^\circ$$

$$s_{\bar{\varepsilon}_o} = \frac{0,223606798}{\sqrt{5}} = 0,10000000001$$

$$\Delta\varepsilon_o = \Delta\varepsilon_{o_d} + \Delta\varepsilon_{o_o} + 3 * s_{\bar{\varepsilon}_o} = 0,5 + 0,25 + 3 * 0,10000000001 = 1,05^\circ \approx 1,1^\circ$$

Wynik doświadczenia:

$$\bar{\varepsilon}_l = (0,0^\circ \mp 0,8^\circ)$$

$$\bar{\varepsilon}_o = (38,1^\circ \mp 1,1^\circ)$$

Wyznaczanie kąta łamiącego pryzmat

$$\varphi = \frac{1}{2}(\bar{\gamma}_l + \bar{\gamma}_p) = \frac{1}{2}(57,2^\circ + 62,9^\circ) = 60,05^\circ \approx 60,1^\circ$$

$$\Delta\varphi_{min} = \frac{1}{2}(-\Delta\gamma_l - \Delta\gamma_p) = \frac{1}{2}(-1,1^\circ - 1,3^\circ) = \frac{1}{2}(-2,4^\circ) = -1,2^\circ$$

$$\Delta\varphi_{max} = \frac{1}{2}(\Delta\gamma_l + \Delta\gamma_p) = \frac{1}{2}(1,1^\circ + 1,3^\circ) = \frac{1}{2}(2,4^\circ) = 1,2^\circ$$

$$\Delta\varphi = \frac{1}{2}(\Delta\varphi_{max} - \Delta\varphi_{min}) = \frac{1}{2}(1,2^\circ - (-1,2^\circ)) = \frac{1}{2}(2,4^\circ) = 1,2^\circ$$

Kąt łamiący pryzmat wynosi:

$$\varphi = (60,1 \mp 1,2) [^\circ]$$

Wyznaczanie kąta najmniejszego odchylenia w pryzmacie

$$\delta = |\bar{\varepsilon}_l - \bar{\varepsilon}_o| = |0^\circ - 38,1^\circ| = 38,1^\circ$$

$$\Delta\delta_{min} = (-\Delta\varepsilon_l - \Delta\varepsilon_o) = (-0,8^\circ - 1,1^\circ) = -1,9^\circ$$

$$\Delta\delta_{max} = (\Delta\varepsilon_l + \Delta\varepsilon_o) = (0,8^\circ + 1,1^\circ) = 1,9^\circ$$

$$\Delta\delta = \frac{1}{2}(\Delta\delta_{max} - \Delta\delta_{min}) = \frac{1}{2}(1,9^\circ - (-1,9^\circ)) = 1,9^\circ$$

Kąt najmniejszego odchylenia wynosi:

$$\delta = (38,1 \mp 1,9) [^\circ]$$

**Wyznaczanie kąta załamania światła**

$$\beta = \frac{1}{2} \varphi = \frac{1}{2} * 60,1^\circ = 30,05^\circ \approx 30,1^\circ$$

$$\Delta\beta_{min} = \frac{1}{2}(-\Delta\varphi) = \frac{1}{2}(-1,2^\circ) = -0,6^\circ$$

$$\Delta\beta_{max} = \frac{1}{2}(\Delta\varphi) = \frac{1}{2}(1,2^\circ) = 0,6^\circ$$

$$\Delta\beta = \frac{1}{2}(\Delta\beta_{max} - \Delta\beta_{min}) = \frac{1}{2}(0,6^\circ - (-0,6^\circ)) = 0,6^\circ$$

Kąt załamania światła wynosi:

$$\beta = (30,1 \pm 0,6) [^\circ]$$

Wyznaczenie kąta padania światła

$$\alpha = \frac{1}{2}(\delta + \varphi) = \frac{1}{2}(38,1^\circ + 60,1^\circ) = 49,1^\circ$$

$$\Delta\alpha_{min} = \frac{1}{2}(-\Delta\delta - \Delta\varphi) = \frac{1}{2}(-1,9^\circ - 1,2^\circ) = -1,55^\circ \approx -1,6^\circ$$

$$\Delta\alpha_{max} = \frac{1}{2}(\Delta\delta + \Delta\varphi) = \frac{1}{2}(1,9^\circ + 1,2^\circ) = 1,55^\circ \approx 1,6^\circ$$

$$\Delta\alpha = \frac{1}{2}(\Delta\alpha_{max} - \Delta\alpha_{min}) = \frac{1}{2}(1,6^\circ - (-1,6^\circ)) = 1,6^\circ$$

Kąt padania światła wynosi:

$$\alpha = (49,1 \pm 1,6) [^\circ]$$

Wyznaczanie współczynnika załamania światła

$$\frac{\Delta n}{n} = |1| \left| \frac{\sin \Delta\alpha}{\sin \alpha} \right| + |-1| \left| \frac{\sin \Delta\beta}{\sin \beta} \right| = \frac{\sin(1,6^\circ)}{\sin(49,1^\circ)} + \frac{\sin(0,6^\circ)}{\sin(30,1^\circ)} = 0,057820968$$

$$n = \frac{\sin \alpha}{\sin \beta} = \frac{\sin(49,1^\circ)}{\sin(30,1^\circ)} = 1,507153 \approx 1,51$$

$$\Delta n = 0,00871 \approx 0,009$$

Współczynnik załamania światła wynosi:

$$n = (1,51 \pm 0,009)$$



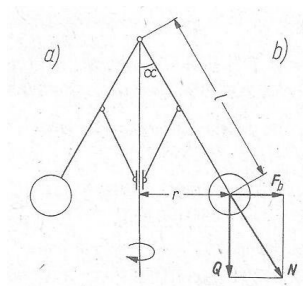
V. Wnioski:

Współczynnik załamania światła, który wyznaczyliśmy jest zgodny z wartością tablicową, bowiem dla badanego pryzmatu wynosił 1,51.

2.4. Zadania

2.4.1. Przykładowe rozwiązanie zadania

Do regulacji prędkości ruchu obrotowego są stosowane regulatory Watta. Składają się one z dwóch kul połączonych z osią obrotu za pomocą przegubowo zawieszonych prętów. Przy zwiększaniu się prędkości obrotowej regulatora – kule pod wpływem wzrastającej siły odśrodkowej bezwładności unoszą się, podnosząc osadzony przesuwnie na osi suwak. Ruch suwaka może być wykorzystany do otwierania i przamykania zaworów w przewodzie parowym lub do innej czynności regulacyjnej. Obliczyć kąt nachylenia ramion regulatora, jeżeli długość prętów, na których zawieszono są kule, wynosi 200 mm, a oś regulatora wykonuje 2 obroty na sekundę².



Dane: $l = 200 \text{ m}$, $f = 2 \text{ Hz}$, $l = 0,2 \text{ m}$, $g = 9,81 \text{ m/s}^2$

Szukane: $\alpha = ?$

Wzory:

$$Q = mg$$

$$v = \frac{2\pi r}{T}$$

$$F_B = \frac{mv^2}{r}$$

$$T = \frac{1}{f}$$



Q – ciężar ciała

F_B – siła bezwładności (siła odśrodkowa)

N – siła naciągu pręta

$$\operatorname{tg} \alpha = \frac{F_B}{Q}$$

$$\sin \alpha = \frac{r}{l} \quad r = l \sin \alpha$$

$$v = 2\pi r f \quad v = 2\pi f l \sin \alpha$$

$$F_B = \frac{m(2\pi f l \sin \alpha)^2}{r}$$

$$F_B = \frac{m(2\pi f l \sin \alpha)^2}{l \sin \alpha}$$

$$F_B = m 4\pi^2 f^2 l \sin \alpha$$

$$\operatorname{tg} \alpha = \frac{m 4\pi^2 f^2 l \sin \alpha}{mg}$$

$$\operatorname{tg} \alpha = \frac{\sin \alpha}{\cos \alpha}$$

$$\frac{\sin \alpha}{\cos \alpha} = \frac{4\pi^2 f^2 l \sin \alpha}{g}$$

$$\frac{1}{\cos \alpha} = \frac{4\pi^2 f^2 l}{g}$$

$$\cos \alpha = \frac{g}{4\pi^2 f^2 l}$$

$$\cos \alpha = \frac{9,81}{4 \cdot 3,14^2 \cdot 2^2 \cdot 0,2} = 0,31$$

Z tablic matematycznym możemy odczytać, że $\alpha = 71^\circ$.

Odpowiedź:

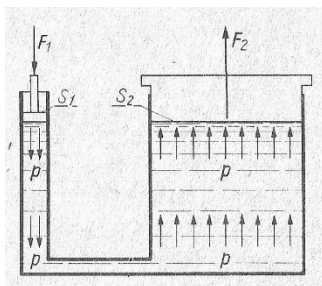
Kąt nachylenia ramion regulatora wynosi 71° .



2.4.2. Zadania

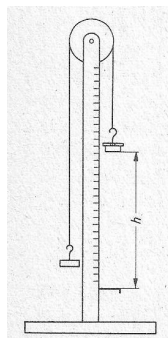
Zadanie 1

Na mniejszy tłok prasy hydraulicznej o średnicy 10 cm działa siła 80 N przesuwając go o 40 cm. Obliczyć siłę nacisku prasy oraz przesunięcie większego tłoka, jeżeli jego średnica wynosi 60 cm. Porównać pracę wykonaną przez obydwa tłoki².



Zadanie 2

Na spadkownicy Atwooda, służącej do pomiaru przyspieszeń, zrównoważono dwa odważniki po 100 g, a następnie obciążono dodatkowo jeden z nich ciężarkiem 10 g. Obliczyć przyspieszenie jakiego dozna układ ciężarków, oraz czas ruchu układu, jeśli wysokości h dodatkowo obciążonego ciężarka od zderzaka ograniczającego jego ruch wynosi 1,61 m².



Zadanie 3

Jak długo należy gotować 0,5 litra wody w czajniku bezprzewodowym o mocy 2200 W, aby doprowadzić ją do wrzenia. Temperatura początkowa wody w czajniku to 24 °C, a ciepło właściwe wody wynosi 4200 J/kg·°C. Jakie natężenie prądu elektrycznego przepływa przez

² Kamiński, W., Kamiński, Z. (2020). Fizyka dla kandydatów na wyższe uczelnie techniczne tom I, PWN



grzałkę, jeżeli czajnik jest podłączony do gniazda w domowej instalacji elektrycznej? Pomiń straty energii.

2.5. Zakończenie

Rozdział ten stanowi jedynie wprowadzenie do świata fizyki, obejmując wybrane zagadnienia teoretyczne. Wszystkie materiały związane z pełnym zakresem przedmiotu otrzymacie Państwo w formie elektronicznej podczas zajęć. W ramach tego rozdziału przedstawiono metody szacowania niepewności pomiarowych oraz zaprezentowano przykładowy protokół dotyczący wyznaczania współczynnika złamania światła n . Dodatkowo zamieszczono również przykład rozwiązania zadania z fizyki.

2.6. Bibliografia

1. Kamiński, W., Kamiński, Z. (2020). Fizyka dla kandydatów na wyższe uczelnie techniczne tom I i II. Wydawnictwo Naukowe PWN.
2. Moebs, W., J. Ling, S., Sanny J. (2018) Fizyka dla szkół wyższych. Tom 1-3. OpenStax (PDF)
3. Orear, J. (2015). Fizyka tom 1/2, Wydawnictwo Naukowo-Techniczne.
4. Holliday, D., Resnick, R., Walker, J. (2015). Podstawy fizyki Tom I-V. Wydawnictwo Naukowe PWN.
5. Jędrzejewski, J., Kruczek, W., Kujawski, A. (2017). Zbiór zadań z fizyki. Tom 1-2. Warszawa: WNT.
6. Różański, S. A., (2014) ZBIÓR ZADAŃ Z FIZYKI Z PRZYKŁADOWYMI ROZWIĄZANIAM. PWSZ w Pile
7. Szuba, S. (2010). Ćwiczenia laboratoryjne z fizyki, Poznań: Wydawnictwo Poznańskiej Księgarni Akademickiej.



Piotr Świta

3. WYTRZYMAŁOŚĆ MATERIAŁÓW W MECHANICE I BUDOWIE MASZYN

3.1. Wstęp

Wytrzymałość materiałów w programie studiów to istotny przedmiot podstawowy mający bezpośrednie powiązanie z zagadnieniami konstrukcyjnymi. Tematyka ta daje podstawy ekonomicznego doboru kształtów i wymiarów, jak również materiału każdej części maszyny lub urządzenia technicznego. Wytrzymałość materiałów to dział mechaniki, gdzie analizowane są odkształcenia, naprężenia i przemieszczenia elementów konstrukcji (maszyny lub innego urządzenia technicznego) wywołane na skutek działania obciążeń. Skutki działania obciążeń nazywane są wyężeniem, stanem odkształcenia i naprężenia lub odpowiedzią materiału. Bazą wyjściową do ustalenia takich informacji jest doświadczenie. W ramach wytrzymałości materiałów przeprowadza się badania doświadczalne, przede wszystkim badania właściwości mechanicznych materiałów. Nie zawsze jednak jest ekonomicznie uzasadnione, aby w każdym przypadku przeprowadzać doświadczenie. Dlatego wykorzystuje się metody analityczne lub numeryczne (np. Metoda Elementów Skończonych, Metoda Elementów Brzegowych). W mechanice rozpatrywane są zarówno ciała sztywne, jak i cała odkształcalne. W przypadku analizy ciał odkształcalnych podstawą omawianego zagadnienia jest teoria sprężystości i teoria plastyczności. Wytrzymałość materiałów służy analizie konstrukcji i jest wykorzystywana do obliczeń związanych z ustalaniem gabarytów elementów konstrukcyjnych. Obliczenia te mają umożliwić ekonomiczne dobranie kształtów, wymiarów i materiału każdej części ustroju konstrukcyjnego, tak aby taki ustrój mógł bezpiecznie spełniać swoją funkcję. Jeżeli uwzględniony zostanie również przewidywany czas użytkowania bezpiecznego wykorzystywania danej maszyny, wówczas będzie to zagadnienie niezawodności. W celu bezpiecznej pracy konstrukcji konieczne jest sprawdzenie trzech warunków: wytrzymałości, sztywności, stateczności.

Warunek wytrzymałości dotyczy wymagania, aby oddziaływania w analizowanym elemencie nie powodowały osiągnięcia wytrzymałości materiału, z którego jest wykonany. Złożone działanie sił na konstrukcję przyczynia się do złożonych stanów naprężenia, a wartościami reprezentatywnymi mogą być wówczas takie, które zostały wyznaczone na podstawie znanych hipotez wytrzymałościowych (np. hipoteza największej energii odkształcenia postaciowego –



hipoteza Hubera, Misesa, Hencky'ego, hipoteza największego naprężenia stycznego – hipoteza Tresca i Guesta).

Warunek sztywności sprowadza się do sprawdzenia wielkości przemieszczeń lub odkształceń analizowanego elementu. Zbyt duża ich wartość może utrudnić lub uniemożliwić prawidłową eksploatację.

Warunek stateczności dotyczy ochrony ustroju konstrukcyjnego przed gwałtowną deformacją kształtu danego elementu konstrukcyjnego.

Aby bliżej wyjaśnić zakres zagadnień działu mechaniki jakim jest wytrzymałość materiałów należy określić założenia dotyczące stanów odkształcenia i naprężenia elementów konstrukcyjnych³. Istotnym jest również aby rzeczywisty ustój konstrukcyjny wyidealizować na potrzeby obliczeniowe w zakresie schematu statycznego, warunków podparcia, połączeń poszczególnych elementów, obciążeń działających na konstrukcję oraz modelu materiału. W zależności od kierunków sił działających na konstrukcję lub kierunków naprężeń rozróżnia się proste lub złożone przypadki wytrzymałościowe. W złożonych stanach naprężeń do oceny wytrzymałości materiału wykorzystuje się hipotezy wytrzymałościowe⁴. Stawia się bowiem hipotezę, że można utworzyć funkcję określającą wyężenie, a jej argumentami są składowe stanu ośrodka ciągłego w danym punkcie (np. składowe stanu naprężenia) i parametry charakteryzujące materiał.

Wytrzymałość materiałów dotyczy również takich zagadnień jak metody energetyczne, stateczność (służąca do sprawdzenia warunku stateczności), elementy dynamiki układów sprężystych (drgań, obciążeń uderowych), zagadnienia nośności granicznej w odniesieniu do układów prętowych⁵, zmęczenie materiału², będące zjawiskiem pękania materiału pod wpływem cyklicznie zmiennych naprężeń. Bardzo istotną grupą zagadnień wytrzymałości materiałów są metody numeryczne do wyznaczania przemieszczeń, naprężeń, odkształceń i obciążeń krytycznych, z których najpopularniejszą z nich jest Metoda Elementów Skończonych⁶. Istotą tej metody jest podział złożonego układu na skończoną liczbę elementów, analiza pojedynczego elementu, a następnie ponowne złożenie wszystkich elementów w celu badania odpowiedzi całego ustroju konstrukcyjnego. Na Metodzie Elementów Skończonych

³ Bibiak-Żochowski M. (red.): *Wytrzymałość materiałów i konstrukcji. Tom I*. Oficyna Politechniki Warszawskiej, Warszawa 2013.

⁴ Dyląg Z., Jakubowicz A., Orłowski Z.: *Wytrzymałość materiałów. Tom 1*. Wydawnictwo Naukowo – Techniczne, Warszawa 1999.

⁵ Wojewódzki W.: *Nośność graniczna konstrukcji prętowych*. Oficyna Wydawnicza Politechniki Warszawskiej, Warszawa 2005.

⁶ Kleiber M.: *Wprowadzenie do Metody Elementów Skończonych*. PWN, Warszawa – Poznań 1986.



opartych jest większość dostępnych komercyjnych programów komputerowych służących analizie konstrukcji.

W niniejszym rozdziale został podany skrót informacji dotyczących sprawdzenia elementarnych warunków wytrzymałościowych, przy czym sprawdzenie to sprowadzono wyłącznie do prętów. Pominięto zakres Statyki jako osobnego działu Mechaniki Ogólnej, który Czytelnik może znaleźć w literaturze⁴. Pełna lista zagadnień wytrzymałości materiałów, wraz z szczegółowymi wyjaśnieniami, znajduje się między innymi w pracach poświęconych literaturze przedmiotu^{2 7 8 9}.

3.2. Schematy statyczne i modele materiałowe

3.2.1. Idealizacja konstrukcji

W celu przeprowadzenia obliczeń statyczno-wytrzymałościowych rzeczywisty ustrój konstrukcyjny należy odwzorować przy pomocy modelu obliczeniowego zwanego schematem statycznym. Taki schemat zawiera odzwierciedlenie geometrii danej konstrukcji oraz warunków podparcia, a także wzajemnych połączeń pomiędzy poszczególnymi elementami. Aby we właściwy sposób zdefiniować schemat statyczny należy badany model zakwalifikować do właściwej grupy. Podstawowe grupy modeli przedstawiają się następująco: pręty, układy prętowe, tarcze, ustroje powierzchniowe (płyty, powłoki), bryły.

Pręt to taki element konstrukcji, którego jeden wymiar (długość) jest wielokrotnie większy od dwóch pozostałych wymiarów (szerokości i grubości) oraz posiada oś prostą lub zakrzywioną np. belka, wał korbowy, ciągnio, łuk, sprężyna. Pręty mogą występować jako pojedyncze elementy lub w zespołach np. kratownice, ramy, ruszty.

Tarcza jest to element konstrukcji, którego jeden wymiar (grubość) jest wielokrotnie mniejszy od dwóch pozostałych (długości i szerokości), a obciążenie występuje w płaszczyźnie środkowej elementu.

Płyta jest to element konstrukcji, którego jeden wymiar (grubość) jest wielokrotnie mniejszy od dwóch pozostałych (długości i szerokości), a obciążenie jest prostopadłe do płaszczyzny środkowej elementu.

⁷ Bibiak-Żochowski M. (red.): *Wytrzymałość materiałów i konstrukcji. Tom I.* Oficyna Politechniki Warszawskiej, Warszawa 2013.

⁸ Bibiak-Żochowski M. (red.): *Wytrzymałość materiałów i konstrukcji. Tom II.* Oficyna Politechniki Warszawskiej, Warszawa 2013.

⁹ Dyląg Z., Jakubowicz A., Orłowski Z.: *Wytrzymałość materiałów. Tom 2.* Wydawnictwo Naukowo-Techniczne, Warszawa 2012.



Powłoka jest to element konstrukcji, którego jeden wymiar (grubość) jest wielokrotnie mniejszy od dwóch pozostałych (długości i szerokości), a obciążanie jest dowolnie skierowane do powierzchni środkowej elementu. Powierzchnia środkowa jest zakrzywiona (jednokrzywiznowa lub dwukrzywiznowa).

Bryła jest to element konstrukcji, którego wszystkie trzy wymiary są tego samego rzędu.

Idealizacja modelu polega również na przyjęciu odpowiednich schematów podparcia i połączeń, które są związane z warunkami ograniczającymi ruch układu mechanicznego wzdłuż poszczególnych kierunków, a więc więzami. Występowanie więzów jest równoznaczne z działaniem sił biernych, które nazywane są reakcjami.

Najczęstszymi sposobami podparcia są: przegub walcowy, przegub kulisty, podpora przegubowa przesuwana, podpora przegubowa nieprzesuwana, zawieszenie na cięgnach wiotkich, oparcie o chropowatą powierzchnię, utwierdzenie całkowite, częściowe utwierdzenie (podatne), podparcie sprężyste.

3.2.2. Idealizacja obciążeń

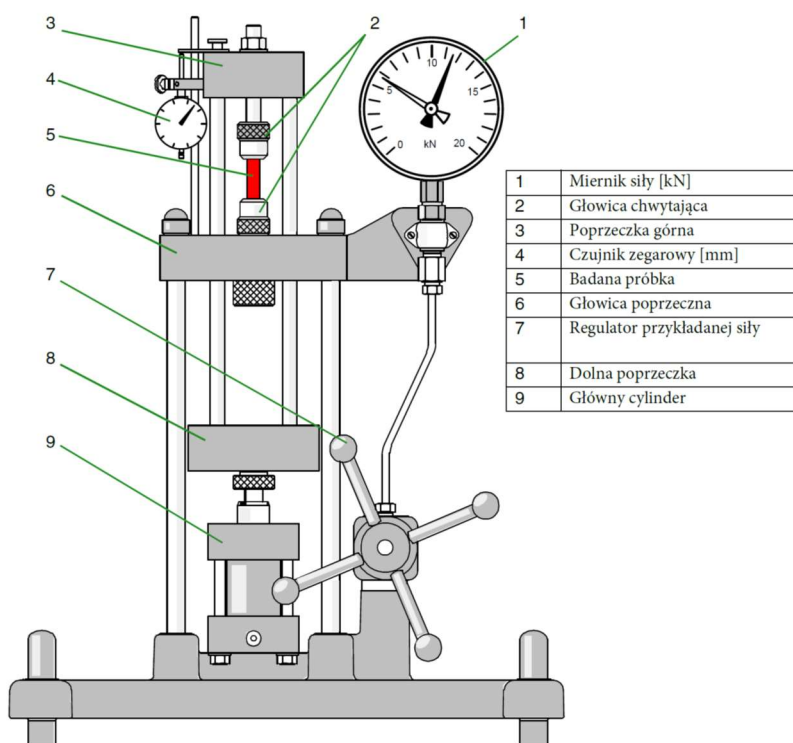
Jeżeli został określony model konstrukcji istotne jest również określenie oddziaływań, a następnie ich właściwa implementacja. Oddziaływania to siły zewnętrzne, ciężar konstrukcji, siły przekazywane przez współpracujące elementy, zmiany temperatury, tarcie, parcie cieczy, opory powietrza, oddziaływania podpór.

Pojęcie siły jest dość szerokie bowiem ze względu na charakter działania i pochodzenie rozróżnia się na stępujące rodzaje sił:

- masowe lub objętościowe, które są proporcjonalne do masy rozłożonej w objętości i działają na wszystkie punkty rozpatrywanego elementu;
- powierzchniowe, powstające przy bezpośrednim zetknięciu się jednego elementu z drugim;
- zewnętrzne, pochodzące od punktów lub elementów nie należących do rozpatrywanego układu mechanicznego;
- wewnętrzne, pochodzące od punktów lub elementów należących do rozpatrywanego układu mechanicznego;
- czynne lub inaczej obciążenia siły zewnętrzne, które mogą wywoływać ruch;
- bierne lub inaczej reakcje powstałe pod wpływem sił czynnych.

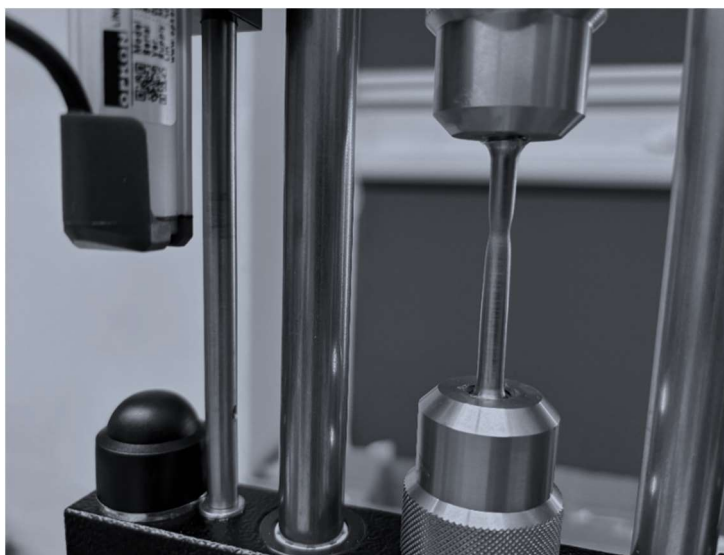
3.2.3. Idealizacja materiału

Puntem wyjścia do idealizacji danego materiału jest badanie laboratoryjne wynikające ze statycznej próby rozciągania pręta. Badanie takie odbywa się przy pomocy maszyny wytrzymałościowej (Rys. 1), która może być zintegrowana z odpowiednią wizualizacją wykresu naprężenie – odkształcenie lub siła – przemieszczenie. W maszynie umieszczony jest pręt wykonany z analizowanego materiału (np. stal, aluminium, miedź, mosiądz, tytan).



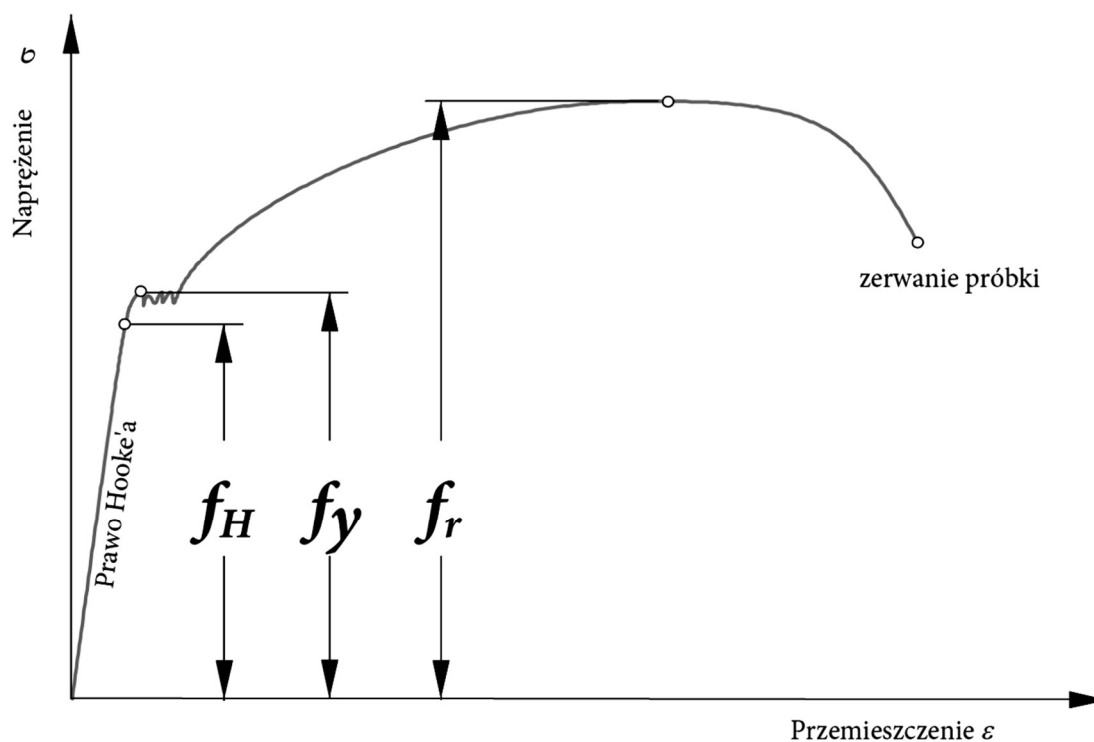
Rys. 1. Maszyna wytrzymałościowa do statycznej próby rozciągania pręta - stanowisko laboratoryjne Akademii Nauk Stosowanych w Koninie (prod. G.U.N.T. Gerätebau GmbH)

W geometrycznie zdefiniowanej próbce wytwarzany jest jednoosiowy stan naprężenia. Ten stan naprężenia jest wytwarzany przez zewnętrzne obciążenie próbki w kierunku podłużnym za pomocą siły ściskającej. Powoduje to, że w przekroju testowym próbki panuje równomierny rozkład naprężeń normalnych. Aby określić wytrzymałość materiału, obciążenie próbki jest powoli i równomiernie zwiększane, aż do zerwania próbki tj. granicy wytrzymałości (Rys. 2).



Rys. 2. Przewężenie rozciąganej próbki stalowej przed zerwaniem - stanowisko laboratoryjne Akademii Nauk Stosowanych w Koninie

To badanie dostarcza wielu informacji dotyczącej charakterystycznych wartości (punktów) na wykresie naprężenie – odkształcenie (Rys. 3), do których należą między innymi wartości naprężeń wykorzystywanych przy formułowaniu warunków wytrzymałościowych danej konstrukcji.



Rys. 3. Wykres naprężenie – odkształcenie otrzymany z badania statycznej próby rozciągania pręta stalowego wraz z charakterystycznymi naprężeniami



Materiał traktuje się zazwyczaj jako ciągły (kontinuum) i jednorodny. Oznacza to, że w dowolnym miejscu właściwości fizyczne są takie same (izotropowość materiału). Do podstawowych parametrów wytrzymałościowych materiału należą:

- granica proporcjonalności proporcjonalności f_H ;
- granica sprężystości f_s ;
- granica plastyczności f_y ;
- granica wytrzymałości f_r ;
- moduł sprężystości podłużnej (moduł Younga) E – wielkość określająca sprężystość materiału przy rozciąganiu i ściskaniu; wartość modułu Younga jest charakterystyczną dla danego materiału i wyraża zależność względnego odkształcenia liniowego ε materiału od naprężenia σ jakie w nim występuje – w zakresie odkształceń sprężystych

$$E = \frac{\sigma}{\varepsilon}; \quad (3.1)$$

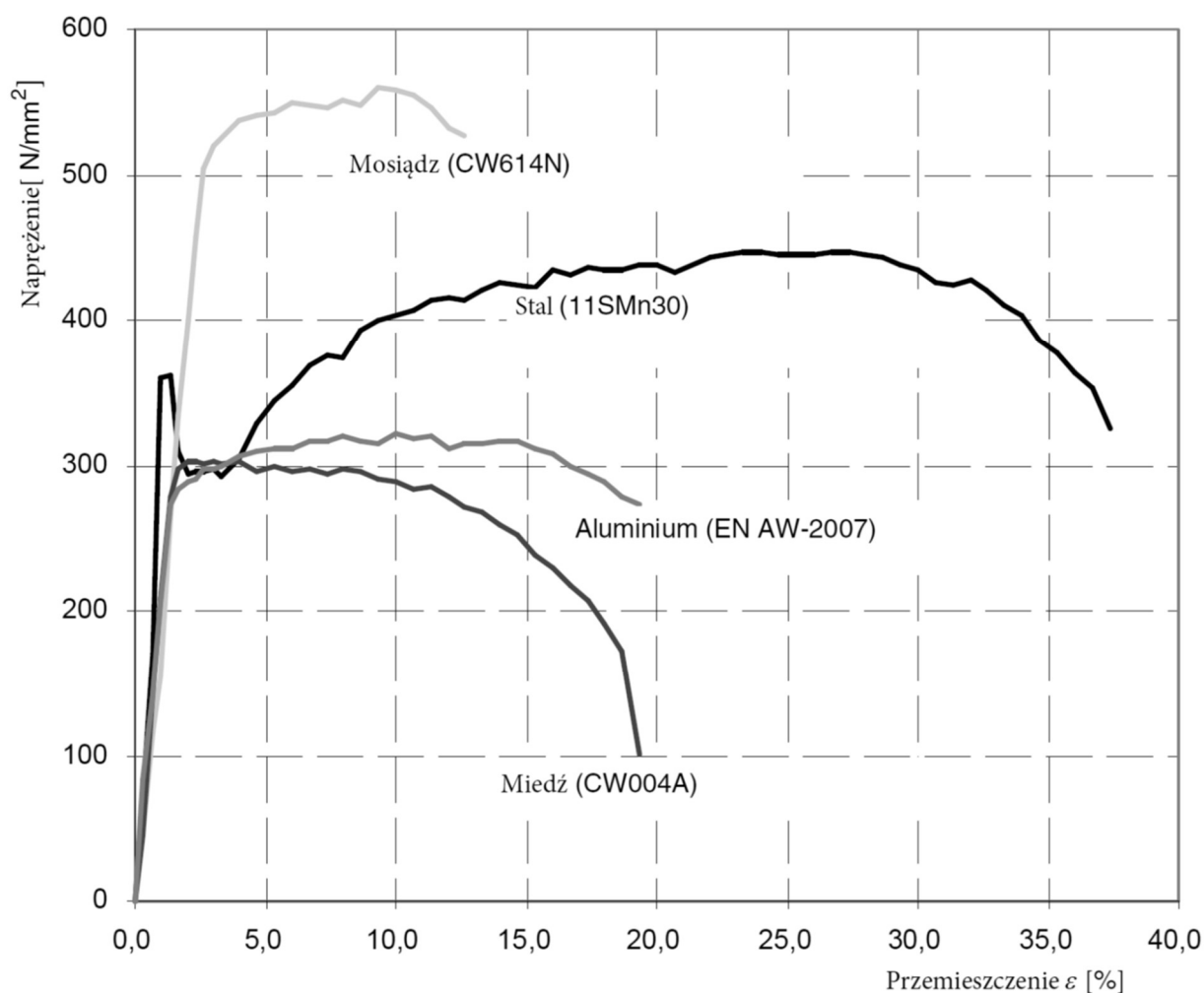
- współczynnik Poissona ν – jest to stosunek względnego odkształcenia prostopadłego do kierunku rozciągania (lub ściskania) do względnego odkształcenia w kierunku działania siły obciążającej;
- współczynnik sprężystości poprzecznej materiału G – inaczej moduł Kirchhoffa lub moduł odkształcalności postaciowej – współczynnik uzależniający odkształcenie postaciowe materiału od naprężenia, jakie w nim występuje; jest to wielkość określająca sprężystość materiału i wyraża się wzorem

$$G = \frac{\tau}{\gamma}, \quad (3.2)$$

gdzie: τ – naprężenie ścinające, γ – odkształcenie postaciowe; współczynnik sprężystości poprzecznej materiału dla materiałów izotropowych bezpośrednio zależy od modułu Younga i współczynnika Poissona:

$$G = \frac{E}{2(1 + \nu)}. \quad (3.2)$$

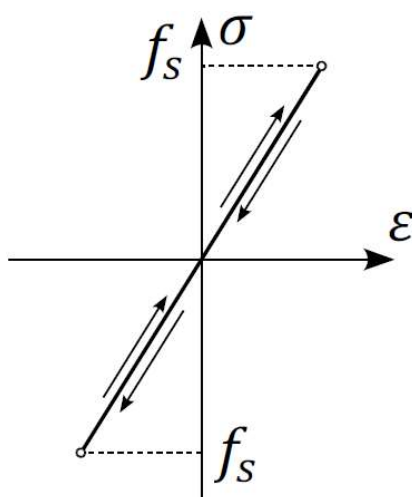
W zależności od analizowanego materiału wykres naprężenie – odkształcenie różni się znacząco (Rys. 4). W każdym jednak przypadku początkowe wartości naprężenia są wprost proporcjonalne do odkształceń. Próbka stalowa ma tendencję do przewężenia przed osiągnięciem granicy wytrzymałości i zerwaniem natomiast dla próbki wykonanej z mosiądzu zerwanie próbki odbywa się gwałtownie jednak przy znacząco większym naprężeniu.



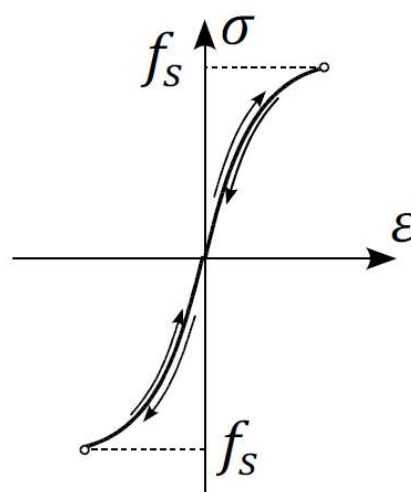
Rys. 4. Wykres naprężenie – odkształcenie otrzymany z badania statycznej próby rozciągania pręta dla różnych materiałów

Znajomość zależności pomiędzy naprężeniem a odkształceniem w warunkach laboratoryjnych daje podstawy do stworzenia modelu materiału opisanego w sposób matematyczny. Istnieje wiele modeli materiału wykorzystywanych w obliczeniach konstrukcyjnych, należy jednak pamiętać o odpowiednim doborze modelu do faktycznego materiału. Model materiału jest uproszczeniem pozwalającym uwzględnić parametry wytrzymałościowe w algorytmach obliczeniowych zarówno otrzymywanych w sposób ścisły jak i numeryczny.

Materiały sprężyste charakteryzują się właściwością, iż po usunięciu obciążenia element powraca do pierwotnego kształtu, a modelem takiego materiału może materiał liniowo – sprężysty (Rys. 5). Materiał ten poddawany zarówno naprężeniom ściskającym jak i rozciągającym, a graniczna wartość naprężeń jest równa granicy sprężystości f_s . Istnieje szereg materiałów, które wymagają zastosowania nieliniowej zależności naprężenie – odkształcenie i najprostszym z nich jest model materiału nieliniowo – sprężystego (Rys. 6).

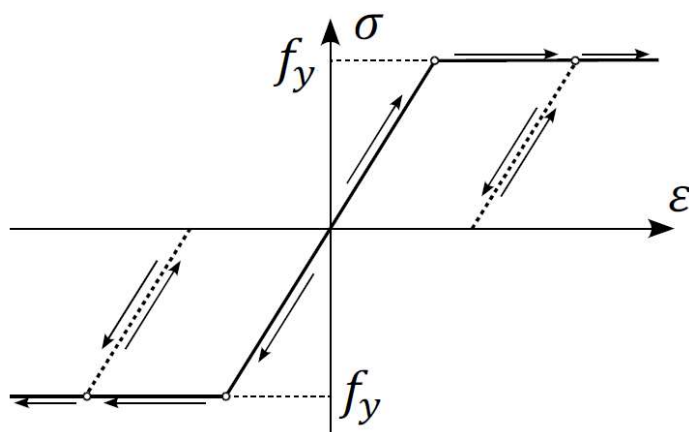


Rys. 5. Model materiału liniowo - sprężystego

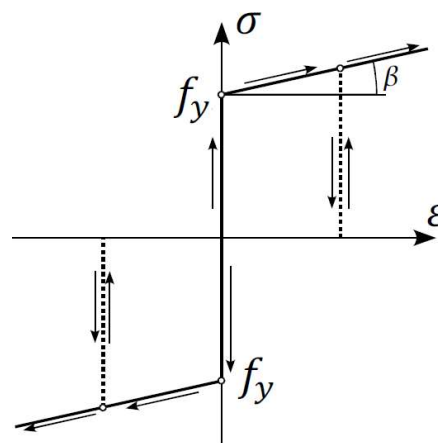


Rys. 6. Model materiału nieliniowo - sprężystego

Materiał plastyczny doznaje trwałych odkształceń, które nazywa się odkształceniami plastycznymi. Szczególnym przypadkiem modelu materiału plastycznego jest model idealnie plastyczny. W modelu tym materiał ulega uplastycznieniu przy ustalonym naprężeniu zastępczym. Najczęściej tym naprężeniem jest granica plastyczności, którą przyjmuje się równą granicy sprężystości. Istnieją dwa rodzaje modelu idealnie plastycznego: model sprężysto idealnie plastyczny (Rys. 7) i model sztywno idealnie plastyczny. W tym modelu całkowite odkształcenie składa się odkształcenia sprężystego ϵ_s oraz odkształcenia plastycznego ϵ_{pl} . Jeżeli odkształcenie sprężyste ϵ_s wynosi zero to materiał jest modelowany modelem sztywno idealnie plastycznym. Istnieją materiały (na przykład stal niskowęglowa), dla których po przekroczeniu granicy plastyczności naprężenia normalne wzrastają - materiał ulega wzmocnieniu. Materiały takie modeluje się za pomocą modeli: sztywno-plastycznego ze wzmocnieniem (rys. 8) oraz sprężysto-plastycznego ze wzmocnieniem.

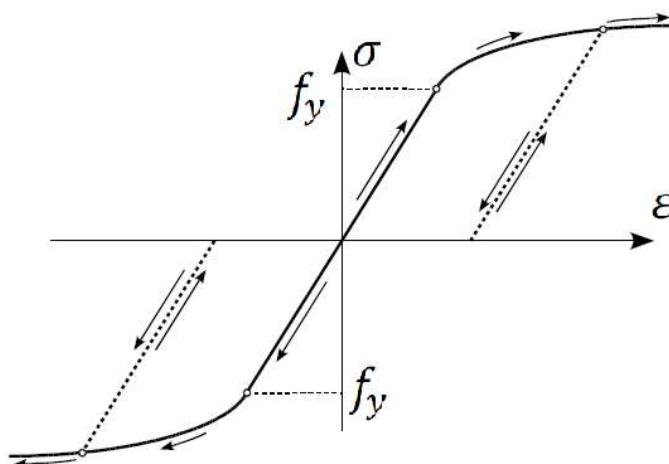


Rys. 7. Model materiału sprężysto idealnie plastycznego



Rys. 8. Model materiału sztywno-plastycznego ze wzmocnieniem

Wzmocnienie może być liniowe lub nieliniowe. Model materiału sztywno-plastycznego ze wzmocnieniem nieliniowym przedstawiono na rysunku 9. Dla modeli z liniowym wzmocnieniem tangens kąta nachylenia prostej wzmocnienia β nazywa się modułem wzmocnienia.



Rys. 9. Model materiału sprężysto-plastycznego ze wzmocnieniem (wzmocnienie nieliniowe)

3.3. Siły wewnętrzne

Ustroje konstrukcyjne podawane są oddziaływaniom, do których należą między innymi siły zewnętrzne w postaci sił powierzchniowych, przyłożonych do powierzchni elementu lub występujące na małych obszarach powierzchni, tzw. siły skupione. Ponadto na konstrukcje działają siły objętościowe wynikające z jej masy $\rho g \Delta V$. Poza siłami zewnętrznymi występują siły wewnętrzne, które są wynikiem wzajemnego wewnętrznego oddziaływania na siebie

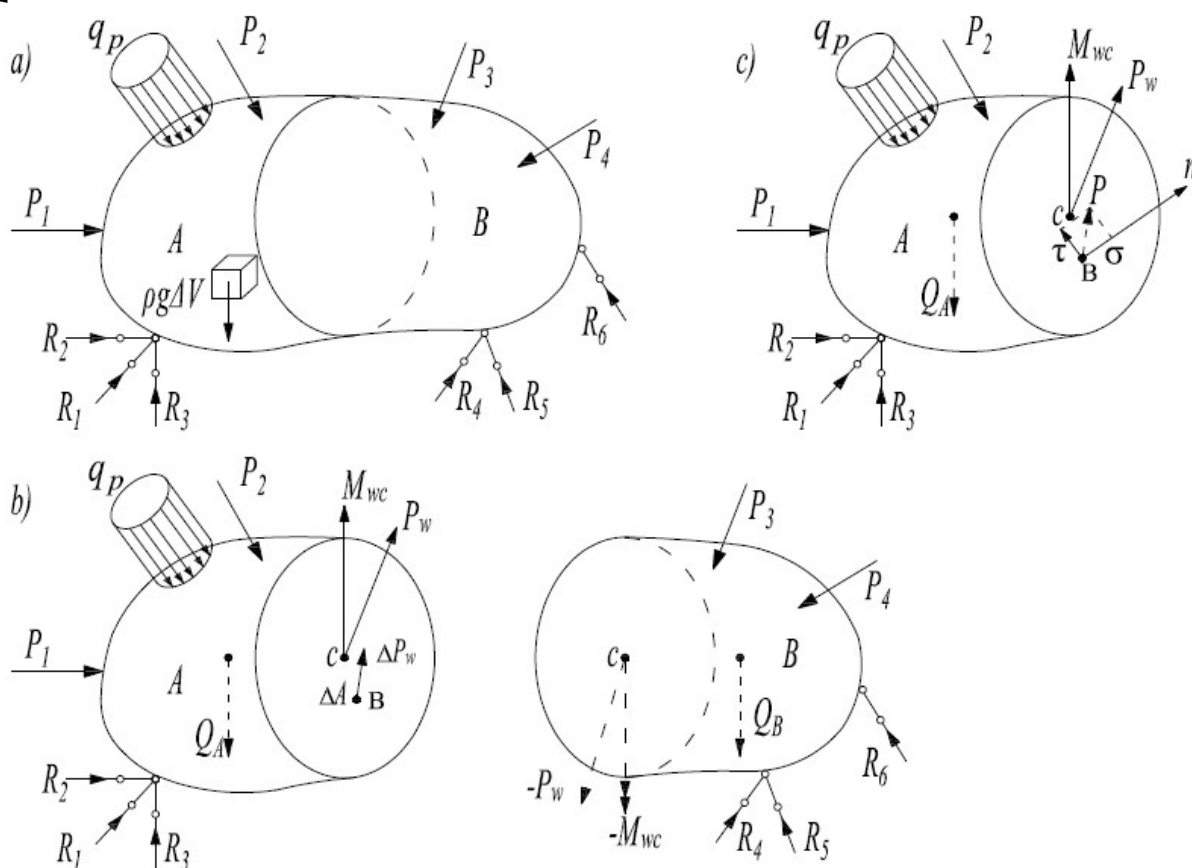


poszczególnych części materiału, należących do danego rozpatrywanego ciała odkształcalnego.

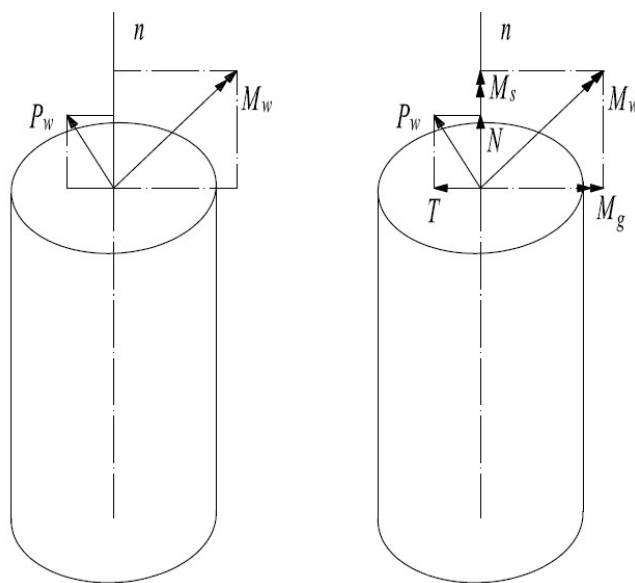
W celu wyznaczenia składowych wektora głównego sił wewnętrznych i wektora momentu rozważane jest ciało o dowolnym kształcie, podparte w sposób zapewniający jego równowagę (Rys. 10). Ciało to znajduje się pod działaniem sił zewnętrznych ($P_1, P_2, \dots, q$) oraz sił biernych czyli reakcji podporowych $R_1, R_2, \dots$

Aby wyznaczyć siły wewnętrzne w dowolnym płaskim przekroju dokonano podziału ciała na dwie części A i B. W celu utrzymania w równowadze części A należy przyłożyć do niej układ sił wewnętrznych tj. wektor główny sił P_w oraz moment M_{wC} wyznaczony względem dowolnie obranego bieguna redukcji C leżącego w płaszczyźnie rozpatrywanego przekroju. Jeżeli część A ma pozostać w równowadze, to siły $P_1, P_2, \dots, R_1, R_2, \dots, P_w$ oraz para sił o momencie M_{wC} muszą spełniać warunki równowagi. Te same warunki równowagi musi spełniać część B, gdzie siły wewnętrzne zostały zredukowane do wektora głównego sił P_w oraz momentu M_{wC} .

Najczęściej stosowanym modelem części konstrukcyjnych jest pręt. Do wyznaczenia sił wewnętrznych w pręcie rozpatruje się przekrój normalny, a środek redukcji przyjmuje się środek ciężkości przekroju (Rys. 11). Wektor główny sił P_w wewnętrznych można rozłożyć na składową N o kierunku prostopadłym do przekroju i składową T w kierunku stycznym do przekroju. Moment główny M_w można rozłożyć na kierunek normalny M_s i kierunek styczny M_g . Uzyskuje się układ uogólnionych sił przekrojowych tj. N siłę podłużną (normalną), T siłę poprzeczną (tnącą), M_s moment skręcający, M_g moment zginający.



Rys. 10. Wektor główny sił P_w oraz wektora moment M_{w0}



Rys. 11. Składowe wektora głównego sił P_w oraz wektora momentu M_{w0}

Jeśli w pręcie wystąpi tylko jedna składowa sił wewnętrznych, to jest prosty przypadek wytrzymałości pręta, do których należą:



- rozciąganie i ściskanie – występuje jedynie siła osiowa N (jeżeli jest zwrócona do wewnątrz rozpatrywanego przekroju wówczas element jest ściskany, jeżeli jest zwrócona na zewnątrz rozpatrywanego przekroju wówczas element jest rozciągany);
- ścinanie - występuje jedynie siła poprzeczna T ;
- zginanie - występuje jedynie moment zginający M_g ;
- skręcanie - występuje jedynie moment skręcający M_s .

3.4. Naprężenia, odkształcenia, przemieszczenia

Naprężenie p w danym punkcie przekroju określa się jako granicę ilorazu różnicowego

$$p = \lim \cdot \frac{\Delta P_w}{\Delta A}, \quad (3.3)$$

gdzie ΔP_w – wektor główny sił, do którego zostały zredukowane siły wewnętrzne; ΔA – wydzielona powierzchnia elementu. Wektor naprężenia p można rozłożyć na składowe:

- składową prostopadłą do powierzchni – naprężenie normalne σ ;
- składową styczną do powierzchni – naprężenie styczne τ .

W zapisie wektorowym naprężenie jest równe

$$p = \sigma + \tau, \quad (3.4)$$

natomiast wartość bezwzględna naprężenia jest równa

$$p = \sqrt{\sigma^2 + \tau^2}, \quad (3.5)$$

gdzie σ i τ to wartości algebraiczne naprężenia normalnego i stycznego.

Przemieszczenia q dowolnego punktu ciała rozumie się jako wektor, którego początkiem jest ten punkt przed odkształceniem ciała, a końcem punkt znajdujący się w nowym położeniu po odkształceniu.

Linioowe odkształcenie względne określa się jako granicę ilorazu $\Delta l/l$, gdy l dąży do zera



$$\varepsilon = \lim_{l \rightarrow 0} \frac{\Delta l}{l}, \quad (3.6)$$

gdzie $\Delta l = l' - l$, natomiast l' jest długością po odkształceniu.

3.5. Założenia wytrzymałości materiałów

Ważnym aspektem omawianego tematu jest przyjęcie dla rozpatrywanego materiału związków między naprężeniami i odkształceniami. Są to równania konstytutywne. Przykładem takich związków, dla materiałów liniowo-sprężystych, jest prawo Hooke'a.

W celu określenia podstawowych związków pomiędzy stanem naprężenia i stanem odkształcenia, w przypadku materiału o liniowych właściwościach sprężystych, można posłużyć się przykładem rozciąganego pręta prostego o stałym przekroju.

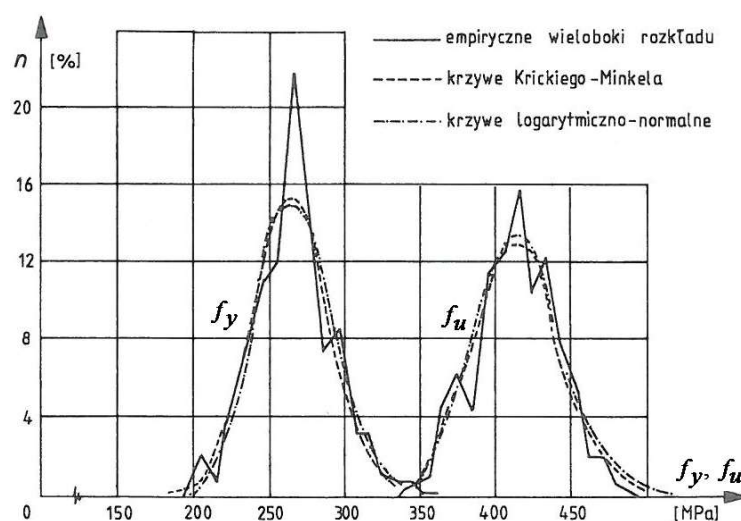
Rozpatrywane zagadnienie można scharakteryzować za pomocą warunków równowagi (jeżeli na nieważki pręt prosty o stałym przekroju działa układ sił lub jedna siła $\mathbf{P}$, która w każdym punkcie przekroju wywołuje siłę rozciągającą $\mathbf{N} = \mathbf{P}$), warunków geometrycznych (określa się wydłużenie względne, wydłużenie całego pręta) i związków fizycznych (prawo Hooke'a).

Podczas rozpatrywania zagadnień statyki przyjmuje się zwykle, że:

- Siły przyłożone do konstrukcji wzrastają od zera do swoich końcowych wartości w sposób ciągły i powolny.
- Działania poszczególnych obciążeń na dany układ sprężysty są od siebie niezależne i w związku z tym wywołane przez nie siły wewnętrzne i odkształcenia można do siebie dodawać (zasada superpozycji).
- W warunkach małych odkształceń i przemieszczeń wartości ich są proporcjonalne do obciążeń (układy liniowo-sprężyste).
- Obowiązuje zasada zeszywnienia.
- Odkształcenia konstrukcji są proporcjonalne do obciążeń i znikają po ich usunięciu.
- Poprzeczne przekroje prętów płaskie przed odkształceniem pozostają płaskie po odkształceniu (hipoteza Bernoulliego).
- Sposób przyłożenia obciążenia ma wpływ na rozkład naprężeń tylko w niewielkim obszarze.

3.6. Badania doświadczalne parametrów wytrzymałościowych materiałów

Materiałem reprezentatywny, w oparciu o który omówiono badania doświadczalne jest stal. W badaniach statystycznych stawia się hipotezę dotyczącą kształtu funkcji rozkładu, a następnie dokonuje się estymacji jego parametrów. W celu weryfikacji poczynionych założeń wykorzystuje się test zgodności założonej hipotetycznie funkcji rozkładu statystycznego z danymi statystycznymi. Do opisu losowego charakteru cech materiałowych stosuje się rozkład logarytmiczno-normalny, Gumbela lub, w przypadku małych populacji, rozkład Gaussa. Nominalna granica plastyczności f_y oznacza wartość minimalną (w sensie probabilistycznym) w odpowiednim przedziale grubości t badanego elementu. Badania rozkładu granicy plastyczności i wytrzymałości polskiej stali zostały opublikowane między innymi w pracy¹⁰ i wykorzystane w pracy¹¹; zgodnie z rysunkiem 12 przedstawione rozkłady te są niesymetryczne.



Rys. 12. Rozkłady granicy plastyczności i wytrzymałości doraźnej⁹

Dla niesymetrycznego rozkładu Gumbela gęstość prawdopodobieństwa ma postać

$$f(f_y) = ke^{-k(f_y-p)}e^{-e^{-k(f_y-p)}} \quad (3.7)$$

¹⁰ Mendera Z.: *Zagadnienia stanów granicznych konstrukcji stalowych*. Zeszyty Naukowe Politechniki Krakowskiej, Kraków 1969;7(33).

¹¹ Biegus A.: *Probabilistyczna Analiza Konstrukcji Stalowych*. PWN, Warszawa 1996.



gdzie parametrami rozkładu są: p - wartość modalna granicy plastyczności stali oraz k – współczynnik dyspersji granicy plastyczności stali.

Wartość średnia i odchylenie standardowe granicy plastyczności wyrażone są jako

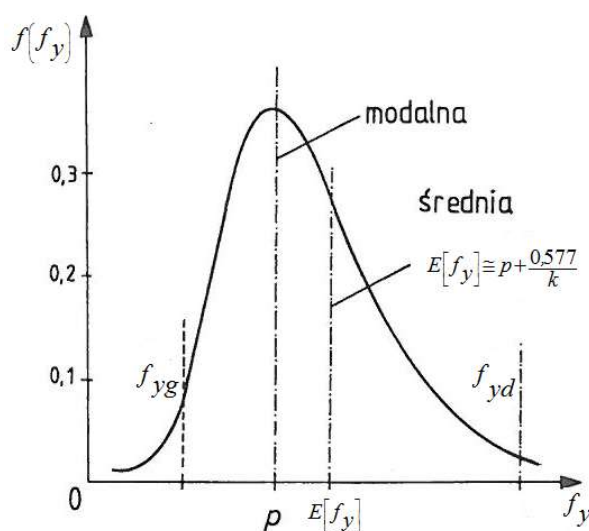
$$E[f_y] \cong p + \frac{0,577}{k}, \quad (3.8)$$

$$\sigma_{f_y} \cong \frac{1,282}{k}. \quad (3.9)$$

Na rysunku 13 został przedstawiony rozkład Gumbela dla granicy plastyczności stali, gdzie określono przedział górnej f_{yg} i dolnej f_{yd} granicy plastyczności, dla 2,5% ryzyka wystąpienia materiału słabszego, gdzie 95% wartości $f_y \in (f_{yg}, f_{yd})$ oraz

$$f_{yg} \cong p - \frac{1,3}{k}, \quad (3.10)$$

$$f_{yd} \cong p + \frac{3,7}{k}. \quad (3.11)$$



Rys. 13. Rozkład Gumbela dla granicy plastyczności stali⁹

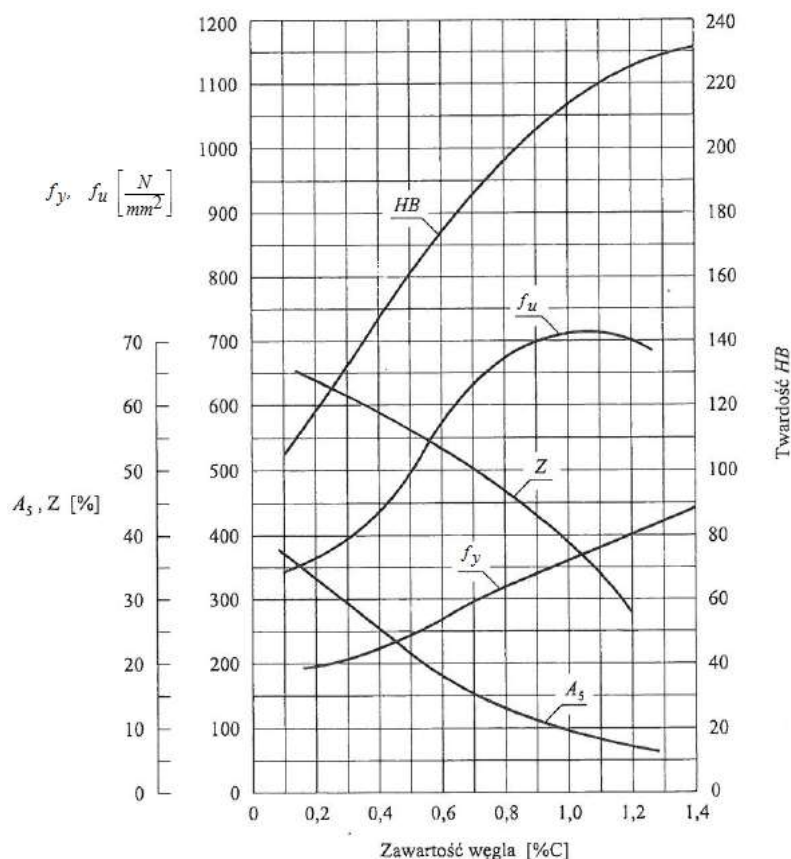
Badania parametrów wytrzymałościowych metali przeprowadza się powszechnie za pomocą statycznej próby rozciągania. Istotnym składnikiem decydującym o właściwościach wytrzymałościowych i technologicznych stali jest zawartość węgla. Na rysunku 14 przedstawiono poglądowo wpływ zawartości węgla na właściwości wytrzymałościowe stali



niestopowej, jak granica plastyczności f_y , wytrzymałość na rozciąganie f_u , twardość HB , wydłużenie całkowite próbki A_5 oraz przewężenie Z . Węgiel w stalach niestopowych zwiększa wytrzymałość stali na rozciąganie i jej twardość, natomiast pogarsza właściwości plastyczne. Losowy charakter parametrów wytrzymałościowych, takich jak granica plastyczności i moduł sprężystości jest między innymi wynikiem imperfekcji strukturalnych materiału, czyli niejednorodnego rozkładu właściwości mechanicznych w obszarze przekroju poprzecznego elementu oraz na długości wyrobu hutniczego. Właściwości materiału zależą od składu chemicznego oraz mikrostruktury w stanie pierwotnym i przerobionym.

Parametry wytrzymałościowe elementów walcowanych są związane z grubością ścianek tych elementów. Im grubszy wyrób hutniczy, tym większa jest niejednorodność struktury w kierunku grubości, gdyż ze wzrostem grubości maleje stopień zgniotu ziarn w środku wyrobu. Mniejsze zróżnicowanie granicy plastyczności stali występuje w wyrobach walcowniczych o cieńszych ściankach. Podobnie jak granica plastyczności, czy granica wytrzymałości, również moduł sprężystości podłużnej cechuje się rozrzutem losowym zależnym od grubości wyrobu.

W większości badania wytrzymałościowe stali wykonywane są na próbkach, które nie są obarczone wstępnymi imperfekcjami technologicznymi, wynikającymi z obróbki mechanicznej, czy też spawania. W przypadku obróbki mechanicznej mamy do czynienia z przecinaniem i ukosowaniem krawędzi elementów. W przypadku przecinania mechanicznego następuje zdeformowanie przecinanego brzegu oraz zmiany strukturalne w części przecinanej; inną formą przecinania jest cięcie tlenem. Wpływ cięcia tlenem na materiał przecinany ujawnia się poprzez nagrzewanie brzegów, co powoduje zmiany strukturalne i w konsekwencji podhartowanie materiału.



Rys. 14. Ilustracja wpływu zawartości węgla na właściwości mechaniczne stali węglowych¹²

Zmiana ta może wywoływać rysy lub pęknięcia na brzegach i dalej rozprzestrzeniać się w materiale, dlatego cięcie tlenem musi być uzależnione do gatunku stali, a tym samym od zawartości węgla, aby ograniczyć to zjawisko. Zróżnicowanie właściwości mechanicznych ma miejsce również w przypadku kształtowania połączeń spawanych. W wyniku spawania powstają zmiany strukturalne wokół spoiny powodujące zmiany strukturalne złącza, a także naprężenia oraz deformacje spawalnicze. Termiczny proces spawania obejmujący okresy nagrzewania i stygnięcia stanowi nieustabilizowany proces termiczny, gdzie pomiędzy stopiwem i materiałem rodzimym powstaje w procesie spawania strefa wpływu ciepła. Jest ona zróżnicowana pod względem struktury na skutek cyklu cieplnego spawania. Wraz ze zmianami struktury występują zmiany granicy plastyczności, wytrzymałości, twardości i modułu sprężystości. W wyniku spawania powstają dodatkowe naprężenia spawalnicze, które są tylko częściowo zmniejszane na skutek obróbki cieplnej; jednak z powodu dużych wymiarów oraz znacznego ciężaru elementów czynność ta bywa pomijana. Wadliwość złączy spawanych można podzielić na szereg grup, z których należy wymienić: pęknięcia, pustki (przestrzeń

¹² Gosowski. B., Kubica E.: *Badania Laboratoryjne z Konstrukcji Metalowych*. Oficyna Wydawnicza Politechniki Wrocławskiej 2001.



w spoinie wypełniona gazem), wtrącenia stałe (obce ciała), przyklejenia i braki przetopu, niezgodności spawalnicze kształtu.

3.7. Warunek wytrzymałości

W zależności od występujących sił wewnętrznych w przekroju poprzecznym pręta, wyznacza się poszczególne wartości naprężenia. Mogą być to następujące przypadki naprężenia:

- a) naprężenia ściskające lub rozciągające – wywołane siłą osiowa (normalna) N ;
- b) zginające i ścinające – w przekroju występuje moment zginający M_g i siła tnąca T ;
- c) skręcające – w przekroju występuje wyłącznie moment zginający M_s .

3.7.1. Ściskanie lub rozciąganie

Naprężenia przy ściskaniu lub rozciąganiu wyznacza się wg wzoru

$$\sigma_i = \frac{N_i}{A_i}, \quad (3.12)$$

gdzie

gdzie N_i – siła normalna (ściskająca lub rozciągająca); A – pole przekroju pręta. Warunek wytrzymałości polega na sprawdzeniu ekstremalnych wartości naprężeń w danym przekroju i porównaniu z wartościami dopuszczalnymi σ_{dop} z przyjętym modelem materiału

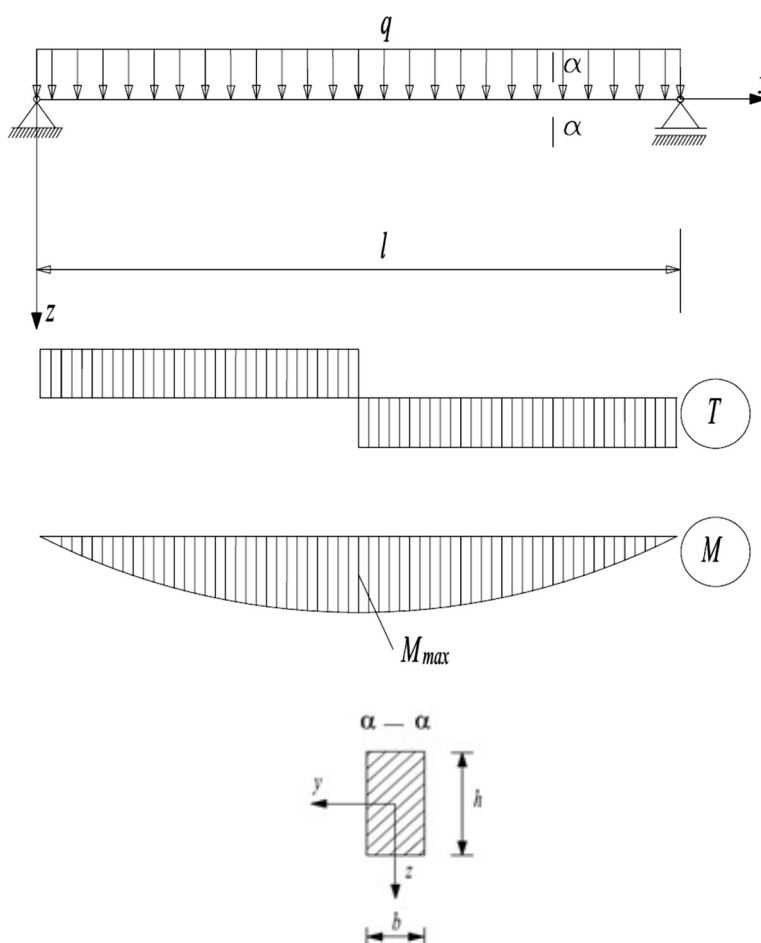
$$\sigma \leq \sigma_{dop}. \quad (3.13)$$

3.7.2. Zginanie i ścinanie

Zginanie i ścinanie zostanie przedstawione na przykładzie belki prostej obciążonej równomiernie obciążeniem q (Rys. 15). Belka ma przekrój prostokątny o wymiarach b i h .

Oprócz schematu statycznego belki przedstawiono wykresy sił wewnętrznych tj. sił poprzecznych T oraz momentu zginającego M .

Element poddany działaniu momentu zginającego ulega jednoczesnemu skróceniu i wydłużeniu w poszczególnych warstwach przekroju. Wykres momentu zginającego został narysowany po stronie rozciąganej, to w warstwie górnej tj. powyżej osi y następuje ściskanie, a poniżej tej osi rozciąganie, natomiast naprężenia na osi y mają wartość równo zero. Oś y możemy zatem nazwać osią obojętną.



Rys. 15. Schemat statyczny belki, rozkład sił wewnętrznych oraz przekrój poprzeczny $\alpha - \alpha$

W celu skorzystania z wzorów na naprężenia przy zginaniu lub ścinaniu należy określić niezbędne charakterystyki geometryczne przekroju. Podstawowe charakterystyki geometryczne wykorzystywane w wzorach wytrzymałościowych przy zginaniu i ścinaniu są następujące²:



- momenty statyczne figury płaskiej jako momenty powierzchniowe pierwszego rzędu względem odpowiednich osi układu współrzędnych tj. S_y i S_z ;
- momenty bezwładności figury płaskiej jako momenty powierzchniowe drugiego rzędu względem odpowiednich osi układu współrzędnych tj. J_y i J_z ;

W przypadku przekroju prostokątnego momenty bezwładności przekroju względem osi układu współrzędnych można zapisać następująco

$$J_y = \frac{b \cdot h^3}{12}, \quad (3.14)$$

$$J_z = \frac{h \cdot b^3}{12}. \quad (3.15)$$

Dla przypadku zginania, w którym występuje jedna składowa momentu zginającego (np. M_y) naprężenia normalne określa się jako

$$\sigma_x = \frac{M_y}{J_y} \cdot z, \quad (3.16)$$

gdzie M_y – moment zginający; J_y – moment bezwładności przekroju względem osi y ; z – odległość rozpatrywanego punktu przekroju do osi obojętnej.

Naprężenia styczne od siły poprzecznej w kierunku osi y i z określa się odpowiednio

$$\tau_{xy} = \frac{T_y \cdot \bar{S}_z}{J_z \cdot b(y)}, \quad \tau_{xz} = \frac{T_z \cdot \bar{S}_y}{J_y \cdot b(z)}, \quad (3.17)$$

gdzie T_y i T_z – siły poprzeczne w kierunku osi y i z ; $\bar{S}_y$ i $\bar{S}_z$ – bezwzględna wartość momentów statycznych względem osi y i z ; $b(y)$ i $b(z)$ – szerokość przekroju w rozpatrywanym punkcie.

Dla przypadku zginania ukośnego, w którym występują dwie składowe momentu zginającego (M_y i M_z) naprężenia normalne określa się jako



$$\sigma_x = \frac{M_y}{J_y} \cdot z - \frac{M_z}{J_z} \cdot y, \quad (3.18)$$

gdzie M_y – moment zginający; J_y – moment bezwładności przekroju względem osi y ; z – odległość rozpatrywanego punktu przekroju do osi obojętnej.

Warunek wytrzymałości polega na sprawdzeniu ekstremalnych wartości naprężeń i porównanie z naprężeniami dopuszczalnymi

$$\sigma_x \leq \sigma_{dop}, \quad (3.19)$$

$$\tau \leq \tau_{dop}. \quad (3.20)$$

3.7.3. Skręcanie

Zagadnienia skręcania należy rozpatrywać w zależności od przekroju. Rozróżnia się czyste skręcanie prętów o przekrojach kolistych i czyste skręcanie prętów o przekrojach niekolistych. Warunek wytrzymałości jest następujący

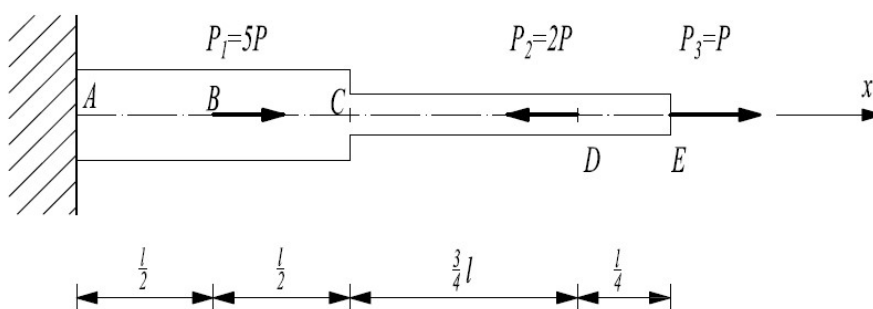
$$\tau_o \leq \tau_{dop}. \quad (3.21)$$

gdzie τ_o – naprężenia styczne przy skręcaniu.

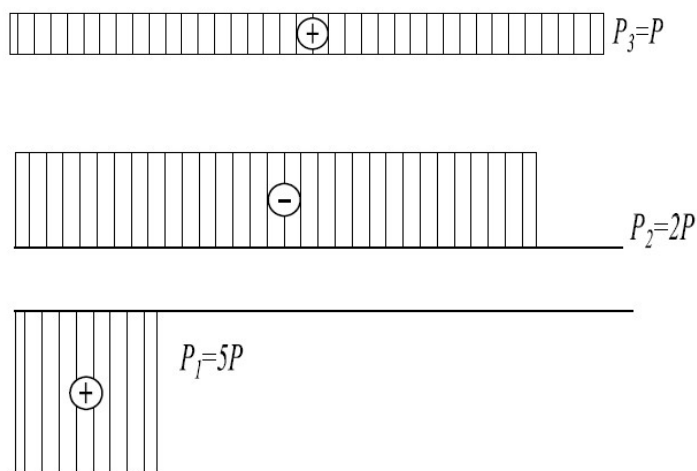
Przykład nr 3.1

Dany jest nieważki pręt, obciążony siłami działającymi wzdłuż jego osi zgodnie z rysunkiem 16. Sprawdzić warunek wytrzymałości oraz całkowite wydłużenie pręta, mając dane naprężenie dopuszczalne $\sigma_{dop} = \frac{4P}{A}$ i długość pręta l . Przekrój poprzeczny pręta jest zmienny skokowo i wynosi $A_{AC} = 2A$, $A_{CE} = A$, moduł sprężystości materiału E .

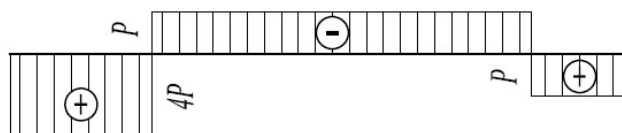
a)



b)



c)



Rys. 16. Schemat statyczny pręta (a); Wykresy naprężeń od poszczególnych sił osiowych (b); Wykres naprężeń dla sumarycznych (c)

Zakłada się, że ekstremalne naprężenia normalne nie przekroczą naprężeń równych granicy proporcjonalności. Wówczas całkowite wydłużenie określone jest wzorem

$$\Delta l = \sum_{i=1}^n \frac{N_i \cdot l_i}{E_i \cdot A_i} .$$

Naprężenia normalne w przekrojach prostopadłych do osi pręta określone są wzorem



$$\sigma_i = \frac{N_i}{A_i} ,$$

gdzie N_i – siła podłużna na danym odcinku, l_i oraz A_i – długość i pole powierzchni przekroju poprzecznego tego odcinka, E_i - moduł sprężystości podłużnej.

Ponieważ element jest wykonany z materiału jednorodnego dla każdego odcinka pręta $E_i = E$.

Wykorzystuje się wykres sił podłużnych od wszystkich obciążeń, a całkowite wydłużenie wynosi

$$\Delta l = \frac{N_{\alpha}^{AB} \cdot l_{AB}}{E \cdot A_{AB}} + \frac{N_{\alpha}^{BC} \cdot l_{BC}}{E \cdot A_{BC}} + \frac{N_{\alpha}^{CD} \cdot l_{CD}}{E \cdot A_{CD}} + \frac{N_{\alpha}^{DE} \cdot l_{DE}}{E \cdot A_{DE}} ,$$

$$\Delta l = \frac{4P \cdot \frac{1}{2}l}{E \cdot 2A} + \frac{-P \cdot \frac{1}{2}l}{E \cdot 2A} + \frac{-P \cdot \frac{3}{4}l}{E \cdot A} + \frac{P \cdot \frac{1}{4}l}{E \cdot A} ,$$

$$\Delta l = \frac{1}{4} \cdot \frac{P \cdot l}{E \cdot A} .$$

Ekstremalne (minimalne) naprężenia normalne wywołane siłą ściskającą wystąpią na odcinku CD i wynoszą

$$\sigma_{min} = \sigma_{CD} = \frac{N_{CD}}{A_{CD}} = \frac{-P}{A} ,$$

natomiast ekstremalne naprężenia wywołane siłą rozciągającą wystąpią na odcinkach AB i DE

$$\sigma_{AB} = \frac{N_{AB}}{A_{AB}} = \frac{4P}{2A} = \frac{2P}{A} , \quad \sigma_{DE} = \frac{N_{DE}}{A_{DE}} = \frac{P}{A} ,$$

gdzie maksymalna ich wartość występuje na odcinku AB tj. $\sigma_{max} = \sigma_{AB}$. Warunek wytrzymałości ma postać

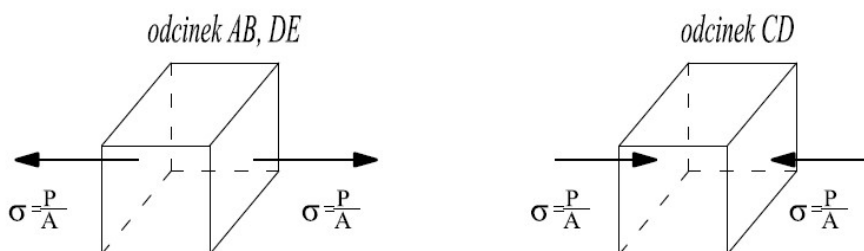


$$\sigma_{max} \leq \sigma_{dop},$$

zatem warunek ten jest spełniony ponieważ

$$\frac{2P}{A} < \frac{4P}{A}.$$

Przy naprężeniach osiowych (ściskających lub rozciągających) ich wartości są stałe w całym przekroju poprzecznym, a ich kierunek jest równoległy do osi sił podłużnych N_i . Jeżeli na pręt działają wyłącznie siły podłużne jest to jednoosiowy stan naprężenia. Elementarny wycinek pręta ograniczony krawędziami prostokątnymi i równoległymi do jego osi x podlega wyłącznie naprężeniom mającym kierunek tej osi.

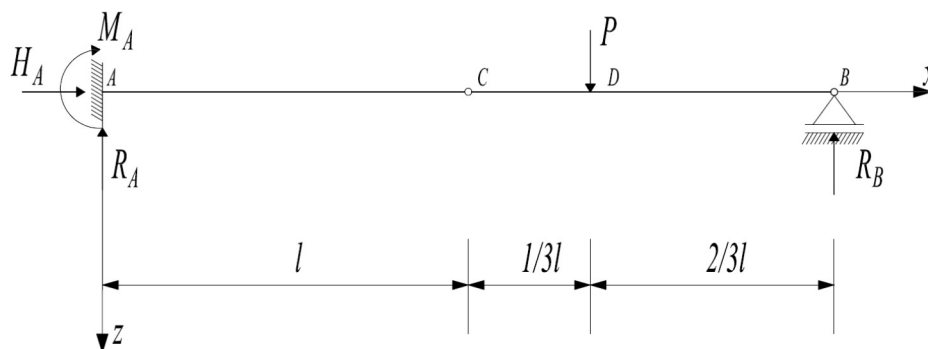


Rys. 17. Naprężenia w elementarnego wycinka pręta

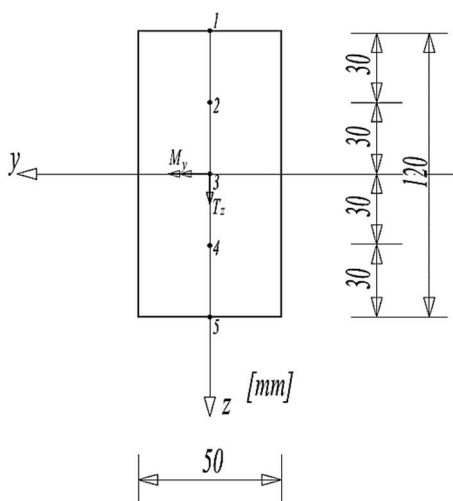
Koniec przykładu.

Przykład nr 3.2

Dana jest belka obciążona obciążeniem skupionym działającym w płaszczyźnie xz zgodnie z rysunkiem 18. Belka posiada przegub w punkcie C. Przekrój belki jest prostokątny o wymiarach $b = 50 \text{ mm}$ i $h = 120 \text{ mm}$ zgodnie z rysunkiem 19. Wyznaczyć wartość naprężenia stycznego i naprężenia normalnego przy zginaniu w poszczególnych punktach przekroju, a także sprawdzić warunek wytrzymałości przy zginaniu. Przyjąć następujące dane: długość pręta $l = 1,5 \text{ m}$, obciążenie $P = 9 \text{ kN}$, $\sigma_{dop} = 100 \frac{\text{N}}{\text{mm}^2}$.



Rys. 18. Schemat statyczny belki z przegubem w punkcie C



Rys. 19. Przekrój poprzeczny belki

Reakcje podpór

W celu rozwiązania zadania należy w pierwszej kolejności wyznaczyć wartości reakcji podporowych z równań równowagi jako

$$1) \sum_{i=1}^n P_{ix} = 0, \quad H_A = 0,$$

$$2) \sum_{i=1}^n P_{iz} = 0, \quad R_A + R_B - P = 0,$$

$$3) \sum_{i=1}^n M_{iA} = 0, \quad -R_B \cdot 2l + P \cdot \frac{4}{3}l + M_A = 0.$$

Z uwagi na występowanie przegubu w punkcie C, możliwe jest zapisanie czwartego warunku ponieważ moment w przegubie jest równy zero, stąd



$$4) M_C = 0, \quad -R_B \cdot l + P \cdot \frac{1}{3}l = 0.$$

Czerty składowe podpór wynoszą

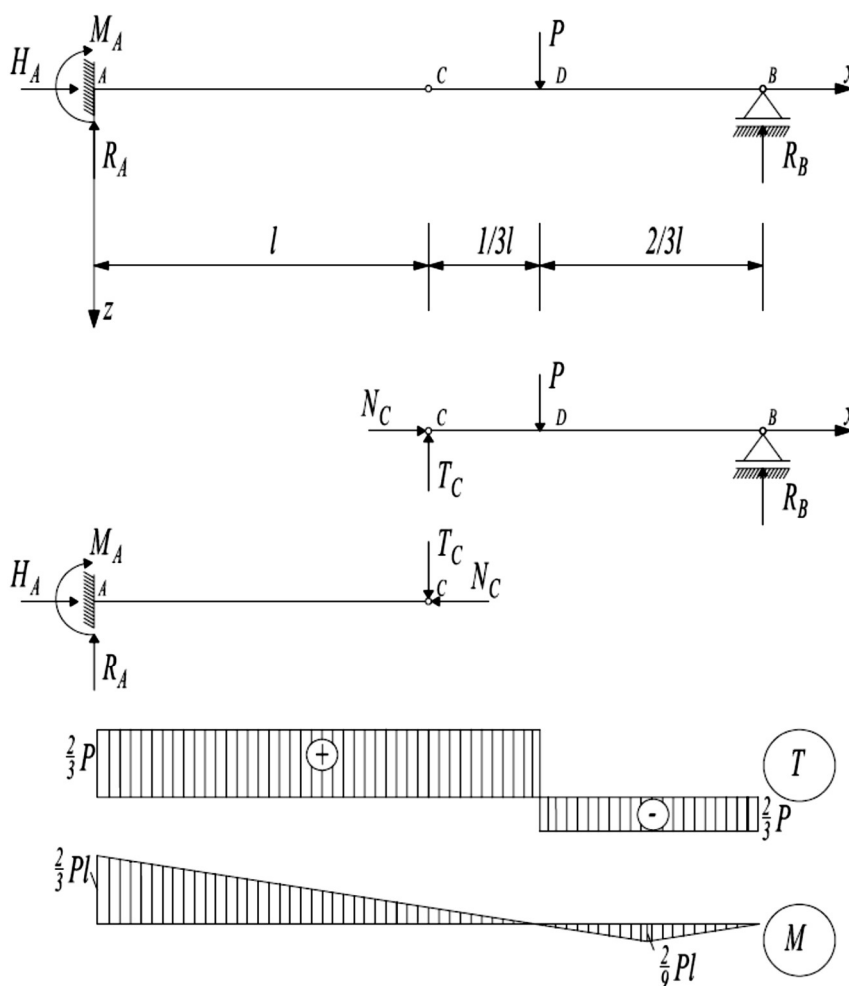
$$R_B = \frac{1}{3}P, \quad M_A = -\frac{2}{3}Pl, \quad R_A = \frac{2}{3}P,$$

Po podstawieniu danych liczbowych otrzymuje się

$$R_B = 3 \text{ kN}, \quad M_A = -9 \text{ kNm}, \quad R_A = 6 \text{ kN}.$$

Siły przekrojowe

Siły wewnętrzne tj. siły tnące jak i momenty zginające zostały przedstawione na poniższych wykresach.



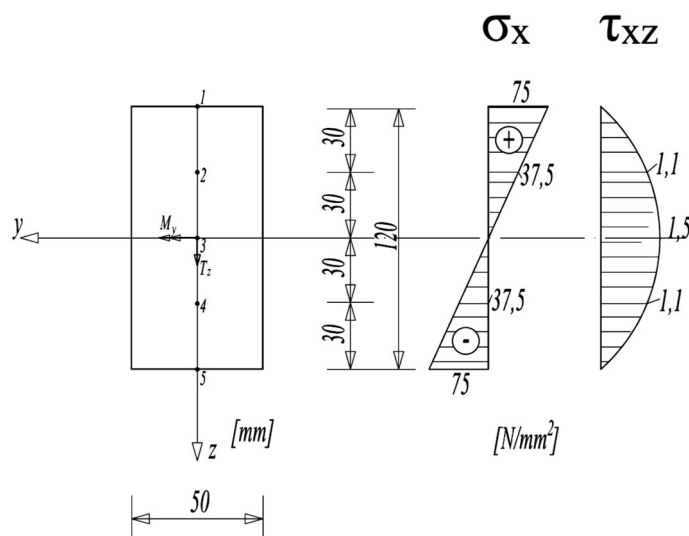
Rys. 20. Wykresy sił wewnętrznych belki – odpowiednio siły tnące T ; momenty zginające M

Ekstremalna wartości siły tnącej wystąpi na odcinku AD i wynosi

$$T_{z,max} = \frac{2}{3}P = 6 \text{ kN}.$$

Ekstremalna wartości momentu zginającego wystąpi w punkcie A, natomiast wartość wynosi

$$M_{y,max} = \frac{2}{3}Pl = 9 \text{ kNm}.$$



Rys. 21. Wykresy naprężeń normalnych σ_x i stycznych τ_{xz}

Moment bezwładności przekroju wyznacza się jako

$$J_y = \frac{b \cdot h^3}{12} = \frac{50 \text{ mm} \cdot (120 \text{ mm})^3}{12} = 7200000 \text{ mm}^4.$$

W celu wyznaczenia naprężenia stycznego należy między innymi określić wartość bezwzględnych momentów statycznych względem osi y

$$\overline{S}_y^1 = \overline{S}_y^5 = 0,$$

$$\overline{S}_y^2 = \overline{S}_y^4 = 50 \text{ mm} \cdot 30 \text{ mm} \cdot 45 \text{ mm} = 67500 \text{ mm}^3,$$

$$\overline{S}_y^3 = 50 \text{ mm} \cdot 60 \text{ mm} \cdot 30 \text{ mm} = 90000 \text{ mm}^3.$$



Szerokość przekroju jest stała i wynosi $b(z) = 50 \text{ mm}$.

Wartości naprężenia stycznego w poszczególnych punktach przekroju wyznaczono ze wzoru (17), a ekstremalna wartość naprężenia stycznego wystąpi w punkcie 3 tj.

$$\tau_{xz}^1 = \frac{T_{z,max} \cdot \overline{S_y^1}}{J_y \cdot b(z)} = \frac{6000 \text{ N} \cdot 0}{7200000 \text{ mm}^4 \cdot 50 \text{ mm}} = 0,$$

$$\tau_{xz}^2 = \frac{T_{z,max} \cdot \overline{S_y^2}}{J_y \cdot b(z)} = \frac{6000 \text{ N} \cdot 67500 \text{ mm}^3}{7200000 \text{ mm}^4 \cdot 50 \text{ mm}} = 1,1 \frac{\text{N}}{\text{mm}^2},$$

$$\tau_{xz}^3 = \frac{T_{z,max} \cdot \overline{S_y^3}}{J_y \cdot b(z)} = \frac{6000 \text{ N} \cdot 90000 \text{ mm}^3}{7200000 \text{ mm}^4 \cdot 50 \text{ mm}} = 1,5 \frac{\text{N}}{\text{mm}^2},$$

$$\tau_{xz}^4 = \tau_{xz}^2.$$

$$\tau_{xz}^5 = \tau_{xz}^1.$$

Wartości naprężenia normalnego w poszczególnych punktach wynoszą

$$\sigma_x^1 = \frac{M_{y,max}}{J_y} \cdot z_1 = \frac{9000000 \text{ Nmm}}{7200000 \text{ mm}^4} \cdot (-60 \text{ mm}) = -75 \frac{\text{N}}{\text{mm}^2},$$

$$\sigma_x^2 = \frac{M_{y,max}}{J_y} \cdot z_2 = \frac{9000000 \text{ Nmm}}{7200000 \text{ mm}^4} \cdot (-30 \text{ mm}) = -37,5 \frac{\text{N}}{\text{mm}^2},$$

$$\sigma_x^3 = \frac{M_{y,max}}{J_y} \cdot z_3 = \frac{9000000 \text{ Nmm}}{7200000 \text{ mm}^4} \cdot 0 = 0,$$

$$\sigma_x^4 = \frac{M_{y,max}}{J_y} \cdot z_4 = \frac{9000000 \text{ Nmm}}{7200000 \text{ mm}^4} \cdot 30 \text{ mm} = 37,5 \frac{\text{N}}{\text{mm}^2},$$

$$\sigma_x^5 = \frac{M_{y,max}}{J_y} \cdot z_5 = \frac{9000000 \text{ Nmm}}{7200000 \text{ mm}^4} \cdot 60 \text{ mm} = 75 \frac{\text{N}}{\text{mm}^2}.$$



Ekstremalne naprężenia normalne wystąpią w punktach najbardziej oddalonych od osi obojętnej, czyli dla współrzędnej $z = |60 \text{ mm}|$, natomiast na osi obojętnej tj. w punkcie 3 wartość naprężenia równa jest zero.

Warunek wytrzymałości na zginanie ma postać $\sigma_x \leq \sigma_{dop}$ i został spełniony ponieważ

$$75 \frac{N}{\text{mm}^2} < 100 \frac{N}{\text{mm}^2}.$$

Koniec przykładu.

3.8. Warunek sztywności

Weryfikacja warunku sztywności ma na celu sprawdzenie wielkości przemieszczeń lub odkształceń i porównanie z wartościami granicznymi. Obliczenia przeprowadza się metodą analityczną (rozwiązania ścisłe) lub metodami numerycznymi (rozwiązania przybliżone). W przypadku przemieszczeń konstrukcji prętowych i rozwiązania analitycznego obliczenia sprowadza się do rozwiązywania równania

$$EJ \frac{d^2 z}{dx^2} = \mp M_g. \quad (3.22)$$

Znak minus lub plus zależy od ustalenia znaku momentu zginającego M_g oraz od orientacji osi układu współrzędnych. Jeżeli wymiary pręta na wzdłuż jego długości będą niezmiennie, to sztywność zginania będzie stała $EJ = \text{const}$ i znalezienie ekstremalnych przemieszczeń pręta uzyska się dzięki dwukrotnemu scałkowaniu równia (3.22). W wyniku całkowania otrzymuje się

$$z = \frac{1}{EJ} \left[\int \left(\int (-M_g) dx + \right) + Cx + D \right], \quad (3.23)$$

gdzie C i D - stałe całkowania, które wyznacza się uwzględniając warunki brzegowe.

Warunek sztywności przyjmuje postać

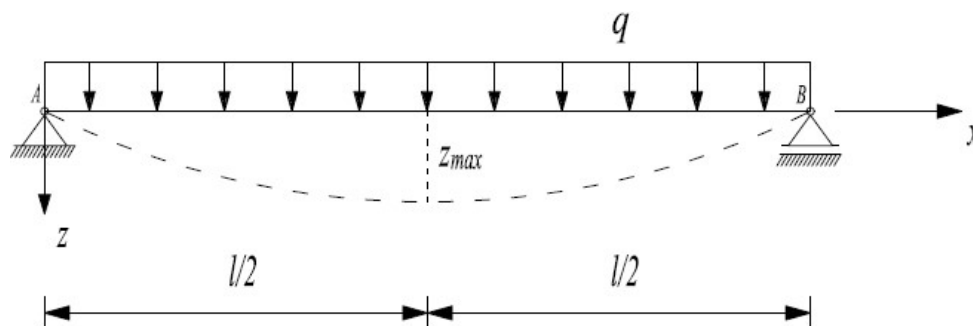
$$z_{max} < z_{dop} , \quad (3.24)$$

gdzie dopuszczalna wartość przemieszczeń z_{dop} jest ustalana w oparciu o warunki prawidłowej eksploatacji rozpatrywanej konstrukcji oraz elementów z nią związanych.

Przykład nr 3.3

Sprawdzić warunek sztywności belek jednoprzęsłowych swobodnie podpartych dla dwóch schematów obciążenia (Rys. 22 i Rys. 23), gdzie przemieszczenie dopuszczalne wynosi $u_{dop} = 10 \text{ mm}$. Dane są: długość $l = 1400 \text{ mm}$, moduł sprężystości podłużnej $E = 210000 \frac{\text{N}}{\text{mm}^2}$, moment bezwładności przekroju $J_y = 1800000 \text{ mm}^4$, siła skupiona $P = 60000 \text{ N}$, obciążenie równomiernie rozłożone $q = 100 \text{ N/mm}$.

Schemat A



Rys. 22. Schemat statyczny belki obciążonej obciążeniem równomiernie rozłożonym

Funkcje przemieszczeń wyznacza się z całkowania równania różniczkowego osi ugiętej np. metodą Clebscha². Po uwzględnieniu warunków brzegowych uzyskuje się równania przemieszczeń belki w funkcji współrzędnej x .

Dla $0 \leq x \leq a$



$$z(x) = \frac{P \cdot b \cdot l \cdot x}{6 \cdot E \cdot J_y} \left(1 - \frac{x^2}{l^2} - \frac{b^2}{l^2} \right),$$

$$z_{x=\frac{l}{2}} = \frac{P \cdot b \cdot l^2}{12 \cdot E \cdot J_y} \left(\frac{3}{4} - \frac{b^2}{l^2} \right),$$

$$a > b.$$

Dla $a \leq x \leq l$

$$z(x) = \frac{P \cdot b \cdot l \cdot x}{6 \cdot E \cdot J_y} \left(1 - \frac{x^2}{l^2} - \frac{b^2}{l^2} + \frac{(x-a)^3}{lbx} \right),$$

$$z_{x=\frac{l}{2}} = \frac{P \cdot b \cdot l^2}{12 \cdot E \cdot J_y} \left(\frac{3}{4} - \frac{b^2}{l^2} + \frac{1}{4} \cdot \frac{(b-a)^3}{bl^2} \right),$$

$$a < b.$$

Maksymalna wartość ugięcia wystąpi dla $x = a = b = \frac{l}{2}$

$$z_{max} = \frac{P \cdot l^3}{48 \cdot E \cdot J_y} = \frac{60000 \text{ N} \cdot (1400 \text{ mm})^3}{48 \cdot 210000 \frac{\text{N}}{\text{mm}^2} \cdot 1800000 \text{ mm}^4},$$

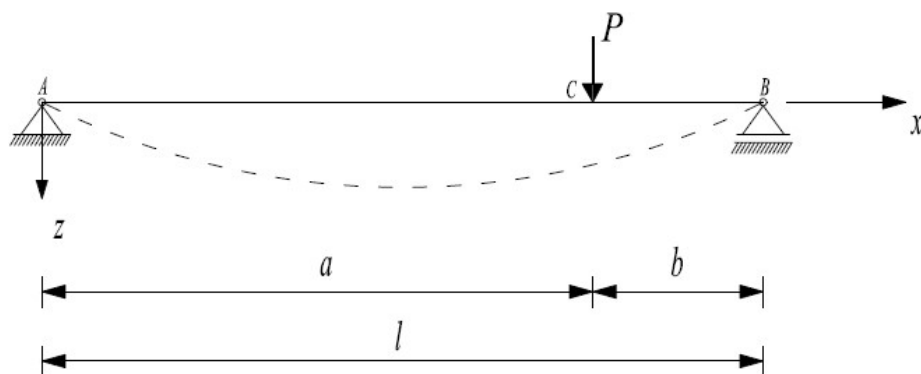
$$z_{max} = 9 \text{ mm}.$$

Warunek sztywności jest spełniony ponieważ

$$z_{max} < z_{dop},$$

$$9 \text{ mm} < 10 \text{ mm}.$$

Schemat rozwiązania B



Rys. 23. Schemat statyczny belki obciążonej siłą skupioną

$$z(x) = \frac{q \cdot l^3 \cdot x}{12 \cdot E \cdot J_y} \left(\frac{1}{2} - \frac{x^2}{l^2} + \frac{1}{2} \cdot \frac{x^3}{l^3} \right),$$

$$z_{max} = z_{x=\frac{l}{2}} = \frac{5}{384} \frac{q \cdot l^4}{E \cdot J_y},$$

$$z_{max} = \frac{5}{384} \cdot \frac{100 \frac{N}{mm} \cdot (1400 \text{ mm})^4}{210000 \frac{N}{mm^2} \cdot 1800000 \text{ mm}^4} = 13 \text{ mm}$$

Warunek sztywności nie jest spełniony ponieważ

$$z_{max} > z_{dop},$$

$$13 \text{ mm} > 10 \text{ mm}.$$

Koniec przykładu.



3.9. Warunek stateczności

3.9.1. Wstęp do stateczności

Za prekursora zagadnień stateczności w obszarze sprężystym uważa się Leonarda Eulera, który określił naprężenia krytyczne jako hiperboliczną funkcję smukłości; jego główne prace zostały opublikowane już w latach 1744-1780¹³.

W odniesieniu do złożonych ustrojów konstrukcyjnych współcześnie dominują numeryczne metody modelowania zjawiska wyboczenia. Wśród kilku najpopularniejszych metod obliczeniowych, znajdujących wykorzystanie w tzw. „solverach” programów komputerowych, należy wymienić przede wszystkim Metodę Elementów Skończonych (MES) i Metodę Różnic Skończonych przedstawioną w pracach^{14 15}.

Większość autorów oprogramowania komputerowego przy rozwiązywaniu zagadnień z zakresu mechaniki skłania się do implementacji MES. Wraz z jej rozwojem stworzona pojawiła się oddzielna grupa zagadnień własnych rozwiązywanych przy pomocy tej metody w odniesieniu do zjawiska wyboczenia. Otrzymywane wektory i wartości własne prowadzące do określania obciążeń oraz sił krytycznych ustroju konstrukcyjnego znajdują potwierdzenie w powszechnych rozwiązaniach analitycznych, zamieszczonych między innymi w pracach S.P. Timoshenki, J.M. Gere¹⁶ i S.P. Timoshenki oraz S. Woinowskiego-Kriegera¹⁷.

Zjawisko wyboczenia jest jedną z form utraty stateczności konstrukcji. Taki stan może nastąpić w przypadku, kiedy siła osiowa P osiągnie pewną wartość krytyczną P_{kr} . Niestateczność definiuje się na ogół jako proces, w którym niewielka zmiana przyczyny (a więc zmiana siły) powoduje dużą zmianę skutku (a więc ugięcia). Punkty krytyczne, czyli punkty na ścieżce równowagi odpowiadające stanom krytycznym konstrukcji zostały przedstawione na rysunku 24. Pierwotnymi punktami krytycznymi są dwa punkty, czyli punkt graniczny (G) i punkt bifurkacji (B). W punkcie granicznym może nastąpić przeskok konstrukcji do nowej konfiguracji równowagi, gdzie występują duże przyrosty przemieszczeń $r(\alpha)$. W punkcie bifurkacji może wystąpić nowa konfiguracja stanu równowagi, inna niż w fazie początkowej. Tak więc utrata stateczności występuje bądź w punkcie bifurkacji, bądź w punkcie granicznym. W sensie ścisłym punkty bifurkacyjne dotyczą modeli idealnych. W układach rzeczywistych,

¹³ Euler L.: Sur La Force Des Colonnes. Men. De l'Acad., 13. Berlin 1757.

¹⁴ Kleiber M.: *Wprowadzenie do Metody Elementów Skończonych*. PWN, Warszawa –Poznań 1986

¹⁵ Pietrzak J., Rakowski G., Wrześniowski G.: *Macierzowa Analiza Konstrukcji*. PWN, Warszawa 1986.

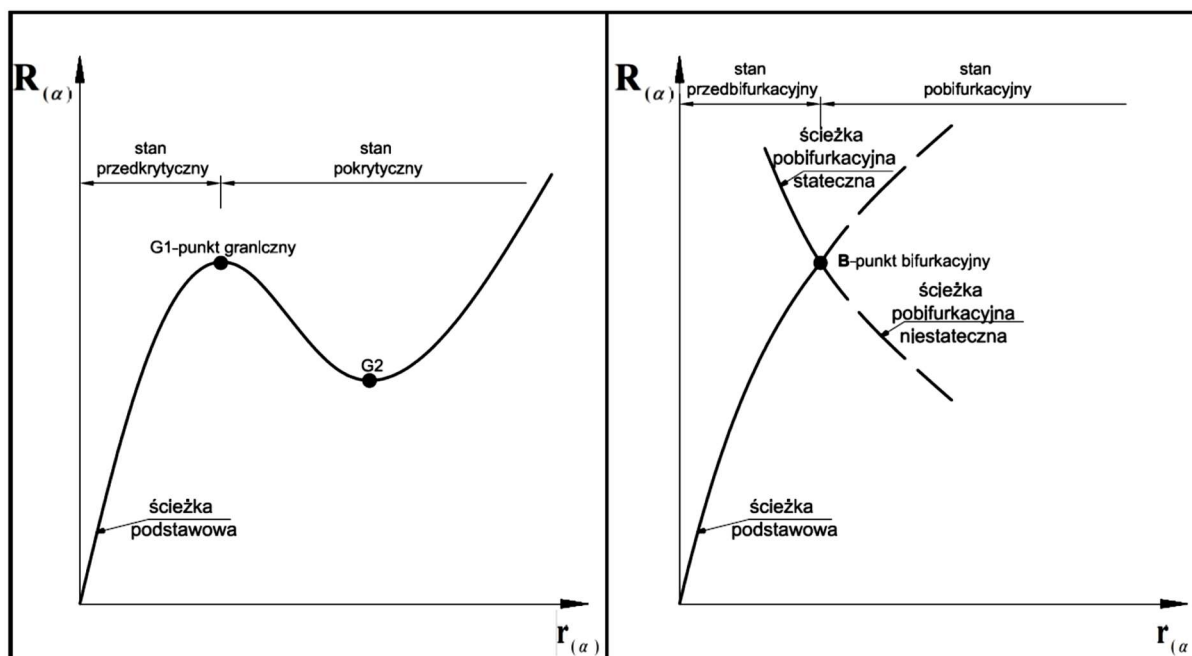
¹⁶ Timoshenko S.P., Gere J.M.: *Teoria Stateczności Sprężystej*. Arkady, Warszawa 1963.

¹⁷ Timoshenko S.P., Woinowsky-Krieger S.: *Teoria Płyt i Powłok*. Arkady, Warszawa 1962.

na skutek zmiennych własności materiałowych, geometrii oraz obciążenia zewnętrznego, punkty bifurkacji równowagi zastępuje się przez punkty maksimum obciążenia⁹.

Jeżeli smukłość elementu jest większa od wartości granicznej, a więc naprężenia w elemencie nie przekraczają granicy proporcjonalności, wówczas wartość krytyczną otrzymuje się w obszarze sprężystym. Poniżej smukłości granicznej, naprężenia mają charakter niesprężysty. Analizą pozwalającą oszacować obciążenie krytyczne wyboczenia sprężystego w stadium przedbifurkacyjnym jest analiza wartości własnych układu.

Dla prętów krępych, których długość jest niewielka w stosunku do wymiarów przekroju – siła krytyczna wywołuje naprężenia bliskie granicy plastyczności, więc sąsiednia postać giętą musi być rozpatrywana w zakresie niesprężystym, a więc wyboczenie jest niesprężyste.



Rys. 24. Ścieżka równowagi konstrukcji

3.9.2. Badania doświadczalne słupów ściskanych osiowo

Istotną rolę w sprawdzeniu poprawności przeprowadzonych analiz numerycznych mają badania eksperymentalne na rzeczywistych modelach, które można podzielić zasadniczo na dwie grupy. Pierwsza grupa dotyczy analiz dla modeli, których smukłość jest większa od smukłości granicznej, wówczas eksperymenty te dotyczą wyboczenia w zakresie sprężystym materiału. Druga grupa to badania elementów w zakresie sprężysto-plastycznym lub plastycznym. Smukłość graniczna ma postać



$$\lambda_{gr} = \pi \cdot \sqrt{\frac{E}{f_H}}, \quad (3.25)$$

gdzie E - moduł sprężystości podłużnej, f_H - granica proporcjonalności.

Doświadczalne wyznaczenie siły krytycznej wyboczenia giętnego w zakresie sprężystym [50] przeprowadza się kilkoma metodami, spośród których należy wymienić przede wszystkim metodę Southwella. Dla klasycznego pręta obustronnie przegubowego równanie linii ugięcia można zapisać w postaci

$$u(z) = \frac{a}{\frac{P_{ycr}}{P} - 1} \cdot \sin \frac{\pi}{l} z, \quad (3.26)$$

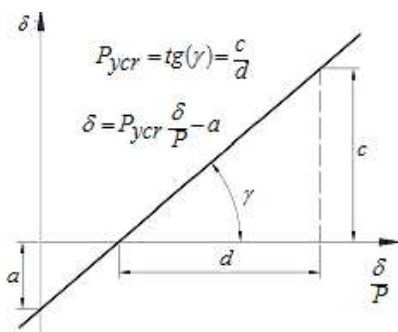
gdzie a to wartość wygięcia początkowego, z - współrzędna opisująca położenie punktu na długości elementu, natomiast l to jego długość. Na podstawie równania otrzymuje się ugięcie w połowie długości pręta $\delta = u(l/2)$ w postaci

$$\delta = P_{ycr} \frac{\delta}{P} - a. \quad (3.27)$$

Równanie (3.27) jest liniowe względem współrzędnych: δ/P i δ , dlatego można wyznaczyć prostą (rys. 25) o współczynniku kierunkowym $\tan(\gamma) = P_{ycr}$. Podczas badań określa się dyskretną postać zależności $\delta_i = f(\delta_i/P_i)$, a następnie aproksymuje otrzymane wyniki metodą regresji liniowej

$$\delta = \eta \frac{\delta}{P} + \nu, \quad (3.28)$$

gdzie parametry η i ν można obliczyć przy pomocy metody najmniejszych kwadratów.



Rys. 25. Wykres zależności $\delta = f(\delta/P)$ ¹⁸

¹⁸ Gosowski. B., Kubica E.: *Badania Laboratoryjne z Konstrukcji Metalowych*. Oficyna Wydawnicza Politechniki Wrocławskiej 2001.



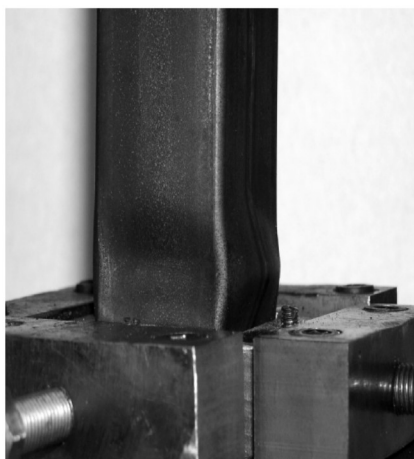
Aproksymacja ta umożliwia wyznaczenie wartości obciążenia krytycznego oraz krzywizny początkowej pręta z zależności

$$P_{ycr} = \eta, \quad a = \nu. \quad (3.29)$$

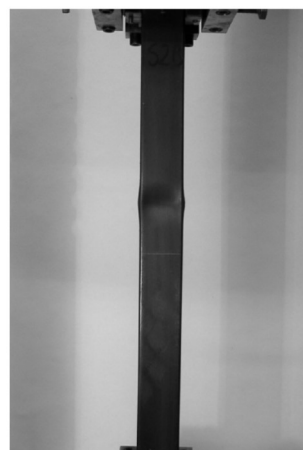
Wyniki analiz numerycznych w większości są zgodne z wynikami znanych rozwiązań analitycznych i znajdują potwierdzenie w badań eksperymentalnych wyboczenia sprężystego. Inny charakter mają analizy w zakresie smukłości $\lambda < \lambda_{gr}$. Złożoność problemu uwidoczniła się już w kilku teoriach dotyczących wyboczenia niesprężystego, do których należy zaliczyć teorię naprężeń krytycznych Tetmajera-Jasińskiego, teorię Engessera-Karmana, Johnsona-Ostenfelda i Engessera-Shanleya¹⁹. W zakresie krępych przekrojów istotną rolę odgrywają badania eksperymentalne dotyczące elementów w odpowiednich zakresach smukłości, wykonanych z konkretnych gatunków i odmian stali, a także posiadające określone losowe imperfekcje wstępne, jak np. naprężenia własne związane z wpływem zgrzewania w przekrojach giętych na zimno. Na losowość zjawiska wyboczenia niesprężystego składa się zwykle szereg czynników, wśród których przede wszystkim należy wymienić smukłość elementu oraz przekrój poprzeczny, co w konsekwencji prowadzi do otrzymania określonej postaci wyboczenia. Powyższa złożoność zjawiska została omówiona między innymi w pracy²⁰, gdzie analizowane były krępe słupy stalowe o przekrojach rurowych czworobocznych. Badania¹⁸ dotyczyły próby ściskania osiowych kształtowników stalowych giętych na zimno, a przyjętym schematem statycznym kształtowników był pręt obustronnie utwierdzony. Próbkki o najmniejszej smukłości uległy wyboczeniu miejscowemu w pobliżu utwierdzenia, natomiast próbki o większych smukłościach uległy wyboczeniu giętnemu, któremu towarzyszy wyboczenie lokalne w połowie wysokości elementu oraz w pobliżu utwierdzenia (Rys. 26-27).

¹⁹ Schanley F.R.: *Inelastic column theory*. Journal of Aeronautical Science 1947;14: 261-267.

²⁰ Glinicka A.: *Doświadczalna analiza wyboczenia niesprężystego kształtowników o przekrojach rurowych czworobocznych*. Instytut Badawczy Dróg i Mostów, Kwartalnik *Drogi i Mosty* 2005;2:5-37.



Rys. 26. Wyboczenie lokalne kształtowników w pobliżu utwierdzenia¹⁸



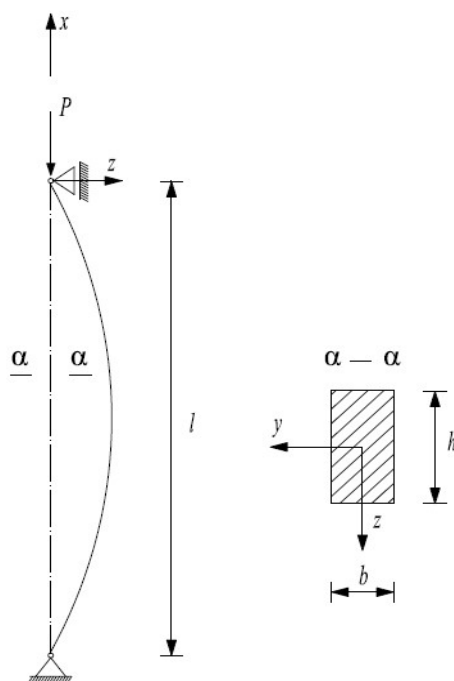
Rys. 27. Wyboczenie lokalne kształtowników w połowie rozpiętości¹⁸

Zagadnienia stateczności badanych kształtowników nie można sprowadzić do jednej płaszczyzny, gdyż w badaniach stwierdzono przestrzenny stan przemieszczenia. W przypadku kształtowników o przekroju prostokątnym wszystkie próbki niezależnie od smukłości zostały zniszczone na skutek wyboczenia giętnego, natomiast przegub plastyczny wytworzył się w połowie wysokości elementu.

Na podstawie przeanalizowanych badań doświadczalnych¹⁸ dotyczących wyboczenia niesprężystego można stwierdzić, iż wyboczenie ma charakter losowy, może mieć formę wyboczenia lokalnego, globalnego wyboczenia giętnego z towarzyszącym mu wyboczeniem lokalnym lub wyboczenia giętnego. Zmiennymi losowymi są tu: smukłość prętowa, klasa i odmiana stali oraz kształt przekroju poprzecznego. Parametry te wpływają na wartość siły krytycznej, postaci wyboczenia oraz wywołane skrócenia i przemieszczenia. W przypadku analizowanych rur obserwuje się najczęściej przeguby plastyczne tworzące się w połowie ich wysokości; wyjątek stanowią rury o bardzo małych smukłościach, gdzie przeguby mogą wystąpić w pobliżu utwierdzeń.

3.9.3. Podstawowe zagadnienia teoretyczne stateczności sprężystej pręta

Dany jest pręt o długości l obustronnie przegubowy, którego schemat statyczny został przedstawiony na rysunku 28. Przekrój pręta jest prostokątny o wymiarach b i h . Pręt wykonany jest z materiału liniowo-sprężystego o module sprężystości podłużnej E .



Rys. 28. Schemat statyczny pręta ściskanego

Siłę krytyczną powodującą wyboczenie ustroju konstrukcyjnego wyznacza się ze wzoru Eulera

$$P_{kr} = \frac{\pi^2 \cdot E \cdot J_{min}}{l_w^2}, \quad (3.30)$$

gdzie J_{min} – minimalny moment bezwładności przekroju względem osi y oraz osi z, wyznaczany jako

$$J_{min} = \min(J_y, J_z),$$

l_w – długość wyboczeniowa zależna od współczynnika długości wyboczeniowej μ . Współczynnik ten jest zależny od warunków podparcia (Rys. 29) i przyjmuje najczęściej następujące wartości:

- w przypadku pręta obustronnie przegubowego $\mu = 1,0$;
- w przypadku pręta jednostronnie przegubowego i utwierdzonego na drugim końcu $\mu = 0,7$;
- w przypadku pręta obustronnie utwierdzonego $\mu = 0,5$;
- w przypadku pręta utwierdzonego na jednym końcu, a drugi koniec jest wolny (wspornik) $\mu = 2,0$.



Rys. 29. Demonstracja przypadków wyboczenia Eulera na stanowisku laboratoryjnym Akademii Nauk Stosowanych w Koninie, prod. G.U.N.T. Gerätebau GmbH (od lewej: pręt obustronnie przegubowy; pręt jednostronnie przegubowy i utwierdzony na drugim końcu; pręt obustronnie utwierdzony; pręt utwierdzony na jednym końcu, a drugi koniec jest wolny)

Długość wyboczeniowa określana jest wzorem

$$l_w = \mu \cdot l. \quad (3.31)$$

Istotnym parametrem geometrycznym pręta jest smukłość określana jako

$$\lambda = \frac{l_w}{i_{min}}, \quad (3.32)$$

gdzie i_{min} – minimalny promień bezwładności przekroju względem osi y oraz osi z wyznaczany jako

$$i_{min} = \min(i_y, i_z),$$

a wzory dla poszczególnych promieni bezwładności są następujące



$$i_z = \sqrt{\frac{J_z}{A}}, \quad i_y = \sqrt{\frac{J_y}{A}}. \quad (3.33)$$

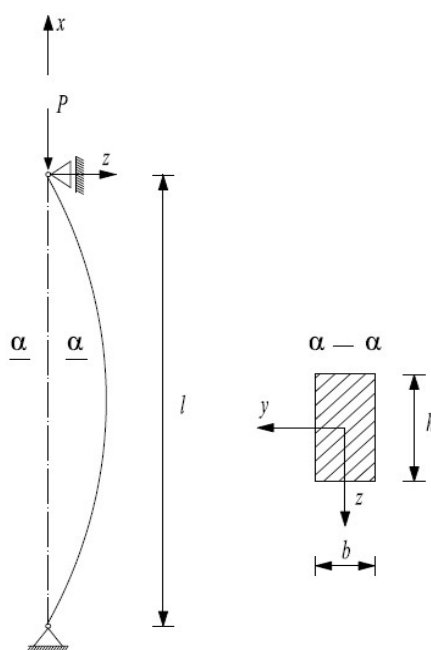
Wzór Eulera (3.30) ma zastosowanie dla prętów, których smukłość spełnia warunek

$$\lambda > \lambda_{gr},$$

gdzie λ_{gr} – smukłość graniczna wg wzoru (3.25).

Przykład nr 3.4

Dany jest pręt stalowy o przekroju prostokątnym obciążony siłą osiową $P = 320000 \text{ N}$ (Rys. 30). Wyznaczyć siłę krytyczną P_{kr} , a następnie sprawdzić warunek stateczności przyjmując następujące dane: moduł sprężystości podłużnej $E = 210000 \frac{\text{N}}{\text{mm}^2}$, granica proporcjonalności $f_H = 180 \frac{\text{N}}{\text{mm}^2}$, szerokość przekroju $b = 40 \text{ mm}$, wysokość przekroju $h = 90 \text{ mm}$, długość pręta $l = 1800 \text{ mm}$. Schemat podparcia w płaszczyźnie xz jest taki sam jak w płaszczyźnie xy .



Rys. 30. Schemat statyczny pręta ściskanego



Ze względu na schemat podparcia pręta współczynnik długości wyboczeniowej wynosi $\mu = 1,00$, a długość wyboczeniowa jest równa

$$l_w = \mu \cdot l = 1,0 \cdot 1800 \text{ mm} = 1800 \text{ mm} .$$

Pole przekroju wynosi

$$A = b \cdot h = 40\text{mm} \cdot 90\text{mm} = 3600 \text{ mm}^2 .$$

Momenty bezwładności przekroju odpowiednio względem osi y oraz osi z wynoszą odpowiednio

$$J_y = \frac{b \cdot h^3}{12} = \frac{40\text{mm} \cdot (90\text{mm})^3}{12} = 2430000 \text{ mm}^4 ,$$

$$J_z = \frac{h \cdot b^3}{12} = \frac{90\text{mm} \cdot (40\text{mm})^3}{12} = 480000 \text{ mm}^4 .$$

Ponieważ $J_{min} = J_z$ czyli $J_z < J_y$, to wyboczenie pręta nastąpi w płaszczyźnie xy . Minimalny promień bezwładności przekroju pręta wynosi

$$i_{min} = i_z = \sqrt{\frac{J_z}{A}} = \sqrt{\frac{480000 \text{ mm}^4}{3600 \text{ mm}^2}} ,$$

$$i_z = 11,55 \text{ mm} .$$

Smukłość pręta wynosi

$$\lambda = \frac{l_w}{i_z} = \frac{1800 \text{ mm}}{11,55 \text{ mm}} ,$$

$$\lambda = 155,84 .$$

Smukłość graniczna wynosi



$$\lambda_{gr} = \pi \cdot \sqrt{\frac{E}{f_H}} = 3,14 \cdot \sqrt{\frac{210000 \frac{N}{mm^2}}{180 \frac{N}{mm^2}}},$$

$$\lambda_{gr} = 107,3.$$

Dla $\lambda > \lambda_{gr}$ wyboczenie jest sprężyste, a siłę krytyczną wyznacza się z wzoru Eulera

$$P_{kr} = \frac{\pi^2 \cdot E \cdot J_{min}}{l_w^2} = \frac{\pi^2 \cdot 210000 \frac{N}{mm^2} \cdot 480000 mm^4}{(1800 mm)^2},$$

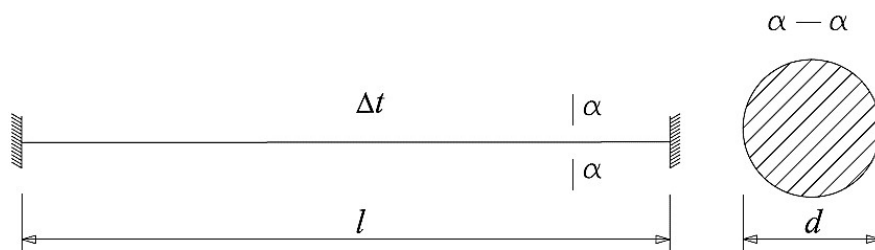
$$P_{kr} = 307054 N.$$

Warunek stateczności jest spełniony ponieważ $P < P_{kr}$.

Koniec przykładu.

Przykład nr 3.5

Stalowy pręt o przekroju kołowym utwierdzony jest na obu końcach w dwóch nieodkształcalnych ściankach (Rys. 31). Zakładając, że pręt ulega równomiernemu podgrzaniu, wyznaczyć wartość przyrostu temperatury Δt_k , przy której pręt utraci stateczność. Dane są: moduł sprężystości podłużnej $E = 210000 \frac{N}{mm^2}$, granica proporcjonalności $f_H = 200 \frac{N}{mm^2}$, współczynnik rozszerzalności cieplnej $\alpha_t = 0,000012 \text{ } ^\circ\text{C}^{-1}$, długość pręta $l = 300 \text{ cm}$, średnica przekroju pręta $d = 5 \text{ cm}$.



Rys. 31. Schemat statyczny pręta



Z uwagi na sposób podparcia pręta współczynnik długości wyboczeniowej wynosi $\mu = 0,5$.

Długość wyboczeniowa równa jest

$$l_w = \mu \cdot l = 0,5 \cdot 300 \text{ cm} = 150 \text{ cm} .$$

Smukłość graniczna wynosi

$$\lambda_{gr} = \pi \cdot \sqrt{\frac{E}{f_H}} = 3,14 \cdot \sqrt{\frac{210000 \frac{N}{mm^2}}{200 \frac{N}{mm^2}}} ,$$

$$\lambda_{gr} = 101,8 .$$

Moment bezwładności przekroju pręta wynosi

$$J_z = J_y = J = \frac{\pi \cdot d^4}{64} = \frac{3,14 \cdot (5\text{cm})^4}{64} ,$$

$$J = 30,7 \text{ cm}^4 .$$

Pole przekroju pręta wynosi

$$A = \frac{\pi \cdot d^2}{4} = \frac{3,14 \cdot (5\text{cm})^2}{4} ,$$

$$A = 19,6 \text{ cm}^2 .$$

Promień bezwładności przekroju pręta wynosi

$$i_z = i_y = i = \sqrt{\frac{J}{A}} = \sqrt{\frac{30,7 \text{ cm}^4}{19,6 \text{ cm}^2}} ,$$

$$i = 1,25 \text{ cm} .$$

Smukłość pręta wynosi



$$\lambda = \frac{l_w}{i} = \frac{150 \text{ cm}}{1,25 \text{ cm}} ,$$

$$\lambda = 120 .$$

Dla $\lambda > \lambda_{gr}$ wyboczenie jest sprężyste.

Podłużna siła wewnętrzna N powstała na skutek podgrzania pręta, obliczana jest z warunku, że pręt nie może zmienić swej długości, dlatego całkowity przyrost długości pręta musi być równy zero, a więc

$$\Delta l = \Delta l_{\Delta t} + \Delta l_N = 0 ,$$

$$\alpha_t \cdot \Delta t \cdot l + \frac{N \cdot l}{E \cdot A} = 0 ,$$

stąd

$$N = E \cdot A \cdot \alpha_t \cdot \Delta t .$$

Utrata stateczności pręta nastąpi jeżeli $N = P_{kr}$, natomiast $\Delta t = \Delta t_k$. Musi zatem zostać spełniony warunek

$$\frac{\pi^2 \cdot E \cdot J}{l_w^2} = E \cdot A \cdot \alpha_t \cdot \Delta t .$$

Po przekształceniu powyższego równania otrzymuje się przyrost temperatury powodujący wyboczenie pręta, który wynosi

$$\Delta t_k = \frac{\pi^2 \cdot J}{l_w^2 \cdot A \cdot \alpha_t} = \frac{\pi^2}{\lambda^2 \cdot \alpha_t} ,$$

$$\Delta t_k = \frac{3,14^2}{120^2 \cdot 0,000012 \text{ } ^\circ\text{C}^{-1}} ,$$

$$\Delta t_k = 57,1 \text{ } ^\circ\text{C} .$$

Koniec przykładu.



3.10. Bibliografia

1. Bibiak-Żochowski M. (red.): *Wytrzymałość materiałów i konstrukcji. Tom I.* Oficyna Politechniki Warszawskiej, Warszawa 2013.
2. Bibiak-Żochowski M. (red.): *Wytrzymałość materiałów i konstrukcji. Tom II.* Oficyna Politechniki Warszawskiej, Warszawa 2013.
3. Biegus A.: *Badania losowej krzywej równowagi granicznej ściskanych blach fałdowych.* Rozprawy Inżynierskie 1988;36:15-27.
4. Biegus A.: *Probabilistyczna Analiza Konstrukcji Stalowych.* PWN, Warszawa 1996.
5. Dyląg Z., Jakubowicz A., Orłoś Z.: *Wytrzymałość materiałów. Tom 1.* Wydawnictwo Naukowo – Techniczne, Warszawa 1999.
6. Dyląg Z., Jakubowicz A., Orłoś Z.: *Wytrzymałość materiałów. Tom 2.* Wydawnictwo Naukowo – Techniczne, Warszawa 2012.
7. Engesser F.: *Über die Knickfestigkeit Gerader Stäbe.* Z. Architektur und Ingenieurwesen 1889.
8. Euler L.: *Sur La Force Des Colonnes.* Men. De l'Acad., 13. Berlin 1757.
9. Glinicka A.: *Wytrzymałość materiałów I.* Oficyna Politechniki Warszawskiej. Warszawa 2011.
10. Glinicka A.: *Doświadczalna analiza wyboczenia niesprężystego kształtowników o przekrojach rurowych czworobocznych.* Instytut Badawczy Dróg i Mostów, *Kwartalnik Drogi i Mosty* 2005;2:5-37.
11. Gosowski B., Kubica E.: *Badania Laboratoryjne z Konstrukcji Metalowych.* Oficyna Wydawnicza Politechniki Wrocławskiej 2001.
12. Grabowski J., Iwanczewska A.: *Zbiór zadań z wytrzymałości materiałów.* Wydawnictwa Politechniki Warszawskiej, Warszawa 1984.
13. Kleiber M.: *Wprowadzenie do Metody Elementów Skończonych.* PWN, Warszawa – Poznań 1986.
14. Misiak J.: *Mechanika ogólna. Tom I. Statyka i kinematyka.* Wydawnictwo Naukowo-Techniczne, Warszawa 1998.
15. Mendera Z.: *Zagadnienia stanów granicznych konstrukcji stalowych.* Zeszyty Naukowe Politechniki Krakowskiej, Kraków 1969;7(33).
16. Pietrzak J., Rakowski G., Wrześniowski G.: *Macierzowa Analiza Konstrukcji.* PWN, Warszawa 1986.



17. Schanley F.R.: *Inelastic column theory*. Journal of Aeronautical Science 1947;14: 261-267.
18. Timoshenko S.P., Gere J.M.: *Teoria Stateczności Sprężystej*. Arkady, Warszawa 1963.
19. Timoshenko S.P., Woinowsky-Krieger S.: *Teoria Płyt i Powłok*. Arkady, Warszawa 1962.
20. Wojewódzki W.: *Nośność graniczna konstrukcji prętowych*. Oficyna Wydawnicza Politechniki Warszawskiej, Warszawa 2005.



Monika Muszyńska

4. METALOZNAWSTWO I OBRÓBKA CIEPLNA

4.1. WSTĘP

Metaloznawstwo oraz obróbka cieplna to niezwykle ważna dziedzina wiedzy technicznej, która odgrywa fundamentalną rolę w projektowaniu i wytwarzaniu nowoczesnych maszyn oraz konstrukcji inżynierskich. Ich znaczenie jest szczególnie istotne dla inżynierów specjalizujących się w mechanice i budowie maszyn, gdzie właściwy dobór i przetwarzanie materiałów mają bezpośredni wpływ na jakość, trwałość oraz bezpieczeństwo projektowanych urządzeń. Osiągnięcie tych celów jest możliwe tylko dzięki głębokiej znajomości właściwości metali, ich struktur, a także metod ich kontrolowanej modyfikacji poprzez obróbkę cieplną. Metale i ich stopy stanowią podstawę wielu konstrukcji i elementów maszyn, ponieważ charakteryzują się szerokim zakresem właściwości mechanicznych, takich jak wytrzymałość, twardość, elastyczność czy odporność na ścieranie. Współczesne wymagania techniczne, zwłaszcza w branżach takich jak motoryzacja, lotnictwo czy przemysł ciężki, stawiają przed konstruktorami coraz większe wyzwania. Aby spełnić te oczekiwania, konieczne jest dokładne zrozumienie mikrostruktury metali oraz wpływu, jaki na nią wywierają różne procesy obróbki cieplnej, takie jak hartowanie, odpuszczanie, wyżarzanie czy nawęglanie.

Obróbka cieplna odgrywa kluczową rolę w kształtowaniu właściwości metali, ponieważ pozwala na kontrolowaną zmianę ich struktury wewnętrznej. Procesy takie jak hartowanie, prowadzące do zwiększenia twardości i wytrzymałości materiału, czy odpuszczanie, mające na celu poprawę jego plastyczności, są niezbędne w produkcji maszyn i konstrukcji narażonych na intensywne obciążenia mechaniczne. Umiejętność doboru właściwych parametrów obróbki cieplnej oraz zrozumienie mechanizmów, które zachodzą w metalach podczas tych procesów, są nieocenione w codziennej pracy inżyniera. Dzięki temu można precyzyjnie kontrolować parametry mechaniczne elementów, co przekłada się na ich niezawodność oraz długowieczność.

Metaloznawstwo jako nauka o strukturze i właściwościach metali pozwala inżynierom na efektywne projektowanie nowych materiałów o wymaganych parametrach użytkowych. Zrozumienie, jak różne pierwiastki stopowe wpływają na właściwości mechaniczne, fizyczne i chemiczne metali, daje możliwość tworzenia materiałów o zróżnicowanych cechach, dostosowanych do specyficznych warunków pracy. Wiedza ta jest kluczowa dla



opracowywania innowacyjnych rozwiązań, które pozwalają na zwiększenie efektywności energetycznej, redukcję masy konstrukcji czy zwiększenie odporności na korozję.

Dla inżyniera mechanika, zwłaszcza w obszarze budowy maszyn, znajomość metaloznawstwa i obróbki cieplnej jest nieodzowna. Każdy etap produkcji maszyn - od projektowania po wykonanie i konserwację - wymaga zrozumienia zachowania materiałów w różnych warunkach eksploatacyjnych. Brak tej wiedzy może prowadzić do niewłaściwego doboru materiałów, co z kolei skutkuje awariami, przedwczesnym zużyciem elementów, a nawet zagrożeniem bezpieczeństwa. Dlatego każda decyzja inżynierska, czy to dotycząca wyboru odpowiedniego stopu, czy zastosowania konkretnej techniki obróbki, musi opierać się na solidnych podstawach teoretycznych i praktycznych z zakresu tych dziedzin.

Obróbka cieplna, mimo że często niewidoczna na pierwszy rzut oka, ma kluczowy wpływ na jakość gotowego produktu. Przykłady takie jak hartowanie powierzchniowe wałów napędowych czy azotowanie przekładni zębatych, pokazują, jak precyzyjna obróbka cieplna może poprawić właściwości mechaniczne kluczowych komponentów maszyn. W branży motoryzacyjnej, lotniczej czy energetycznej, gdzie niezawodność i trwałość są absolutnie krytyczne, umiejętne wykorzystanie procesów obróbki cieplnej pozwala na znaczne zwiększenie wydajności i bezpieczeństwa pracy urządzeń (Dobrzański 2002).

4.2. BUDOWA METALI

Metale w stanie stałym mają strukturę krystaliczną, co oznacza, że atomy w ich wnętrzu są uporządkowane w regularny, powtarzalny sposób. Ten układ atomów tworzy tzw. sieć krystaliczną. W metalach najczęściej spotyka się trzy podstawowe rodzaje sieci krystalicznych:

- Struktura regularna przestrzennie centrowana (RPC), (BCC - *Body-Centered Cubic*), w tej strukturze atomy metalu są ułożone w sposób, gdzie jeden atom znajduje się w centrum sześcianu, a pozostałe atomy na jego wierzchołkach. Przykładami metali o tej strukturze są m.in. żelazo (w temperaturze poniżej 912°C), chrom i molibden. RPC charakteryzuje się relatywnie luźnym upakowaniem atomów, co wpływa na właściwości mechaniczne materiału, takie jak twardość i kruchość.
- Struktura regularna ściennie centrowana (RSC), (FCC - *Face-Centered Cubic*), w strukturze RSC atomy są ułożone na wierzchołkach sześcianu oraz na środkach jego ścian. Ten rodzaj sieci krystalicznej występuje m.in. w aluminium, miedzi i żelazie (powyżej 912°C). RSC jest bardziej zwarte niż RPC, co skutkuje większą



plastycznością materiału. Metale o tej strukturze są zazwyczaj bardziej kowalne i ciągliwe.

- Struktura heksagonalna zwartouporządkowana HZ, (HCP - *Hexagonal Close-Packed*) ta struktura charakteryzuje się gęstym upakowaniem atomów w heksagonalnym układzie. Przykładami metali o tej strukturze są m.in. tytan, magnez i cynk. Ze względu na specyficzne ułożenie atomów, metale o strukturze HZ mogą wykazywać niższą plastyczność w porównaniu do RSC, ale wciąż są odporne na odkształcenia (Przybyłowicz 2003).

4.2.1 Wiązania metaliczne

Unikalną cechą metali jest sposób, w jaki atomy są ze sobą połączone, czyli wiązanie metaliczne. W tym typie wiązania atomy metalu oddają swoje zewnętrzne elektrony, tworząc tzw. „chmurę elektronową” wokół rdzeni atomowych. Elektrony te, nazywane elektronami przewodnictwa, są swobodne i mogą poruszać się przez całą strukturę krystaliczną metalu. To zjawisko jest kluczowe dla wyjaśnienia wielu charakterystycznych właściwości metali, takich jak:

- Przewodnictwo elektryczne - swobodne elektrony w metalu umożliwiają łatwy przepływ prądu elektrycznego. Kiedy przykładamy różnicę potencjałów (napięcie), elektrony przemieszczają się w kierunku dodatniej elektrody, co prowadzi do przepływu prądu.
- Przewodnictwo cieplne - dzięki ruchowi swobodnych elektronów, metale są również dobrymi przewodnikami ciepła. Ciepło w metalu przenoszone jest zarówno przez wibracje atomów (fonony), jak i przez swobodne elektrony, które przekazują energię cieplną.
- Plastyczność i kowalność - swobodne elektrony w metalu pozwalają na to, by atomy przesuwały się względem siebie bez zrywania wiązań. Dzięki temu metale mogą być odkształcane plastycznie, co umożliwia ich formowanie poprzez kucie, walcowanie czy tłoczenie.



4.2.2. Defekty w strukturze krystalicznej

Metale mają regularną strukturę krystaliczną, jednakże w praktyce nigdy nie jest ona idealna. W rzeczywistych materiałach występują różnego rodzaju defekty, które mają istotny wpływ na ich właściwości mechaniczne. Do najważniejszych występujących defektów należą:

- Defekty punktowe - to lokalne zaburzenia w sieci krystalicznej, takie jak luki atomowe (brakujące atomy) lub atomy obce w sieci. Mogą one wpływać na wytrzymałość, plastyczność i inne właściwości materiału.
- Defekty liniowe (dyslokacje) - dyslokacje to defekty w postaci linii, wzdłuż których nastąpiło zaburzenie układu atomów. Obecność dyslokacji ułatwia plastyczne odkształcanie się metalu, co czyni go bardziej kowalnym i ciągliwym. Kontrola dyslokacji jest kluczowym elementem w procesach wzmacniania metali, takich jak hartowanie czy kucie.
- Defekty powierzchniowe - to zaburzenia na granicach ziaren, czyli miejscach, gdzie kończy się jeden kryształ, a zaczyna drugi o innym ułożeniu atomów. Im mniejsze ziarna, tym materiał ma większą wytrzymałość, co wynika z większej liczby granic ziaren, które hamują ruch dyslokacji (Przybyłowicz 2003).

4.2.3. Właściwości mechaniczne i fizyczne metali

Budowa metali bezpośrednio wpływa na ich właściwości mechaniczne i fizyczne, które decydują o ich zastosowaniach w przemyśle. Do najważniejszych właściwości należą:

- Wytrzymałość na rozciąganie - zależy od struktury krystalicznej i obecności dyslokacji, metale mogą być wzmacniane poprzez odpowiednie procesy czy obróbkę, np. hartowanie czy walcowanie.
- Twardość - zdolność materiału do opierania się zarysowaniom i odkształceniom, twardość metalu zależy od jego składu chemicznego i sposobu obróbki.
- Plastyczność - zdolność metalu do trwałych odkształceń bez pękania, wiąże się z możliwością przesuwania się dyslokacji w sieci krystalicznej.
- Przewodnictwo elektryczne i cieplne - jak już wcześniej wspomniano, metale są doskonałymi przewodnikami dzięki obecności swobodnych elektronów.



4.3. PRODUKCJA METALI

Produkcja metali to wieloetapowy proces, który obejmuje wydobycie surowców mineralnych, ich przekształcanie oraz dalszą obróbkę w celu uzyskania czystego metalu o odpowiednich właściwościach. Różnorodność procesów stosowanych do wydobycia i przetwarzania metali zależy od specyfiki surowca, jak i od końcowego przeznaczenia produktu. Wraz z postępem technologicznym rozwijają się nowe, bardziej efektywne i ekologiczne metody produkcji.

4.3.1. Wydobycie rud metali

Produkcja metali rozpoczyna się od wydobycia odpowiednich surowców mineralnych, które zawierają metale w postaci związków chemicznych, najczęściej jako rudy. Proces wydobycia rud odbywa się zarówno metodami odkrywkowymi, jak i podziemnymi, w zależności od głębokości występowania złóż. Wydobycie odkrywkowe polega na usuwaniu warstw ziemi pokrywających złoża i jest stosowane w przypadku płytko położonych rud. Kopalnie głębinowe, stosowane do bardziej skomplikowanych złóż, wymagają zaawansowanych technologii wydobywczych i stosowania specjalistycznego sprzętu, takiego jak maszyny drążące tunele.

Jednym z przykładów jest wydobycie rud żelaza, które jest surowcem do produkcji stali. W przypadku miedzi stosuje się zarówno metody odkrywkowe, jak i podziemne, a złoża tego metalu są eksploatowane na całym świecie, w tym w Chile, Australii i Polsce.

4.3.2. Procesy wzbogacania rud

Po wydobyciu rudy muszą być poddane procesom wzbogacania, które mają na celu usunięcie niepożądanych składników mineralnych oraz zwiększenie zawartości metalu w materiale. Metody wzbogacania rud różnią się w zależności od ich składu chemicznego i mineralogicznego.

- Flotacja: jest to jedna z najczęściej stosowanych metod wzbogacania rud, szczególnie w przypadku rud miedzi. Proces ten polega na zmieszaniu zmielonej rudy z wodą i odpowiednimi odczynnikami, które umożliwiają oddzielenie metalu od skały płonnej. Pęcherzyki powietrza unoszą cząstki metalu na powierzchnię, skąd mogą być łatwo usunięte.



- Separacja magnetyczna: w przypadku rud żelaza, wykorzystuje się ich właściwości magnetyczne, aby oddzielić metal od reszty materiału. Separatory magnetyczne przyciągają cząstki żelaza, umożliwiając ich wzbogacenie.

Procesy te są niezwykle ważne, ponieważ obniżają koszty dalszej obróbki rud i poprawiają efektywność redukcji metali w następnych etapach produkcji (Mazurkiewicz i in. 2003).

4.3.3. Redukcja rud

Kolejnym krokiem w produkcji metali jest redukcja, która polega na usunięciu tlenu z tlenków metali, co umożliwia uzyskanie czystego metalu. Istnieją różne metody redukcji, z których każda jest dostosowana do rodzaju rudy i specyfiki metalu.

- **Redukcja w piecach hutniczych** - jest jedną z najstarszych i najczęściej stosowanych metod otrzymywania metali, szczególnie w przypadku produkcji żelaza i stali. Wielki piec to olbrzymie urządzenie, w którym rudy żelaza (np. hematyt lub magnetyt) są redukowane za pomocą koksu, co prowadzi do otrzymania surówki – żelaza z domieszką węgla. Surówka ta jest następnie przetwarzana w procesie stalowniczym na stal.

Proces ten obejmuje następujące reakcje chemiczne:

- redukcja tlenku węgla (CO), powstałego w wyniku spalania koksu, który reaguje z tlenkami żelaza
- powstanie czystego żelaza oraz tlenków węgla (CO₂), które są odprowadzane jako gaz odpadowy.

Tego rodzaju redukcja stosowana jest również w produkcji innych metali, takich jak cynk i ołów, jednak specyfika każdego z tych procesów może się różnić.

- **Elektroliza** - to metoda stosowana głównie w przypadku metali, które trudno jest zredukować chemicznie, jak np. aluminium. Proces ten polega na przepuszczaniu prądu elektrycznego przez stopioną rudę lub roztwór wodny zawierający metal, co powoduje rozkład chemiczny związku i wytrącanie czystego metalu na elektrodach. Jednym z najlepszych przykładów jest elektroliza tlenku glinu (Al₂O₃), który jest surowcem do produkcji aluminium. Jest on rozpuszczany w stopionej kriolicie, a następnie poddawany elektrolizie. Na katodzie osadza się aluminium, natomiast na anodzie wydziela się tlen. Proces ten jest również stosowany do oczyszczania metali, takich jak



miedź, złoto czy srebro, gdzie pozwala na uzyskanie metali o bardzo wysokiej czystości, niezbędnych w przemyśle elektronicznym i jubilerskim (Mazurkiewicz i in. 2003).

4.3.4. Rafinacja i oczyszczanie

Otrzymane w wyniku redukcji metale często zawierają domieszki, które należy usunąć, aby uzyskać metal o wysokiej czystości. Proces rafinacji, zwany także oczyszczaniem, jest kluczowy, zwłaszcza w przypadku metali szlachetnych i metali o wysokich wymaganiach jakościowych.

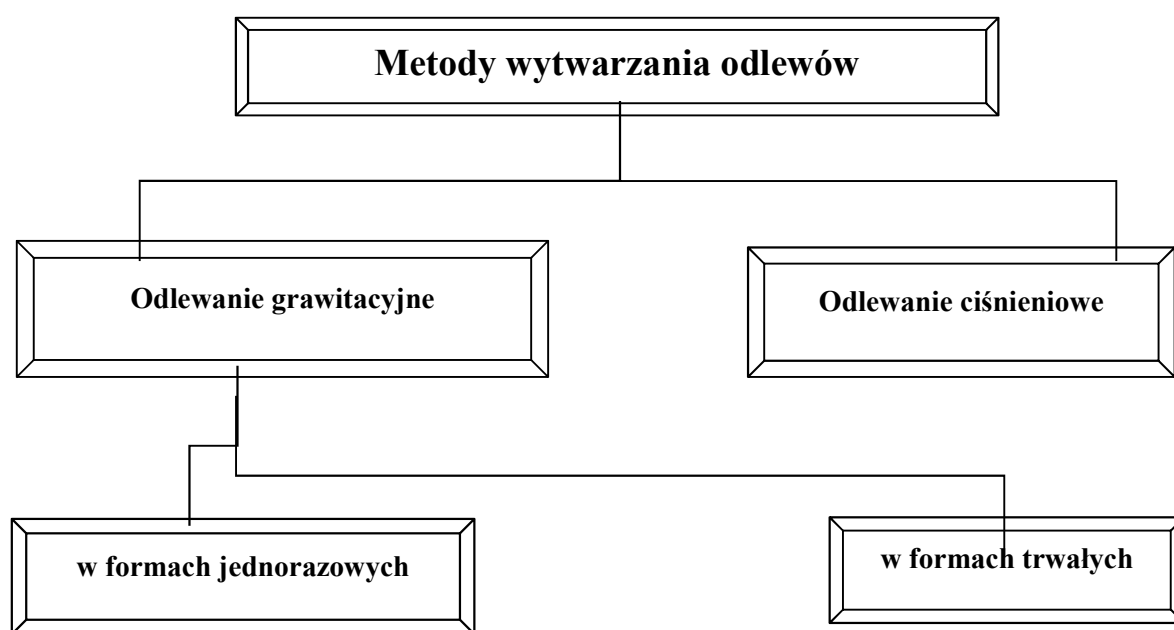
- **Rafinacja miedzi** - w przemyśle miedziowym powszechnie stosowaną metodą rafinacji jest elektroliza. Surowa miedź, uzyskana w procesach hutniczych, służy jako anoda, a czysta miedź osadza się na katodzie. W procesie tym następuje usunięcie zanieczyszczeń, takich jak ołów, cynk czy żelazo. Podobne metody stosuje się do oczyszczania innych metali, takich jak srebro czy złoto.
- Destylacja w przypadku metali o niskiej temperaturze wrzenia, takich jak cynk czy rtęć, stosuje się destylację. Metale te są podgrzewane do temperatury, w której zaczynają parować, a następnie pary te są skraplane, co pozwala na uzyskanie czystego metalu.

4.3.5. Formowanie i obróbka

Po uzyskaniu metali w postaci czystej lub rafinowanej, konieczne jest ich dalsze przetwarzanie, aby nadać im pożądane właściwości mechaniczne oraz kształty. Procesy te obejmują:

- **Walcowanie**: stosowane głównie w produkcji blach i drutów, polega na formowaniu metalu poprzez jego wielokrotne przepuszczanie między obracającymi się wałkami.
- **Kucie**: polega na kształtowaniu metalu poprzez uderzenia młota lub nacisk pras mechanicznych. Kucie umożliwia produkcję wytrzymałych elementów konstrukcyjnych, takich jak osie, wały czy koła zębate.
- **Odewanie**: proces odlewania polega na wlewaniu stopionego metalu do form, gdzie zastyga, przyjmując kształt pożądanego produktu. Odewanie jest często stosowane w produkcji silników, elementów konstrukcyjnych i części maszyn.

Odlewanie jest jednym z najstarszych i najpowszechniej stosowanych procesów produkcji metali, który polega na wlewaniu ciekłego metalu do formy, gdzie następnie zastyga, tworząc określony kształt. Proces ten umożliwia produkcję części metalowych o skomplikowanych kształtach, które trudno byłoby uzyskać za pomocą innych metod (Braszczyński 1989). Istnieje wiele różnych technik odlewania, a każda z nich jest dostosowana do specyficznych wymagań związanych z typem metalu, jego właściwościami oraz końcowym przeznaczeniem wyrobów. Poniżej opisano najpopularniejsze sposoby odlewania (Rys.1).



Rys. 1. Techniki wytwarzania form – odlewów

źródło: opracowanie własne na podstawie:

Górny Z.: *Odlewnicze stopy metali nieżelaznych* (1992)

- **Odlewanie piaskowe (w formie)** - jest jedną z najczęściej stosowanych metod, zwłaszcza w produkcji dużych i ciężkich elementów. Proces ten polega na wlewaniu ciekłego metalu do formy wykonanej z piasku kwarcowego związanego spoiwem, które utrzymuje jej kształt. Formy piaskowe są jednorazowe, co oznacza, że po zastygnięciu odlewu są niszczone, aby wydobyć gotowy produkt.

Zalety odlewania piaskowego:

- niskie koszty formy, co sprawia, że metoda ta jest opłacalna nawet przy produkcji pojedynczych elementów,
- możliwość tworzenia odlewów o skomplikowanych kształtach i dużych wymiarach.



Wady:

- gorsza jakość powierzchni odlewów w porównaniu z innymi metodami,
- niższa dokładność wymiarowa, co wymaga dodatkowej obróbki wykańczającej.
- **Odlewanie ciśnieniowe** - polega na wtryskiwaniu ciekłego metalu pod wysokim ciśnieniem do metalowej formy (matrycy). Jest to jedna z najnowocześniejszych i najbardziej precyzyjnych metod odlewania, stosowana głównie do produkcji masowej elementów o wysokiej jakości powierzchni i dokładności wymiarowej.

Zalety odlewania ciśnieniowego:

- bardzo wysoka dokładność wymiarowa i gładka powierzchnia odlewów,
- szybkość produkcji, co czyni tę metodę opłacalną przy dużych seriach produkcyjnych,
- możliwość automatyzacji procesu, co zwiększa wydajność.

Wady:

- wysokie koszty narzędzi i form, co sprawia, że metoda ta jest opłacalna głównie przy masowej produkcji,
- ograniczenia dotyczące wielkości odlewów, zazwyczaj stosowana do małych i średnich elementów.
- **Odlewanie kokilowe (grawitacyjne)** - zwane również odlewaniem grawitacyjnym, polega na wlewaniu ciekłego metalu do trwałej formy (kokili) wykonanej z metalu. Metal zastygając w formie, przybiera jej kształt. Kokile mogą być używane wielokrotnie, co czyni tę metodę bardziej opłacalną niż odlewanie piaskowe przy średnich i dużych seriach produkcji.

Zalety odlewania kokilowego:

- wyższa jakość powierzchni i dokładność wymiarowa niż w przypadku odlewania piaskowego,
- możliwość wielokrotnego użycia kokili, co obniża koszty produkcji przy większych seriach.

Wady:

- ograniczenia dotyczące kształtu i złożoności odlewów w porównaniu z odlewaniem piaskowym,
- wyższe koszty formy początkowej niż w przypadku odlewania piaskowego.
- **Odlewanie ciągłe** - to proces, w którym ciekły metal jest wylewany do formy,



a następnie schładza się i zastygając, jest ciągle wyciągany z formy w postaci długich prętów, rur lub innych kształtów. Metoda ta jest szeroko stosowana do produkcji stalowych i aluminiowych półfabrykatów, takich jak blachy, druty i profile.

Zalety odlewania ciągłego:

- wysoka wydajność i możliwość produkcji dużych ilości materiału w sposób ciągły,
- mniejsze zużycie energii i materiałów w porównaniu z tradycyjnymi metodami odlewania.

Wady:

- ograniczenia w zakresie kształtów, które można uzyskać tą metodą,
- wymaga zaawansowanego sprzętu i technologii, co zwiększa koszty początkowe inwestycji (Górny 1992).

4.4. PODZIAŁ METALI

Metale stanowią podstawową grupę materiałów wykorzystywanych w wielu dziedzinach przemysłu, inżynierii, budownictwa i technologii. Aby lepiej zrozumieć ich różnorodność oraz zastosowanie, dokonuje się podziału metali na różne kategorie w zależności od ich właściwości chemicznych, fizycznych, a także składu. Podstawowy podział obejmuje metale **żelazne i nieżelazne**, a dodatkowo uwzględnia także metale szlachetne, metale ziem rzadkich oraz stopy metali.

4.4.1. Metale żelazne

Metale żelazne, zwane również metalami ferromagnetycznymi, to grupa metali, w których głównym składnikiem chemicznym jest żelazo (Fe). Metale te charakteryzują się dobrą wytrzymałością mechaniczną, wysoką odpornością na ściskanie, a także zdolnością do przewodzenia ciepła i elektryczności. Żelazo i jego stopy są szeroko stosowane w budownictwie, motoryzacji, produkcji maszyn oraz narzędzi.

- Stal - jest najważniejszym i najpowszechniej używanym metalem żelaznym. Stanowi stop żelaza z węglem (w ilości do ok. 2%), często z dodatkiem innych pierwiastków, takich jak mangan, chrom, nikiel czy molibden, które poprawiają jej właściwości. Stal występuje w różnych rodzajach, które mają zastosowanie w zależności od potrzeb:



- Stal węglowa: najczęściej stosowana, zawiera od 0,02% do 2,11% węgla. Znajduje zastosowanie w produkcji konstrukcji stalowych, maszyn i narzędzi.
- Stal nierdzewna: dzięki zawartości chromu (co najmniej 10,5%) ma wysoką odporność na korozję i jest szeroko stosowana w przemyśle spożywczym, chemicznym oraz w budownictwie.
- Stal narzędziowa: Charakteryzuje się wysoką twardością i odpornością na zużycie, co czyni ją idealną do produkcji narzędzi.
- Żeliwo - to stop żelaza z węglem (powyżej 2,11%), często z domieszką krzemu i manganu. Jest to materiał kruchy, ale jednocześnie charakteryzuje się wysoką odpornością na ściskanie i odpornością na zużycie. Żeliwo ma zastosowanie w produkcji rur, elementów maszyn oraz urządzeń grzewczych.
- Stopy żelaza – żelazo może być mieszane z innymi metalami w celu tworzenia stopów o określonych właściwościach. Przykłady takich stopów to stal narzędziowa, stale sprężynowe czy stale konstrukcyjne, które są używane w wielu gałęziach przemysłu, w tym w motoryzacji i lotnictwie.

4.4.2. Metale nieżelazne

Metale nieżelazne to grupa metali, które nie zawierają żelaza lub zawierają je w niewielkich ilościach. Charakteryzują się różnorodnymi właściwościami, takimi jak odporność na korozję, lekkość, wysoka przewodność elektryczna oraz łatwość formowania. W tej grupie wyróżniamy kilka podkategorii, takich jak metale lekkie, metale ciężkie, a także metale szlachetne.

- Metale lekkie - charakteryzują się niską gęstością i są szczególnie cenione w przemyśle lotniczym, motoryzacyjnym oraz w produkcji sprzętu sportowego. Najważniejsze metale z tej grupy przedstawiono poniżej:
 - Aluminium (Al): posiada wysoką odporność na korozję, jest lekkie i ma dobrą przewodność cieplną oraz elektryczną. Stosowane jest w konstrukcjach samolotów, samochodów, puszkach do napojów oraz oknach i drzwiach.
 - Tytan (Ti): lekki, wytrzymały, odporny na korozję i wysokie temperatury ma zastosowanie w przemyśle lotniczym, wojskowym oraz medycznym, zwłaszcza w implantach i protezach.



- Magnez (Mg): jeden z najlżejszych metali strukturalnych, stosowany w produkcji lekkich konstrukcji, a także w stopach z aluminium, co zwiększa ich wytrzymałość.
- Metale ciężkie - charakteryzują się większą gęstością niż metale lekkie i znajdują zastosowanie w przemyśle ciężkim oraz elektrotechnice. Najważniejsze z nich to:
 - Miedź (Cu): posiada doskonałą przewodność elektryczną i cieplną, co sprawia, że jest niezastąpiona w przewodach elektrycznych, elektronice, a także w przemyśle grzewczym, posiada także odporność na korozję.
 - Ołów (Pb): dzięki swojej dużej gęstości i odporności na korozję, znajduje zastosowanie w akumulatorach samochodowych, osłonach przed promieniowaniem i w budownictwie.
 - Cynk (Zn): używany głównie do pokrywania stali w procesie cynkowania, co zapobiega jej korozji. Cynk jest również składnikiem stopów, takich jak mosiądz (stop cynku z miedzią).
- Metale szlachetne - to grupa metali charakteryzujących się dużą odpornością na korozję i utlenianie, a także wysoką wartością ekonomiczną. Stosowane są głównie w jubilerstwie, elektronice oraz przemyśle chemicznym. Do najważniejszych metali szlachetnych należą:
 - Złoto (Au): posiada doskonałe właściwości przewodzące, jest odporne na utlenianie i korozję, co czyni je idealnym materiałem do produkcji biżuterii oraz elektroniki.
 - Srebro (Ag): najlepszy przewodnik elektryczności wśród metali, stosowane w elektronice, medycynie oraz jubilerstwie.
 - Platyna (Pt): stosowana w katalizatorach samochodowych, przemyśle chemicznym oraz jubilerstwie, platyna jest niezwykle odporna na wysokie temperatury i korozję.

4.4.3. Metale ziem rzadkich

Metale ziem rzadkich to grupa 17 pierwiastków chemicznych, które mają specyficzne właściwości magnetyczne, optyczne oraz katalityczne. Choć ich nazwa sugeruje, że są one rzadkie, w rzeczywistości występują one w przyrodzie, ale są rozproszone i trudne do wydobycia w czystej formie. Metale ziem rzadkich mają kluczowe znaczenie w nowoczesnych technologiach, takich jak:

- Neodym (Nd): używany do produkcji silnych magnesów stosowanych w silnikach elektrycznych, turbinach wiatrowych i głośnikach.



- Lantan (La): stosowany w akumulatorach, katalizatorach i optyce.
- Europ (Eu): używany w produkcji fosforów w ekranach telewizyjnych oraz oświetleniu fluorescencyjnym czy energii jądrowej jako absorber neutronów.

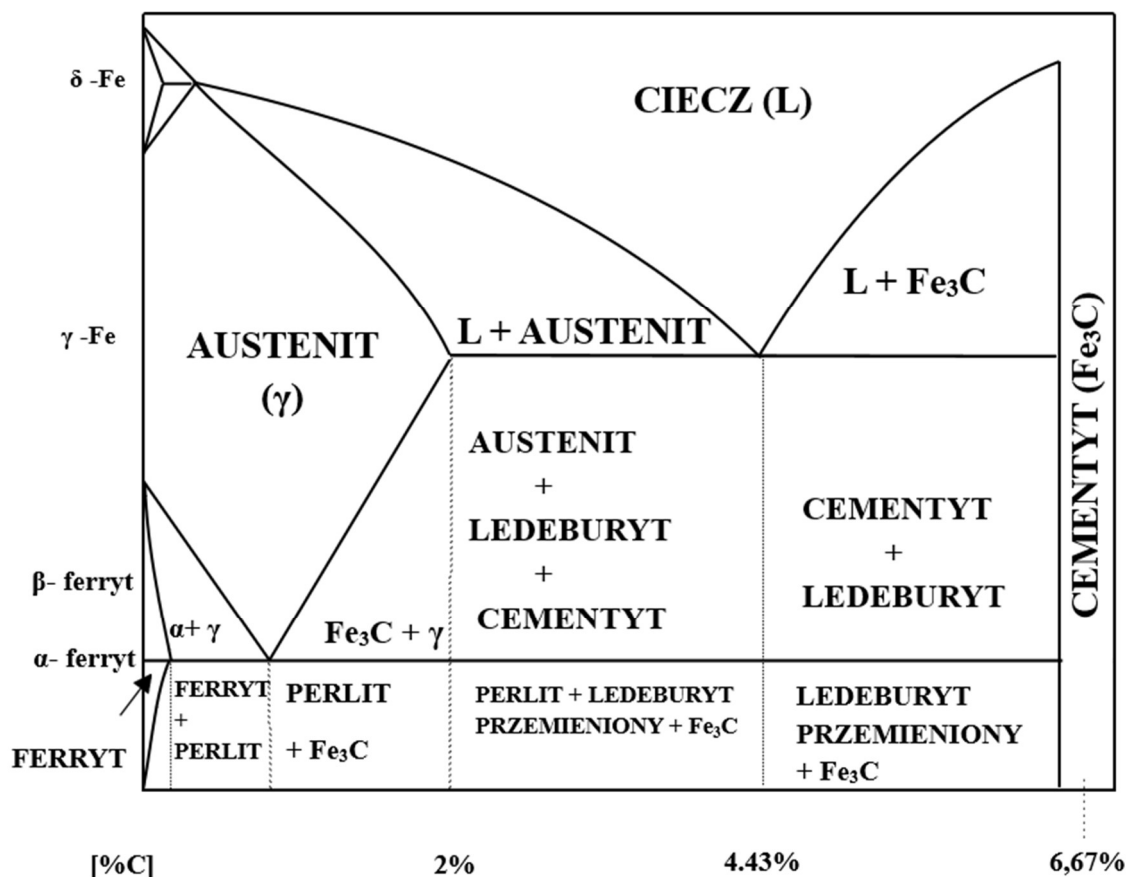
4.4.4. Stopy metali

Stopy metali to materiały, które składają się z co najmniej dwóch pierwiastków, gdzie przynajmniej jednym jest metal. Stopy są tworzone w celu uzyskania określonych właściwości, takich jak wytrzymałość, twardość, odporność na korozję czy zdolność do przewodzenia ciepła i prądu. Najważniejsze przykłady stopów to m.in.:

- Stal: opisana wcześniej, będąca stopem żelaza i węgla.
- Mosiądz: stop miedzi i cynku, szeroko stosowany w armaturze, instrumentach muzycznych oraz w dekoracjach.
- Brąz: stop miedzi z cyną, używany w produkcji elementów mechanicznych, jak łożyska, oraz w sztuce.

4.5. WYKRES ŻELAZO - WĘGIEL

Wykres żelazo-węgiel (znany również jako diagram fazowy żelazo-węgiel), przedstawiony na Rys. 2, to podstawowe narzędzie używane w metalurgii do zrozumienia przemian fazowych w stopach żelaza z węglem.



Rys. 2. Wykres fazowy żelazo – węgiel

źródło: opracowanie własne na podstawie: Dobrzański L.A.,
Podstawy nauki o materiałach i metaloznawstwo, 2002

Umożliwia on przewidywanie, jakie struktury krystaliczne (fazy) będą występować w stali i żelowie w zależności od temperatury oraz zawartości węgla. Wiedza ta jest kluczowa w projektowaniu procesów obróbki cieplnej, takich jak hartowanie, odpuszczanie czy wyżarzanie, ponieważ różne struktury krystaliczne mają wpływ na właściwości mechaniczne stopów. Wykres żelazo-węgiel obejmuje kilka istotnych obszarów i punktów, które mają kluczowe znaczenie dla zrozumienia przemian fazowych:

- Ferryt to faza o strukturze regularnej przestrzennie centrowanej (RPC), która jest stabilna w temperaturach poniżej 912°C. Ferryt zawiera bardzo małe ilości węgla (do 0,02%) i jest stosunkowo miękki oraz plastyczny. Występuje w stali niskowęglowej oraz w stalach o bardzo niskiej zawartości węgla.
- **Austenit** - to faza, która występuje w stali w wyższych temperaturach (powyżej 727°C) i ma strukturę regularną ściennie centrowaną (RSC). Austenit może rozpuszczać więcej



węgla niż ferryt (do 2,11%), co sprawia, że ma większą wytrzymałość i plastyczność. W procesie obróbki cieplnej stali austenit może przekształcać się w inne fazy, takie jak martenzyt lub perlit, co jest zależne od szybkości chłodzenia.

- **Perlit** - to mieszanina ferrytu i cementytu, która tworzy się w wyniku powolnego chłodzenia austenitu. Ma strukturę lamelarną (warstwową) i występuje w stali o zawartości węgla około 0,8%. Perlit łączy w sobie cechy obu faz: plastyczność ferrytu oraz twardość cementytu. Jest to jedna z najważniejszych struktur w stalach o średniej zawartości węgla.
- **Cementyt (Fe_3C)** - czyli węglik żelaza III, to twarda i krucha faza, która występuje w stopach żelaza z węglem. Ma bardzo dużą zawartość węgla (6,67%) i odgrywa kluczową rolę w kształtowaniu właściwości mechanicznych stali. Cementyt tworzy się w wyniku przesycenia austenitu węglem oraz podczas chłodzenia stali i żeliwa.
- **Martenzyt** - to faza, która powstaje w wyniku bardzo szybkiego chłodzenia (hartowania) austenitu. Ma strukturę tetragonalną, która jest bardzo twarda i krucha. Stal hartowana, w której występuje martenzyt, charakteryzuje się dużą wytrzymałością, jednak jest ona bardzo krucha, dlatego często poddaje się ją dodatkowej obróbce cieplnej (np. odpuszczaniu) w celu poprawy plastyczności.

Liniowe punkty wykresu: eutektyka i eutektoid

- Punkt eutektyczny znajduje się na wykresie w temperaturze 1147°C przy zawartości węgla 4,3%. To tutaj zachodzi przemiana ciecz $\rightarrow$ austenit + cementyt (w przypadku żeliwa).
- Punkt eutektoidalny znajduje się przy temperaturze 727°C i zawartości węgla 0,8%, w tym punkcie austenit przekształca się w mieszaninę ferrytu i perlitu.

Zastosowanie wykresu żelazo-węgiel: jest to niezbędne narzędzie dla inżynierów i metalurgów, ponieważ umożliwia przewidywanie przemian fazowych i odpowiednie dostosowanie procesów obróbki cieplnej, aby uzyskać pożądane właściwości mechaniczne stali. W zależności od zawartości węgla oraz temperatury obróbki cieplnej, można modyfikować strukturę stopu, co pozwala na tworzenie materiałów o różnych parametrach wytrzymałościowych, twardości, elastyczności i odporności na ścieranie.



4.6. PODZIAŁ STALI

Podział stali może być dokonany według różnych kryteriów, takich jak: zawartość węgla, skład chemiczny, struktura, sposób obróbki cieplnej oraz zastosowanie.

4.6.1. Podział ze względu na zawartość węgla

- Stale niskowęglowe (miękkie): zawierają do 0,25% węgla. Charakteryzują się dobrą plastycznością i są łatwe w obróbce plastycznej na zimno. Znajdują zastosowanie m.in. w przemyśle motoryzacyjnym oraz budowlanym, w produkcji blach, rur i prętów.
- Stale średniowęglowe: ich zawartość węgla wynosi od 0,25% do 0,60%. Posiadają lepszą wytrzymałość niż stale niskowęglowe, ale ich plastyczność jest mniejsza. Znajdują zastosowanie w produkcji narzędzi oraz części maszyn o wyższej wytrzymałości, takich jak wały, korbowody czy sprężyny.
- Stale wysokowęglowe (twarde): zawierają powyżej 0,60% węgla. Są bardzo twarde i odporne na zużycie, ale trudne w obróbce plastycznej. Wykorzystywane są m.in. do produkcji narzędzi tnących, łożysk oraz sprężyn.

4.6.2. Podział według składu chemicznego

- Stale niestopowe (węglowe): zawierają głównie żelazo i węgiel, bez istotnych dodatków stopowych. Ich właściwości zależą przede wszystkim od zawartości węgla. Znajdują zastosowanie w budownictwie, motoryzacji oraz jako materiały konstrukcyjne.
- Stale stopowe: oprócz węgla, zawierają również inne pierwiastki stopowe, takie jak chrom, nikiel, molibden, wanad, które nadają im specyficzne właściwości. Mogą być odporne na korozję, wysokie temperatury lub zmęczenie materiału. Przykłady stali stopowych to:
 - Stale nierdzewne: zawierają co najmniej 11% chromu, co nadaje im odporność na korozję. Stosowane są w produkcji naczyń, sprzętu medycznego, rur oraz konstrukcji narażonych na działanie czynników atmosferycznych.
 - Stale narzędziowe: wzbogacane o takie pierwiastki jak wolfram, chrom czy wanad, są używane do produkcji narzędzi o wysokiej odporności na zużycie.



4.6.3. Podział według struktury

- Stale ferrytyczne: charakteryzują się obecnością ferrytu w strukturze, co daje im dobrą odporność na korozję i plastyczność, jednak ich wytrzymałość jest relatywnie niska. Stosowane są głównie w przemyśle chemicznym i naftowym.
- Stale martenzytyczne: otrzymywane w wyniku obróbki cieplnej, mają strukturę martenzytu, która zapewnia im dużą twardość. Wykorzystywane są w produkcji narzędzi i elementów maszyn wymagających dużej wytrzymałości.
- Stale austenityczne: zawierają nikiel i chrom, co powoduje, że są odporne na korozję, mają dużą wytrzymałość i dobrą plastyczność. Często wykorzystywane w przemyśle spożywczym, farmaceutycznym oraz przy produkcji wyrobów precyzyjnych.

4.6.4. Stale szybko tnące

Stal szybko tnąca (oznaczana skrótem HSS - *High Speed Steel*) jest szczególną odmianą stali narzędziowej, która charakteryzuje się zdolnością do zachowania swoich właściwości, nawet w bardzo wysokich temperaturach. Dzieje się tak dzięki obecności pierwiastków stopowych, takich jak wolfram, molibden, wanad oraz kobalt, które poprawiają jej odporność na wysokie temperatury i ścieranie.

Charakterystyka:

- Właściwości: wyjątkowo odporna na ścieranie, zachowuje twardość nawet w temperaturze powyżej 600°C. Ma bardzo wysoką wytrzymałość i trwałość.
- Zastosowanie: stale szybko tnące są powszechnie używane do produkcji narzędzi skrawających, takich jak frezy, wiertła, piły i narzędzia tokarskie, które muszą pracować z dużymi prędkościami obróbki i przy dużym obciążeniu termicznym.

4.6.5. Stale żarowytrzymałe

Stale żarowytrzymałe to specjalne stale stopowe, które wykazują dużą odporność na działanie wysokich temperatur i obciążeń mechanicznych. Głównymi pierwiastkami stopowymi odpowiedzialnymi za te właściwości są chrom, nikiel, molibden i wolfram. Te stale są przystosowane do pracy w temperaturach powyżej 500°C, nie tracąc swoich właściwości wytrzymałościowych.

Charakterystyka:



- Właściwości: wysoka odporność na pękanie, korozję i utlenianie w wysokich temperaturach. Zachowują dobrą wytrzymałość mechaniczną przy długotrwałej ekspozycji na ciepło.
- Zastosowanie: stale żarowytrzymałe są stosowane w przemyśle energetycznym, lotniczym oraz motoryzacyjnym. Wykorzystywane są do produkcji elementów turbin, pieców przemysłowych, wymienników ciepła oraz kotłów, które muszą wytrzymać działanie ekstremalnych temperatur i dużych obciążeń.

4.6.6. Podział według zastosowania stali:

- Stale konstrukcyjne: używane do budowy konstrukcji maszyn, budynków i innych dużych obiektów inżynierskich. Ich cechą charakterystyczną jest dobra wytrzymałość na obciążenia oraz plastyczność.
- Stale narzędziowe: przeznaczone do produkcji narzędzi tnących, form, matryc, które muszą być odporne na zużycie i wysokie temperatury. Mają wysoką twardość i są stosowane w przemyśle metalurgicznym i motoryzacyjnym.
- Stale odporne na korozję: dzięki dodatkom stopowym, takim jak chrom, są odporne na działanie środowisk agresywnych, np. kwasów czy soli. Wykorzystywane są w przemyśle chemicznym, morskim oraz farmaceutycznym.

4.7. ŻELIWO

Żeliwo to stop żelaza z węglem, którego zawartość węgla wynosi zazwyczaj powyżej 2,14%. W przeciwieństwie do stali, gdzie węgiel jest rozpuszczony w roztworze żelaznym, w żeliwie występuje on w postaci wolnej (grafit) lub związanej (cementyt). Charakteryzuje się ono dużą twardością i kruchością, ale również doskonałą leżnością, co czyni je jednym z najczęściej stosowanych materiałów odlewniczych. Dzięki swoim specyficznym właściwościom, żeliwo znajduje zastosowanie w wielu dziedzinach przemysłu, takich jak budowa maszyn, motoryzacja czy przemysł ciężki.

Podział żeliwa opiera się na jego strukturze, sposobie krystalizacji oraz na formie, w jakiej występuje węgiel w stopie. Wyróżniamy głównie żeliwo szare, białe, sferoidalne i ciągliwe, z których każde posiada specyficzne cechy i zastosowania.



4.7.1. Podział żeliwa

Żeliwo szare to najpowszechniej używany rodzaj żeliwa, w którym węgiel występuje w postaci grafitu. Cechą charakterystyczną tego materiału jest jego szary kolor na powierzchni przełomu, co wynika z obecności grafitu. Ze względu na łatwość w odlewaniu i dobrą obrabialność, jest szeroko stosowane w przemyśle.

Charakterystyka:

- **Struktura:** węgiel w postaci płatków grafitu w osnowie metalicznej (ferryt lub perlit).
- **Właściwości:** dobrze tłumi drgania, jest odporny na ścieranie i łatwo się obrabia, jego wadą jest kruchość i niska wytrzymałość na rozciąganie.
- **Zastosowanie:** żeliwo szare stosuje się do produkcji odlewów maszynowych, części silników, bloków silnikowych, obudów skrzyni biegów, a także elementów rur kanalizacyjnych.

Żeliwo białe większość węgla w nim zawarta jest związana w postaci cementytu (Fe_3C), co sprawia, że materiał ten nie ma widocznych płatków grafitu, a przełom odlewu ma biały połysk. Żeliwo białe jest bardzo twarde i kruche, przez co jest trudne w obróbce, ale charakteryzuje się wysoką odpornością na zużycie.

Charakterystyka:

- **Struktura:** węgiel występuje głównie w postaci węglika żelaza (cementytu), bez widocznego grafitu.
- **Właściwości:** jest twarde, bardzo kruche i trudne do obrabiania mechanicznego. Jest odporne na ścieranie, jednakże słabo tłumi drgania.
- **Zastosowanie:** wykorzystywane do produkcji elementów narażonych na wysokie ścieranie, takich jak części młynów, walce hutnicze, bębny hamulcowe oraz narzędzia tnące.

Żeliwo sferoidalne (żeliwo z grafitem kulkowym) jest modyfikowaną wersją żeliwa szarego, w której grafit przybiera postać kulek, a nie płatków. Taka struktura poprawia właściwości mechaniczne tego żeliwa, zwłaszcza jego wytrzymałość i odporność na pękanie. W celu uzyskania kulistej formy grafitu, do stopu dodaje się niewielkie ilości magnezu lub ceru (Rączka i Sakwa 1986):.

Charakterystyka:

- **Struktura:** węgiel w postaci sferoidalnych cząstek grafitu równomiernie rozmieszczonych w osnowie ferrytowo-perlitowej.



- **Właściwości:** jest bardziej wytrzymałe i elastyczne niż żeliwo szare, charakteryzuje się większą odpornością na pękanie i lepszą wytrzymałością na rozciąganie. Posiada także lepsze właściwości plastyczne.
- **Zastosowanie:** wykorzystywane do produkcji wałów korbowych, zaworów, rur ciśnieniowych, elementów zawieszonych samochodowych oraz części maszyn pracujących pod dużymi obciążeniami.

Żeliwo ciągliwe powstaje poprzez odpowiednią obróbkę cieplną żeliwa białego. Proces ten powoduje wytrącanie węgla w postaci drobnych cząstek grafitu, co sprawia, że żeliwo staje się bardziej plastyczne i ciągliwe, a jednocześnie zachowuje część twardości żeliwa białego.

Charakterystyka:

- **Struktura:** grafit występuje w formie drobnych, kulistych wtrąceń, co poprawia jego plastyczność i wytrzymałość mechaniczną.
- **Właściwości:** wysoka plastyczność, dobra odporność na rozciąganie i udarność. Żeliwo ciągliwe jest bardziej wytrzymałe na rozciąganie niż żeliwo szare.
- **Zastosowanie:** stosowane w produkcji elementów wymagających wysokiej plastyczności i wytrzymałości, takich jak łączniki rurowe, osie, ramy pojazdów oraz części maszyn precyzyjnych.

4.7.2. Właściwości i zastosowanie żeliwa

Żeliwo charakteryzuje się następującymi właściwościami:

- **Twardość i odporność na ścieranie:** dzięki zawartości cementytu (w żeliwie białym) lub grafitu (w żeliwie szarym), żeliwo jest wytrzymałe na ścieranie, co czyni je idealnym do zastosowań narażonych na intensywną eksploatację.
- **Kruchość:** większość żeliw, szczególnie białe i szare, charakteryzuje się niską plastycznością i dużą kruchością.
- **Dobra lejność:** żeliwo łatwo formuje się w skomplikowane kształty podczas odlewania, co czyni je idealnym materiałem do produkcji odlewów.
- **Tłumienie drgań:** żeliwo, zwłaszcza szare, ma doskonałą zdolność tłumienia drgań, co czyni je użytecznym w produkcji maszyn pracujących w warunkach dużych obciążeń dynamicznych.

Żeliwo, dzięki swojej różnorodności, znalazło zastosowanie w wielu dziedzinach przemysłu:



- **Budownictwo:** żeliwo stosowane jest w produkcji elementów konstrukcyjnych, takich jak rury kanalizacyjne, pokrywy studzienek, kraty, a także w armaturze wodociągowej.
- **Motoryzacja:** żeliwo sferoidalne oraz ciągliwe wykorzystywane jest do produkcji części samochodowych, takich jak zawieszenia, wały korbowe, obudowy silników i inne elementy narażone na duże obciążenia.
- **Przemysł maszynowy:** żeliwo szare znajduje zastosowanie w produkcji podstaw i korpusów maszyn, gdzie istotne jest tłumienie drgań. Żeliwo białe z kolei jest wykorzystywane w częściach maszyn narażonych na duże ścieranie, takich jak młyny i walce.
- **Energetyka:** żeliwo, szczególnie sferoidalne, używane jest do produkcji rur ciśnieniowych, kotłów i innych elementów instalacji, gdzie wymagane są dobre właściwości mechaniczne przy pracy w wysokich temperaturach i ciśnieniach.

4.8. OBRÓBKĄ CIEPLNĄ I CIEPLNO-CHEMICZNĄ METALI

Obróbka cieplna i cieplno-chemiczna to jedne z kluczowych procesów stosowanych w inżynierii materiałowej, które umożliwiają zmianę właściwości fizycznych, mechanicznych i chemicznych metali. Głównie cele tych procesów obejmują zwiększenie twardości, odporności na zużycie, poprawę wytrzymałości oraz odporności na korozję. Dzięki zastosowaniu obróbki cieplnej możliwe jest dostosowanie właściwości materiału do specyficznych wymagań eksploatacyjnych. Procesy cieplno-chemiczne dodatkowo wzbogacają obróbkę cieplną poprzez modyfikację składu chemicznego warstwy powierzchniowej, co prowadzi do uzyskania bardziej pożądanego właściwości na powierzchni materiału.

Obróbka cieplna polega na precyzyjnym sterowaniu temperaturą, czasem trwania poszczególnych faz oraz szybkością chłodzenia. Procesy te opierają się na termodynamicznych przemianach w materiale, co pozwala na kontrolowane zmiany struktury krystalicznej i mikrostruktury stopu. Z kolei obróbka cieplno-chemiczna, oprócz zmian temperaturowych, wprowadza do materiału aktywne chemicznie pierwiastki, takie jak węgiel, azot czy bor, co prowadzi do modyfikacji właściwości warstwy wierzchniej (WW).



4.8.1. Procesy obróbki cieplnej:

Wyżarzanie to jeden z najstarszych i najczęściej stosowanych procesów obróbki cieplnej. Polega on na nagrzeniu metalu do odpowiedniej temperatury, utrzymywaniu tej temperatury przez określony czas, a następnie powolnym chłodzeniu. Celem wyżarzania jest redukcja naprężeń wewnętrznych, poprawa plastyczności oraz zmiana struktury krystalicznej. W zależności od zastosowanego procesu, wyróżnia się różne rodzaje wyżarzania:

- **Wyżarzanie pełne** - polega na nagrzeniu materiału do temperatury wyższej niż temperatura rekrytalizacji i jego powolnym chłodzeniu, dzięki temu procesowi uzyskuje się materiał o niskim poziomie naprężeń i jednorodnej mikrostrukturze.
- **Wyżarzanie normalizujące** - odbywa się przez nagrzenie metalu do temperatury nieco powyżej temperatury przemiany fazowej i chłodzeniu w powietrzu, proces ten prowadzi do uzyskania drobnoziarnistej struktury i poprawy właściwości mechanicznych, takich jak wytrzymałość i twardość.
- **Wyżarzanie rekrytalizujące** - stosowane w celu usunięcia efektów zgniotu po obróbce plastycznej na zimno, polega na nagrzeniu materiału do temperatury rekrytalizacji, co powoduje wzrost nowych ziaren.

Hartowanie jest procesem, który ma na celu zwiększenie twardości i wytrzymałości metali. Polega na szybkim schładzaniu metalu z wysokiej temperatury po uprzednim nagrzeniu go do odpowiedniego zakresu temperatury (zwykle powyżej temperatury przemiany fazowej). Proces ten prowadzi do przemiany austenitu w martenzyt, co znacząco zwiększa twardość materiału, hartowanie stosuje się głównie dla stali, jednak proces ten można stosować również dla innych stopów metali. Ważnym elementem hartowania jest wybór odpowiedniego medium chłodzącego, zastosowanie wody, oleju lub powietrza wpływa na szybkość chłodzenia, a tym samym na ostateczne właściwości materiału. Istnieje również proces zwany hartowaniem powierzchniowym, w którym szybkie schłodzenie dotyczy jedynie warstwy wierzchniej materiału, podczas gdy wnętrze pozostaje mniej zmienione.

Odpuszczanie to proces stosowany po hartowaniu w celu zredukowania naprężeń i poprawy udarności materiału. Polega on na nagrzeniu metalu do temperatury niższej niż temperatura hartowania, utrzymywaniu go w tej temperaturze przez określony czas, a następnie powolnym chłodzeniu, w zależności od temperatury odpuszczania można uzyskać różne kombinacje twardości i plastyczności. Odpuszczanie niskie prowadzi do zachowania wysokiej twardości, natomiast odpuszczanie wysokie poprawia plastyczność kosztem pewnej utraty twardości.



Starzenie - proces starzenia odnosi się głównie do stopów aluminium, miedzi oraz stali nierdzewnych i ma na celu poprawę ich właściwości mechanicznych, takich jak twardość i wytrzymałość. Starzenie może być naturalne (przebiegające w temperaturze otoczenia) lub sztuczne (przyspieszane w wyższych temperaturach). W wyniku starzenia dochodzi do wydzielania się faz wtórnych, co prowadzi do utwardzenia materiału.

4.8.2. Obróbka cieplno-chemiczna

Obróbka cieplno-chemiczna jest procesem bardziej zaawansowanym niż tradycyjna obróbka cieplna, oprócz poddawania metalu zmianom temperaturowym, modyfikowane są także jego właściwości chemiczne. Do materiału wprowadzane są różne pierwiastki chemiczne, takie jak węgiel, azot, bor czy chrom, co prowadzi do zmiany składu chemicznego warstwy powierzchniowej metalu.

Nasycanie węglem - nawęglanie - to proces cieplno-chemiczny, w którym powierzchnia metalu jest nasycana węglem, polega on na umieszczeniu metalowych elementów w atmosferze bogatej w węgiel (np. gaz, proszek lub ciecz) w podwyższonej temperaturze. Dzięki temu procesowi powierzchnia materiału staje się znacznie twardsza, podczas gdy jego wnętrze zachowuje pierwotną plastyczność. Nawęglanie jest szczególnie popularne w przypadku stali niskowęglowych, które po obróbce uzyskują doskonałe właściwości wytrzymałościowe i twardość na powierzchni.

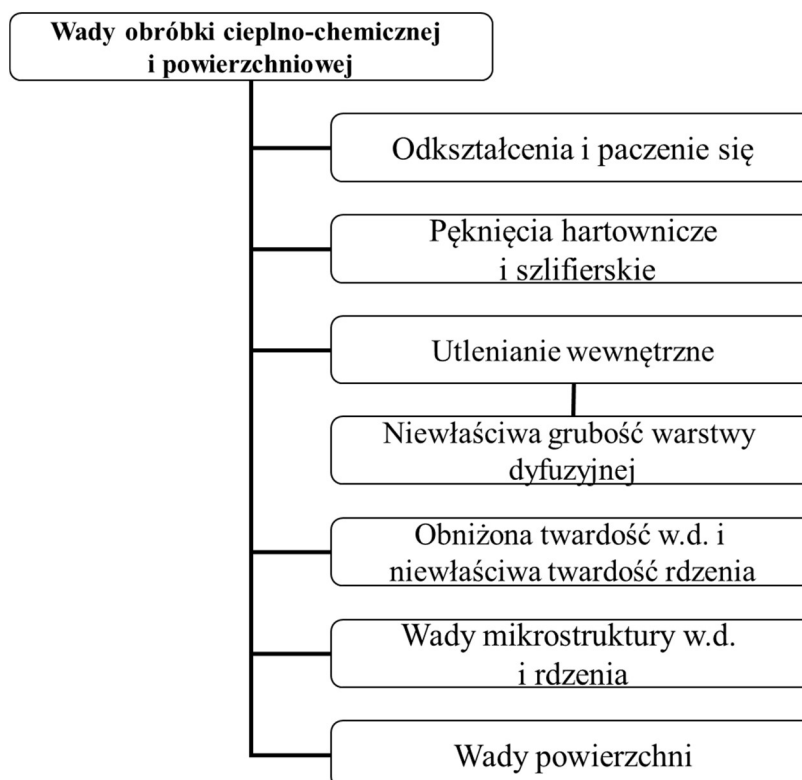
Nasycanie azotem - azotowanie jest procesem, w którym powierzchnia metalu nasycana jest azotem. Odbywa się w atmosferze amoniaku lub w środowisku gazowym zawierającym azot w podwyższonej temperaturze, proces ten prowadzi do powstania bardzo twardej warstwy powierzchniowej, która jest odporna na zużycie i korozję. Azotowanie stosowane jest głównie do stali stopowych, narzędziowych oraz części maszyn, które wymagają wysokiej odporności na ścieranie.

Cynkowanie, chromowanie i borowanie – to inne formy obróbki cieplno-chemicznej, które mają na celu zwiększenie odporności metalu na korozję oraz zużycie. Cynkowanie polega na nasyceniu powierzchni metalu cynkiem, co zabezpiecza go przed korozją atmosferyczną. Chromowanie polega na wprowadzeniu chromu do powierzchni, co nadaje jej wysoką odporność na ścieranie i korozję. Borowanie natomiast zwiększa twardość i odporność na zużycie poprzez wprowadzenie boru do powierzchni materiału.



4.8.3. Zastosowanie obróbki cieplnej i cieplno-chemicznej

Procesy obróbki cieplnej i cieplno-chemicznej znajdują szerokie zastosowanie w przemyśle, w szczególności w produkcji narzędzi, elementów maszyn, części samochodowych i lotniczych. Dzięki nim możliwe jest dostosowanie właściwości materiałów do specyficznych wymagań, takich jak wysoka twardość, odporność na zużycie, wytrzymałość na zmęczenie czy odporność na korozję. Procesy te mają także kluczowe znaczenie w przemyśle energetycznym, gdzie trwałość i wytrzymałość materiałów są priorytetowe. Obróbka cieplno-chemiczna źle przeprowadzona może powodować różnego rodzaju wady i zmiany powierzchni materiału (Rys. 3).



Rys. 3. Wady obróbki cieplno-chemicznej
źródło: opracowanie własne na podstawie: Dobrzański L.A.,
Podstawy nauki o materiałach i metaloznawstwo, 2000

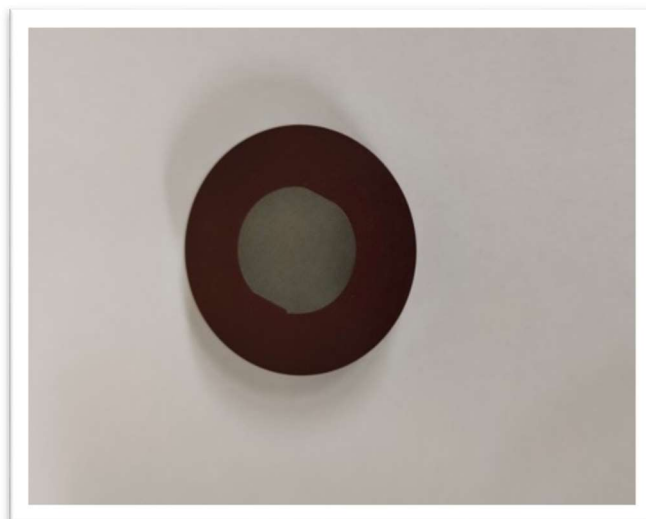
4.9. BADANIA METALOGRAFICZNE

Badania mikroskopowe polegają na analizie powierzchni próbek, zarówno tych naturalnych, jak i odpowiednio przygotowanych przekrojów (zglądów), przy użyciu powiększeń od 40 do 1500x. Dzięki temu można:



- ocenić mikrostrukturę materiału, biorąc pod uwagę rodzaj, rozmieszczenie i wielkość poszczególnych składników,
- określić strukturę i grubość różnych warstw, takich jak hartowane, dyfuzyjne czy galwaniczne,
- zidentyfikować wady, np. wtrącenia niemetaliczne, mikropęknięcia czy korozję międzykrystaliczną,
- przybliżenie zawartości węgla w stalach niestopowych poprzez analizę udziału różnych faz (Przybyłowicz 2011).

Do tych obserwacji używa się mikroskopów metalograficznych, które pracują na zasadzie odbicia światła od nieprzezroczystych powierzchni próbek -zgładów, (Rys.4), co odróżnia je od mikroskopów biologicznych, które wymagają przezroczystości próbek dla światła widzialnego.



Rys. 4. Zgład metalograficzny przygotowany do badania
źródło: opracowanie własne – fotografia z kolekcji autorki

Przygotowanie zgładów metalograficznych obejmuje kilka etapów:

- Przygotowanie próbek: proces zaczyna się od wycięcia fragmentu reprezentatywnego dla badanego materiału, a następnie szlifowania, polerowania i trawienia.
- Wycinanie próbek: wybór odpowiedniego miejsca do pobrania próbki wymaga doświadczenia. Kluczowe jest, aby proces ten nie zmienił struktury materiału, np. przez zastosowanie narzędzi wysokoobrotowych bez chłodzenia. Należy unikać wpływu ciepła podczas cięcia, np. przy użyciu palnika.



- Szlifowanie: odbywa się na specjalnych szlifierkach, przy użyciu papieru ściernego o różnych gradacjach (od 120 do 600), zmieniając kierunek szlifowania o 90° po każdej zmianie papieru. Małe próbki lub te o delikatnych krawędziach umieszcza się w uchwytych lub inkluduje w żywicach. Dla cienkich warstw stosuje się czasem zgłady skośne, co pozwala uzyskać dokładniejsze dane.
- Polerowanie: w celu uzyskania lustrzanej powierzchni zgładu, stosuje się polerowanie mechaniczne, chemiczne lub elektrolityczne. Polerowanie mechaniczne, najczęściej stosowane, usuwa rysy i powierzchniową warstwę zgniotu, jednak czasem wymaga ponownego trawienia. Polerowanie elektrolityczne działa poprzez anodowe rozpuszczanie próbki, co pozwala uzyskać gładką powierzchnię, ale jest najlepiej dopasowane do stopów jednofazowych.
- Trawienie: polega na aplikacji odpowiednich substancji chemicznych, które reagują z powierzchnią próbki, odsłaniając mikrostrukturę. W zależności od badanego materiału i celu analizy, dobierane są różne odczynniki oraz parametry trawienia, takie jak czas, temperatura i stężenie. Jeśli efekt trawienia nie jest zadowalający, można powtórzyć proces, modyfikując jego warunki.

Każdy etap ma na celu jak najlepsze przygotowanie próbki do analizy mikroskopowej, co pozwala na uzyskanie precyzyjnych wyników dotyczących struktury materiału (Kubiński 2019): .

4.9.1. Mikroskop metalograficzny, zasada działania

Mikroskop metalograficzny (Rys. 5) to precyzyjne urządzenie optyczne, które jest wykorzystywane do obserwacji nieprzezroczystych materiałów, takich jak metale, stopy, ceramika czy materiały kompozytowe. Jego główną cechą jest zdolność do analizowania powierzchni próbek, które nie przepuszczają światła, co odróżnia go od mikroskopów biologicznych, które wymagają przezroczystości próbek.



Rys. 5. Mikroskop metalograficzny Delta Optical MET-1000 TRF,
źródło: opracowanie własne, fotografia z kolekcji autorki

Elementy budowy mikroskopu metalograficznego:

- Źródło światła: mikroskopy metalograficzne wykorzystują światło odbite od powierzchni próbki, dlatego niezbędne jest zastosowanie intensywnego i równomiernego źródła światła. Zwykle stosuje się oświetlenie halogenowe lub LED, które zapewnia odpowiednią jasność i stabilność strumienia świetlnego, oświetlenie może być regulowane, aby dostosować intensywność do potrzeb obserwacji.
- Układ optyczny: mikroskopy metalograficzne wyposażone są w zestaw soczewek tworzących układ optyczny, który składa się z:
 - Obiektywów - są to soczewki znajdujące się najbliżej próbki, odpowiadające za uzyskanie obrazu. W mikroskopach metalograficznych stosuje się obiektywy o różnym powiększeniu, np. 5x, 10x, 20x, 50x, a nawet 100x, obiektywy te są umieszczone w obrotowej głowicy, co pozwala na ich szybkie przełączanie.
 - Okularów - stanowią one soczewki, przez które użytkownik obserwuje powiększony obraz próbki. Mikroskopy metalograficzne mogą być wyposażone w okulary o powiększeniu od 10x do 20x, łącznie z obiektywem tworzą końcowe powiększenie obserwowanego obrazu.
 - Tuby optycznej - to część mikroskopu, w której obraz jest formowany i przekazywany od obiektywu do okularu.



- Stolik mikroskopowy - próbka jest umieszczana na stoliku mikroskopowym, który może być ręcznie lub automatycznie regulowany w poziomie i pionie. Precyzyjne ustawienie próbki umożliwia obserwację wybranego fragmentu powierzchni w różnych powiększeniach. Stolik mikroskopu metalograficznego może być wyposażony w systemy do mocowania próbek, aby zapobiec ich przesuwaniu podczas obserwacji.
- Mikrometr - pozwala na dokładne ustawienie odległości pomiędzy obiektywem a próbką, co jest kluczowe do uzyskania ostrego obrazu. Mikrometry umożliwiają regulację zarówno w osi pionowej (zgrubna i precyzyjna ostrość), jak i w osi poziomej (ruch próbki w kierunku x-y).
- Kondensor - jest to element odpowiadający za skupienie światła na próbce, co pozwala na równomierne oświetlenie powierzchni materiału. Dzięki temu uzyskuje się wyraźny i kontrastowy obraz badanej struktury.

Dodatkowe elementy:

- Kamera cyfrowa - nowoczesne mikroskopy metalograficzne mogą być wyposażone w kamery, które umożliwiają rejestrację obrazów w wysokiej rozdzielczości. Taki system pozwala na dokumentowanie wyników badań oraz ich dalszą analizę w programach komputerowych.
- System komputerowy - niektóre mikroskopy metalograficzne mogą być połączone z oprogramowaniem do analizy obrazu, co ułatwia obliczenia dotyczące wielkości ziaren, porowatości czy rozkładu faz w materiale.

Zasada działania mikroskopu metalograficznego : mikroskopy metalograficzne działają na zasadzie obserwacji światła odbitego od powierzchni próbki. W przeciwieństwie do mikroskopów biologicznych, które analizują światło przechodzące przez przezroczyste obiekty, mikroskopy metalograficzne wymagają zastosowania systemu oświetlenia, który oświetla próbkę od góry. Światło pada na powierzchnię próbki, odbija się od niej i wraca do obiektywu mikroskopu, gdzie jest przetwarzane na obraz.

Obiektyw zbiera światło odbite od próbki i tworzy powiększony obraz mikrostruktury materiału. Następnie obraz ten jest przekazywany przez system soczewek i tubę optyczną do okularów, gdzie jest powiększany po raz kolejny. Użytkownik może oglądać obraz bezpośrednio przez okulary lub rejestrować go za pomocą kamery cyfrowej. Dlaczego stosuje się mikroskopy metalograficzne do badania struktury materiałów? Oto kilka kluczowych powodów, dla których są stosowane w analizie strukturalnej:



- Analiza mikrostruktury: mikroskopy metalograficzne umożliwiają szczegółową ocenę mikrostruktury materiałów. Dzięki dużym powiększeniom (do 1500x) można zidentyfikować różne fazy i składniki, które tworzą materiał, co ma kluczowe znaczenie w inżynierii materiałowej. Struktura metalu, układ ziaren, obecność wtrąceń i inne detale mogą znacząco wpływać na właściwości mechaniczne i wytrzymałościowe materiału.
- Ocena jakości materiału: metalograficzne badania są kluczowe w ocenie jakości i jednorodności materiałów, dzięki mikroskopom można wykryć defekty takie jak mikropęknięcia, wtrącenia niemetaliczne czy porowatość, które mogą obniżać wytrzymałość materiału i jego odporność na korozję. Mikroskopy te są również używane do badania struktury materiałów po różnych procesach obróbki, takich jak hartowanie czy galwanizacja.
- Badanie procesów przemysłowych - mikroskopy metalograficzne są niezbędne w przemyśle metalurgicznym do analizy wpływu różnych procesów technologicznych na strukturę materiałów. Obróbka cieplna, spawanie, hartowanie czy formowanie mają bezpośredni wpływ na mikrostrukturę metalu, a mikroskopowe badania pozwalają na ocenę, czy procesy te zostały przeprowadzone prawidłowo.
- Zastosowanie w analizie awarii – w przypadku awarii konstrukcji metalowych mikroskopy metalograficzne są używane do analizy przyczyn uszkodzeń. Badanie mikrostruktury uszkodzonego materiału pozwala na zidentyfikowanie, czy awaria była spowodowana przez zmęczenie materiału, korozję, mikropęknięcia, czy inne czynniki.
- Przybliżona ocena składu chemicznego - na podstawie obserwacji struktury metali, zwłaszcza stali, możliwa jest przybliżona ocena zawartości pierwiastków, takich jak węgiel. W zależności od obecności i rozmieszczenia poszczególnych faz, mikroskopowa analiza pozwala na oszacowanie zawartości węgla w stopie, co ma znaczenie w kontrolowaniu właściwości mechanicznych stali.

4.10. BADANIA METALOGRAFICZNE – PRZYKŁADY ZADAŃ

Badania mikrostruktury materiałów, w tym metali i stopów, pozwalają na dokładną ocenę ich budowy wewnętrznej. Proces analizy obejmuje szereg kroków, od przygotowania mikroskopu metalograficznego po rejestrowanie wyników obserwacji. Poniższa instrukcja zawiera zarys procesu badania mikrostrukturalnego.



Przygotowanie mikroskopu metalograficznego:

- Uruchomienie mikroskopu: na początku należy włączyć mikroskop i upewnić się, że źródło światła jest poprawnie ustawione. W mikroskopach metalograficznych najczęściej stosuje się światło odbite od powierzchni próbki, dlatego oświetlenie musi być dostosowane do jej specyfiki.
- Wybór obiektywu: mikroskopy metalograficzne oferują szeroki zakres powiększeń, od niskich (50x, 100x) po wyższe (200x, 500x, a nawet 1000x). Zaleca się rozpoczęcie od mniejszych powiększeń w celu zlokalizowania interesującego obszaru, a następnie przejście na wyższe, aby zobaczyć więcej szczegółów.

Regulacja ostrości i uzyskiwanie obrazu:

- Umieszczenie próbki: próbkę, czyli zgląd metalograficzny, należy precyzyjnie ułożyć na stoliku mikroskopu, mechanizmy regulacyjne stolika umożliwiają precyzyjne ustawienie pozycji próbki.
- Ostrość obrazu: pierwszym krokiem jest ustawienie zgrubnej ostrości za pomocą większego pokrętki, a następnie bardziej precyzyjne dostrojenie obrazu przy użyciu pokrętki mikro. W przypadku większych powiększeń wymagana jest szczególna ostrożność przy regulacji.
- Dostosowanie oświetlenia: kluczowe jest także ustawienie odpowiedniej intensywności światła, zbyt silne światło może prowadzić do przeświecienia, natomiast zbyt słabe uniemożliwi dostrzeżenie szczegółów struktury.

Wybór powiększenia:

- Powiększenie początkowe (50x-100x): na początku badania zaleca się niższe powiększenie, aby uzyskać ogólny obraz struktury materiału.
- Większe powiększenia (200x, 500x): po wstępnej analizie można przejść na wyższe powiększenia, aby zobaczyć bardziej szczegółowe aspekty, takie jak ziarna, wtrącenia czy mikropęknięcia, powiększenia powyżej 500x umożliwiają analizę najdrobniejszych elementów struktury.

Analiza mikrostruktury:

- Identyfikacja faz i składników: mikroskop pozwala na obserwację różnych faz i składników struktury, ziarna mogą różnić się kolorami i rozmiarami, co pozwala wnioskować o ich składzie chemicznym i orientacji krystalograficznej.



- Granice ziaren: widoczne granice ziaren to linie oddzielające poszczególne obszary, granice te mogą być bardziej wyraźne dzięki trawieniu, które uwypukla różnice między nimi.
- Obserwacja wad: warto zwrócić uwagę na defekty takie jak pęknięcia, wtrącenia czy porowatość, ciemniejsze obszary mogą wskazywać na obecność zanieczyszczeń lub mikropęknięć.
- Opis struktury: na podstawie uzyskanych obserwacji należy opisać strukturę materiału, można określić, czy ziarna są równomiernie rozmieszczone, czy widoczne są wady oraz czy mikrostruktura jest jednofazowa lub wielofazowa.

Przykłady opisu mikrostruktury:

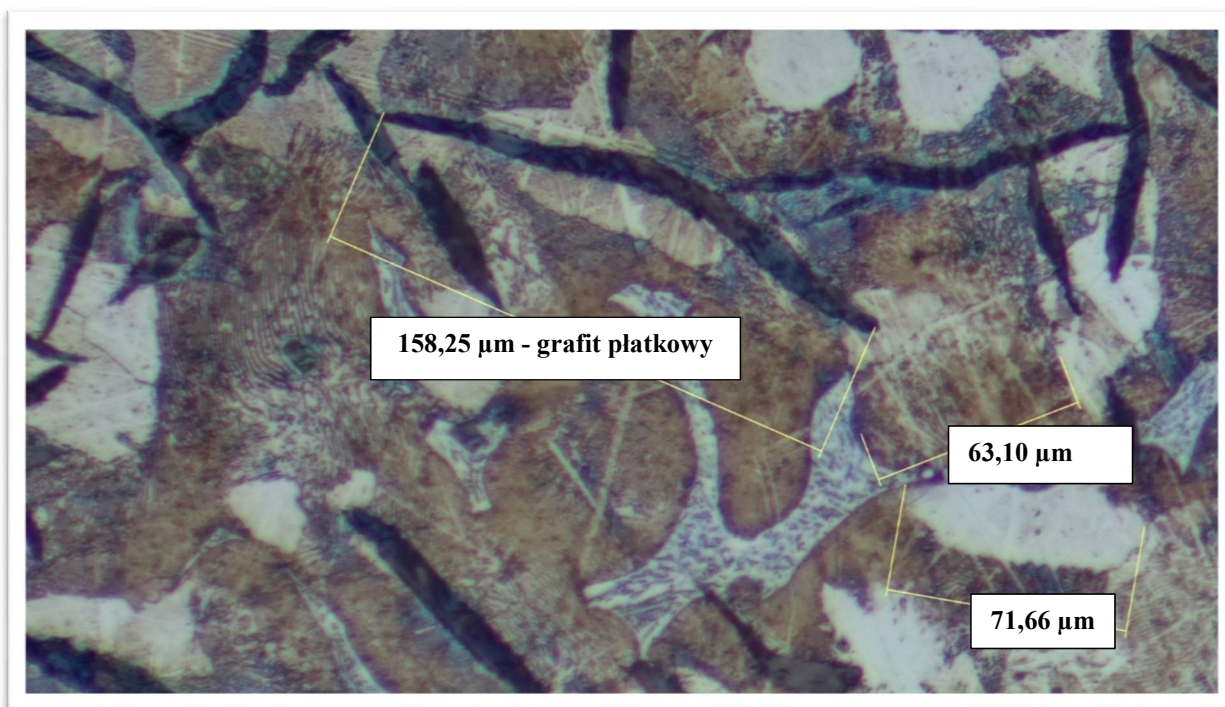
- Struktura jednofazowa: obserwowane są jednolite ziarna ferrytu, które różnią się orientacją krystalograficzną, ale nie zawierają wyraźnych defektów.
- Struktura wielofazowa: w badanym materiale widać zróżnicowane fazy, w tym ferryt i perlit, a także obecność wtrąceń niemetalicznych.
- Wady: mikrostruktura ujawnia liczne wady, takie jak pęknięcia i wtrącenia, które mogą świadczyć o nieprawidłowościach produkcyjnych.

Dokumentacja wyników:

- Notatki: wszelkie obserwacje dotyczące struktury, rozmiaru ziaren, obecności defektów należy dokładnie zapisać, warto także zwrócić uwagę na wpływ obróbki termicznej na strukturę materiału.
- Fotografie: zaleca się wykonanie zdjęć badanej struktury, aby później móc przeprowadzić dodatkową analizę lub porównać wyniki z innymi próbkami.

Podsumowanie badania:

Na podstawie zebranych informacji dokonuje się ostatecznej oceny materiału, uwzględniając jakość próbki, obecność defektów oraz zgodność z normami, dokładna interpretacja wyników badań metalograficznych dostarcza cennych informacji o właściwościach materiału, co jest kluczowe w dalszej analizie jego zastosowań przemysłowych.



Rys. 6. Żeliwo szare – obraz mikroskopowy
źródło: opracowanie własne, fotografia z kolekcji autorki

Mikrostruktura żeliwa szarego charakteryzuje się obecnością grafitu w postaci płatków oraz osnowy metalicznej, która może składać się z ferrytu, perlitu lub ich mieszaniny. Wygląd mikrostruktury żeliwa szarego jest ściśle związany z jego składem chemicznym, warunkami odlewania oraz obróbki cieplnej. Główne składniki mikrostruktury żeliwa szarego to:

- **Płatki grafitu:** grafit w żeliwie szarym występuje w formie płatków, co jest jego najbardziej charakterystyczną cechą. Płatki te są rozproszone w osnowie metalicznej i wpływają na właściwości mechaniczne żeliwa, takie jak kruchość oraz niska wytrzymałość na rozciąganie. Płatki grafitu pełnią funkcję naturalnych dyslokacji, które mogą osłabiać materiał, gdyż są to miejsca koncentracji naprężeń, jednak ich obecność nadaje żeliwu dobre właściwości tłumiące drgania oraz dobrą skrawalność.
- **Kształt i wielkość płatków:** zależą od warunków chłodzenia i składu chemicznego. Wolne chłodzenie sprzyja powstawaniu większych płatków grafitu, podczas gdy szybsze chłodzenie prowadzi do ich drobniejszej struktury.
- **Rozmieszczenie płatków:** wpływa na właściwości mechaniczne żeliwa, nierównomierne rozmieszczenie płatków może prowadzić do defektów strukturalnych, takich jak lokalne osłabienie materiału.



Osnowa metaliczna: osnowa w żeliwie szarym może składać się z ferrytu, perlitu lub ich kombinacji, w zależności od zawartości węgla, krzemu i innych dodatków stopowych oraz od szybkości chłodzenia po odlaniu.

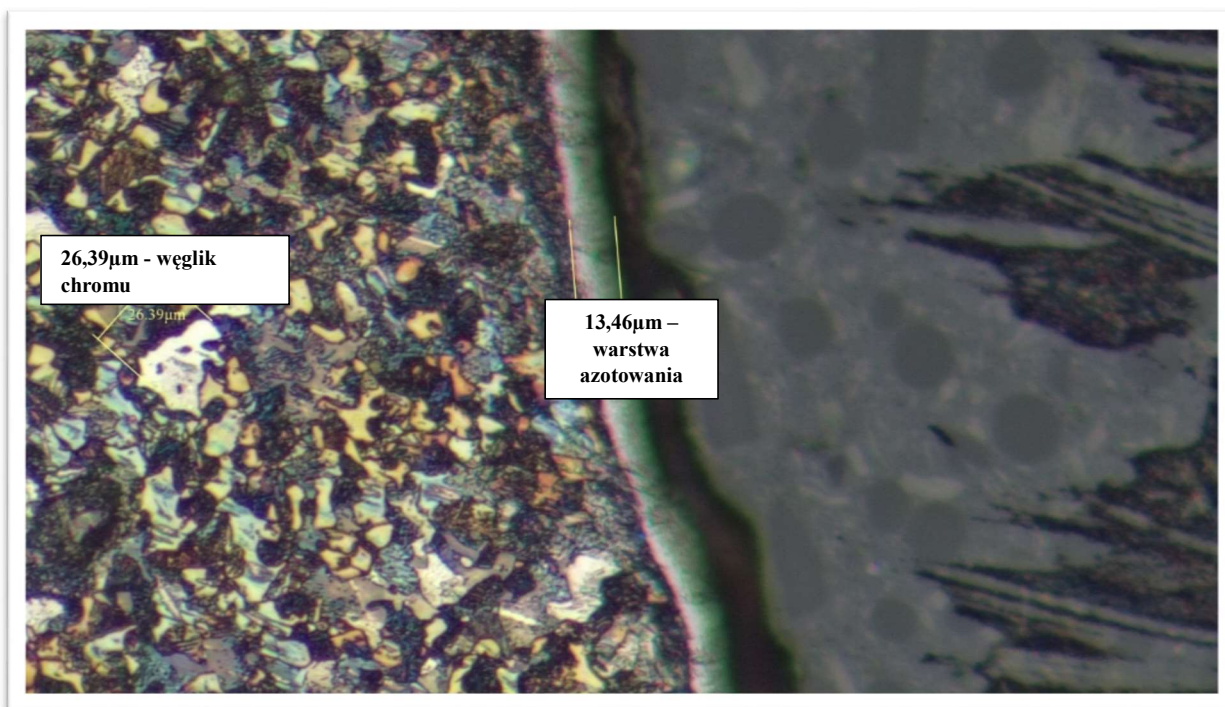
- **Ferryt:** to miękka faza żelazo-węglowa, o strukturze regularnej kubicznej przestrzennie centrowanej (RPC), jest plastyczny i ma niską wytrzymałość mechaniczną, obecność dużej ilości ferrytu w osnowie sprawia, że żeliwo staje się bardziej plastyczne, lecz mniej wytrzymałe.
- **Perlit:** jest to mieszanina ferrytu i cementytu, która tworzy lamelarną strukturę o wysokiej twardości i wytrzymałości, perlit zwiększa twardość żeliwa, poprawiając jego odporność na zużycie, ale jednocześnie zmniejsza plastyczność materiału.

Fazy węglkowe (cementyt): w przypadku żeliwa szarego z wyższą zawartością węgla lub szybko chłodzonego mogą pojawiać się węgliki (cementyt - Fe_3C). Cementyt jest twardą i kruchą fazą, która zwiększa twardość materiału, ale zmniejsza jego skrawalność i podatność na obróbkę.

Grafityzacja: proces grafityzacji, który zachodzi podczas krystalizacji, polega na wydzielaniu się węgla w postaci grafitu, zamiast tworzenia twardego cementytu (Fe_3C), żeliwo szare różni się od żeliwa białego właśnie procesem grafityzacji - w żeliwie białym grafit nie występuje, a węgiel pozostaje w postaci cementytu, co nadaje mu znacznie większą twardość i kruchość.

Wpływ składu chemicznego:

- **Krzem (Si):** dodatek krzemu sprzyja grafityzacji, co prowadzi do tworzenia płatków grafitu kosztem cementytu, wyższa zawartość krzemu zwiększa ilość ferrytu w osnowie i poprawia plastyczność materiału.
- **Węgiel (C):** zawartość węgla wpływa na ilość grafitu oraz perlitu w osnowie, wyższa zawartość węgla sprzyja tworzeniu większej ilości grafitu.
- **Dodatki stopowe:** dodatki takie jak mangan (Mn), fosfor (P) czy chrom (Cr) mogą wpływać na twardość i wytrzymałość żeliwa, zmieniając ilość i rozkład faz w mikrostrukturze.



Rys. 7. Stal 100Cr6 – po azotowaniu gazowym -zdjęcie mikroskopowe
źródło: opracowanie własne, fotografia z kolekcji autorki

Po azotowaniu struktura powierzchniowa stali zmienia się znacząco w porównaniu z jej stanem przed procesem. Możemy wyróżnić następujące strefy:

Warstwa związków azotowych (tzw. biała warstwa): jest to cienka, ale twarda warstwa powstała bezpośrednio na powierzchni stali. Składa się głównie z:

- **Azotków żelaza (Fe_4N - faza ϵ i $\text{Fe}_2\text{-3N}$ - faza γ'):** azotki te tworzą twardą, odporną na zużycie i korozję warstwę. Biała warstwa ma dużą twardość (ok. 1000-1200 HV), co zapewnia doskonałą odporność na ścieranie.
- **Azotków chromu:** ponieważ stal 100Cr6 zawiera chrom, azot reaguje także z chromem, tworząc azotki chromu (CrN), te związki dodatkowo zwiększają twardość i odporność na korozję.

Strefa dyfuzyjna: znajduje się pod warstwą związków azotowych i jest to warstwa, w której azot wnika do struktury stali, ale nie tworzy wyraźnych związków chemicznych, w tej strefie dochodzi do:

- **Dyfuzji azotu:** azot przenika w głąb struktury stali, tworząc lokalne przesycenie atomów azotu w osnowie martenzytycznej.



- **Wzrostu twardości:** wprowadzenie azotu prowadzi do umocnienia osnowy poprzez tzw. mechanizm wzmocnienia roztworowego, azot będąc małym atomem, powoduje napięcia w sieci krystalicznej, co podnosi twardość.
- **Powstawania węgliko-azotków:** zamiast typowych węglików chromu (Cr_3C_2), mogą powstawać węgliko-azotki ($\text{Cr}_3(\text{C}_x\text{N})$), które są bardziej stabilne w wyższych temperaturach i charakteryzują się większą twardością.

Osnowa wewnętrzna: poza strefą dyfuzyjną struktura stali pozostaje w dużej mierze niezmieniona, w osnowie tej wciąż dominuje martenzyt oraz węgliki chromu, jednakże strefa ta nie jest tak twarda jak warstwy powierzchniowe.

Zmiany właściwości po azotowaniu: azotowanie gazowe wprowadza szereg zmian w mikrostrukturze, które prowadzą do poprawy właściwości mechanicznych stali 100Cr6 m.in.:

- **Zwiększona twardość powierzchni:** warstwa azotkowa i strefa dyfuzyjna znacznie podnoszą twardość powierzchni, co zwiększa odporność na ścieranie.
- **Odporność na korozję:** obecność azotków żelaza i chromu na powierzchni poprawia odporność na korozję, szczególnie w środowiskach agresywnych.
- **Zwiększona trwałość zmęczeniowa:** dzięki warstwie azotkowej i umocnieniu powierzchni materiału, stal staje się bardziej odporna na cykliczne obciążenia, co przekłada się na wyższą trwałość zmęczeniową.
- **Zachowanie właściwości rdzenia:** azotowanie zmienia głównie właściwości powierzchniowe stali, pozostawiając rdzeń stosunkowo miękki i bardziej plastyczny, co pozwala materiałowi na pochłanianie obciążeń dynamicznych.

4.11. BIBLIOGRAFIA

1. Braszczyński J.: (1989): Teoria procesów odlewniczych, PWN, Warszawa
2. Dobrzański L. A. (2006): Materiały inżynierskie i projektowanie materiałowe. Podstawy nauki o materiałach i metaloznawstwo. WNT, Warszawa
3. Dobrzański L.A. (2002), Podstawy nauki o materiałach i metaloznawstwo, WNT, Warszawa
4. Górny Z. (1992): Odlewnicze stopy metali nieżelaznych, WNT Warszawa
5. Kubiński W. (2019): Wybrane metody badania materiałów, PWN, Warszawa
6. Mazurkiewicz J. i in. (2003), Podstawy technologii przetwórstwa metali.
7. Przybyłowicz K. (2003): Metaloznawstwo, WNT, Warszawa



8. Przybyłowicz K. (2011): Metody badania tworzyw metalicznych, Wydawnictwo Politechniki Świętokrzyskiej, Kielce
9. Rączka J., Sakwa W. (1986): Żeliwo, Poradnik inżyniera. Odlewnictwo, WNT, Warszawa, s.1-201



5. GRAFIKA INŻYNIERSKA W MECHANICE I BUDOWIE MASZYN

5.1. Zagadnienia teoretyczne z przykładami

Rysunek jest jedną z form wypowiedziania się i wzajemnego porozumiewania się ludzi, a zarazem najbardziej uniwersalnym środkiem wyrażania i przekazywania myśli, gdyż nie wymaga znajomości języka.

Specjalnym rodzajem rysunku jest rysunek techniczny, służący do porozumiewania się ludzi zatrudnionych w produkcji: konstruktorów oraz bezpośrednich wykonawców. Dzięki międzynarodowemu ujednoliceniu formy oraz stosowaniu uproszczeń rysunek techniczny jest czytelny dla każdego pracownika technicznego, bez konieczności znajomości języka. Umożliwia to wymianę osiągnięć technicznych między krajami oraz korzystanie z wszelkich źródeł informacji technicznej.

Rysunek techniczny musi być łatwo i jednoznacznie rozumiany, jest on wykonywany zgodnie z ustalonymi zasadami i przepisami, wynikającymi z państwowych i międzynarodowych norm oraz zaleceń. W poszczególnych krajach normalizacją zajmują się Komitety Normalizacyjne. W Polsce prace normalizacyjne prowadzi Polski Komitet Normalizacji, który zgodnie z międzynarodowymi zaleceniami opracowuje Polskie Normy (PN).

Urządzenia elektryczne i elektroniczne są wykonywane w Polsce na podstawie dokumentacji technicznych zgodnych z polskimi normami rysunku technicznego maszynowego i rysunku technicznego elektrycznego. Rysunek techniczny maszynowy za pomocą rzutowania i wymiarowania przedstawia budowę urządzenia oraz kształt i wymiary składowych części mechanicznych, zwanych częściami maszyn (w skrócie częściami). Rysunek techniczny elektryczny za pomocą znormalizowanych symboli graficznych przedstawia urządzenie lub jego elementy składowe, linie zaś symbolizują połączenia elektryczne między nimi²¹.

Rosnące wciąż wraz z rozwojem i postępem technicznym w przemyśle maszynowym - wymagania dotyczące dokumentacji technicznej powodują konieczność stosowania różnych

²¹ Paprocki K. *Rysunek techniczny*. Wydawnictwa Szkolne i Pedagogiczne 1996. s.6



rodzajów i odmian rysunków technicznych maszynowych, o cechach odpowiadających celowi, do jakiego są przeznaczone.

Podział rysunków technicznych maszynowych można przeprowadzać według wielu mniej lub bardziej ważnych kryteriów. Mnogość tych kryteriów spowodowała, że wszystkie dotychczasowe próby opracowania pełnej klasyfikacji rysunków maszynowych miały raczej charakter teoretyczny i nie doprowadziły do poprawnego rozwiązania tego zagadnienia.

Dla praktyki przemysłowej istotne znaczenie ma nie tyle klasyfikacja rysunków, ile uporządkowanie nazw najczęściej spotykanych rodzajów rysunków, w celu uniknięcia ewentualnych nieporozumień. Dlatego w normach ustalono i zdefiniowano terminy stosowane w dokumentacji technicznej wyrobów, dotyczące rysunków technicznych we wszystkich dziedzinach zastosowania²².

5.2. Formaty rysunkowe i linie rysunkowe

Formatem zasadniczym arkusza jest format A4 o wymiarach 210x297 mm. Formaty A3, A2, A1 i A0 powstają przez zwielokrotnienie formatu A4, przy czym format A3 - 2A4, format A2 - 2A3 itd. (tab. 1). Formaty od A4 do A0 noszą nazwę formatów podstawowych²³.

Tab. 1. Wymiary i kształt arkuszy rysunkowych

Format	Wymiar arkusza (mm)
A0	841 x 1189
A1	594 x 841
A2	420 x 594
A3	297 x 420
A4	210 x 297

Linia, według normy, to obiekt geometryczny, którego długość jest większa niż połowa grubości. Obiekt graficzny, którego długość jest mniejsza lub równa połowie grubości, nazywa się **kropką**.

Zapis na rysunku powstaje w wyniku użycia linii różnego rodzaju i grubości. Dobra znajomość tego zagadnienia to klucz do poprawnego wykonania i odczytania rysunku technicznego.

Jakie sprawy i problemy reguluje norma *Linie rysunkowe*?

- Określa rodzaje, nazwy i budowę linii.

²² Dobrzański T. *Rysunek techniczny maszynowy*, Wydawnictwo PWN, 2005., s. 9

²³ Dobrzański T. *Rysunek techniczny maszynowy*, Wydawnictwo PWN, 2005., s. 11



- Ustala obowiązujące grubości linii oraz zasady ich wyboru do danego opracowania rysunkowego.
- Opisuje budowę linii nieciągłych oraz zasady ich rysowania.
- Określa jednoznacznie zastosowanie linii.

W rysunku technicznym stosuje się następujące rodzaje linii (tabl. 2). Każda z nich ma opisany przez normę odpowiadający jej numer, graficzną budowę oraz nazwę²⁴.

Tab. 2. Rodzaje linii rysunkowych oraz ich zastosowanie

Lp.	Rodzaj linii	Linia - budowa	Odmiana grubości	Podstawowe zastosowanie
1.	Ciągła		cienka	• linie wymiarowe • pomocnicze linie wymiarowe • linie odniesienia • linie kreskowania przekrojów
			gruba	• widoczne zarysy widoków i przekrojów • ślady płaszczyzn przekrojów • zarysy kładów przesuniętych • obramowanie rysunku
			bardzo gruba	• połączenia klejone i lutowane
2.	Falista		cienka	• urwania i przerwania rzutów • linia oddzielająca widok od przekroju
3.	Zygzakowa		cienka	• urwania i przerwania rzutów
4.	Kreskowa		cienka	• niewidoczne zarysy przedmiotu
5.	Punktowa		cienka	• się symetrii • koła i linie podziałowe
			gruba	• powierzchnie podlegające obróbce cieplnej • powleczenia
6.	Dwupunktowa		cienka	• linie gięcia na rozwinięciach • skrajne położenia ruchomych części
7.	Wielopunktowa		cienka	• ma zastosowanie na rysunku budowlanym i w kartografii

5.3. Rzutowanie prostokątne, przekroje i wymiarowanie

Rzuty aksonometryczne odzwierciedlają przedmiot w sposób poglądowy, wyraźny i czytelny, również dla człowieka nie znającego zasad rysunku technicznego. Sporządzenie rysunku w rzutach aksonometrycznych, szczególnie rysunku przedmiotu o złożonych kształtach, jest bardzo pracochłonne, wymaga czasu i sporych umiejętności. Z tych m.in. powodów w technice mają zastosowanie rysunku wykonane według innych reguł. Metodą rzutowania najczęściej stosowaną w rysunku technicznym jest **rzutowanie prostokątne**²⁵.

²⁴ Lewandowski T., *Rysunek techniczny dla mechaników*, WSiP, 2018, s. 21

²⁵ Lewandowski T., *Rysunek techniczny dla mechaników*, WSiP, 2018, s. 62-63



Rzutowanie prostokątne metoda europejską E polega na wyznaczaniu rzutów prostokątnych przedmiotu na wzajemnie prostopadłych rzutniach, przy założeniu, że przedmiot rzutowany znajduje się między obserwatorem i rzutnią.

Jeżeli umieścimy przedmiot wewnątrz wyobraźalnego prostopadłościanu, którego wszystkie ściany są rzutniami, i wyznaczymy na tych rzutniach rzuty prostokątne przedmiotu wg metody E, to po rozwinięciu ścian prostopadłościanu w sposób pokazany na otrzymamy układ rzutów tego przedmiotu pokazany na rys. 1.

Poszczególne rzuty mają następujące nazwy (rys. 1):

rzut w kierunku A - rzut z przodu (rzut główny),

rzut w kierunku B - rzut z góry,

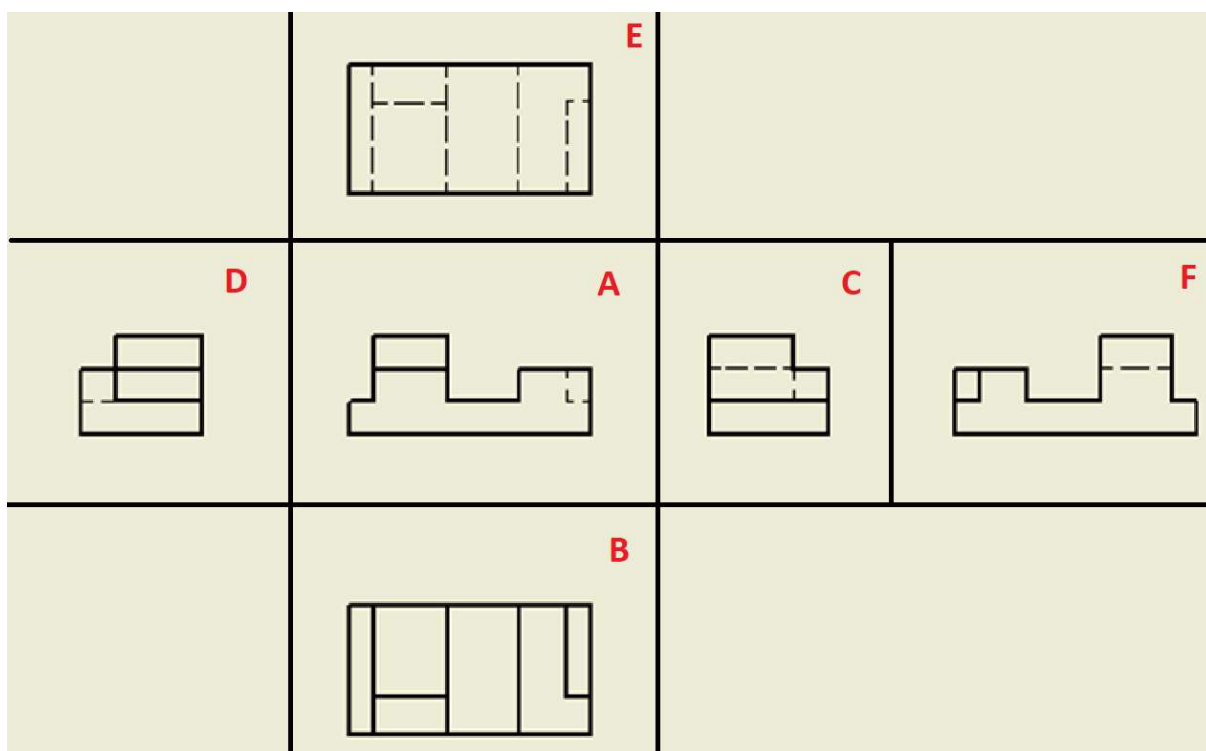
rzut w kierunku C - rzut od lewej strony,

rzut w kierunku D - rzut od prawej strony,

rzut w kierunku E - rzut z dołu,

rzut w kierunku F - rzut z tyłu.

Rzutowanie metodą europejski E obowiązuje w Polsce i w wielu innych krajach²⁶.



Rys. 1. Rzutowanie metodą europejski E

²⁶ Dobrzański T. Rysunek techniczny maszynowy, Wydawnictwo PWN, 2005, s. 32



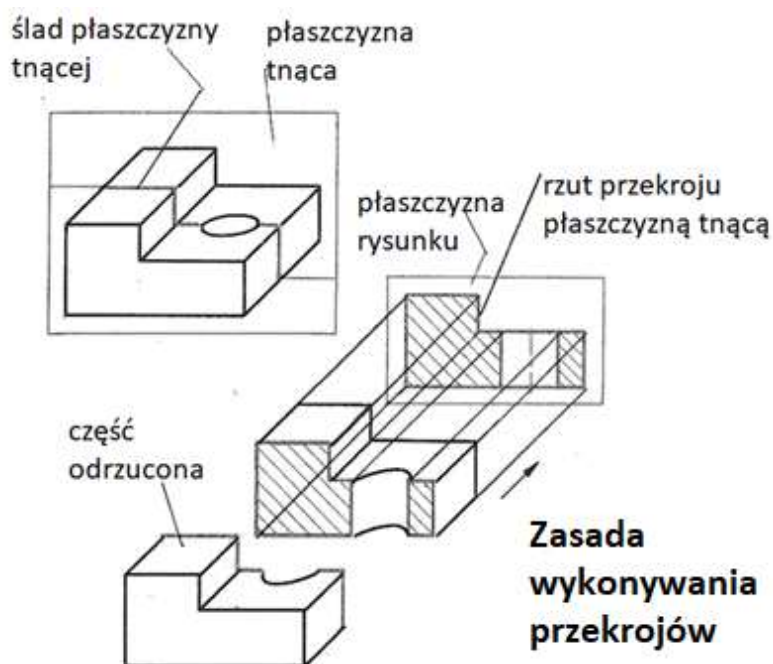
Rzutowanie prostokątne metodą amerykańską A – rzutowanie tą metodą różni się od metody E tym, że rzutnia znajduje się między obserwatorem a przedmiotem rysowanym, co powoduje, że w układzie rzutów wg metody A (rys. 2) niektóre rzuty są jak gdyby poprzestawiane w porównaniu z układem wg metody E (rzuty B z E i C z D). Rzutowanie metodą A jest stosowane głównie w krajach anglosaskich.

Przy rysowaniu przedmiotów w rzutach prostokątnych należy stosować następujące zasady:

- Liczba rzutów powinna być ograniczona do minimum niezbędnego do jednoznacznego przedstawienia kształtów przedmiotu i wymiarowania go. Wszystkie sześć rzutów, jak na rys. 1, rysuje się tylko wtedy, gdy przedmiot ma skomplikowaną budowę. W ogromnej większości przypadków wystarczają trzy rzuty (najczęściej A, B i C), dwa lub jeden (rzut główny, który nie może być pominięty na żadnym rysunku, niezależnie od liczby rzutów). Przy rzutowaniu metodą E przekroje umieszcza się albo na miejscach odpowiednich widoków, gdy te ostatnie nie są potrzebne, albo na dowolnych wolnych miejscach na arkuszu. Natomiast jeżeli jednego z widoków nie można umieścić na arkuszu zgodnie z metodą rzutowania E, to można go przesunąć równolegle na dowolne miejsce na arkuszu. Gdy któryś z rzutów (oprócz głównego) musi być z jakiegoś powodu obrócony o pewien kąt w stosunku do swego właściwego położenia, to nad takim rzutem należy umieścić znak $\odot$.
- Przedmiot rysowany powinien być tak ustawiony wewnątrz wyobraźnego prostopadłościanu, aby większość jego charakterystycznych płaszczyzn i osi była równoległa lub prostopadła do rzutni, gdyż ułatwia to rysowanie i wymiarowanie.
- Rzut główny, zarówno rysunku złożeniowego, jak i rysunku pojedynczej części maszynowej, powinien - jeśli to jest możliwe - przedstawiać przedmiot w położeniu, jakie ma on zajmować w rzeczywistości (tzw. położenie użytkowe), widziany od strony uwidaczniającej - najwięcej jego cech charakterystycznych. Od zasady tej dopuszcza się następujące odstępstwa:
 - ✓ długie przedmioty, których położenie użytkowe jest pionowe, można rysować w położeniu poziomym, przy czym dolną część przedmiotu umieszcza się z prawej strony rzutu,
 - ✓ przedmioty, których położenie użytkowe nie jest ani poziome, ani pionowe, oraz przedmioty, które przyjmują różne położenia podczas Użytkowania (np. korbowody), rysuje się w położeniu poziomym lub pionowym.

- Na rysunku wykonawczym przedmiot przedstawia się najczęściej w położeniu, jakie zajmuje podczas obróbki nadającej mu najwięcej kształtów charakterystycznych, a więc np. wrzeciono wiertarki - w położeniu poziomym (położenie podczas toczenia wrzeciona), natomiast na rysunkach operacyjnych i zabiegowych (w dokumentacji technologicznej) - w położeniu, jakie przedmiot ma zajmować podczas konkretnej operacji czy zabiegu.
- Widoki rozwinięte przedmiotów rysuje się w celu pokazania budowy przedmiotów walcowych i stożkowych oraz przedmiotów wyginanych z blachy²⁷.

Przekroje – rzuty przedmiotu w postaci widoków często nie dają pełnego wyobrażenia o jego kształcie, zwłaszcza gdy ma on złożoną budowę, wewnętrzną. Niewidoczne krawędzie przedmiotu zaznaczano za pomocą linii kreskowej cienkiej, zmniejsza to jednak czytelność rysunku. W celu dokładnego pokazania wewnętrznego kształtu przedmiotu stosuje się przekroje. Zasady wykonywania przekrojów zgodnie z normą przedstawiono na rys. 2²⁸.



Rys. 2. Zasady wykonywania przekroju

²⁷ Dobrzański T. Rysunek techniczny maszynowy, Wydawnictwo PWN, 2005, s. 33

²⁸ Paprocki K. *Rysunek techniczny*. Wydawnictwa Szkolne i Pedagogiczne 1996. s. 40



Przekrój powstaje przez przecięcie przedmiotu wyobraźną płaszczyzną tnącą i odrzucenie części przedmiotu znajdującej się przed płaszczyzną tnącą. Następnie wykonuje się rzut przeciętego przedmiotu na płaszczyznę rysunku.

Oznaczanie i kreskowanie przekrojów:

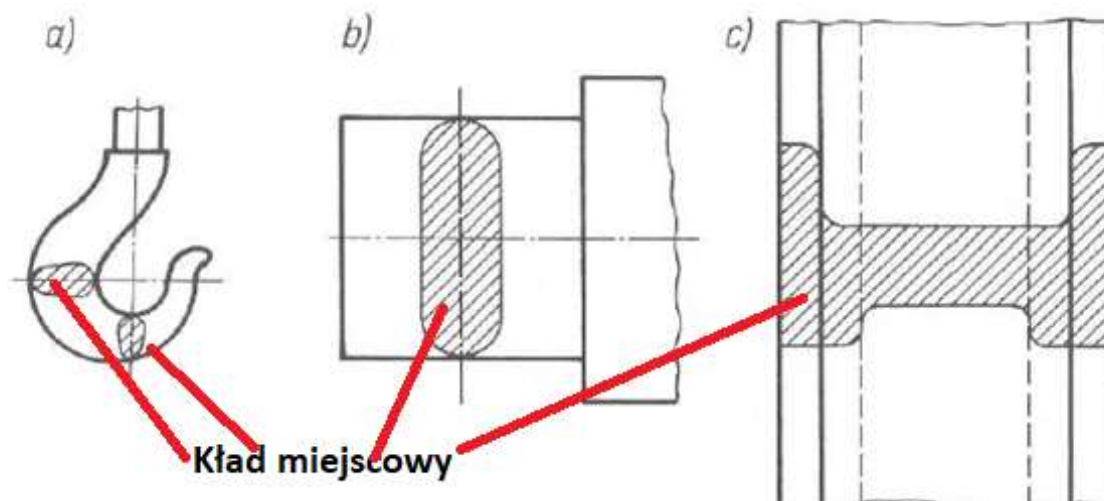
- a) usytuowanie przekroju względem rzutu głównego zależy od położenia płaszczyzny przekroju i jej oznaczenia przez odcinki linii grubej, strzałki i litery;
- b) jeśli przekrój jest jednoznaczny, to można go nie oznaczać;
- c) jeśli należy wykonać kilka przekrojów przedmiotu, to oznacza się je kolejnymi literami wielkimi (z wyjątkiem liter: I, O, R, Q, X);
- d) przekrój jest rzutem przeciętego przedmiotu, muszą więc być na nim zaznaczone wszystkie szczegóły przedmiotu, leżące poza płaszczyzną tnącą;
- e) płaszczyznę przekroju kreskuje się liniami ciągłymi cienkimi, nachylnymi pod kątem 45° do głównych krawędzi przedmiotu. Linie kreskowania powinny przebiegać przez cały obszar płaszczyzny przekroju przy takim samym nachyleniu i odległości t , mimo przerw w polu przekroju. W rysunkach o średniej wielkości można przyjąć $t = 2\text{mm}$. Dla kilku przekrojów tego samego przedmiotu kreskowanie powinno być zgodne co do kierunku, nachylenia i odległości;
- f) pola przekroju załamane pod kątem 45° można kreskować pod kątem 30° ;
- g) wąskie pola przekroju mogą być zaczernione;
- h) na rysunkach złożeniowych kreskowanie powierzchni przedmiotu stykających się części powinno różnić się kierunkiem lub przynajmniej odległością²⁹.

Kłady

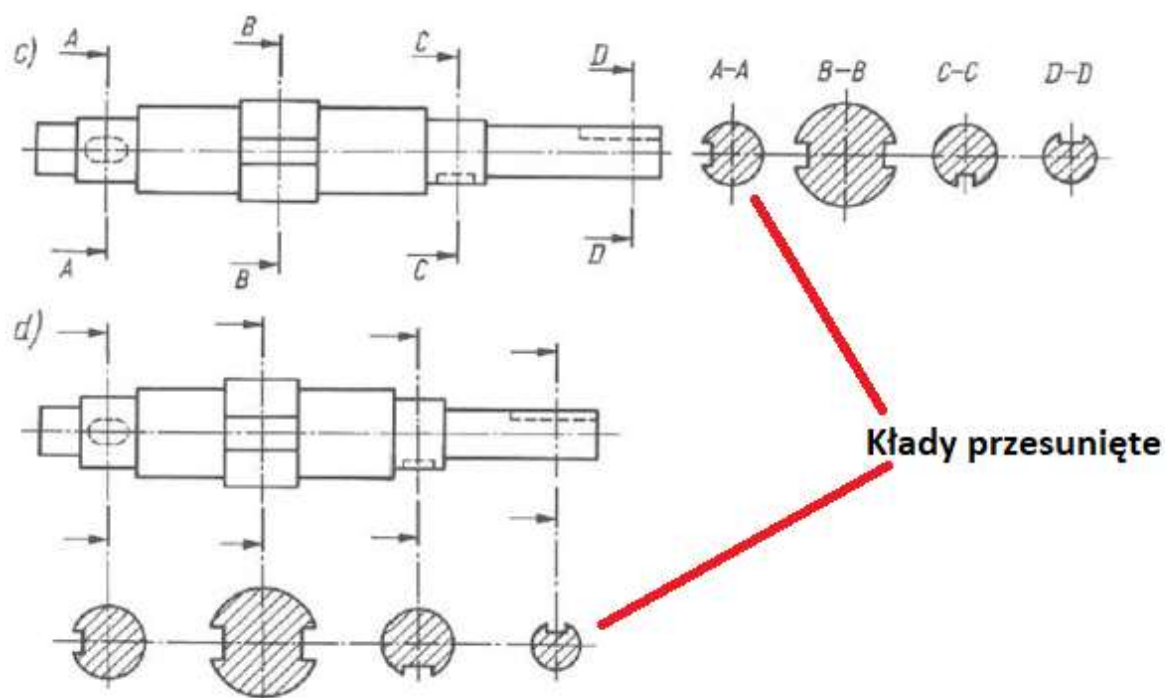
Kład jest to zarys figury płaskiej leżącej w płaszczyźnie poprzecznego przekroju przedmiotu, obrócony wraz z tą płaszczyzną o 90° i położony na widoku przedmiotu (kład miejscowy - rys. 3), lub poza jego zarysem (kład przesunięty - rys. 4). Kierunek obrotu płaszczyzny przekroju wraz z kładem powinien być zgodny z kierunkiem patrzenia na przedmiot od strony prawej lub od dołu, a więc płaszczyznę tę należy obracać w lewo lub do góry, w zależności od jej położenia. Kłady miejscowe wolno rysować tylko wtedy, gdy nie zaciemniają rysunku; rysuje się je liniami cienkimi, zaś kłady przesunięte - liniami grubymi, jak zwykle rzuty. Jeżeli płaszczyzna przekroju przechodzi przez oś otworu walcowego lub stożkowego, to kład

²⁹ Paprocki K. *Rysunek techniczny*. Wydawnictwa Szkolne i Pedagogiczne 1996. s. 41

uzupełnia się widokiem krawędzi otworów leżących za płaszczyzną przekroju (aby uniknąć „rozpadnięcia się” kładu na dwie odrębne części). We wszystkich innych przypadkach, w których kład składałby się z dwóch lub więcej części oddzielnych; należy rysować nie kład, lecz przekrój³⁰.



Rys. 3. Kład miejscowy



Rys. 4. Kład przesunięty

³⁰ Dobrzański T. Rysunek techniczny maszynowy, Wydawnictwo PWN, 2005, s. 38

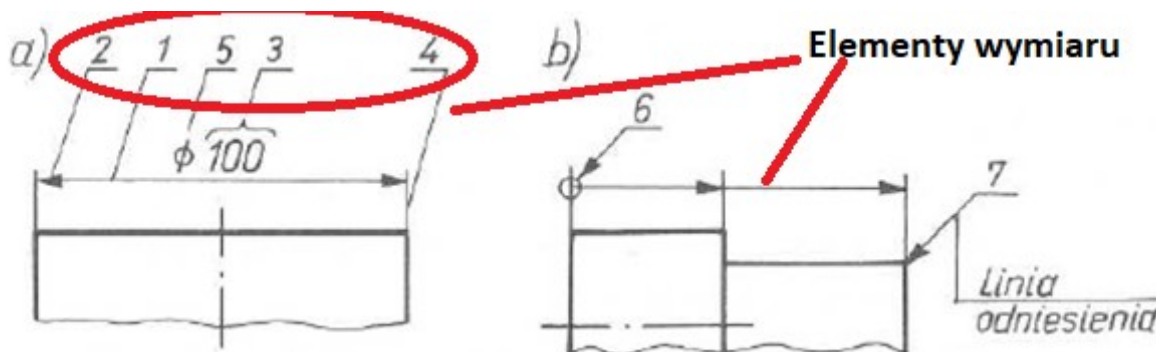
Wymiarowanie

Rzuty przedmiotu odwzorowują jego budowę i kształt, nie stanowią jednak informacji wystarczającej do jego wykonania. Wykonawca, poza kształtem i budową, musi ponadto znać wymiary poszczególnych elementów geometrycznych narysowanego przedmiotu. Wymiarowanie, czyli podawanie wymiarów na widokach, przekrojach i kładach, podobnie jak zasady rzutowania, jest objęte normalizacją. Z tego powodu wymiarowanie nie może być dowolne, przypadkowe i wykonane według indywidualnych pomysłów autora rysunku.

Wymiar rysunkowy składa się z kilku elementów graficznych (rys. 6):

- linii wymiarowej, znaku ograniczenia linii wymiarowej, liczby wymiarowej,
- pomocniczej linii wymiarowej, znaku wymiarowego (rys. 6 a),
- oznaczenia początku linii wymiarowej oraz linii odniesienia (rys. 6 b).

Wymienione elementy nie zawsze występują równocześnie, ale każdy z nich musi spełniać określone wymagania graficzne.



Rys. 6 a, b. Elementy wymiaru rysunkowego

Linie wymiarowe należy rysować:

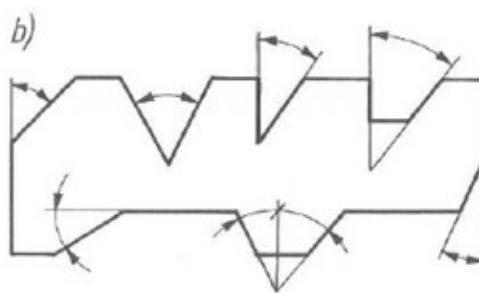
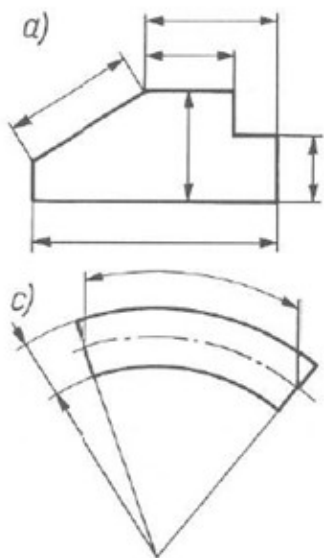
1. jako linie ciągłe cienkie zakończone znakami ograniczenia:
 - przy wymiarowaniu odcinka prostoliniowego - jako równoległe do tego odcinka (rys. 7a),
 - przy wymiarowaniu kąta - jako łuk okręgu zatoczonego z wierzchołka wymiarowanego kąta (rys. 7b),
 - przy wymiarowaniu łuku - współśrodkowo z wymiarowanym łukiem (rys. 7c),



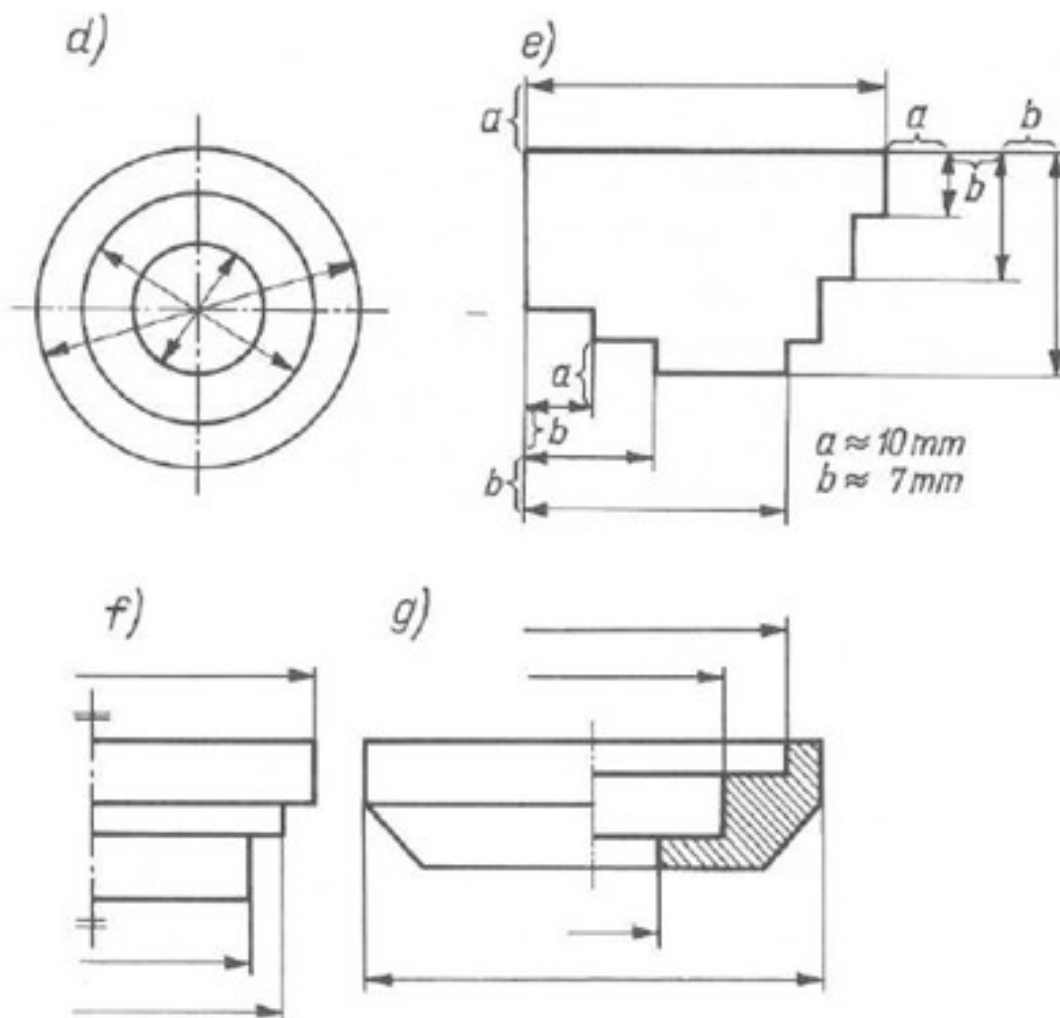
- przy wymiarowaniu odległości między łukami współśrodkowymi - promieniowo do wymiarowanych łuków (rys. 7c),
 - przy wymiarowaniu średnicy okręgu - jako odcinki łączące dwa punkty okręgu i przechodzące przez środek okręgu (rys. 2d).
2. w odległości nie mniejszej niż 10 mm od linii zarysu przedmiotu i 7 mm od równoległej osi wymiarowej (rys. 2e).
 3. w przypadkach szczególnych - jako urwane w odległości $2 \div 10$ mm poza środkiem okręgu lub osią symetrii (rys. 2f, g).

Linie wymiarowe nie mogą:

1. wzajemnie się przecinać, z wyjątkiem linii wymiarowych okręgów kół współśrodkowych,
2. spełniać żadnych innych funkcji na rysunku poza określaniem wymiaru,
3. Jeśli to możliwe, nie powinny się przecinać z pomocniczymi liniami wymiarowymi i liniami odniesienia,
4. leżeć na przedłużeniu osi symetrii pomocniczych linii wymiarowych oraz linii zarysu przedmiotu,
5. leżeć na jednej prostej (przy wymiarowaniu promieni łuków okręgów współśrodkowych)³¹.



³¹ Lewandowski T., *Rysunek techniczny dla mechaników*, WSiP, 2018, s. 122



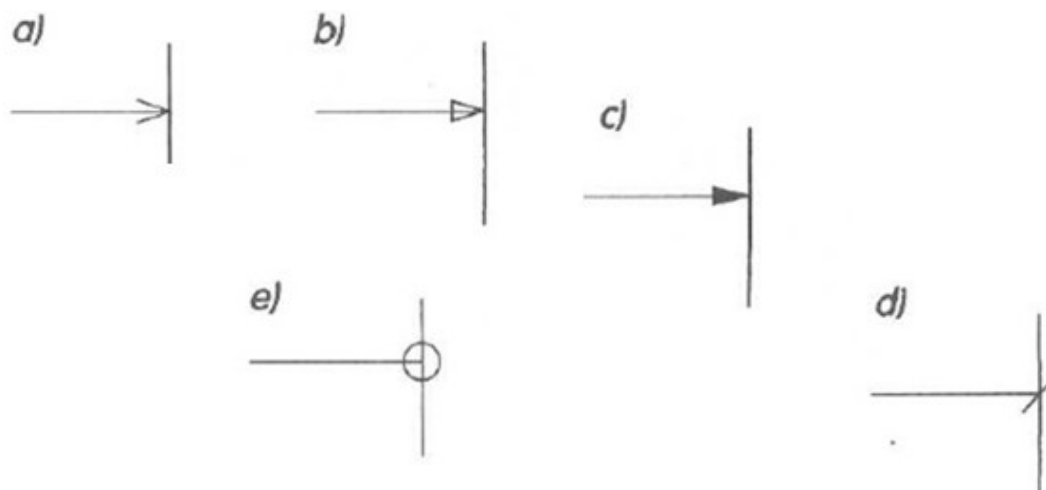
Rys. 7 a, b, c, d, e f, g. Rysowanie linii wymiarowych

Znaki ograniczenia linii wymiarowych - grot rysuje się krótkimi, cienkimi liniami tworzącymi ostrze. Grot może być otwarty, zamknięty lub zamknięty i zaczerniony, natomiast ostrze grotu może mieć dowolny kąt rozwarcia, zawarty w przedziale 15 do 90⁰ (rys. 8 a,b,c). Wielkość grotu powinna być proporcjonalna do wielkości rysunku.

W zasadzie ostrza grotów powinny dotykać od wewnątrz linii, między którymi wymiar ma być podany. W braku miejsca grotki można umieszczać na zewnątrz tych linii, na przedłużeniach linii wymiarowej.

Dopuszczalne jest zastępowanie grotów cienkimi (lub o grubości linii cyfr wymiarowych) kreskami o długości co najmniej 3,5 mm, nachylonymi pod kątem 45⁰ do linii wymiarowych lub kropkami o średnicy ok. 1 mm (rys. 8 e,d). ³².

³² Dobrzański T. Rysunek techniczny maszynowy, Wydawnictwo PWN, 2005, s. 38



Rys. 8. Znaki ograniczenia linii wymiarowych

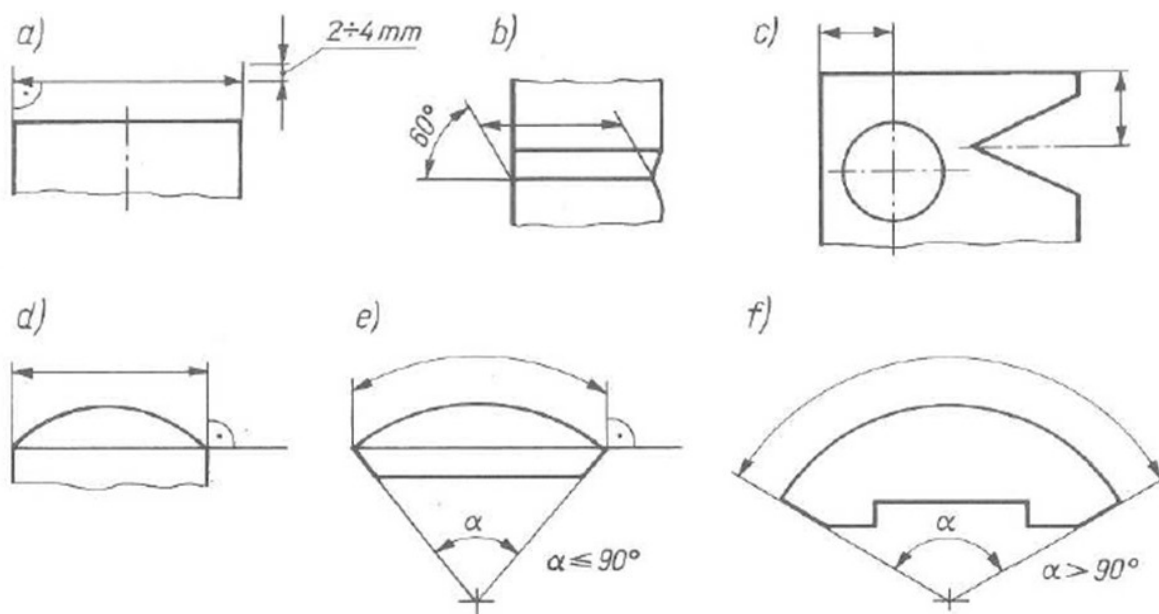
Liczby wymiarowe na rysunkach technicznych maszynowych wymiary liniowe (długościowe) podaje się w milimetrach, przy czym oznaczenie „mm” pomija się, nawet gdy liczba wymiarowa podana jest z dokładnością do trzech znaków dziesiętnych za przecinkiem. Jeżeli konieczne jest podanie wymiaru w innych jednostkach, to za liczbą wymiarową należy umieścić oznaczenie jednostki miary, np. 10 cm. Przy konstruowaniu zaleca się przyjmować - o ile to możliwe - wartości liczbowe wymiarów (przynajmniej ważniejszych) z ciągów wymiarów normalnych.

Pomocnicze linie wymiarowe (rys. 9) należy rysować:

- jako linie ciągłe cienkie, przeciągnięte 2-4mm poza odpowiadające im linie wymiarowe (rys. 9a),
- prostopadłe do odpowiadających wymiarów liniowych (rys. 9a), a w szczególnych przypadkach jako dwie ukośne linie równoległe (rys. 9b),
- jako przedłużenia osi symetrii, gdy zachodzi taka potrzeba (rys. 9c),
- prostopadłe do cięciwy łuku przy wymiarowaniu tej cięciwy (rys. 4d),
- prostopadłe do cięciwy łuku przy wymiarowaniu łuku opartego na kącie wierzchołkowym nie większym niż 90° (rys. 4e),
- promieniowo - przy wymiarowaniu łuku opartego na kącie wierzchołkowym większym od 90° (rys. 4f).

Należy unikać wzajemnego przecinania się pomocniczych linii wymiarowych oraz ich prowadzenia równoległe do linii kreskowania przekrojów³³.

³³ Dobrzański T. Rysunek techniczny maszynowy, Wydawnictwo PWN, 2005, s. 43



Rys. 9 a, b, c, d, e, f. Rysowanie pomocniczych linii wymiarowych

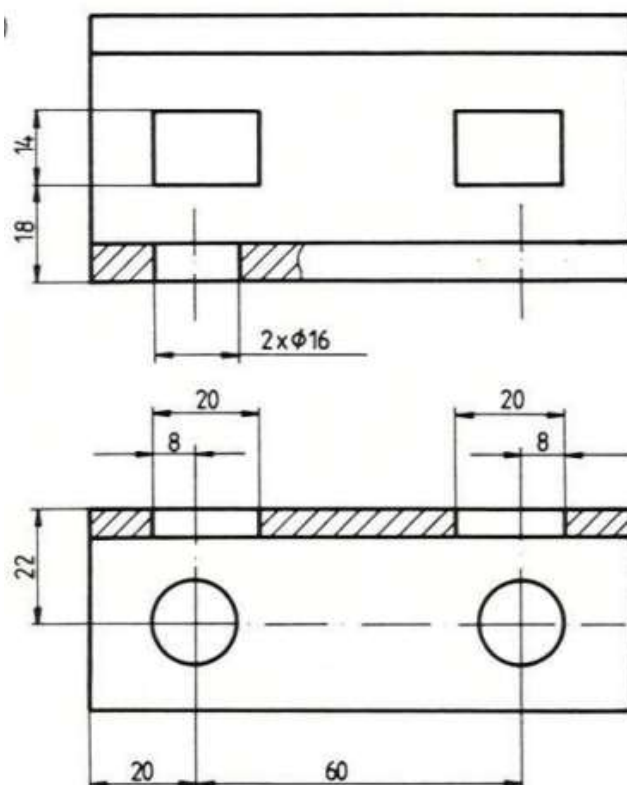
Zasady ogólne wymiarowania, nazywane niekiedy porządkowymi, są bardziej szczegółowym uściśleniem zasady jednoznaczności, niesprzeczności i zupełności w odniesieniu do układu wymiarów. Podają one ograniczenia, którym podlega układ wszystkich wymiarów zawartych na rysunku i z tego powodu muszą być bezwzględnie przestrzegane.

Jednoznaczność, niesprzeczność i zupełność zapisu konstrukcji osiąga się, gdy: zapis zawiera tylko wymiary konieczne, pominięto wymiary oczywiste dla przyjętego układu wymiarów, nie powtórzono żadnego wymiaru i nie zamknięto łańcucha wymiarowego.

Zasada wymiarów koniecznych stanowi, że na rysunku należy podać tylko te wymiary, które są niezbędne do opisu geometrycznych cech konstrukcyjnych przedmiotu w określonym stanie jego wykonania. Układ wymiarów musi być zatem podporządkowany rodzajowi rysunku i jego przeznaczeniu. Inne wymiary są potrzebne do wykonania surowego odlewu lub odkuwki, inne do wykonania części gotowej, a jeszcze inne do zmontowania całego zespołu lub wyrobu.

Wyodrębnienie poszczególnych grup wymiarów odpowiadających kolejnym etapom tworzenia przedmiotu i przedstawienie ich na oddzielnych rysunkach poprawia czytelność zapisu, ale wymaga sporządzenia bardziej rozbudowanej dokumentacji rysunkowej. Z tego też względu taki sposób postępowania jest stosowany głównie podczas przygotowania produkcji seryjnej lub masowej, gdzie kosztem zwiększonych nakładów na opracowanie dokumentacji uzyskuje się poprawę organizacji procesów wytwarzania.

Każdy kolejny rysunek przedmiotu, dla którego wykonano odrębne rysunki odlewu lub odkuwki, rysunki zabiegowe, czynnościowe, montażowe itp., zawiera tylko wymiary potrzebne do jego dalszej obróbki, a więc wymiary powierzchni obrabianych i wymiary określające położenie tych powierzchni względem powierzchni ukształtowanych w poprzednim procesie technologicznym (rys. 10). Na rysunkach przedmiotów o nieskomplikowanej konstrukcji lub wytwarzanych jednostkowo podaje się równocześnie wszystkie wymiary powierzchni obrabianych i nieobrabianych. O tym, którą cechę konstrukcyjną opisują poszczególne układy wymiarów, wnioskuje się wówczas na podstawie znajomości technik wytwarzania.



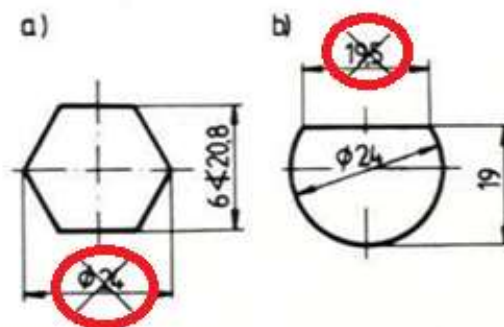
Rys. 10. Zasada pomijania wymiarów oczywistych

Zasada pomijania wymiarów oczywistych (rys. 11) ściśle wynika z zasady wymiarów koniecznych. Skoro bowiem na rysunku powinny być podane wyłącznie wymiary konieczne, to nie należy podawać wymiarów oczywistych, takich jak kątów 0° (180°) między prostymi równoległymi i kąta 90° między prostymi prostopadłymi. Jednak wymiar kąta 90° można pominąć tylko wówczas, gdy ma on charakter wymiaru swobodnego, który na rysunku występuje bez odchyłek, lub gdy dokładność kąta prostego ustalono za pomocą tolerancji prostopadłości. W każdym innym przypadku tolerowania kąta prostego musi być podany na rysunku jego wymiar nominalny (90°). Wymiarami oczywistymi mogą być również inne



wartości kątów, jak również wymiary liniowe, ale tylko takie, które nie muszą być tolerowane.

W przeciwnym bowiem przypadku muszą być one zawsze opisane liczbowo.



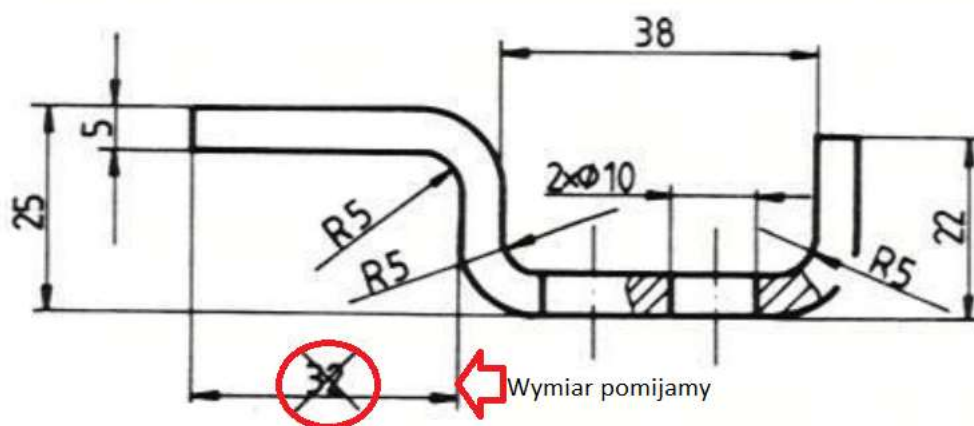
Rys. 11. Przykłady pomijania wymiarów oczywistych

Zasada niepowtarzania wymiarów stanowi, że wymiar może być podany na rysunku tylko jeden raz, niezależnie od liczby rzutów czy liczby arkuszy, na których przedmiot jest przedstawiony.

Podczas wymiarowania, szczególnie zaś gdy przedmiot jest przedstawiony na kilku arkuszach, można bardzo łatwo popełnić błąd opisując tę samą cechę konstrukcyjną dwoma różnymi wartościami liczbowymi. Doprowadzi to do wadliwego wykonania przedmiotu lub wadliwej jego oceny, jeżeli nie zostanie w porę zauważone.

Inną przyczyną, dla której trzeba bezwzględnie przestrzegać zasady niepowtarzania wymiarów jest okoliczność wprowadzania zmian. Jeżeli brak pewności, że wymiary nie są powtórzone, zachodzi konieczność odszukania wszystkich powtórzeń, których liczba najczęściej nie jest znana. Przeoczenie jednego z zapisów powoduje, że zaistnieją różne informacje o tej samej cesze konstrukcyjnej. Konsekwencją tego będzie wadliwy przedmiot. Występowanie wymiaru tylko jeden raz gwarantuje, że dokonana zmiana będzie jednoznaczna i niesprzeczna.

Zasada niezamykania łańcucha wymiarowego (rys. 12) stanowi, że wymiar należy pominąć, jeżeli można go obliczyć jako wypadkową pozostałych. W innym przypadku wystąpi pośrednie powtórzenie wymiarów. Sytuacja taka nazywa się potocznie „przewymiarowaniem”, które utrudnia, a niekiedy wręcz uniemożliwia wykonanie przedmiotu zgodnie z wymaganymi dokładnościami.



Rys. 12. Zasada niezamykania łańcucha wymiarowego

Zakaz zamykania łańcucha wymiarowego dotyczy wszystkich wymiarów, niezależnie od miejsca (rzutu) ich umieszczenia, jeżeli odnoszą się one do tej samej cechy konstrukcyjnej lub cech związanych ze sobą.

Z szeregu wymiarów tworzących zamknięte ogniwo pomija się zawsze najmniej ważny. Jednak gdy podanie wymiaru zamykającego jest uzasadnione, to należy umieścić go w nawiasach okrągłych i wówczas - jako wymiar pomocniczy - nie zamknie łańcucha³⁴.

Zasady wymiarowania wynikające z potrzeb konstrukcyjnych i technologicznych.

Wymiarować można od baz konstrukcyjnych, obróbkowych lub pomiarowych. Wymiarowanie od baz konstrukcyjnych stosuje się, gdy zależy na podaniu na rysunku tych wymiarów, które mają bezpośredni wpływ na działanie wymiarowanej części maszynowej w zespole, do którego należy, i na montaż. Są to wymiary mające wpływ na położenie części w zespole, na wymagane luzy i wciski itd.

Do zalet wymiarowania od baz konstrukcyjnych należą:

- krótkie łańcuchy wymiarowe, co ułatwia analizę wymiarową całego zespołu i wyrobu oraz zwiększa dokładność wykonania części,
- niezmienność w znacznym stopniu rysunku części, gdyż zmian technologii nie mają wpływu na wymiarowanie.

Wadą wymiarowania od baz konstrukcyjnych jest właśnie oderwanie się od technologii, co zmusza często w przygotowaniu produkcji do przeliczania wymiarów, zacieśniania tolerancji

³⁴ Boder A., Dudziak M., *Zapis konstrukcji*, Wydawnictwo Politechniki Poznańskiej, Poznań 1995, s. 272



itd. Jeżeli przedmiot nie ma powierzchni ważniejszych od innych, to zwykle wymiaruje się go od dwóch końców, co często pokrywa się zwymiarowaniem od baz obróbkowych.

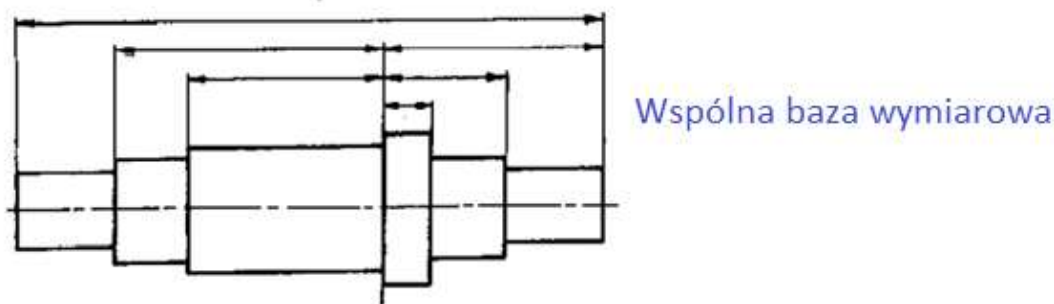
Na rysunkach części współpracujących za bazę wymiarową konstrukcyjną należy przyjmować powierzchnię styku tych części po zmontowaniu. Jest to tzw. wymiarowanie od wspólnych baz wymiarowych.

Wymiarowanie od baz obróbkowych stosuje się, gdy zależy przede wszystkim na uproszczeniu procesu technologicznego. Dzięki wymiarowaniu od baz obróbkowych osiąga się:

- a) łatwiejsze uzyskanie dokładnych wymiarów części obrabianej,
- b) możliwość zaplanowania przez technologa przebiegu obróbki bez potrzeby przeliczania wymiarów i zacieśniania tolerancji, dzięki czemu tolerancje wymiarów podane na rysunkach mogą być w pełni wykorzystane, co zmniejsza koszty produkcji,
- c) możliwość takiego ustawiania części do obróbki, że budowa uchwytów obróbkowych staje się prostsza i tańsza.

Bazy konstrukcyjne i obróbkowe często nie pokrywają się, w każdym przypadku ich niezgodności należałoby rozstrzygnąć czy wymiarować od jednych, czy od drugich. W praktyce jednak rzadko stosuje się wymiarowanie tylko od baz konstrukcyjnych lub tylko od obróbkowych. Zwykle podaje się wymiary najważniejsze (wchodzące w łańcuchy wymiarowe mechanizmów) od baz konstrukcyjnych i toleruje dokładnie, pozostałe zaś wymiary podaje się od baz obróbkowych, z większymi tolerancjami lub bez tolerancji.

Oczywiście, sprawa się upraszcza, gdy bazy obróbkowe pokrywają się z konstrukcyjnymi. Przy wymiarowaniu od baz obróbkowych wymiary odnoszące się do jednej operacji należy, o ile to możliwe, podawać na jednym rzucie i to grupując je po jednej stronie rzutu³⁵.



Rys. 13. Wymiarowanie części współpracujących od wspólnej bazy wymiarowej

³⁵ Dobrzański T. Rysunek techniczny maszynowy, Wydawnictwo PWN, 2005, s. 62



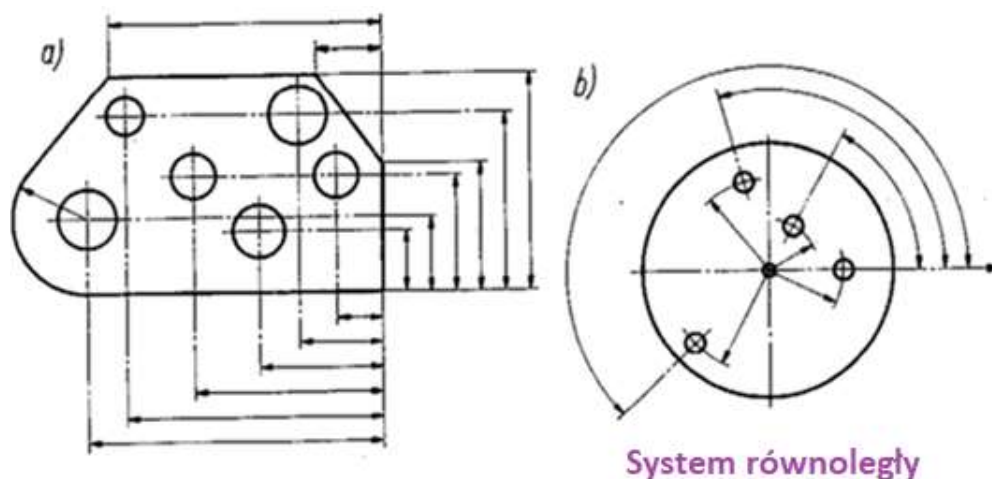
Wymiarowanie równoległe, szeregowe i mieszane

Wymiary biegnące w jednym kierunku można podawać na rysunkach w trzech układach: równoległym, szeregowym lub mieszanym.

Wymiarowanie w układzie równoległym (rys. 14) polega na podawaniu wszystkich wymiarów równoległych od jednej bazy (powierzchni lub linii). W płaskim prostokątnym układzie współrzędnych baz tych, wzajemnie prostopadłych, jest dwie; tak samo dwie bazy występują w płaskim biegunowym układzie współrzędnych. Natomiast w układach przestrzennych występują trzy wzajemnie prostopadłe bazy.

Przy wymiarowaniu w układzie równoległym dokładność każdego wymiaru uzyskana w wyniku obróbki zależy tylko od dokładności samej obróbki, a nie zależy w ogóle od dokładności innych wymiarów przedmiotu. Dlatego też ten sposób wymiarowania stosuje się, gdy zależy na uzyskaniu dokładnego położenia pewnej ilości powierzchni przedmiotu od wybranej uprzednio bazy.

Poza tym wymiarowanie równoległe ma tę zaletę, że dowolny wymiar przedmiotu nie podany na rysunku można obliczyć jako wypadkowy tylko dwóch podanych wymiarów. Wymiarowania równoległego nie należy jednak stosować, gdy dokładne mają być właśnie wymiary wynikające jako wypadkowe przy takim wymiarowaniu, gdyż zmusza to do znacznego zmniejszania tolerancji wymiarów, z których wynika wymiar wypadkowy.

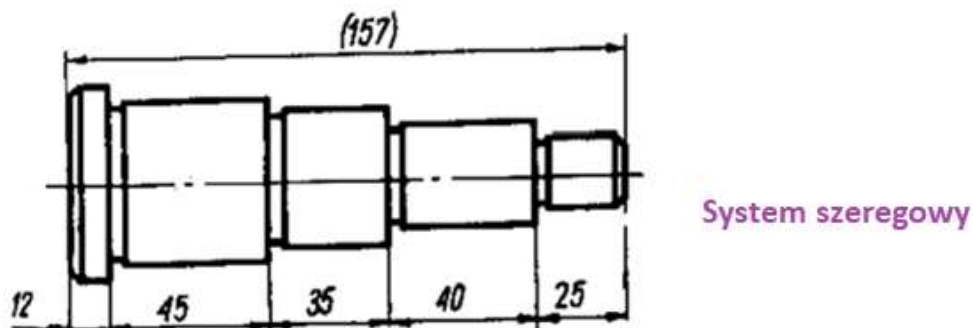


Rys. 14. Wymiarowanie w systemie równoległym

Wymiarowanie w układzie szeregowym polega na wpisywaniu wymiarów równoległych jeden za drugim (rys. 15). Ten sposób wymiarowania stosuje się gdy zależy na dokładności wzajemnego położenia sąsiednich elementów przedmiotu, a nie na dokładnym ich położeniu

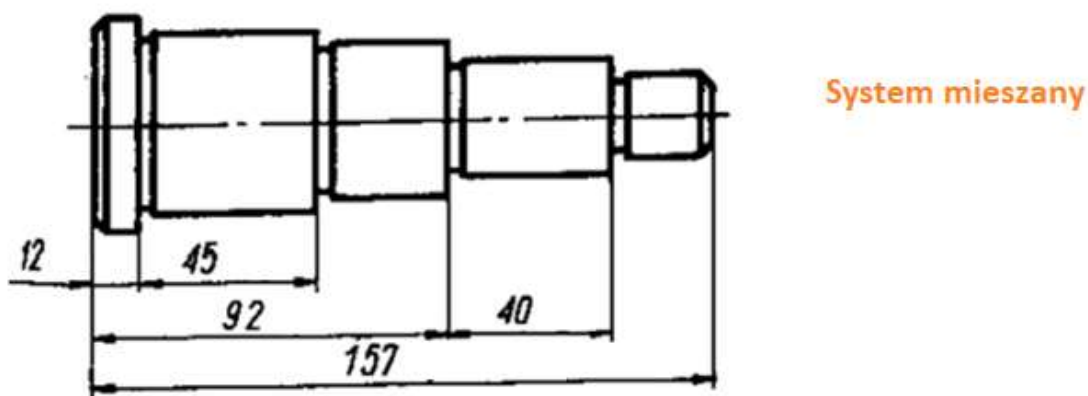


względem jednej bazy. W ten sposób wymiaruje się często przedmioty, które mają być obrabiane zespołem narzędzi pracujących jednocześnie (np. równoczesne wiercenie wielu otworów przy użyciu głowicy wielowrzecionowej).



Rys. 15. Wymiarowanie w systemie szeregowym

Wymiarowanie w układzie mieszanym (rys. 16) jest połączeniem obu sposobów omówionych wyżej i łączy zalety obu tych sposobów. Przy wymiarowaniu mieszanym położenie tych powierzchni, które powinny się znajdować w ściśle określonych odległościach od pewnej bazy, wymiaruje się od tej bazy, zaś położenia pozostałych powierzchni względem poprzednich lub między sobą określa się krótkimi łańcuchami wymiarowymi, czyli wymiaruje szeregowo. Dzięki takiemu wymiarowaniu, wszystkie ważne wymiary przedmiotu mogą być na rysunku bezpośrednio podane, a zatem i bezpośrednio sprawdzone.



Rys. 16. Wymiarowanie w systemie mieszanym

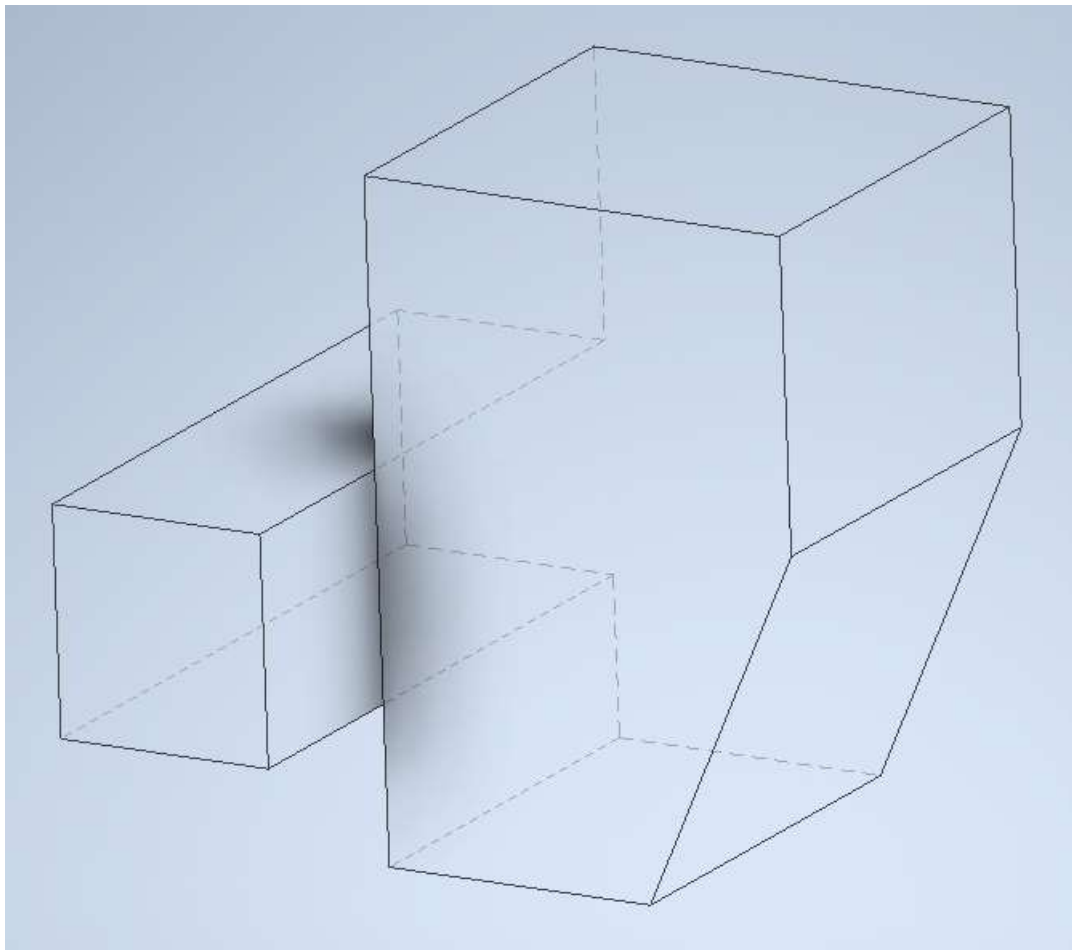
Przed wyborem sposobu wymiarowania należy każdorazowo zbadać, które wymiary części maszynowej są ważne ze względu na jej przyszłe działanie, oraz zorientować się co do



przypuszczalnego przebiegu procesu technologicznego tej części. Najczęściej stosuje się wymiarowanie mieszane³⁶.

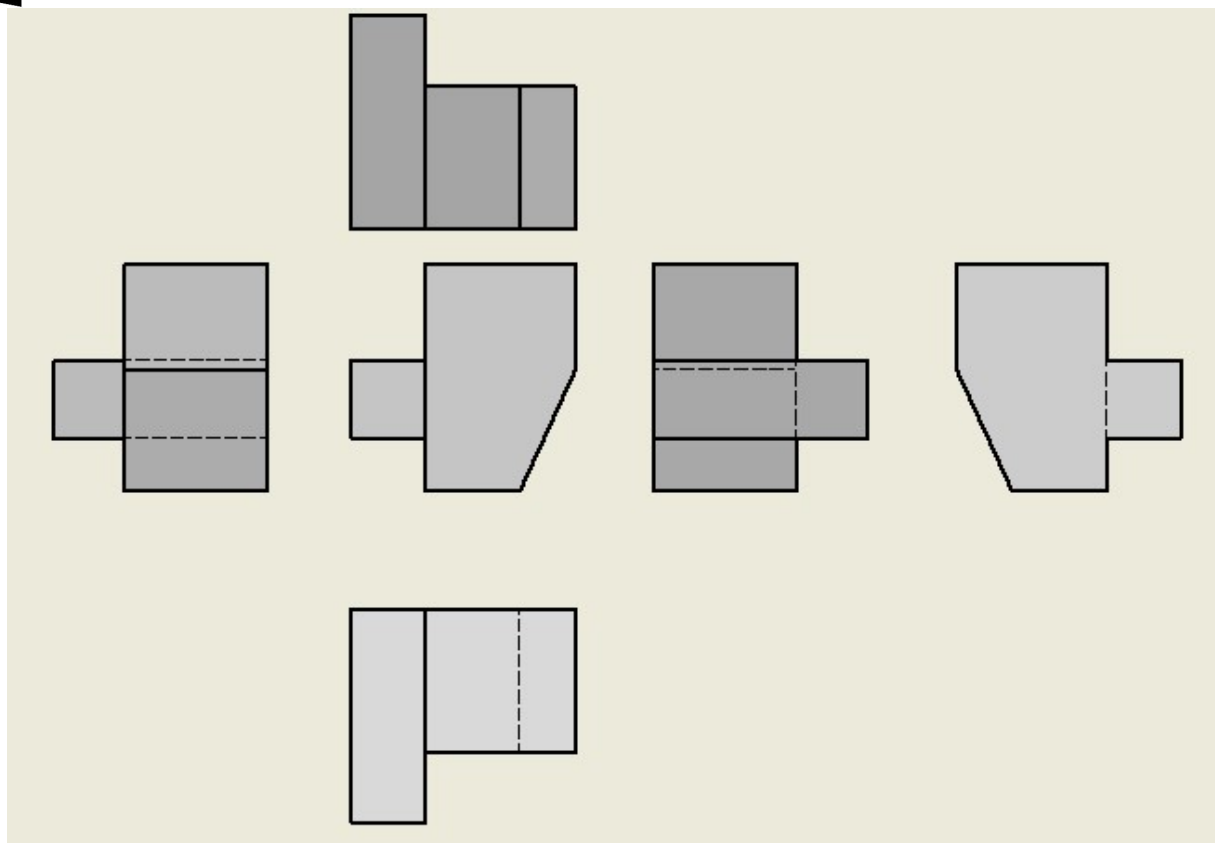
5.4. Zadania

Zadanie numer 1. Proszę o wykonanie rzutowanie europejskie dla poniższego elementu

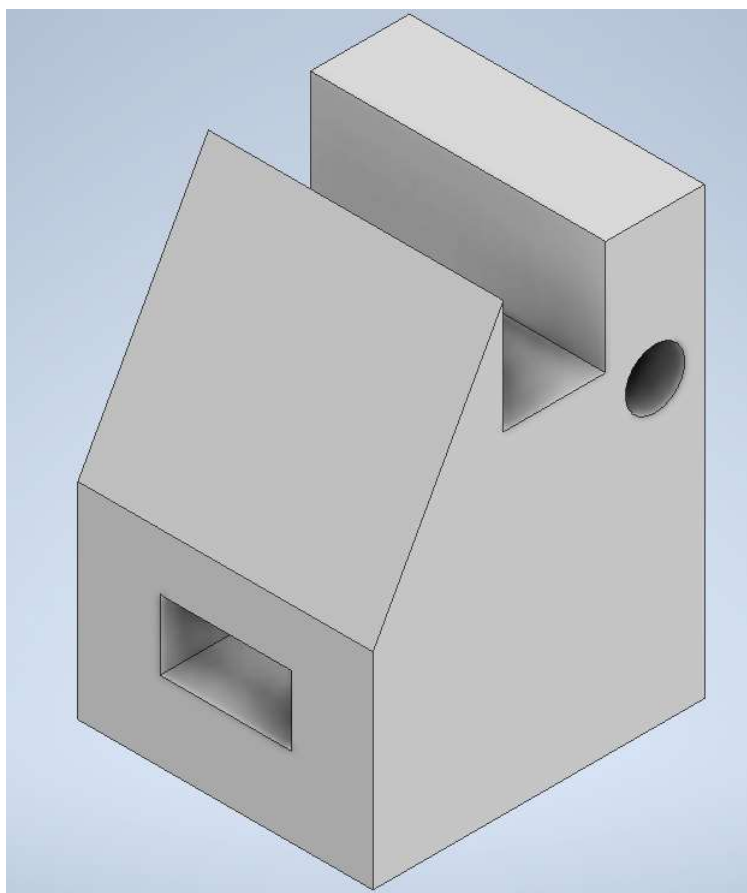


Rozwiązanie dla zadania numer 1

³⁶ Dobrzański T. Rysunek techniczny maszynowy, Wydawnictwo PWN, 2005, s. 57

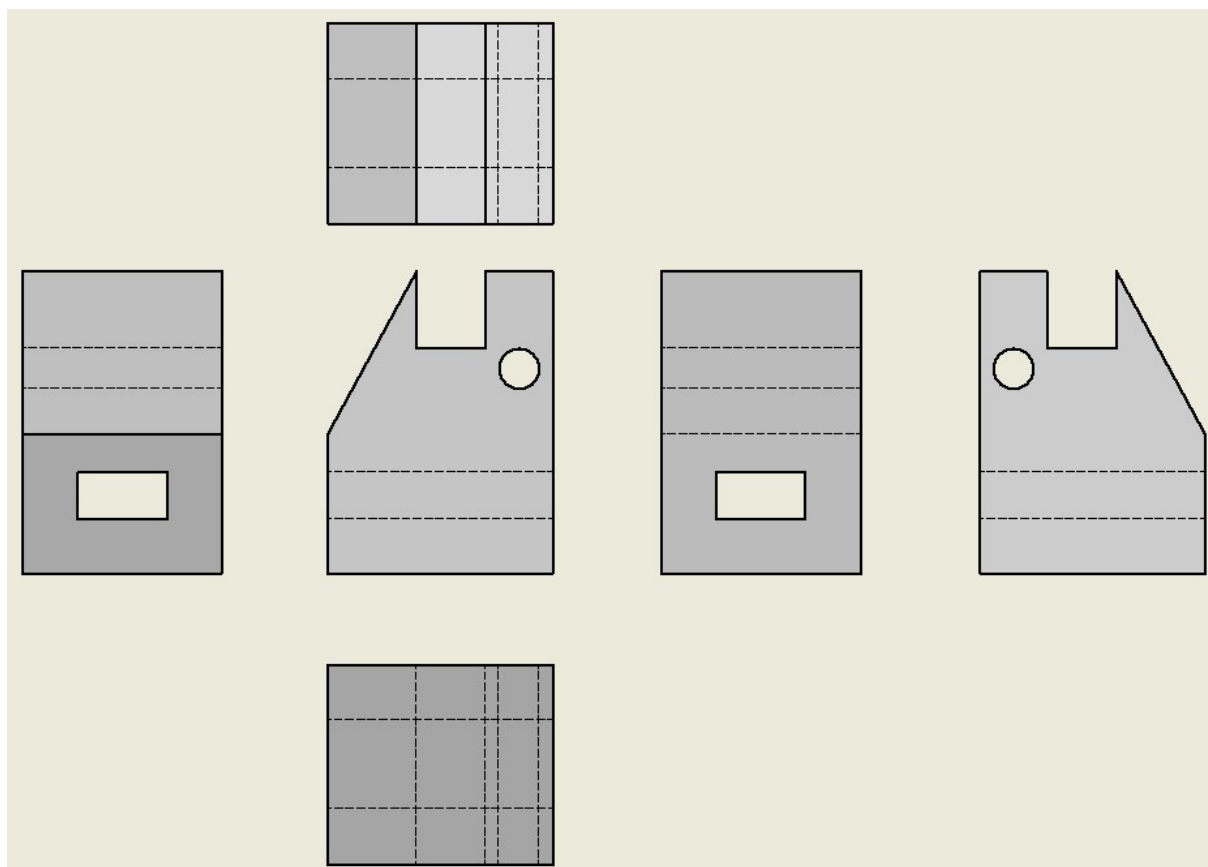


Zadanie numer 2. Proszę o wykonanie rzutowanie europejskie dla poniższego elementu

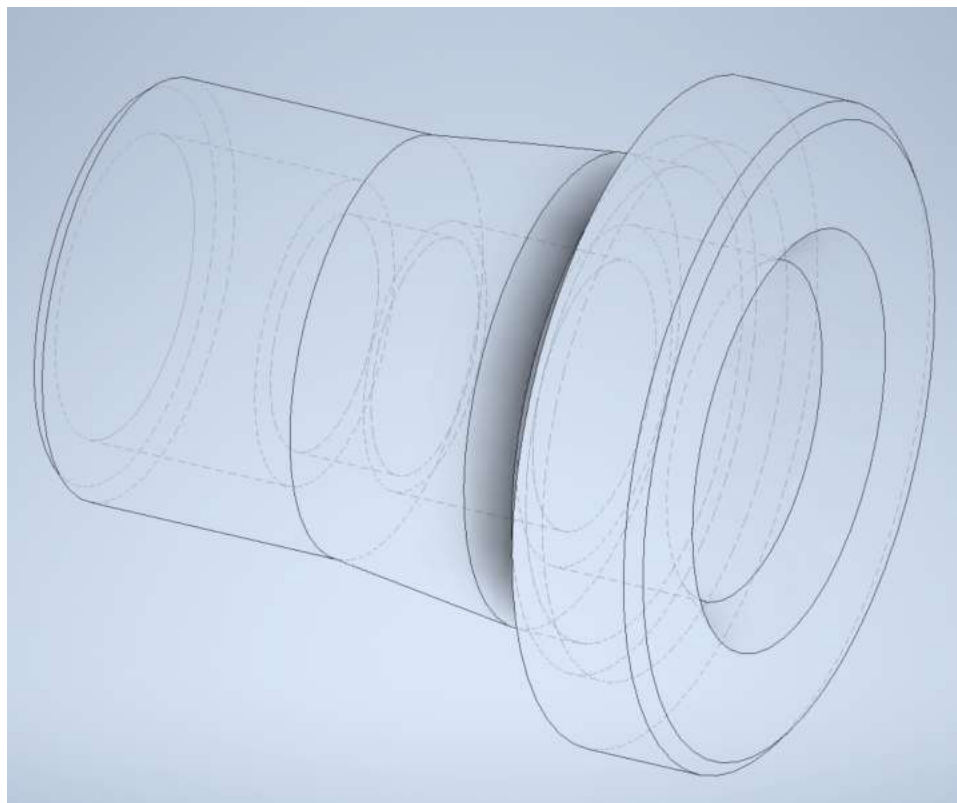




Rozwiązanie dla zadania numer 2

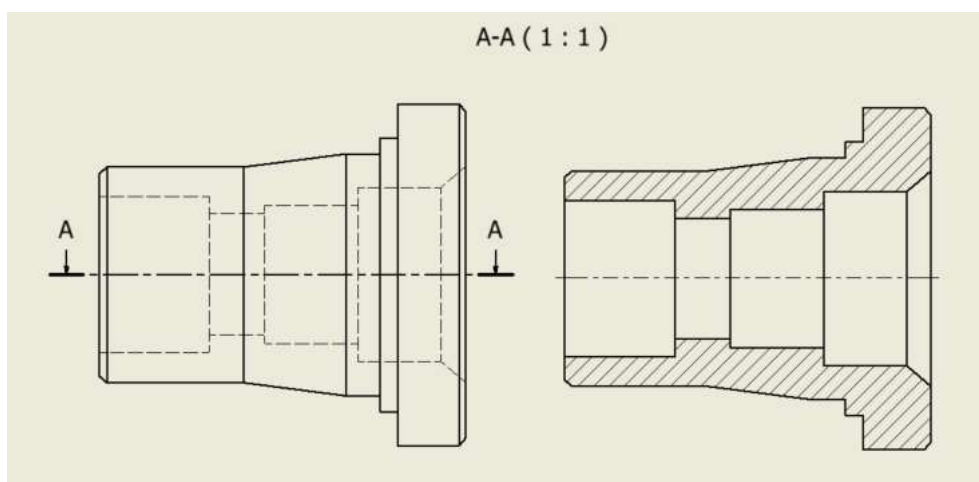


Zadanie numer 3. Proszę o wykonanie przekroju dla poniższego elementu

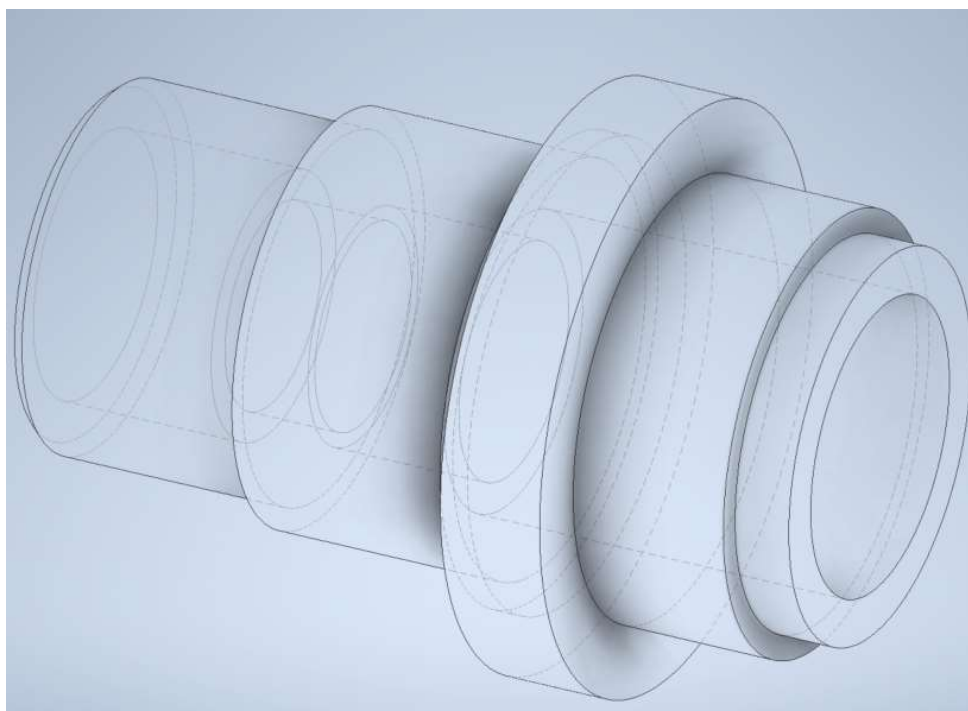




Rozwiązanie zadania numer 3



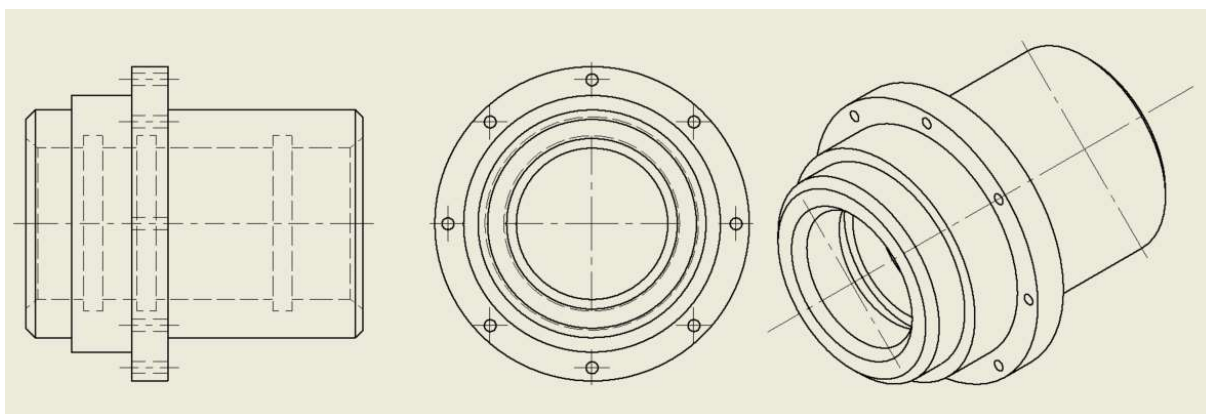
Zadanie numer 4. Proszę o wykonanie przekroju dla poniższego elementu



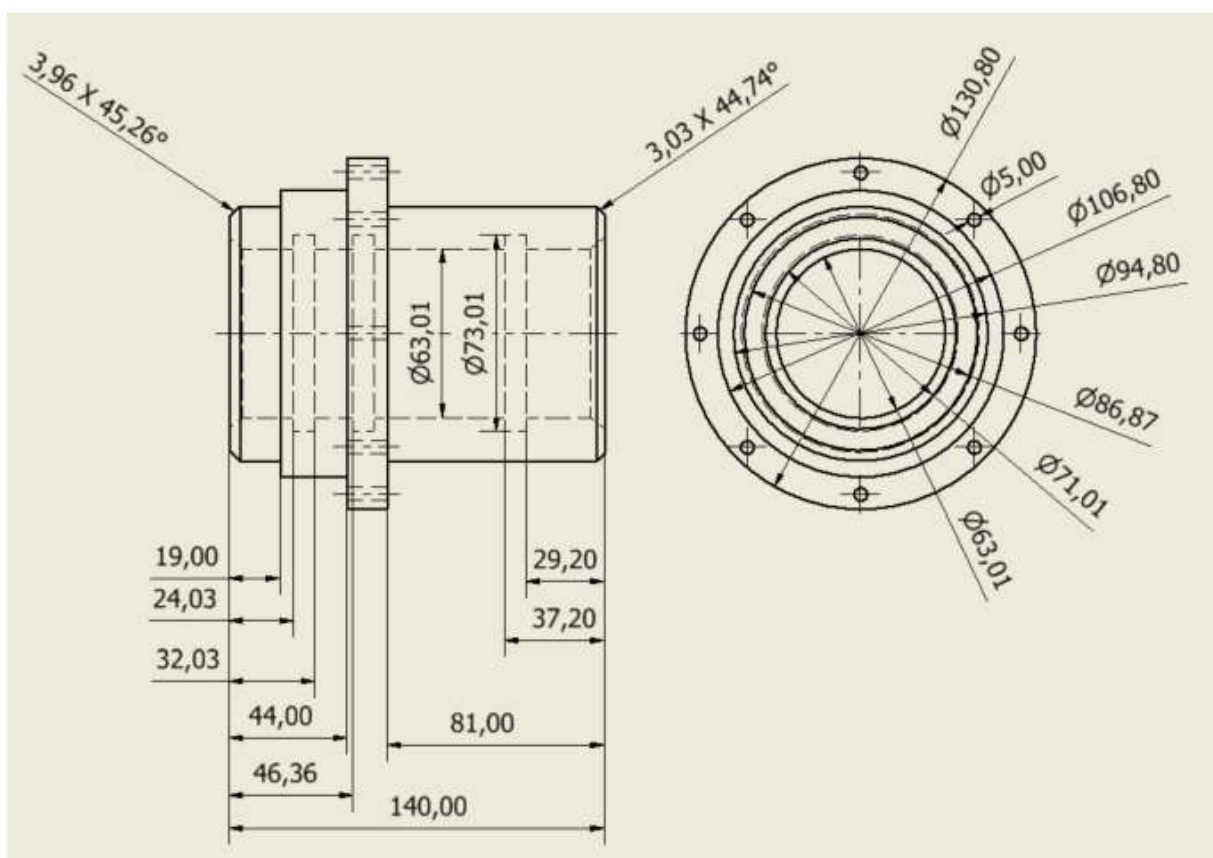
The image displays three orthographic views of a mechanical component, likely a flange or a connector. The views are arranged horizontally: a front view on the left, a top view in the center, and an isometric view on the right. The front view shows a central cylindrical body with a flange on the left and a hexagonal base on the right. The top view shows a circular flange with six bolt holes and a central circular opening. The isometric view provides a 3D perspective of the component, showing the flange, the central body, and the hexagonal base.



Zadanie numer 6. Proszę o wykonanie wymiarowania dla poniższego elementu



Rozwiązanie zadania numer 6



5.5. Bibliografia

1. Boder A., Dudziak M., *Zapis konstrukcji*, Wydawnictwo Politechniki Poznańskiej, Poznań 1995.
2. Dobrzański T., *Rysunek techniczny maszynowy*, Wydawnictwo PWN, 2005.
3. Lewandowski T., *Rysunek techniczny dla mechaników*, WSiP, 2018.
4. Paprocki K., *Rysunek techniczny*. Wydawnictwa Szkolne i Pedagogiczne 1996.