



Fundusze Europejskie
dla Wielkopolski

Dofinansowane przez
Unię Europejską



SAMORZĄD
WOJEWÓDZTWA
WIELKOPOLSKIEGO



Automatyka i robotyka

**Skrypt dla studentów
studiów niestacjonarnych**

Rok II

redakcja

Robert Cieślak

Konin 2025

Tytuł

Automatyka i robotyka
Skrypt dla studentów studiów niestacjonarnych
Rok II

Redakcja naukowa

Robert Cieślak

Autorzy

Tomasz Kapłon
Kamil Łodygowski
Andrzej Milecki
Robert Rogaczewski

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

ELEKTROTECHNIKA I ELEKTRONIKA (Tomasz Kapłon)	5
1. Prąd stały	5
2. Prąd przemienny	17
3. Magnetyzm i układy magnetyczne	27
4. Układy trójfazowe	35
5. Wybrane maszyny i urządzenia elektryczne	38
6. Elektronika	45
Bibliografia	59
MONITOROWANIE I DIAGNOSTYKA (Kamil Łodygowski)	60
1. Diagnostyka	60
2. Niezawodność maszyn	64
3. Diagnostyka wibroakustyczna	68
4. Badania nieniszczące	70
5. Czujniki w monitorowaniu maszyn	77
Bibliografia	80
PODSTAWY AUTOMATYKI I ROBOTYKI (Tomasz Kapłon)	82
1. Podstawy automatyki	82
2. Podstawy robotyki	122
Bibliografia	133
PROJEKTOWANIE LINII ZAUTOMATYZOWANYCH (Andrzej Milecki)	135
Wprowadzenie	135
1. Podstawy automatyzacji urządzeń	136
2. Wybrane elementy stosowane w automatyzacji	141
3. Czujniki analogowe	146
4. Napędy i elementy wykonawcze	150



5. Silniki trójfazowe	153
6. Podstawy sterowników przemysłowych typu PLC	157
7. Programowanie sterowników PLC.....	164
8. Zasilanie i połączenia urządzeń zautomatyzowanych.....	168
9. Podstawy doboru elementów i projektowania linii zautomatyzowanych	171
Bibliografia.....	179
PLANOWANIE I STEROWANIE PRODUKCJĄ (Robert Rogaczewski)	180
Wstęp.....	180
1. Planowanie produkcji	181
2. Sterowanie produkcją	198
Zakończenie	209
Bibliografia.....	209



Tomasz Kapłon

ELEKTROTECHNIKA I ELEKTRONIKA

1. Prąd stały

Pojęcia podstawowe

Prądem elektrycznym nazywamy uporządkowany przepływ ładunków elektrycznych przez ośrodek przewodzący. Jeżeli wielkość tego przepływu oraz jego kierunek są stałe to określamy go jako prąd elektryczny stały. W przypadku, kiedy te wartości ulegają zmianie mówimy o prądzie elektrycznym przemiennym. Specyficzna odmianą prądu przemiennego, szeroko stosowana w elektrotechnice jest prąd sinusoidalnie zmienny, który będzie omawiany w dalszej części.

Wielkość przepływu ładunków elektrycznych przez przewodnik jest nazywana **natężeniem prądu elektrycznego**, określa ona wartość ładunków elektrycznych przepływających przez ośrodek przewodzący (przewód) w określonym czasie. Jednostką natężenia prądu elektrycznego jest **amper [A]**. Jeżeli ten przepływ jest stały oznaczamy natężenie prądu elektrycznego literą I , oraz definiujemy poniższym wzorem

$$I = \frac{Q}{t} \quad (1)$$

gdzie:

I - natężenie prądu elektrycznego

Q – wartość przepływającego ładunku elektrycznego

t – czas przepływu

Czynnikiem powodującym przepływ prądu elektrycznego przez przewód jest różnica potencjałów elektrycznych na jego końcach. Tą różnicę potencjałów **nazywamy napięciem elektrycznym** i oznaczamy literą U . Jednostką napięcia elektrycznego jest wolt [V].

W trakcie przepływu prądu elektrycznego, nośniki ładunków natrafiają nierzadko na przeszkody zaburzające ich swobodny przepływ. Przykładowo elektrony na swojej drodze spotykają jądra atomów, przez co kierunek ich ruchu zostaje zaburzony. Utrudnienia te spowalniają przepływ ładunków. Wielkość tych utrudnień jest określana przez tak zwany **opór elektryczny** nazywany także **rezystancją**. Symbolem rezystancji jest litera R , a jednostką jest om [Ω - duża litera omega].

Prąd elektryczny może płynąć przez zamkniętą drogę nazywaną obwodem elektrycznym. Przyjmuje się, że kierunkiem dodatnim przepływu prądu elektrycznego jest kierunek od

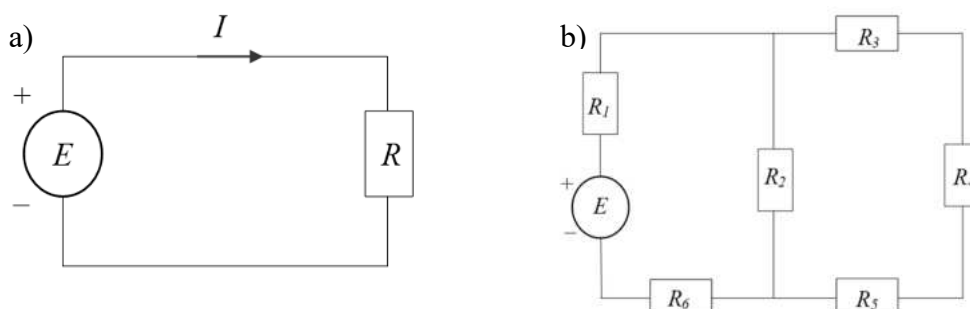


wyższego potencjału do niższego. Obwód elektryczny składa się z co najmniej z trzech elementów:

- źródła energii elektrycznej - jest to aktywny element obwodu elektrycznego
- odbiornika elektrycznego - jest to element bierny obwodu elektrycznego
- przewodów łączących.

Wyróżniamy obwody:

- proste (nierozgałęzione) - istnieje w nich tylko jedna droga przepływu prądu (rys. 1a)
- złożone (rozgałęzione) - zawierająco najmniej dwie niezależne i zamknięte drogi przepływu prądu (rys. 1b).



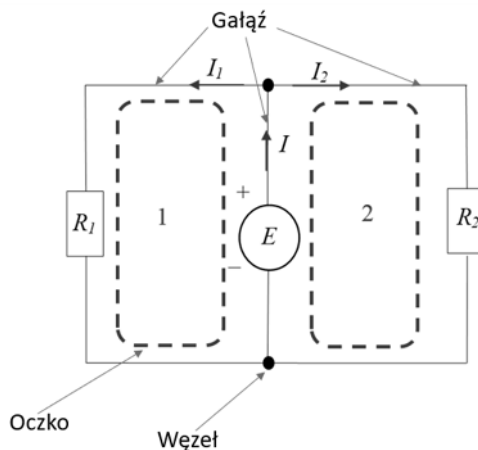
Rys. 1. Obwód prosty (a) oraz złożony (b) E – źródło napięcia, R – rezystancja poszczególnych odbiorników

W obwodzie elektrycznym można wyróżnić następujące elementy (rys. 2):

Węzeł - punkt połączenia co najmniej trzech przewodów

Gałąź - łączy dwa sąsiednie węzły i może zawierać jeden lub więcej odbiorników

Oczko - zbiór połączonych ze sobą elementów tworzących zamkniętą drogę dla przepływu prądu, taką że po usunięciu któregośkolwiek elementu ze zbioru pozostałe elementy nie tworzą drogi zamkniętej.



Rys. 2. Elementy obwodu elektrycznego



Podstawowe prawa dla obwodów prądu stałego

Zależność wiążąca ze sobą wartości natężenia prądu, napięcia elektrycznego oraz rezystancji jest nazywana **prawem Ohma**. Ma ona następującą postać:

$$I = \frac{U}{R} \quad (2)$$

gdzie:

I – natężenie prądu elektrycznego

U – napięcie prądu elektrycznego

R – opór elektryczny

Wartość rezystancji przewodu elektrycznego jest zależna od właściwości materiału, z którego został wykonany, a także od jego parametrów geometrycznych. Rezystancja rośnie wraz z długością przewodu. Wynika to z faktu, że im dłuższa droga przepływu prądu tym większe jest prawdopodobieństwo, że na swojej drodze elektron spotka przeszkodę. Natomiast im większe pole powierzchni przekroju przewodu tym niższa rezystancja- wynika to z faktu, że ładunki mają szerszą drogę, przez którą mogą swobodniej przepływać. Wielkością materiałową określającą jak dobrym przewodnikiem jest dany materiał jest **rezystywność** inaczej **opór właściwy**. Jego symbolem jest ρ (mała litera ro), a jednostką $\Omega \cdot m$. Zależność na opór elektryczny ma następującą postać:

$$R = \rho \frac{l}{S} \quad (3)$$

gdzie:

R - rezystancja

ρ – rezystywność

l – długość przewodu

S – pole powierzchni przekroju przewodu

Tabela 1 przedstawia zbiór wartości rezystywności dla metali powszechnie stosowanych w technice.



Tabela 1. Rezystywność wybranych materiałów

Materiał	Rezystywność w temperaturze 20°C
	$\Omega \cdot m$
Srebro	$1,62 \cdot 10^{-8}$
Złoto	$2,44 \cdot 10^{-8}$
Miedź	$1,75 \cdot 10^{-8}$
Aluminium	$2,88 \cdot 10^{-8}$
Żelazo	$9,6 \cdot 10^{-8}$
Nikiel	$7,3 \cdot 10^{-8}$

Na wartość rezystancji wpływ ma także temperatura przewodu. W przypadku metali rezystancja wzrasta wraz ze wzrostem temperatury, w przypadku półprzewodników można znaleźć takie, których rezystancja spada. Zależność opisującą zależność między temperaturą a rezystancją danego elementu przedstawia poniższe równanie.

$$R_T = R_{20}(1 + \alpha \Delta T) \quad (4)$$

gdzie:

R_T – rezystancja elementu w danej temperaturze

R_{20} – rezystancja elementu w temperaturze odniesienia (20°C)

α – temperaturowy współczynnik rezystancji ($\frac{1}{K}$)

ΔT – różnica między temperaturą elementu a temperaturą odniesienia (w kelwinach)

W tabeli 2 zebrano wartości temperaturowych współczynników rezystancji dla wybranych materiałów.

Tabela 2. Temperaturowe współczynniki rezystancji dla wybranych materiałów

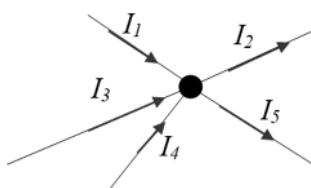
Materiał	Temperaturowy współczynnik rezystancji
	K^{-1}
Srebro	$4,1 \cdot 10^{-3}$
Złoto	$3,6 \cdot 10^{-3}$
Miedź	$3,9 \cdot 10^{-3}$
Aluminium	$4,4 \cdot 10^{-3}$
Żelazo	$6,5 \cdot 10^{-3}$
Nikiel	$6,0 \cdot 10^{-3}$



W obliczeniach obwodów elektrycznych oprócz prawa Ohma zastosowanie mają także dwa prawa Kirchhoffa sformułowane na podstawie bilansu energetycznego.

I prawo Kirchhoffa dotyczy węzłów obwodu elektrycznego. W każdym węźle obwodu elektrycznego suma natężeń prądów wpływających do węzła równa się sumie natężeń prądów wypływających z węzła. W przypadku przedstawionym na rysunku 3:

$$I_1 + I_3 + I_4 = I_2 + I_5 \quad (5)$$



Rys. 3. Prądy wpływające i wypływające z węzła

Jeżeli przyjmie się że prądy wpływające są dodatnie a wypływające ujemne to wzór ogólny przyjmuje postać:

$$\sum_{k=1}^n I_k = 0 \quad (6)$$

gdzie:

I_k – prąd z danej gałęzi łączącej się z węzłem

n – liczba gałęzi łączących się w danym węźle

II prawo Kirchhoffa dotyczy zależności między siłami elektromotorycznymi i spadkami napięć w oczku. W dowolnym oczku obwodu elektrycznego suma algebraiczna sił elektromotorycznych (napięć źródłowych) jest równa sumie algebraicznej spadków napięć na rezystancjach tego oczka. Dla przykładowego oczka na rysunku 4 można zapisać:

$$E_1 - E_3 + E_2 = U_1 - U_4 - U_3 + U_2 \quad (7)$$

gdzie:

E_i – napięcia źródłowe w oczku

U_i – spadki napięcia na poszczególnych odbiornikach

a po zastosowaniu prawa Ohma i podstawieniu w miejscu poszczególnych spadków napięć U_i odpowiednich iloczynów I_i oraz R_i :

$$E_1 - E_3 + E_2 = R_1 I_1 - R_4 I_4 - R_3 I_3 + R_2 I_2 \quad (8)$$



$$P = R \cdot I^2 = \frac{U^2}{R} \quad (12)$$

P - moc

R - rezystancja

U – napięcie elektryczne

I - natężenie prądu elektrycznego

Ponieważ jednostki podstawowe układu SI dżule i waty są zbyt małe do stosowania w licznikach w rzeczywistych instalacjach elektrycznych, często stosuje się jednostki pochodne

Do określenia mocy:

1 kW (kilowat) = 1000 W

1 MW (megawat) = 1 000 000 W

Natomiast do opisanie energii:

1 kWh (kilowatogodzinę) = 1000 W · 3600 s = 3,6 · 10⁶ J.

Źródła energii elektrycznej

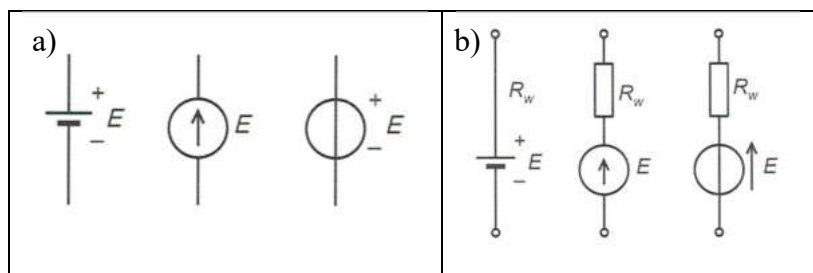
Ze względu na sposób dostarczania energii elektrycznej źródła można podzielić na dwa rodzaje:

- źródła napięciowe
- źródła prądowe.

Źródłem napięciowym (idealnym) nazywamy element, na którego zaciskach występuje napięcie elektryczne o nieziennej wartości bez względu na obciążenie jakie jest podłączone do jego zacisków. Tego typu obiekt nie jest możliwy do realizacji fizycznej, ponieważ każde źródło posiada pewną rezystancję wewnętrzną. W związku z tym definiuje się także rzeczywiste źródło napięciowe. W najprostszym modelu źródła rzeczywistego przyjmuje się, iż ma ono dwa parametry: siłę elektromotoryczną (napięcie źródłowe) oraz rezystancję wewnętrzną. Rzeczywiste źródła napięciowe można podzielić na:

- źródła chemiczne - baterie, akumulatory
- źródła elektromagnetyczne - prądnice prądu stałego i zmiennego
- źródła termoelektryczne - termoogniwa, termopary
- źródła fotoelektryczne - tzw. fotoogniwa.

Symbole stosowane dla źródeł napięciowych przedstawia rysunek 5.



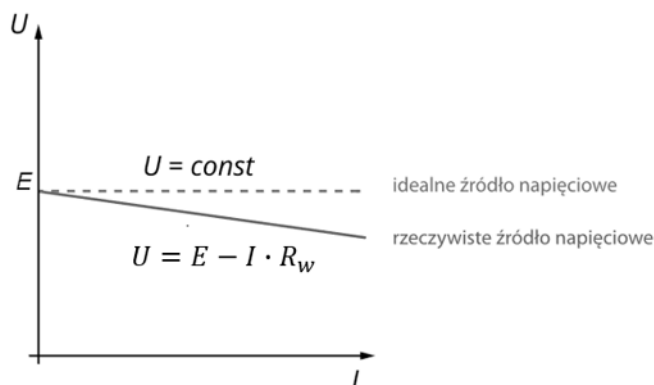
Rysunek 5. Symbole źródła napięcia a - idealne, b - rzeczywiste

Źródło napięciowe może znajdować się w trzech stanach pracy, określanych jako stan jałowy, stan zwarcia oraz stan obciążenia. **Stan jałowy** występuje, gdy obwód elektryczny jest przerwany, między zaciskami źródła występuje nieskończenie duża rezystancja, wskutek czego nie występuje przepływ prądu. Wówczas napięcie mierzone na zaciskach rzeczywistego źródła napięciowego jest największe. **Stan zwarcia** występuje, gdy zaciski źródła zostaną ze sobą bezpośrednio zwarte, rezystancja między nimi wynosi wówczas 0. W takim przypadku wartość natężenia prądu płynącego między zaciskami osiągnie wartość maksymalną, w przypadku rzeczywistego źródła napięciowego ograniczoną jedynie jego rezystancją wewnętrzną. **Stan obciążenia** jest typowym stanem pracy źródła napięciowego. Wówczas między zaciskami źródła jest podłączony odbiornik o określonej rezystancji. W takim przypadku wartość prądu płynącego w obwodzie, oraz napięcia mierzonego na zaciskach źródła rzeczywistego zależy od rezystancji odbiornika oraz rezystancji wewnętrznej źródła. Rysunek 6 przedstawia zbiorczo stany pracy źródła napięciowego.

Stan jałowy ($I_0 = 0$; $R_0 = \infty$)	Stan zwarcia ($I_0 = I_Z$; $R_0 = 0$) $I_Z = \frac{E}{R_w}$	Stan obciążenia $U = E - R_w \cdot I$ $I = \frac{E}{R + R_w}$

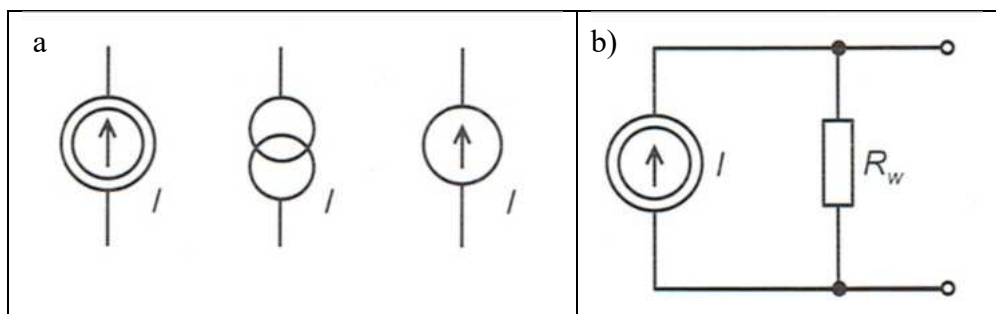
Rysunek 6. Stany pracy źródła napięciowego rzeczywistego

Na rysunku 7 przedstawiono charakterystykę prądowo napięciową idealnego źródła napięciowego i rzeczywistego.

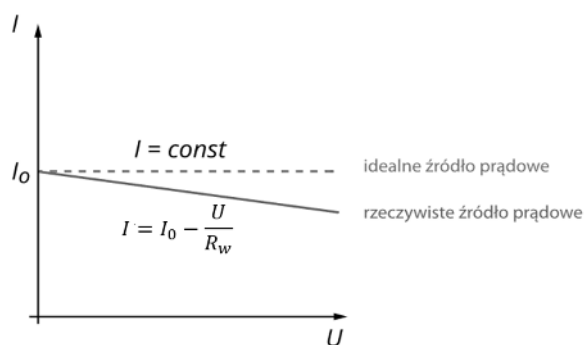


Rysunek 7. Charakterystyka prądowo-napięciowa idealnego i rzeczywistego źródła napięciowego

Źródłem prądowym (idealnym) nazywamy element, który wymusza pomiędzy swoimi zaciskami prąd elektryczny o niezmiennym natężeniu bez względu na obciążenie, jakie podłączymy do tego źródła. W rzeczywistym źródle napięciowym występuje jednakże także rezystancja wewnętrzna, wskutek czego natężenie prądu także zależy od obciążenia. Rzeczywiste źródło prądowe modeluje się jako idealne źródło prądowe połączone równolegle z rezystancją wewnętrzną. Rysunek 7a przedstawia symbole źródła prądowego, natomiast rysunek 8 charakterystykę prądowo napięciową idealnego i rzeczywistego źródła prądowego.



Rys. 7a. Symbole źródła prądowego a) idealnego, b) – rzeczywistego



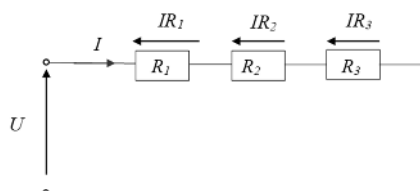
Rys. 8. Charakterystyka prądowo-napięciowa idealnego i rzeczywistego źródła prądowego

Łączenie rezystorów i źródeł prądu elektrycznego

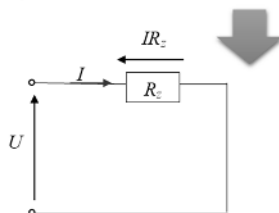
W trakcie obliczeń obwodów elektrycznych występuje potrzeba zastępowania rezystancji poszczególnych odbiorników tak zwaną rezystancją zastępczą R_Z . Rezystancja zastępcza jest to więc rezystancja jaką miałby rezystor równoważny kilku (myślowo) zastępowanym rezystorom. Sposób obliczenia rezystancji zastępczej zależy od połączenia rezystorów.

Przykładowe **połączenie szeregowe rezystorów** zostało przedstawione na rysunku 9. Przy takim połączeniu przez wszystkie rezystory przepływa prąd o takim samym natężeniu.

a)



b)



Rys. 9. a - połączenie szeregowe rezystorów, b - obwód równoważny

W przypadku przedstawionym na rysunku 9a można więc zapisać:

$$U = IR_1 + IR_2 + IR_3 \quad (13)$$

W przypadku obwodu równoważnego z rysunku 9b:

$$U = IR_Z \quad (14)$$

Na podstawie przyrównania zależności na napięcie U :

$$IR_Z = IR_1 + IR_2 + IR_3 \quad (15)$$



$$R_Z = R_1 + R_2 + R_3 \quad (16)$$

Dla przypadku uogólnionego można zapisać:

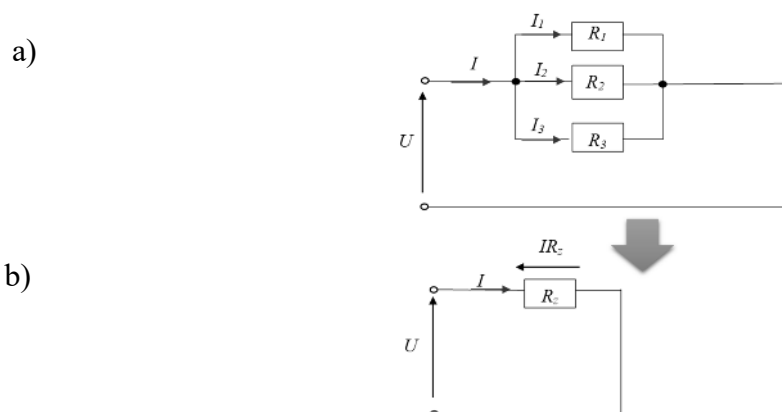
$$R_Z = \sum_{k=1}^n R_k \quad (17)$$

gdzie:

R_Z – rezystancja zastępcza

R_k – rezystancja poszczególnych odbiorników połączonych szeregowo

Przykładowe **połączenie równoległe rezystorów** zostało przedstawione na rysunku 10. Przy takim połączeniu na wszystkich rezystorach występuje spadek napięcia o takiej samej wartości.



Rys. 10. a - połączenie równoległe rezystorów, b - obwód równoważny

W przypadku przedstawionym na rysunku 10a można zapisać:

$$I = I_1 + I_2 + I_3 \quad (18)$$

$$I_1 = \frac{U}{R_1} \quad (19)$$

$$I_2 = \frac{U}{R_2} \quad (20)$$

$$I_3 = \frac{U}{R_3} \quad (21)$$

W przypadku obwodu równoważnego z rysunku 10b:

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

Na podstawie przyrównania zależności na natężenie prądu I :

$$\frac{U}{R_Z} = \frac{U}{R_1} + \frac{U}{R_2} + \frac{U}{R_3} \quad (23)$$

$$\frac{1}{R_Z} = \frac{1}{R_1} + \frac{1}{R_2} + \frac{1}{R_3} \quad (24)$$



Dla przypadku uogólnionego można zapisać:

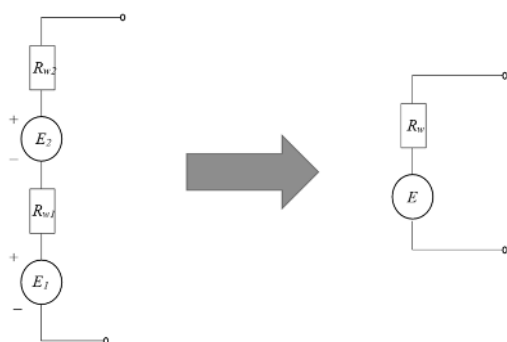
$$\frac{1}{R_Z} = \frac{1}{R_1} + \frac{1}{R_2} + \dots + \frac{1}{R_n} \quad (25)$$

gdzie:

R_Z – rezystancja zastępcza

R_k – rezystancja poszczególnych odbiorników połączonych szeregowo

W przypadku łączenia szeregowego źródeł napięcia, jak na rysunku 11 zastępcze napięcie źródłowe jest sumą napięć poszczególnych źródeł.



Rys. 11. Połączenie szeregowe źródeł napięciowych

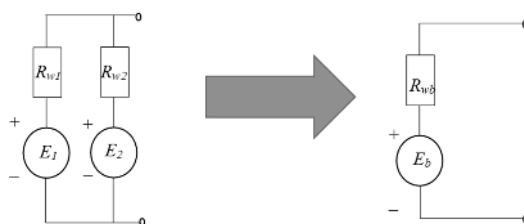
$$E = E_1 + E_2 \quad (26)$$

$$R_w = R_{w1} + R_{w2} \quad (27)$$

$$E = E_1 = E_2 \quad (28)$$

$$R_{wb} = \frac{R_w}{2} \quad (29)$$

Jeżeli następuje łączenie takich samych źródeł napięciowych równolegle jak na rysunku 12, napięcie źródłowe pozostaje niezmienione, zmniejsza się natomiast rezystancja wewnętrzna. Dzięki temu można uzyskać większy maksymalny prąd na wyjściu. W przypadku połączenia źródeł napięciowych o różnych napięciach źródłowych, silniejsze źródło zacznie ładować słabsze i jeżeli to nie jest do tego dostosowane może zostać uszkodzone.



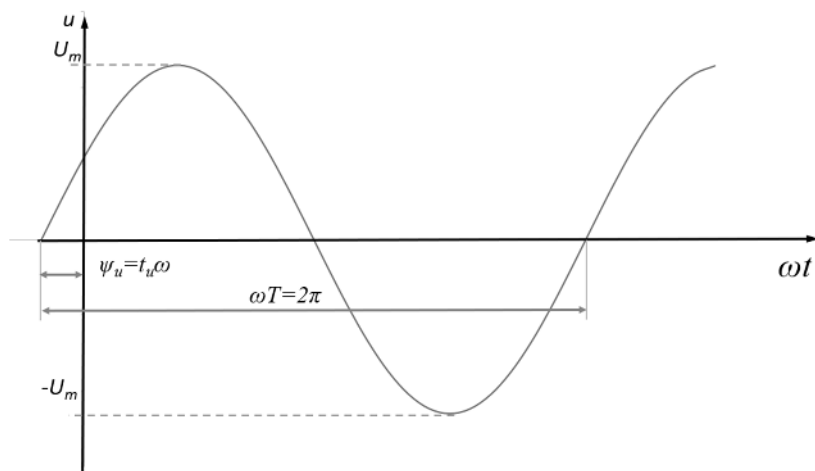
Rys. 12. Połączenie równoległe źródeł napięciowych

Wymagane: $E_1 = E_2$ oraz $R_{w1} = R_{w2} = R_w$



2. Prąd przemienny

Zależnie od zmiany kierunku i wartości przepływu prądu dzieli się go na prąd stały oraz zmienny. Jeżeli zmiana napięcia elektrycznego i natężenia prądu ma charakter sinusoidalny nazywamy go prądem sinusoidalnie zmiennym. Powszechnie w skrócie mówi się o nim jako o prądzie przemiennym. Przykładowy przebieg napięcia prądu sinusoidalnego oraz parametry opisujące go przedstawiono na rysunku 13.



Rysunek 13. Przykładowy przebieg napięcia prądu sinusoidalnie zmiennego

Przebieg napięcia można opisać wzorem:

$$u = U_m \sin(\omega t + \psi_u) \quad (30)$$

W wzorze tym:

u – jest to napięcie chwilowe (V)

U_m – jest to amplituda, czyli największa wartość napięcia chwilowego (V)

T – okres prądu sinusoidalnego (s)

ω – pulsacja (rad/s)

ψ_u – początkowe przesunięcie fazowe

Wyróżnić można także częstotliwość prądu sinusoidalnego f , której jednostką są herce (Hz), będącą odwrotnością okresu.

Przebieg natężenia prądu elektrycznego także ma charakter sinusoidalny. W zależności od właściwości danego obwodu elektrycznego przebiegi sinusoidalne napięcia i natężenia prądu mogą być ze sobą zgodne w fazie, albo przesunięte o pewien kąt.



Wartości średnie i skuteczne prądu przemiennego

Ponieważ w przypadku prądów sinusoidalnie zmiennych wartości natężeń prądów oraz napięć są zmienne w czasie wykonywanie dla nich obliczeń może być skomplikowane. Z tego powodu dla ich uproszczenia wprowadza się równoważne wartości stałe w czasie. Zależnie od celu obliczeń posługujemy się wartościami średnimi albo skutecznymi.

W przypadku rozważań dotyczących ładunku elektrycznego posługujemy się wartościami średnimi. Wartość średnia prądu sinusoidalnego jest to wartość zastępczego prądu stałego, który w czasie równym połowie okresu przenosi taki sam ładunek jak prąd sinusoidalny.

W przypadku wartości średniej natężenia prądu można ją wyprowadzić w sposób następujący

$$I_{sr} \frac{T}{2} = \int_0^{T/2} i dt \quad (31)$$

$$I_{sr} = \frac{2}{T} \int_0^{T/2} i dt \quad (32)$$

$$I_{sr} = \frac{2}{T} \int_0^{T/2} I_m \sin \omega t dt = \left[\frac{2I_m}{\omega T} (-\cos \omega t) \right]_0^{T/2} = \frac{2}{\pi} I_m \approx 0,637 I_m \quad (33)$$

gdzie:

I_{sr} – wartość średnia natężenia prądu

I_m – amplituda natężenia prądu

W przypadku wartości średniej napięcia, podobnie:

$$U_{sr} = \frac{2}{\pi} U_m \approx 0,637 U_m \quad (34)$$

W przypadku rozważań energetycznych prąd sinusoidalny zastępuje się równoważnym prądem stałym o natężeniu I albo napięciu U , który wydzielilby na rezystorze R taką samą ilość energii w postaci ciepła jak prąd sinusoidalny płynący przez taki sam okres czasu T . Przy czym wartości I oraz U nazywamy wartościami skutecznymi. Zapisujemy je za pomocą dużych liter. Zależność na wartość skuteczną natężenia prądu sinusoidalnego można wyprowadzić w sposób następujący:

$$RI^2T = \int_0^T Ri^2 dt \quad (35)$$

$$I = \sqrt{\frac{1}{T} \int_0^T i^2 dt}$$

(36)

$$I = \sqrt{\frac{1}{T} \int_0^T I_m^2 \sin^2 \omega t dt} = \sqrt{\frac{I_m^2}{T} \int_0^T (1 - \cos 2\omega t) dt} = \sqrt{\frac{I_m^2}{2}} = \frac{I_m}{\sqrt{2}} \quad (37)$$

W przypadku wartości skutecznej napięcia podobnie:



$$U = \frac{U_m}{\sqrt{2}} \quad (38)$$

Do zapisu wartości prądu elektrycznego przemiennego można zastosować także liczby zespolone. Zapis zespolony wartości skutecznej wykorzystuje związek między zapisem wielkości sinusoidalnie zmiennym a zapisem wykładniczym liczby zespolonej. Przykładowo:

$$\underline{I} = I e^{j\varphi} \quad (39)$$

gdzie:

I – wartość skuteczna

φ – argument równy fazie początkowej

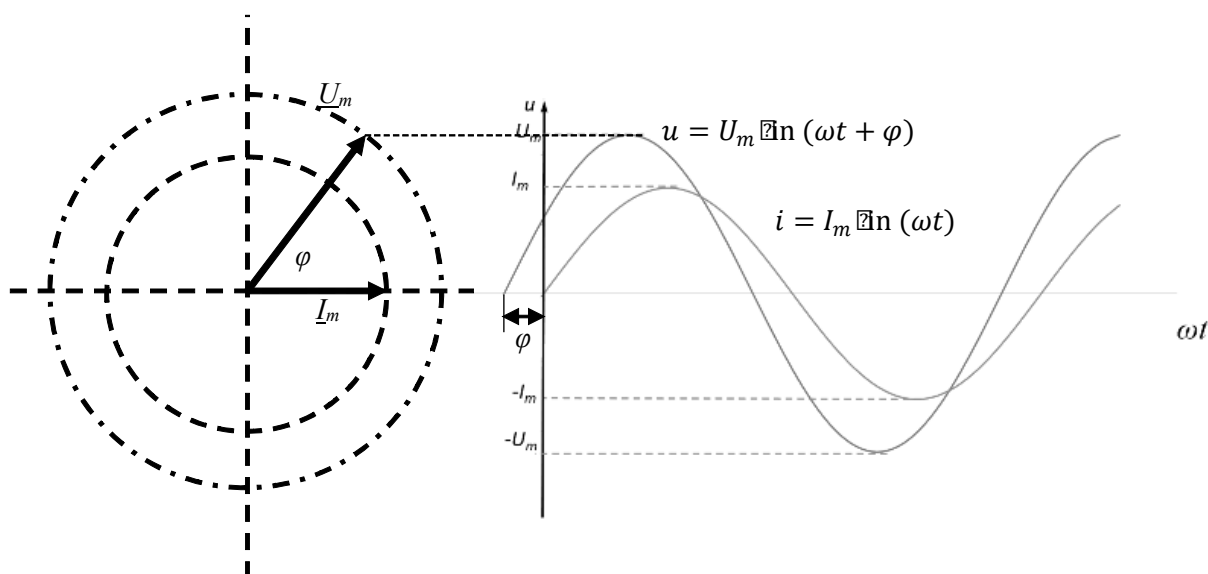
j – liczba urojona definiowana jako $j^2 = -1$

Co można rozisać:

$$\underline{I} = I e^{j\varphi} = I \cos(\varphi) + j I \sin(\varphi) \quad (40)$$

Wykresy wskazowe prądów i napięć

Wartości zmienne prądu sinusoidalnego często dużo wygodniej jest przedstawiać jako tak zwane wektory wirujące zamiast w postaci przebiegu czasowego. W przypadku takiego wektora jego moduł (długość) jest równy amplitudzie wielkości sinusoidalnej (natężenia albo napięcia). Kąt natomiast jaki tworzy wektor z półosią odciętych jest równy fazie początkowej. Wykresy wskazowe sporządza się także dla wartości skutecznych - wówczas moduły wektorów są równe wartościom skutecznym. Przykładowy wykres wskazowy przedstawia rysunek 14.



Rysunek 14. Przykład wektora wskazowego dla natężenia prądu i napięcia przesuniętego o kąt φ



Moc oraz energia w obwodach prądu przemiennego

W przypadku prądu stałego wartości napięcia i natężenia są stałe, z tego powodu ich iloczyn określający moc był także wartością stałą. Ponieważ w przypadku prądu przemiennego wartości napięcia i natężenia są zmienne w czasie, również moc obliczoną jako ich iloczyn określa się jako moc chwilową – p . Małe litery wskazują że mówimy o wartościach chwilowych.

$$p = u \cdot i \quad (41)$$

Dla lepszego zobrazowania mocy wygodniej niż za pomocą wartości chwilowych jest posługiwać się wartościami skutecznymi napięcia i natężenia. W przypadku obwodów prądu przemiennego sinusoidalnego wyróżniamy trzy rodzaje mocy: moc czynną, moc bierną oraz moc pozorną.

Moc czynna jest mocą zamienianą na faktyczną pracę użyteczną, podobne jak w przypadku prądu stałego jej symbolem jest P , a jednostką wat (W). Jej wartość jest iloczynem wartości skutecznych napięcia i natężenia prądu oraz współczynnika mocy. Współczynnik mocy jest wartością cosinusa kąta przesunięcia fazowego między przebiegiem napięcia i natężenia prądu – φ .

$$P = UI \cos \varphi \quad (42)$$

Moc bierna – Q służy między innymi wytworzeniu pól magnetycznych, nie jest zamieniana na pracę użyteczną. Jej jednostką jest war (war). Jest określona przez zależność:

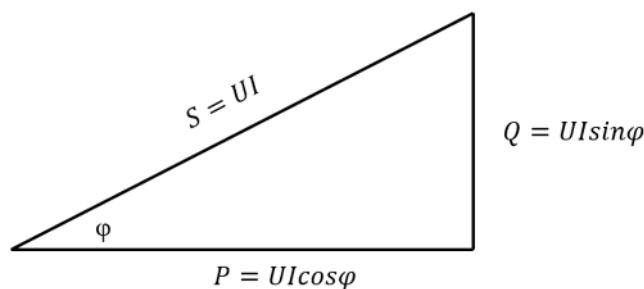
$$Q = UI \sin \varphi \quad (43)$$

Wyróżniamy moc bierną indukcyjną Q_L , które przyjmuje się za dodatnie, oraz moc pozorną pojemnościową Q_C umownie ujemną. W obwodach prądu zmiennego moce pozorne indukcyjne i pojemnościowe mogą nawzajem się skompensować.

Moc pozorna – S jest to największa możliwa wartość mocy czynnej w obwodzie przy $\cos \varphi = 1$. Jest to iloczyn wartości skutecznych napięcia i prądu. Jej jednostką jest woltoamper (V·A).

$$S = UI \quad (44)$$

Relacje między poszczególnymi mocami można zilustrować za pomocą tak zwanego trójkąta mocy, przedstawionego na rysunku 15.



Rysunek 15. Trójkąt mocy

Na podstawie zależności między bokami trójkąta prostokątnego można zapisać:

$$S = \sqrt{P^2 + Q^2} \quad (45)$$

Uwzględniając że moc bierna wydzielana w obwodzie może być zarówno dodatnia, jak i ujemna:

$$S = \sqrt{P^2 + (Q_L - Q_C)^2} \quad (46)$$

Energia czynna - A_{cz} w układach prądu przemiennego jest to energia zamieniana na pracę mechaniczną lub ciepło. Jej wartość jest równa pomnożeniu mocy biernej przez czas

$$A_{cz} = P \cdot t \quad (47)$$

Istnieje także pojęcie energii biernej, która jest równa mocy czynnej pomnożonej przez czas

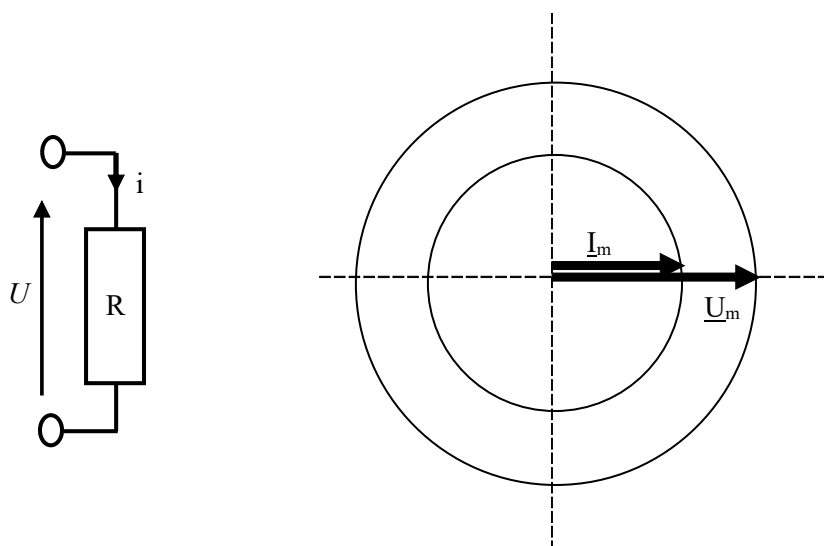
$$A_b = Q \cdot t \quad (48)$$

Drobni odbiorcy energii elektrycznej płacą tylko za pobieraną moc czynną. Dużi natomiast są wyposażeni także w liczniki mocy biernej i za jej odbiór także płacą.

Elementy RLC obwodu prądu przemiennego

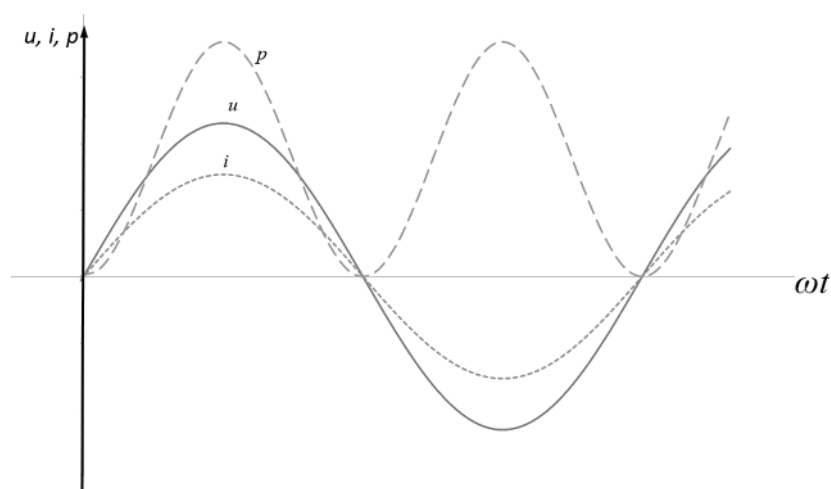
W obwodach prądu przemiennego wyróżniamy elementy rezystancyjne, indukcyjne i pojemnościowe. W swojej idealnej modelowej postaci każdy cechuje się specyficzną właściwością rezystancją, indukcyjnością, oraz pojemnością. Rzeczywiste elementy posiadają jednakże także właściwości pasożytnicze, to oznacza, że przykładowo rezystor oprócz rezystancji posiada także pewną indukcyjność i pojemność. Są one jednak zazwyczaj znacznie mniejsze od podstawowego parametru.

Elementy rezystancyjne nazywane są opornikami albo rezystorami. Ich symbol przedstawia rysunek 16. Oznaczane są symbolem R, a ich podstawowym parametrem jest rezystancja, czyli opór czynny.



Rys. 16. Rezystor i wykres wskazowy przebiegu napięcia i prądu

Cechuje je to, że wydzielą się na nich tylko moc czynna oraz przebiegi prądu i napięcia są ze sobą zgodne w fazie – nie występuje na nich przesunięcie przebiegu napięcia i natężenia prądu. Brak przesunięcia między natężeniem prądu i napięciem widać na wykresie ich przebiegu na rysunku 17.

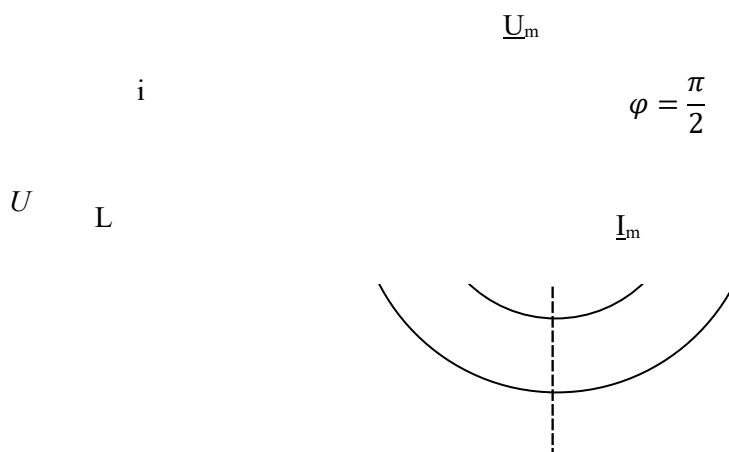


Rys. 17. Zmiany w czasie na rezystorze napięcia, prądu i mocy chwilowej

Elementami indukcyjnymi są najczęściej różnego rodzaju cewki. Symbol cewki przedstawia rysunek 18. Na cewce wydzielą się moc bierna o charakterze indukcyjnym, dla której przyjmuje się, że ma znak dodatni. Podstawowym parametrem cewki jest jej indukcyjność oznaczona symbolem L . Jednostką indukcyjności są henry (H). W przypadku łączenia cewek w sposób szeregowy ich indukcyjność zastępcza jest sumą indukcyjności połączonych



elementów. Natomiast w przypadku łączenia równoległego odwrotność indukcyjności zastępczej jest sumą odwrotności indukcyjności połączonych cewek.

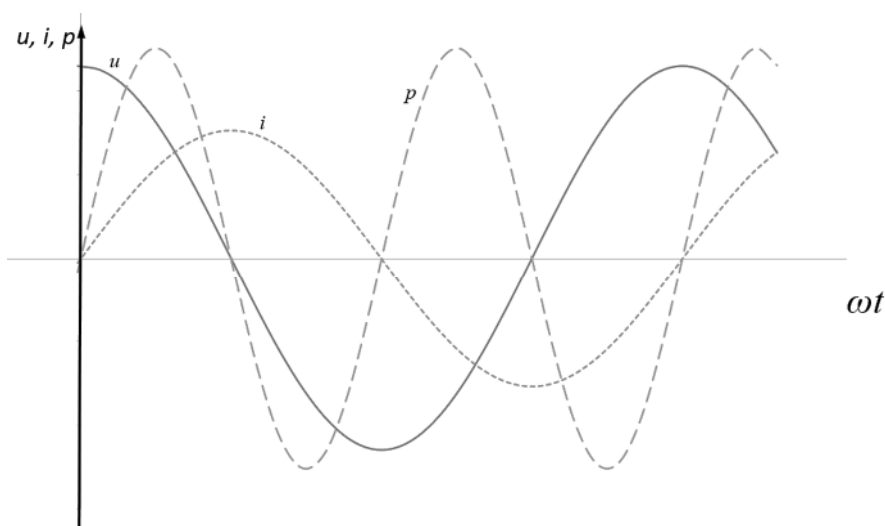


Rys. 18. Symbol cewki oraz jej wykres wskazowy

W elementach indukcyjnych występuje opór bierny indukcyjny zależny od indukcyjności cewki oraz od częstotliwości prądu płynącego przez cewkę. Opisuje go poniższy wzór:

$$X_L = 2\pi fL \quad (49)$$

W przypadku cewki przebieg napięcia wyprzedza przebieg prądu o 90° . Co przedstawia rysunek 19.

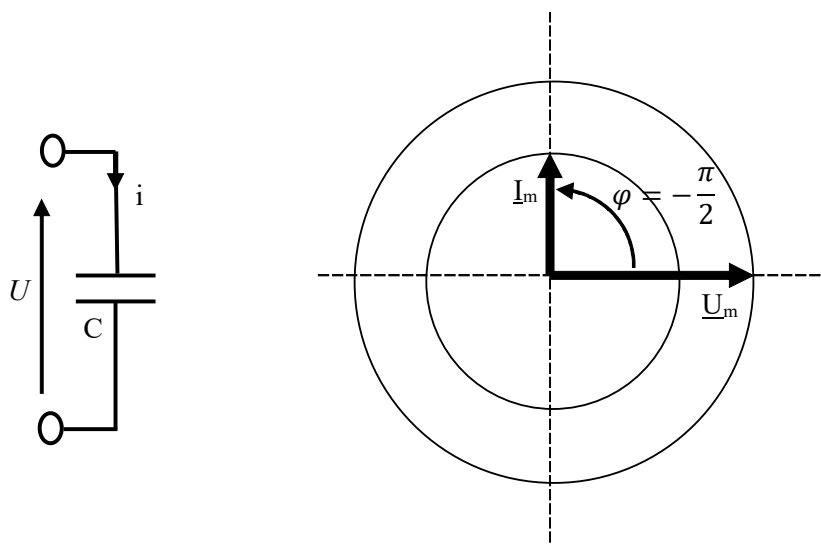


Rys. 19. Przebieg w czasie napięcia, prądu i mocy chwilowej na cewce

Elementami pojemnościowymi są najczęściej różnego rodzaju kondensatory. Symbol kondensatora przedstawia rysunek 20. Na kondensatorze wydzielą się moc bierna o charakterze pojemnościowym, dla której przyjmuje się, że ma znak ujemny. Podstawowym parametrem



kondensatora jest jego pojemność oznaczona symbolem C . Jednostką pojemności są farady (F). Przypadku kondensatorów przy ichłączeniu równoległym pojemność zastępcza równa się sumie pojemności kondensatorów połączonych. Natomiast przy połączeniu szeregowym odwrotność pojemności zastępczej równa się sumie odwrotności pojemności połączonych kondensatorów.

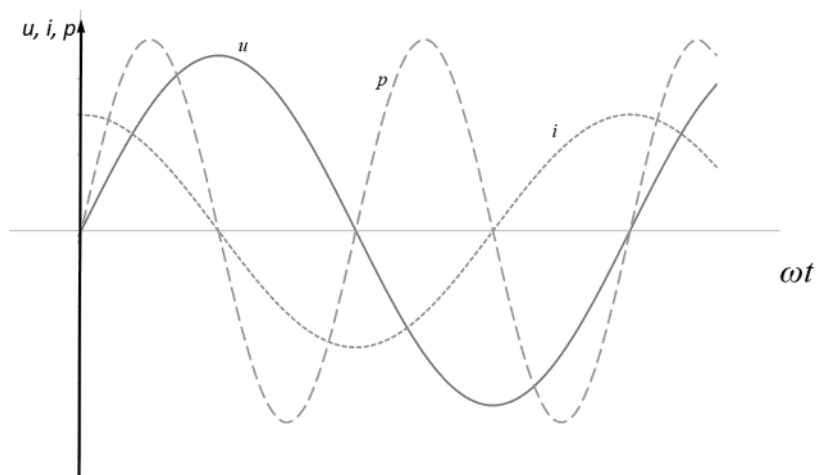


Rys. 20. Symbol kondensatora oraz jego wykres wskazowy

W elementach pojemnościowych występuje opór bierny pojemnościowy zależny od pojemności kondensatora oraz od częstotliwości prądu płynącego przez cewkę. Opisuje go poniższy wzór:

$$X_C = \frac{1}{2\pi f C} \quad (50)$$

W przypadku kondensatora przebieg napięcia jest opóźniony względem prądu o 90° . Co przedstawia rysunek 21.

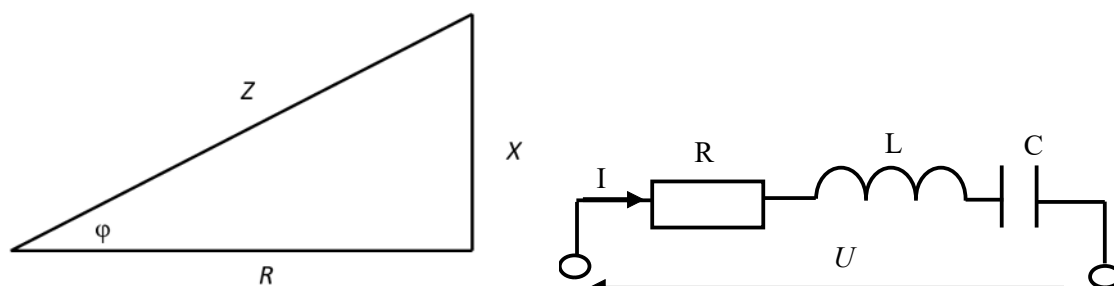


Rys. 21. Przebieg w czasie napięcia, prądu i mocy chwilowej na kondensatorze

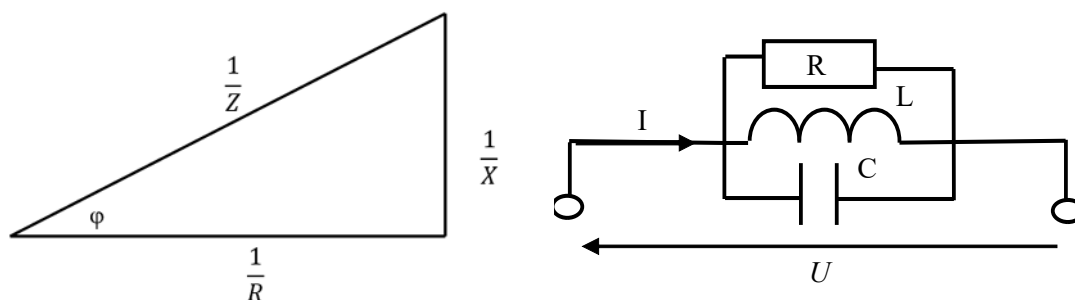


Opór w obwodach prądu przemiennego

W przypadku obwodów prądu przemiennego możemy wyróżnić podobnie jak z mocą opór czynny - rezystancję (R), opór bierny - reaktancję (X) oraz opór pozorny - impedancję (Z). W obliczeniach umownie przyjmuje się opór bierny indukcyjny za dodatni, a pojemnościowy za ujemny. Relacje między poszczególnymi oporami można przedstawić za pomocą trójkąta prostokątnego, z rozróżnieniem czy elementy są połączone szeregowo czy równoległe. Rysunek 22 ilustruje połączenie szeregowe, a 23 równoległe.



Rys. 22. Trójkąt oporów dla elementów RLC połączonych szeregowo



Rys. 23. Trójkąt oporów dla elementów RLC połączonych równoległe

Prawa Kirchhoffa w obwodach prądu sinusoidalnie zmiennego

Prawa Kirchhoffa są spełnione dla obwodów sinusoidalnie zmiennych dla wartości chwilowych oraz dla wartości skutecznych zespolonych.

I prawo Kirchhoffa:

$$\sum_{k=1}^n i_k = 0 \quad (51)$$

$$\sum_{k=1}^n I_k = 0 \quad (52)$$

II prawo Kirchhoffa:

$$\sum_{l=1}^m e_l = \sum_{k=1}^n u_{Rk} + u_{Lk} + u_{Ck} \quad (53)$$

$$\sum_{l=1}^m \underline{E}_l = \sum_{k=1}^n \underline{U}_{Rk} + \underline{U}_{Lk} + \underline{U}_{Ck} \quad (54)$$

Rezonans elektryczny

W obwodach prądu sinusoidalnego impedancja, a zatem także wartości prądów oraz ich przesunięć fazowych w stosunku do napięcia zależą od reaktancji indukcyjnej



i pojemnościowej, ponieważ we wzorze określającym wartość impedancji Z występuje różnica reaktancji pojemnościowej i indukcyjnej,

$$Z = \sqrt{R^2 + (X_L - X_C)^2} \quad (55)$$

może dojść do ich wzajemnego zniesienia pod warunkiem:

$$X_L = X_C \quad (56)$$

Prąd wpływający do obwodu będzie taki jakby w obwodzie była tylko rezystancję, nie będzie przesunięcia fazowego między napięciem i prądem.

Wyróżnia się:

- Rezonans dla obwodu szeregowego- rezonans napięć
- Rezonans dla obwodu równoległego- rezonans prądów

Rezonans napięć występuje w obwodzie szeregowym przy spełnieniu warunku rezonansu

$$\underline{U}_L + \underline{U}_C = 0 \quad (57)$$

Po podstawieniu do warunku rezonansu, wyrażenia na wartość reaktancji indukcyjnej i pojemnościowej

$$X_L = X_C \quad (57)$$

$$2\pi fL = \frac{1}{2\pi fC} \quad (58)$$

Otrzymać można wartość częstotliwości rezonansowej

$$f_0 = \frac{1}{2\pi\sqrt{LC}} \quad (60)$$

Rezonans prądów występuje w obwodzie równoległe połączonych elementów R , L , C po spełnieniu warunku rezonansu

$$\underline{I}_L + \underline{I}_C = 0 \quad (61)$$

Wzór określający częstotliwość rezonansową jest taki sam jak w przypadku rezonansu napięć

Poprawa współczynnika mocy

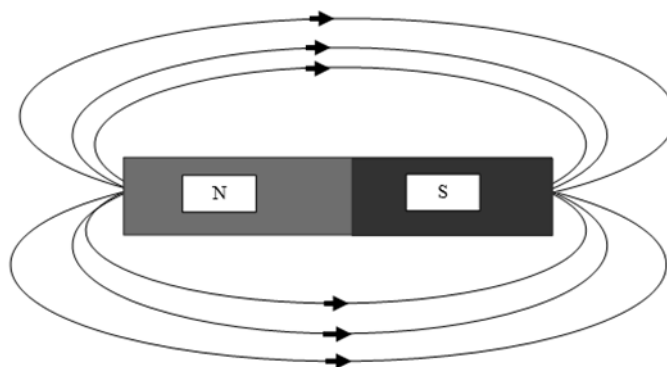
Mała wartość współczynnik mocy zwiększa straty przy przesyle energii. Ponadto przedsiębiorcy płacą dodatkowe opłaty za pobierana moc bierną. Z tych powodów dąży się do kompensacji mocy biernej, aby minimalizować odbiór z sieci mocy biernej. Ponieważ zazwyczaj u odbiorców przemysłowych moc bierna ma charakter indukcyjny, ze względu na występowanie silników elektrycznych, kompensacja mocy biernej polega na dołączaniu dodatkowej pojemności do układu. Kompensacja może być wykonywana indywidualnie przy maszynach elektrycznych, albo grupowo.



3. Magnetyzm i układy magnetyczne

Pole magnetyczne i magnetyzm

Już w starożytności obserwowano oddziaływania magnetyczne. Zauważono wtedy, że kawałki rudy żelaza, albo magnetytu, mogą przyciągać drobinki żelaza. W średniowieczu zaobserwowano, że za pomocą naturalnych magnesów można nadawać właściwości magnetyczne kawałkom stali przez ich jednokierunkowe pocieranie i powstałe tak magnesy mogą zachować swoje właściwości przez długi czas. Zaobserwowano także, że opiłki stalowe w pobliżu magnesu układają się wzdłuż linii w pobliżu magnesu. W XIX zaobserwowano, że prąd płynący w przewodniku może oddziaływać na igłę magnetyczną znajdującą się w pobliżu. Stwierdzono także, że pole magnetyczne w pobliżu przewodu ma postać koncentrycznych okręgów leżących w płaszczyźnie prostopadłej do osi przewodu, przez który płynie prąd. Przestrzeń, w której występują siły magnetyczne nazywa się polem magnetycznym. Jest ono przedstawiane graficznie w postaci linii pola. Doświadczalnie stwierdzono, że linie pola magnetycznego są zawsze zamknięte oraz że magnes ma zawsze dwa bieguny. Dodatni nazwano północnym - N, a ujemny południowym - S. Umownie przyjęto także, że linie magnetyczne wychodzą z bieguna północnego i wchodzą do południowego, tak jak na rysunku 24.



Rys. 24. Pole magnetyczne

Pole magnetyczne jest przedstawiane za pomocą linii pola magnetycznego. Sumę wszystkich linii pola magnetycznego przechodzących przez określony przekrój nazywa się strumieniem magnetycznym – Φ . Strumień magnetyczny jest wielkością skalarną, a jego jednostką jest woltosekunda [$V \cdot s$], zwana także weberem [Wb]. Istnieje także pojęcie gęstość strumienia magnetycznego, czyli liczby linii pola przypadających na jednostkę powierzchni (S), nazywa



się indukcją magnetyczną – B . Jest ona wielkością wektorową. Jej jednostką jest $\frac{Wb}{m^2}$, czyli tesla [T]. Indukcja magnetyczna może być zdefiniowana wzorem:

$$B = \frac{d\Phi}{dS} \quad (62)$$

W związku z czym strumień magnetyczny można zapisać jako:

$$\Phi = \int_S B dS \quad (63)$$

W przypadku równomiernego pola magnetycznego:

$$B = \frac{\Phi}{S} \quad (64)$$

Prawo Biota i Savarta pozwala określić w dowolnym punkcie przestrzeni indukcję pola magnetycznego obchodząc od odcinka przewodnika, przez który przepływa prąd elektryczny.

$$dB = \frac{\mu dl \sin \alpha}{4\pi r^2} \quad (65)$$

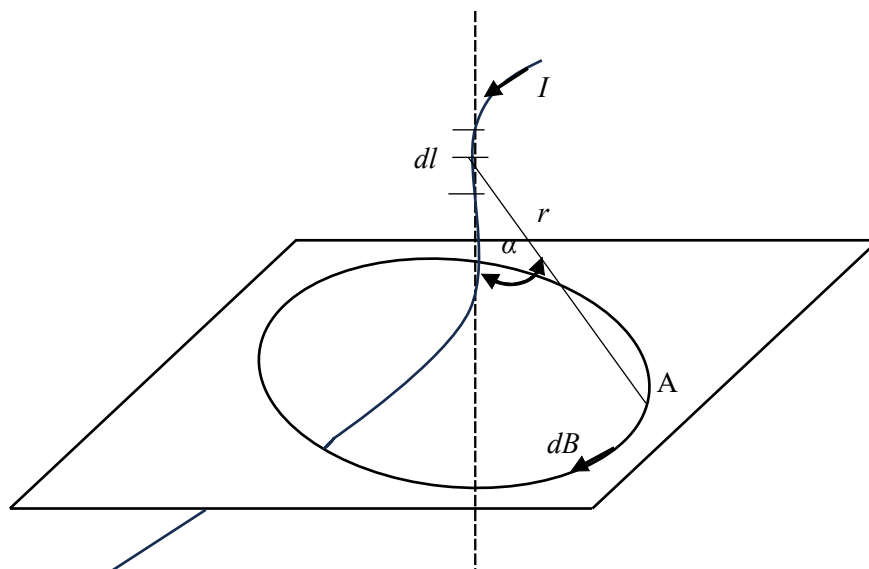
Gdzie zgodnie z rysunkiem 25:

dB – cząstka indukcji pochodząca od odcinka przewodu dl , przez który płynie prąd o natężeniu I

μ - przenikalność magnetyczna bezwzględna środowiska

r – odległość punktu A, w którym oblicza się indukcję magnetyczną od elementu dl

α – kąt między osią elementu dl a odcinkiem łączącym element dl z punktem A



Rys. 25. Ilustracja do prawa Biota i Savara

Rozwiązanie powyższego równania dla przypadku bardzo długiego prostoliniowego przewodu przyjmuje postać:



$$B = \frac{\mu I}{2\pi r} \quad (66)$$

Gdzie:

B – wartość indukcji w ośrodku o przenikalności magnetycznej μ w odległości r od przewodu, przez który płynie prąd I .

Natomiast wartość indukcji w środku kołowego przewodnika w postaci okrągłej ramki przyjmuje postać:

$$B = \frac{\mu I}{2r} \quad (67)$$

W powyższych wzorach można zauważyć, że wartość indukcji magnetycznej zależy od wartości współczynnika opisującego właściwości magnetyczne ośrodka - przenikalności magnetycznej bezwzględnej środowiska - μ .

Wartość przenikalności magnetycznej bezwzględnej określa się jako krotność przenikalności magnetycznej próżni - μ_0 .

$$\mu = \mu_0 \mu_r \quad (68)$$

gdzie:

$\mu_0 = 4\pi \cdot 10^{-7}$ H/m – przenikalność magnetyczna próżni

μ_r – przenikalność magnetyczna względna środowiska, jest ona wartością bezwymiarową.

Pole magnetyczne może być także opisywane inną jednostką wektorową – natężeniem pola magnetycznego- H . Jej wartość zależy od ukształtowania obwodów elektrycznych i prądów płynących przez nie, ale nie zależy od właściwości środowiska. Jednostką natężenia pola magnetycznego jest A/m. Zależność między indukcją magnetyczną B i natężeniem pola magnetycznego można opisać wzorem:

$$H = \frac{B}{\mu} \quad (69)$$

gdzie:

μ - przenikalność magnetyczna [H/m] - wielkość określająca właściwości magnetyczne danego ośrodka

Właściwości magnetyczne materiałów

W ujęciu mikroskopowym każde środowisko można postrzegać jako próżnię, w której znajdują się atomy składające się z jąder i krążących wokół nich elektronów. Elektrony te na skutek swojego ruchu tworzą prądy okrężne, które z kolei wytwarzają własne pole magnetyczne. Tak więc krążąc po swoich orbitach elektrony tworzą elementarne dipole magnetyczne. To



wewnętrzne pole magnetyczne może oddziaływać w różny sposób z zewnętrznym polem magnetycznym, zależnie od rodzaju substancji.

W przypadku tak zwanych diamagnetyków pole magnetyczne pole prądów elementarnych przeciwdziała zewnętrznemu polu magnetycznemu. Na skutek tego wypadkowa indukcja magnetyczna jest mniejsza niż w próżni. Względna przenikalność magnetyczna - μ_r takich materiałów jest mniejsza od jedności:

$$\mu_r < 1 \quad (70)$$

Przykładowymi diamagnetykami są np. miedź, złoto, gazy szlachetne.

W przypadku kiedy pole magnetyczne prądów elementarnych większości atomów współdziała z zewnętrznym polem magnetycznym, to wypadkowa indukcja jest większa niż w próżni. Takie materiały nazywamy paramagnetykami. Ich względna przenikalność magnetyczna jest większa od jedności.

$$\mu_r > 1 \quad (71)$$

Przykładowym paramagnetykiem jest aluminium.

Specyficznymi materiałami są tak zwane ferromagnetyki. W ich przypadku zewnętrzne pole magnetyczne współdziała z polem magnetycznym prądów elementarnych. Indukcja magnetyczna jest wielokrotnie większa niż w próżni. Przenikalność magnetyczna względna dla tych materiałów jest wielokrotnie większa od jedności, może dochodzić nawet do kilku tysięcy.

$$\mu_r \gg 1 \quad (72)$$

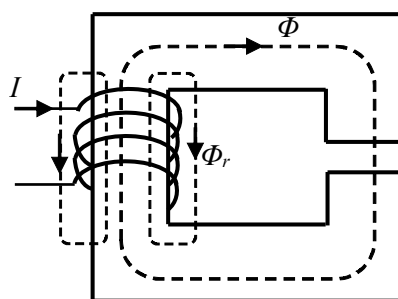
Ponadto nie jest to wartość stała, na jej wartość oprócz rodzaju materiału wpływa także wartość natężenia pola magnetycznego. Zależność μ_r od wartości natężenia pola magnetycznego, nie jest także prostoliniowa i zależy także od samej historii namagnesowywania- charakterystyka ta wykazuje silną histerezę. Przykładowymi ferromagnetykami są żelazo, nikiel, kobalt.

Obwód magnetyczny

Kiedy przez cewkę popłynie prąd elektryczny powstanie strumień magnetyczny o pewnej wartości. Jeżeli w cewce tej umieści się materiał ferromagnetyczny ukształtowany tak, by przynajmniej część pola magnetycznego mogła się w nim zamknąć to powstanie strumień magnetyczny o wartości znacznie większej, niż gdyby strumień magnetyczny zamykał się tylko w powietrzu. Wynika to ze znacznie większej przenikalności magnetycznej ferromagnetyku i wynikającego z niej oporu magnetycznego nazywanego także **reluktancją**. Te właściwości materiałów ferromagnetycznych wykorzystuje się w wielu urządzeniach do tworzenia



odpowiednio dużych i odpowiednio skierowanych strumieni magnetycznych. Zespół elementów służących do wytworzenia strumienia magnetycznego i jego pokierowaniem nazywa się obwodem magnetycznym. Na obwód magnetyczny składają się zazwyczaj elementy służące wytworzeniu i wzbudzeniu pola magnetycznego- mogą to być cewki elektryczne, albo magnesy trwałe, elementy ferromagnetyczne służące pokierowaniu strumienia oraz szczeliny powietrzne. Szczeliny powietrzne są występującymi z konieczności przerwami w materiale ferromagnetycznym. Na rysunku (rysunek 26) poniżej pokazano przykładowy obwód magnetyczny z szczeliną powietrzną. Zaznaczono także, że część strumienia magnetycznego wytworzonego przez cewkę zamyka się w powietrzu, jest to tak zwany strumień rozproszenia.



Rys. 26. Obwód magnetyczny

Dla obwodów magnetycznych duże znaczenie ma tak zwane prawo przepływu. Głosi ono, że wzdłuż drogi zamkniętej suma iloczynów natężenia pola magnetycznego i długości odcinka, wzdłuż którego natężenie pola magnetycznego nie ulega zmianie równa się sumie przepływu prądów obejmowanych przez tę drogę zamkniętą:

$$H_1 l_1 + H_2 l_2 + \dots H_n l_n = I_1 z_1 + I_2 z_2 + \dots I_m z_m \quad (73)$$

gdzie:

H_i – stałe natężenie pola magnetycznego wzdłuż i-tego odcinka

l_i – długość i-tego odcinka

I_j – natężenie prądu przepływającego przez j-tą cewkę

z_j – ilość zwojów w j-tej cewce.

W przypadku cewki, w której występuje równomierne pole magnetyczne, wytworzone przez jedną cewkę;

$$Hl = Iz \quad (74)$$

Iloczyn Iz nazywa się często przepływem, wzbudnością albo siłą magnetomotoryczną i oznacza Θ .



$$\theta = Iz \quad (75)$$

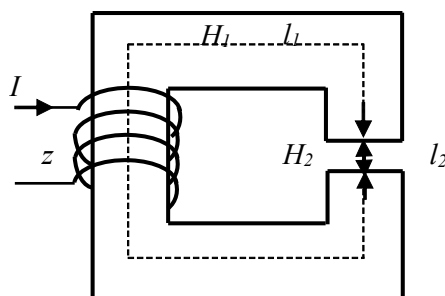
Jednostką są ampéry albo amperozwoje. Prawo przepływu można więc zapisać także w poniższej postaci:

$$\sum_{k=1}^n H_k l_k = \theta \quad (76)$$

Iloczyn $H_k l_k$ nazywa się napięciem magnetycznym.

Dla obwodu magnetycznego z poniższego rysunku (rysunek 27) równanie ułożone na podstawie prawa przepływu przybierze postać:

$$H_1 l_1 + H_2 l_2 = Iz \quad (77)$$



Rys. 27. Obwód magnetyczny ze szczeliną powietrzną

Opór magnetyczny danego odcinka obwodu magnetycznego można określić na podstawie wzoru:

$$R_{mk} = \frac{l_k}{\mu_k S_k} \quad (78)$$

gdzie:

R_{mk} – opór magnetyczny (reluktancja) danego odcinka magnetowodu

l_k – długość danego odcinka

S_k – pole przekroju danego odcinka

μ_k – przenikalność magnetyczna materiału danego odcinka obwodu magnetycznego

Całkowity opór magnetyczny danego obwodu jest sumą oporów magnetycznych poszczególnych odcinków.

Dla obwodu magnetycznego formułuje się prawo Ohma wiążące strumień magnetyczny, siłę magnetomotoryczną oraz opór magnetyczny:

$$\Phi = \frac{\theta}{\sum_{i=1}^n R_{mi}} \quad (79)$$

gdzie:



Φ – strumień magnetyczny [Wb]

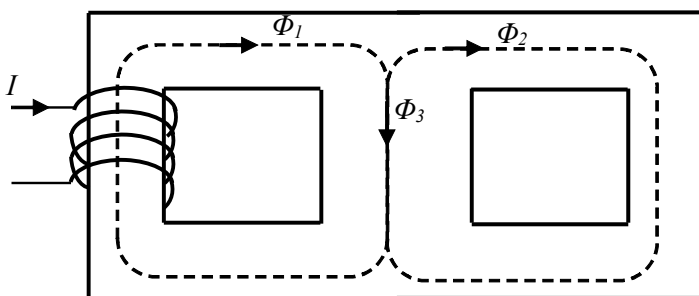
Θ – siła magnetomotoryczna $\Theta = I z$ [A]

R_m – opór magnetyczny [Ω]

Prawa Kirchhoffa dla obwodów magnetycznych

Dla obwodów magnetycznych, podobnie jak dla obwodów elektrycznych, formułuje się prawa Kirchhoffa. Pierwsze prawo głosi, że suma algebraiczna strumieni magnetycznych w węźle obwodu magnetycznego jest równa zero. Przykład takiego węzła widać na rysunku 28.

$$\sum_{k=1}^n \Phi_k = 0 \quad (80)$$



Rys. 28. Obwód magnetyczny z węzłem, w którym rozdziela się strumienie

Drugie głosi, że w oczku obwodu magnetycznego suma spadków napięć magnetycznych jest równa sumie sił magnetomotorycznych:

$$\sum_{k=1}^n \Theta_k = \sum_{k=1}^n H_k \Delta l_k \quad (81)$$

gdzie:

Θ_k - siła magnetomotoryczna

$H_k \Delta l_k$ - spadek napięcia magnetycznego

Indukcyjność własna

Sumę strumieni magnetycznych przenikających poszczególne zwoje cewki nazwa się strumieniem skojarzonym z cewką i oznacza Ψ . Jeżeli strumienie są jednakowe, a liczba zwojów cewki wynosi z , to:

$$\Psi = z\Phi \quad (82)$$

W przypadku strumienia skojarzonego Ψ z cewką o stałej przenikalności wywołany przez prąd I :

$$\Psi = LI \quad (83)$$

Gdzie L określa się jako indukcyjność własną cewki, której jednostką są henry [H].

W przypadku długiej cewki jej indukcyjność własna wynosi:



$$L = \frac{\mu z^2}{l} S \quad (84)$$

gdzie:

μ - przenikalność magnetyczna ośrodka

l - długość cewki

z – ilość zwojów

S - pole powierzchni wewnątrz zwoju

Zjawisko indukcji magnetycznej

Zjawisko indukcji magnetycznej polega na powstaniu w przewodniku lub uzwojeniu siły elektromotorycznej przy jakiegokolwiek zmianie strumienia skojarzonego z tym obwodem. Wartość siły elektromotorycznej e indukowanej w obwodzie wskutek zmian strumienia magnetycznego jest proporcjonalna do szybkości zmian strumienia magnetycznego skojarzonego

$$e = -\frac{d\Psi}{dt} \quad (85)$$

Znak minus jest wyrazem reguły Lenza, mówiącej że zwrot indukowanej siły elektromotorycznej jest taki że wywołany przez nią prąd przeciwstawia się poprzez wytworzone przezeń pole magnetyczne zmianom strumienia magnetycznego które go wywołują. Dokonując podstawienia $\Psi = L i$ otrzymujemy:

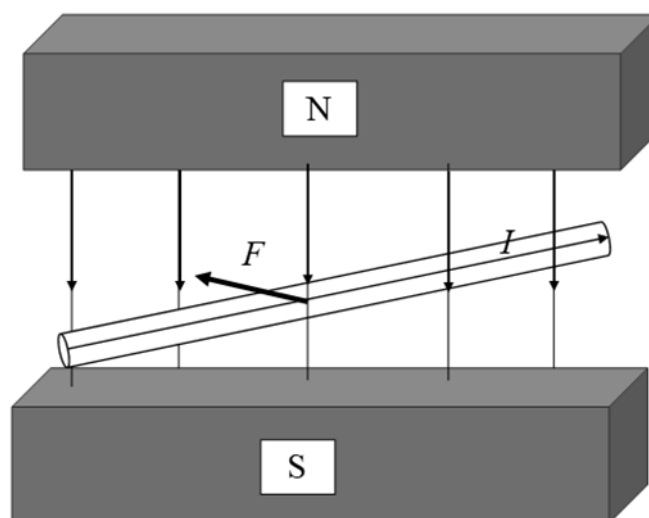
$$e = -L \frac{di}{dt} \quad (86)$$

Siła elektrodynamiczna

W przypadku przewodu znajdującego się w polu magnetycznym, przez który przepływa prąd elektryczny działa na niego tak zwana siła elektrodynamiczna. Jej wartość zgodnie z prawem Ampera zależy od indukcji pola magnetycznego B , wartości natężenia prądu płynącego przez przewód oraz długości przewodu w polu magnetycznym zgodnie z poniższym wzorem:

$$F = B I l \sin(l, B) \quad (87)$$

Gdzie $\sin(l, B)$ jest wartością sinusa kąta pomiędzy osią przewodu, a linią pola magnetycznego przechodzącego przez tę oś. Zwrot siły określa reguła lewej ręki, czyli jeżeli linie pola magnetycznego będą wchodzić w otwartą dłoń, cztery wyprostowane palce będą ułożone wzdłuż przewodu zgodnie ze zwrotem płynącego prądu to wyprostowany kciuk wskaże zwrot siły. Prawo to ilustruje rysunek 29.



Rys. 29. Siła elektrodynamiczna powstająca w przewodzie w polu magnetycznym

4. Układy trójfazowe

Obwód trójfazowy jest to obwód zawierający trzy sprzężone źródła napięcia sinusoidalnego, mające tę samą częstotliwość, których przebiegi czasowe są przesunięte względem siebie w fazie o kąt 120° , czyli $2\pi/3$. Napięcie trójfazowe wytwarza się w prądnicach trójfazowych. Takie prądnice zawierają trzy uzwojenia, w których indukują się trzy siły elektromotoryczne przesunięte między sobą w fazie o 120° . Oznaczając te siły elektromotoryczne jako e_a , e_b , e_c oraz przyjmując początkowe przesunięcie fazowe $\psi=0$, można zapisać:

$$e_a = E_m \sin(\omega t) \quad (88)$$

$$e_b = E_m \sin(\omega t - 2\pi/3) \quad (89)$$

$$e_c = E_m \sin(\omega t - 4\pi/3) \quad (90)$$

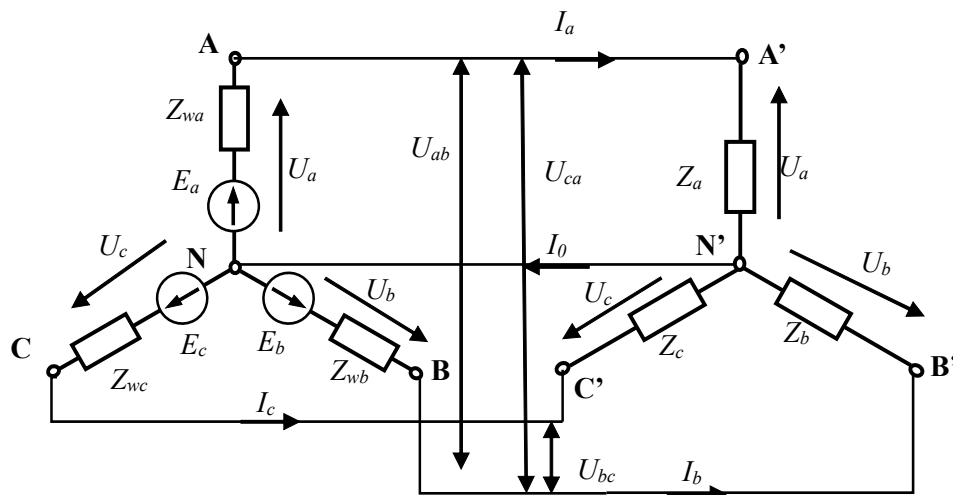
Jeżeli trzy uzwojenia prądnicy są stosowane jako niezależne źródła napięcia to otrzymamy **układ trójfazowy nieskojarzony** - takie rozwiązanie nie znalazło jednak zastosowania w praktyce. W praktyce stosuje się układy skojarzone otrzymane przez odpowiednie połączenie uzwojeń. Wyróżnia się:

- układ trójfazowy połączony w gwiazdę
- układ trójfazowy połączony w trójkąt

Obwody skojarzone w gwiazdę

Układ, w którym końce faz prądnic lub odbiornika będą połączone razem nazywamy połączonym w gwiazdę. W układzie takim wspólny węzeł nazywa się punktem neutralnym.

Jeżeli punkt neutralny prądnicy jest połączony z punktem neutralnym odbiornika to taki układ nazywa się trójfazowym czteroprzewodowym, w przeciwnym wypadku jest to układ trójfazowy trzyprzewodowy. Przewody łączące początki faz prądnicy z odbiornikiem nazywa się przewodami fazowymi. Rysunek 30 przedstawia obwód skojarzony w gwiazdę.



Rys. 30. Obwód skojarzony w gwiazdę

W układzie tym występują napięcia:

- międzyfazowe (albo międzyprzewodowe) - występujące między przewodami fazowymi: U_{ab}, U_{bc}, U_{ca}
- napięcia fazowe między przewodami fazowymi a neutralnym: U_a, U_b, U_c

W przypadku układu symetrycznego, cechującego się równymi modułami sił elektromotorycznych na każdej gałęzi przesuniętymi między sobą w fazie o 120° oraz równymi impedancjami. Wartości modułów napięć międzyprzewodowych są większe $\sqrt{3}$ -krotnie od napięć fazowych. Ponadto prąd płynący linią łączącą punkty neutralne wynosi 0. Dla układu symetrycznego:

$$I = I_f \quad (91)$$

$$U = \sqrt{3}U_f \quad (92)$$

Obwody skojarzone w trójkąt

W układzie, w którym koniec jednej fazy prądnicy jest połączony z początkiem drugiej nazywamy układem trójfazowym połączonym w trójkąt. W układzie trójkątowym napięcia międzyfazowe są równe napięciom fazowym prądnicy i odbiornika. Obwód skojarzony w trójkąt przedstawia rysunek 31. Występują w tym układzie dwa rodzaje prądów:

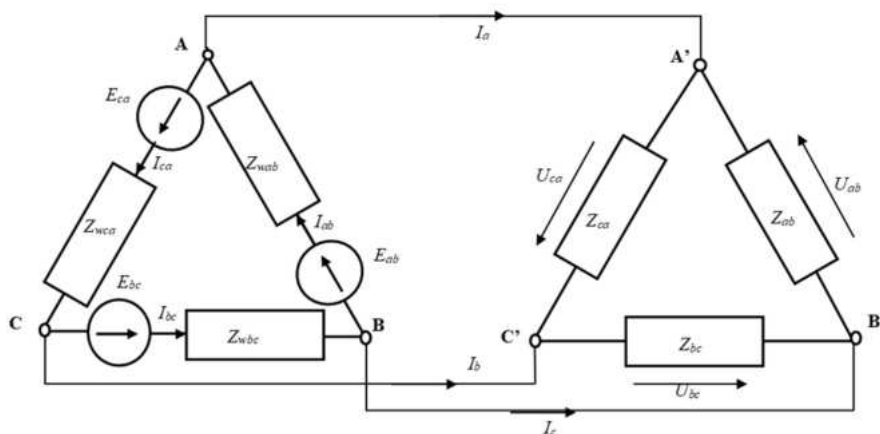
- Prądy przewodowe płynące między prądnicą i odbiornikiem



$$I_a = I_b = I_c = I \quad (93)$$

- Prądy fazowe płynące przez fazy prądnicy lub odbiornika

$$I_{ab} = I_{bc} = I_{ca} = I_f \quad (94)$$



Rys. 31. Obwód skojarzony w trójkąt

Dla układu symetrycznego:

$$U_{ab} = U_{bc} = U_{ca} = U = U_f \quad (95)$$

$$I = \sqrt{3}I_f \quad (96)$$

Moc w obwodach trójfazowych

Moc czynna wytwarzana przez źródło trójfazowe jest równa sumie mocy poszczególnych faz

$$P = P_a + P_b + P_c = U_a I_a \cos \varphi_a + U_b I_b \cos \varphi_b + U_c I_c \cos \varphi_c \quad (105)$$

Dla układu symetrycznego

$$P = 3U_f I_f \cos \varphi_f \quad (97)$$

Gdzie:

U_f - napięcie fazowe

I_f - prąd

Ponieważ łatwiej jest zmierzyć prądy przewodowe i napięcie międzyfazowe:

$$P = \sqrt{3}UI \cos \varphi_f \quad (98)$$

Gdzie:

U - napięcie międzyfazowe



I - prąd przewodowy

φ - kąt przesunięcia między napięciem fazowym a prądem fazowym

Moc bierną i pozorną określają podobne wzory:

$$Q = \sqrt{3}UI \sin \varphi_f \quad (99)$$

$$S = \sqrt{3}UI \quad (100)$$

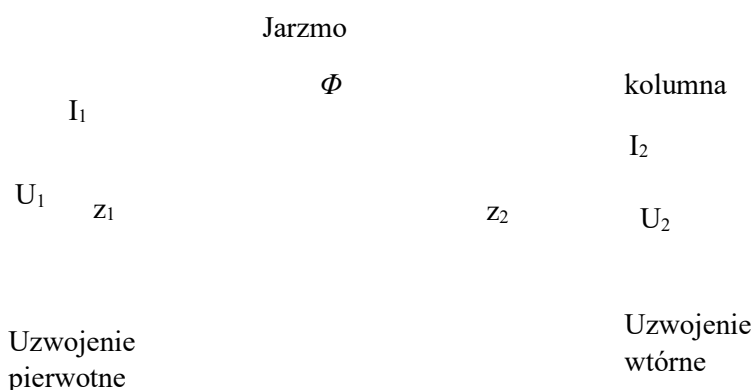
5. Wybrane maszyny i urządzenia elektryczne

Transformator

Transformator jest urządzeniem elektrycznym służącym do przekazywania energii elektrycznej z jednego obwodu do drugiego za pośrednictwem pola magnetycznego. W zależności od środowiska, w jakim zamyka się wytworzony przez uzwojenie główny strumień magnetyczny, rozróżnia się transformatory powietrzne i transformatory z rdzeniem ferromagnetycznym. W elektroenergetyce transformatory stosuje się do przetwarzania energii prądu przemiennego o jednym napięciu na energię prądu przemiennego o innym napięciu.

Transformator taki jak na rysunku 32, składa się z rdzenia, w którym zamyka się strumień magnetyczny oraz uzwojeń.

Uzwojenie, do którego doprowadza się energię nazywa się pierwotnym, natomiast uzwojenie, którym odprowadza się energię określa się wtórnym.



Rys. 32. Budowa transformatora

Ze względu na liczbę uzwojeń rozróżnia się transformatory dwuuzwojeniowe, wielouzwojeniowe oraz jednouzwojeniowe, czyli tzw. autotransformatory.



W przypadku typowego transformatora na rdzeniu z materiału ferromagnetycznego, który stanowi zamkniętą drogę dla strumienia magnetycznego nawinięte są dwa uzwojenia odizolowane elektrycznie od siebie i rdzenia. Uzwojenie pierwotne posiada z_1 zwojów, a wtórne z_2 zwojów.

Podczas pracy transformatora do uzwojenia pierwotnego podłączane jest napięcie sinusoidalne o napięciu u_1 . To napięcie powoduje przepływ przez uzwojenie pierwotne prądu sinusoidalnego, który natomiast powoduje powstanie w rdzeniu ferromagnetycznym zmiennego strumienia magnetycznego Φ . Ten Zmienny strumień magnetyczny sprzęgający oba uzwojenia indukuje w nich siły elektromotoryczne.

Dla prądu sinusoidalnego wartości skuteczne tych sił elektromotorycznych wynoszą:

$$E_1 = \frac{2\pi f z_1 \Phi_m}{\sqrt{2}} = 4,44 z_1 f \Phi_m \quad (101)$$

$$E_2 = \frac{2\pi f z_2 \Phi_m}{\sqrt{2}} = 4,44 z_2 f \Phi_m \quad (102)$$

A ich wzajemny stosunek, nazywany przekładnią transformatora:

$$\vartheta = \frac{E_1}{E_2} = \frac{z_1}{z_2} \quad (103)$$

Transformator może znajdować się w trzech różnych stanach pracy.

Stan jałowy – stan, w którym uzwojenia pierwotne połączone jest ze źródłem napięcia sinusoidalnego, a uzwojenie wtórne jest otwarte. Dla tego stanu można przyjąć:

$$\underline{U}_1 = \underline{E}_1 \quad (104)$$

$$\vartheta = \frac{E_1}{E_2} = \frac{U_1}{U_2} \quad (105)$$

Stan obciążenia – typowy stan pracy, w którym do uzwojenia pierwotnego jest podłączone napięcie sinusoidalne a do wtórnego odbiornik

Stan zwarcia – jest to stan, w którym zaciski wtórne są zwarte, a zasilanie pierwotne jest zasilane napięciem. W praktyce jest to stan awaryjny. Doprowadza się do niego celowo jedynie przy pomiarach rezystancji i reaktancji uzwojeń. Ponieważ rezystancje uzwojeń i reaktancje rozproszeniowe są niskie przez uzwojenia w tym stanie popłynie duży prąd, co w wyniku przemiany energii elektrycznej w ciepło w krótkim czasie może doprowadzić do stopienia izolacji uzwojeń.

Transformatory znajdują zastosowanie przykładowo:

- W elektrowniach transformują energię podwyższając napięcie, co umożliwia zmniejszenie strat na przesyle przez linie energetyczne. W stacjach transformatorowych



służą natomiast do obniżenia napięcia do poziomu umożliwiającego podłączenie odbiorców.

- Stosowane są przy pomiarach wysokich napięć, albo dużych prądów sinusoidalnych.
- Stosuje się też transformatory spawalnicze, prostownikowe, autotransformatory.

W trakcie pracy transformatora występują w nim straty energii. Są to głównie straty tak zwane histerezy oraz związane z prądami wirowymi. Straty histerezy są związane z przemagnesowywaniem materiału rdzenia, z tego powodu stosuje się materiały o małej pętli histerezy, czyli magnetycznie miękkie. Straty związane z powstawaniem prądów wirowych są natomiast redukowane przez wykonywanie transformatorów z cienkich blaszek izolowanych od siebie. Blaszki te wykonywane są z materiałów o dużej rezystywności, typowo z stali krzemowej.

Prądnica i silnik prądu stałego

Wziętość maszyn prądu stałego może pracować jako silnik albo generator przy tej samej budowie. Taka maszyna składa się z trzech podstawowych elementów:

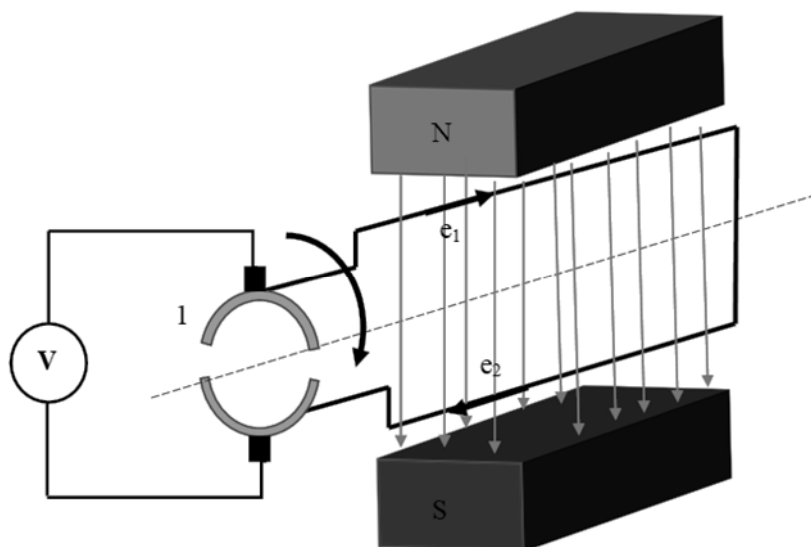
- magneśnicy
- wirnika
- komutatora

Magneśnica jest nieruchomym stojanem składającym się z jarzma oraz biegunów magnetycznych, którymi są elektromagnesy albo magnesy stałe. W przypadku elektromagnesów ich uzwojenie nazywa się uzwojeniem wzbudzenia albo wzbudnikiem.

Wirnik jest obracającym się element, z rowkami, w których znajduje się uzwojenie wirnika nazywane twornikiem.

Komutator jest częścią maszyny umożliwiającą zmianę kierunku przepływu prądu. Łączy uzwojenie twornika z zaciskami wyjściowymi maszyny. W przypadku komutatora mechanicznego składa się z półpięścieni połączonych z uzwojeniem twornika, ślizgających się po przewodzących szczotkach, które łączą się z zaciskami wyjściowymi maszyny.

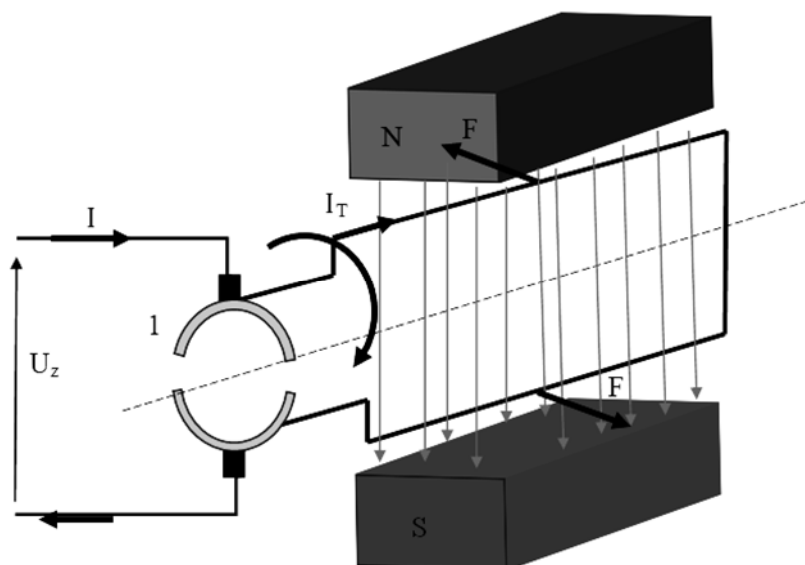
Uproszczony schemat przedstawiający działanie prądnicy prądu stałego przedstawia poniższy rysunek 33. Umieszczone w obudowie-stojanie elektromagnesy (albo magnesy stałe) wytwarzają stałe pole magnetyczne. W polu tym umieszczona jest obrotowa ramka utworzona z dwóch równoległych przewodów.



Rys. 33. Ilustracja działania prądnicy prądu stałego, 1 - komutator

Obrót ramki powoduje, że boki ramki przecinają linię sił pola magnetycznego, na skutek indukują się w nich siły elektromotoryczne e_1 i e_2 , a wypadkowa wartość siły elektromotorycznej (SEM) ma wartość $e_w = e_1 + e_2$. Ponieważ kąt przecinania linii pola magnetycznego zmienia się, dlatego indukowana SEM zmienia się od 0 V w położeniu poziomym ramki do maksymalnej w położeniu pionowym. Indukowana SEM jest odprowadzana do zacisków wyjściowych za pośrednictwem komutatora, współpracującego z przewodzącymi szczotkami. Komutator pełni także rolę mechanicznego prostownika powodując, że na wyjściu otrzymuje się napięcie jednokierunkowe tętniące. W rzeczywistym generatorze na wirującym wale rozmieszczonych byłoby więcej takich ramek, dzięki czemu minimalizowany byłby efekt tętnienia.

W przypadku silnika prądu stałego z magnesami trwałymi jego budowa jest taka sama jak w przypadku prądnicy, jednakże do zacisków wyjściowych jest podłączone źródło napięcie stałego doprowadzające prąd, na skutek czego wirnik zaczyna się obracać. Ilustrację działania silnika prądu stałego przedstawia poniższy rysunek 34.

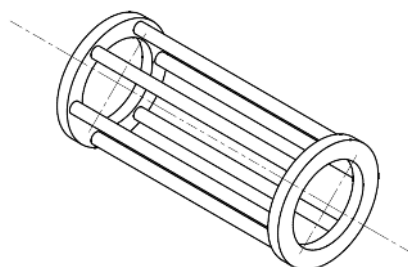


Rys. 34. Zasada działania silnika prądu stałego z magnesami trwałymi

W przewodzie umieszczonym w polu magnetycznym, przez który płynie prąd powstaje siła elektrodynamiczna. W przypadku ramki, przez którą płynie prąd, powstają dwie siły o zwrotach zgodnych z zasadą lewej ręki, tworzące moment obrotowy obracający nią. Jego wartość zależy od chwilowego położenia kąтового ramki. Kiedy płaszczyzna ramki pokrywa się z kierunkiem sił pola magnetycznego jest on największy, kiedy jest ona prostopadłą wynosi on zero. Dzięki połączeniu przez komutator, co pół obrotu zmienia się kierunek przepływu prądu w ramce, dzięki czemu wirnik obraca się w jednym kierunku. W przypadku prawdziwych silników na wale umieszczanych jest wiele uzwojeń, dzięki czemu silnik utrzymuje równy moment i prędkość w całym zakresie obrotu.

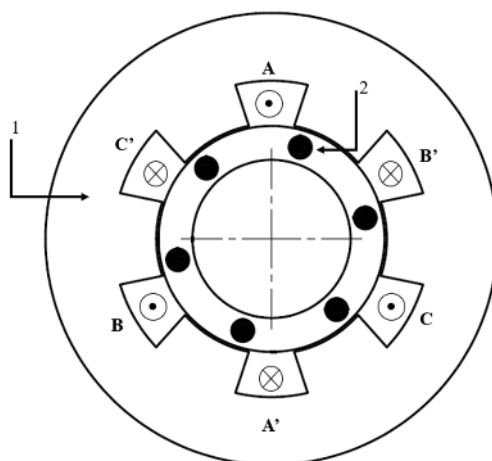
Silnik prądu przemiennego asynchroniczny trójfazowy

Asynchroniczny trójfazowy silnik prądu przemiennego jest jednym z najbardziej rozpowszechnionych silników w przemyśle. Silnik ten składa się z stojana i wirnika. Na wirniku umieszczone jest uzwojenie, które nie jest podłączane do zasilania z zewnątrz. Uzwojenie może mieć formę pierścieniowego albo klatkowego. Klatkowe rozwiązanie jest obecnie powszechniejsze. Określenie klatkowe, pochodzi od charakterystycznego kształtu tego uzwojenia przypominającego klatkę. Składa się ono bowiem z prętów połączonych ze sobą pierścieniami przy podstawach. Rysunek 35 przedstawia takie uzwojenie.



Rys. 35. Uzwojenie klatkowe silnika asynchronicznego

W stojanie, wewnątrz którego obraca się wirnik, znajdują się żłobki, w których umieszcza się uzwojenia. Każde z tych uzwojeń ma taką samą liczbę zwojów i są one przesunięte względem siebie co 120° . Schemat stojana przedstawia rysunek 36. Poszczególne uzwojenia są zasilane przez osobne fazy prądu trójfazowego.



Rys. 36. Stojan silnika asynchronicznego trójfazowego z umieszczonym w środku wirnikiem,
1 - stojan, 2 - wirnik, AA', BB', CC' - uzwojenia stojana

Jeżeli końcówki uzwojeń zostaną połączone w gwiazdę lub trójkąt i podłączy się je do sieci trójfazowej symetrycznej, to przez uzwojenia popłyną prądy i_a , i_b , i_c . Każda z faz stojana przy przepływie tych prądów wytworzy oscylujące pole magnetyczne. Wypadkowe pole magnetyczne będzie natomiast sumą geometryczną wektorów indukcji magnetycznej poszczególnych cewek. Powstający wektor wypadkowy będzie się obracał, a czas jego pełnego obrotu będzie równy okresowi prądu trójfazowego zasilającego silnik.

Powstające wirujące pole magnetyczne będzie przecinać pręty początkowo nieruchomego wirnika indukuje w nich siłę elektromotoryczną. Pod wpływem tej zaindukowanej siły elektromotorycznej w prętach wirnika zaczniesz płynąć prąd o znacznej wartości. Ponieważ pręty, przez które płynie prąd, znajdują się w polu magnetycznym zaczniesz działać na nie siła elektrodynamiczna. Na skutek powstania tych sił powstaje moment, powodujący obrót wirnika.



Wirnik silnika obraca się z kierunkiem zgodnym z kierunkiem wirowania pola magnetycznego, ale z niższą prędkością, gdyż w przypadku zrównania się prędkości wirnika i prędkości wirowania pola magnetycznego ustałoby indukowanie się prądu w prętach, a tym samym nastąpiłby zanik momentu obrotowego. Opóźnienie prędkości obrotowej silnika asynchronicznego w stosunku do prędkości synchronicznej nazywa się poślizgiem - s . Jest zdefiniowany poniższym wzorem:

$$s = \frac{n_1 - n}{n_1} \quad (106)$$

gdzie:

s - poślizg

n_1 - prędkość synchroniczna

n - prędkość obrotowa silnika

Prędkość synchroniczną określa wzór:

$$n_1 = \frac{60f_1}{p} \left[\frac{\text{obr}}{\text{min}} \right] \quad (107)$$

gdzie:

n_1 - prędkość synchroniczna

f_1 - częstotliwość prądu zasilającego [Hz]

p - liczba par biegunów

Ostatecznie wzór na prędkość obrotową silnika asynchronicznego przyjmuje postać:

$$n = n_1(1 - s) = \frac{60f}{p}(1 - s) \quad (108)$$

Jak można zauważyć prędkość silnika obrotowego zależy od poślizgu, liczby par biegunów oraz częstotliwości prądu zasilającego. Współcześnie sterowanie prędkością obrotową silników asynchronicznych trójfazowych odbywa się głównie poprzez zmianę częstotliwości, za pomocą urządzeń zwanych przemiennikami częstotliwości.

Uzwojenia silnika mogą być połączone mogą być połączone w gwiazdę albo trójkąt. Sposób połączenia z odpowiadającym im standardowym sposobem oznaczenia zacisków przedstawiono w tabeli nr 3. Przy połączeniu w trójkąt silnik może osiągnąć większy moment, ale przy poborze większego prądu, niż w przypadku połączenia w gwiazdę. Pobór szczególnie dużego prądu występuje przy rozruchu, z tego powodu niekiedy korzystniej jest załączać silnik przy połączeniu w gwiazdę, a następnie przełączyć na połączenie w trójkąt. Z tego powodu często znajdują zastosowanie przełączniki gwiazda-trójkąt.



Tabela 3. Połączenie w gwiazdę i trójkąt uzwojeń silnika asynchronicznego trójfazowego

gwiazda		
trójkąt		

6. Elektronika

Półprzewodniki

Półprzewodnikami są materiały, których rezystywność jest znacznie większa niż rezystywność przewodników, ale znacznie mniejsza niż rezystywność dielektryków. Podstawowymi materiałami zaliczającymi się do tej grupy, stosowanymi w elektronice są krzem i german.

W przypadku tych atomów pasmo walencyjne jest oddzielone od pasma przewodnictwa przerwą energetyczną. Szerokość tej przerwy jest tak duża, że w normalnych warunkach tylko nieliczne elektrony walencyjne atomów krzemu lub germanu mogą zgromadzić niezbędną energię do pokonania przerwy energetycznej i przejść do pasma przewodnictwa.

Półprzewodniki jednorodne posiadają niewiele elektronów swobodnych, co objawia się dużą rezystancją właściwą materiału półprzewodnikowego. Z tego powodu stosuje się domieszkowanie.

Domieszkowanie polega na wprowadzeniu do struktury kryształu dodatkowych atomów pierwiastka, który nie wchodzi w skład półprzewodnika naturalnego. Ponieważ w wiązaniach



wewnątrzatomowych bierze udział ustalona liczba elektronów zamiana któregoś z jonów na atom domieszki powoduje wystąpienie nadmiaru lub niedoboru elektronów.

Wprowadzenie domieszki produkującej niedobór elektronów (w stosunku do ilości ci niezbędnej do stworzenia wiązań) powoduje powstanie półprzewodnika typu p, natomiast wprowadzenie domieszki produkującej nadmiar elektronów (w stosunku do ilości niezbędnej do stworzenia wiązań) powoduje powstanie półprzewodnika typu n.

Złączeniem p-n nazywane jest złącze utworzone przez dwa półprzewodniki niesamoistne o różnych typach przewodnictwa: p oraz n. W obszarze typu n występują nośniki większościowe ujemne (elektrony) oraz unieruchomione w siatce krystalicznej atomy domieszek (donory). Analogicznie, w obszarze typu p, występują nośniki większościowe o ładunku elektrycznym dodatnim (dziury) oraz atomy domieszek (akceptory).

W stanie niespolaryzowanym, czyli gdy z zewnątrz nie jest przyłożone żadne pole elektryczne, w pobliżu styku obszarów p i n swobodne nośniki większościowe przemieszczają się, co spowodowane jest różnicą koncentracji nośników. Gdy elektrony przemieszczają się do obszaru typu p, natomiast dziury do obszaru typu n dochodzi do rekombinacji z nośnikami większościowymi, które nie przeszły na drugą stronę złącza.

W efekcie rekombinacji w pobliżu złącza powstaje tak zwana warstwa zubożona, nazywana także warstwą zaporową. Nieruchomy ładunek dodatni po stronie N hamuje przepływ dziur z obszaru P, natomiast ładunek ujemny po stronie P hamuje przepływ elektronów z obszaru N, na skutek czego przepływ nośników większościowych praktycznie ustaje. Powstaje tak zwana bariera potencjału.

Bariera potencjałów uniemożliwia przepływ ładunków przez złącze p-n. Jej działanie może zostać zneutralizowane poprzez dołączenie napięcia zewnętrznego o wartości równej napięciu bariery potencjałów i przeciwnym kierunku polaryzacji. Jeżeli napięcie zewnętrzne przekroczy wartość progową, wytworzone pole elektryczne dostarczy swobodnym nośnikom ładunków energię wystarczającą do przekroczenia złącza i przez złącze p-n popłynie prąd elektryczny utworzony przez nośniki większościowe.

Jeżeli natomiast do złącza podłączone zostanie napięcie zewnętrzne o polaryzacji zgodnej z napięciem bariery potencjałów wytworzone zewnętrzne pole elektryczne wzmocni odpychające działanie bariery potencjałów na nośniki ładunków swobodnych. Wówczas przez złącze będzie płynął tylko niewielki prąd utworzony przez mniejszościowe nośniki ładunku



Elementy oparte na półprzewodnikach jednorodnych

W elektronice najczęściej stosuje się elementy, w których zastosowano różnego rodzaju złącza półprzewodników typu p i n. Jednak występują także elementy oparte na jednorodnych półprzewodnikach.

Termistory – są to elementy półprzewodnikowe, których rezystancja zależy od temperatury.

Występują w kilku typach:

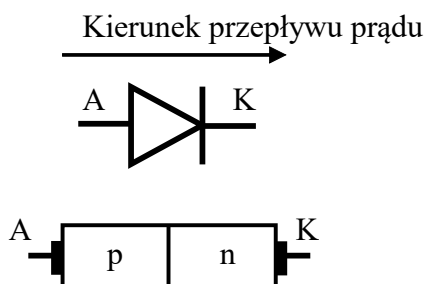
- PTC (*positive temperature coefficient*) – ich rezystancja zwiększa się wraz ze wzrostem temperatury
- Przez złącze będzie płynął tylko niewielki prąd utworzony przez mniejszościowe nośniki ładunku
- Przez złącze będzie płynął tylko niewielki prąd utworzony przez mniejszościowe nośniki ładunku

Warystory – są elementami o rezystancji zmieniającej się nieliniowo pod wpływem występującego na nim napięcia

Hallotrony – są elementami półprzewodnikowymi wytwarzającymi sygnał napięciowy pod wpływem zewnętrznego pola magnetycznego. Znalazły zastosowanie w czujnikach wykrywających pole magnetyczne.

Diody

Najprostszym elementem elektronicznym wykorzystującym złącze p-n jest dioda półprzewodnikowa. Ma właściwość jednokierunkowego przewodzenia prądu i po włączeniu do obwodu zasilanego napięciem przemiennym działa jak prostownik. Symbol diody, oraz schemat jej budowy przedstawia rysunek 37.



Rys. 37 Symbol diody i jej schematyczna budowa

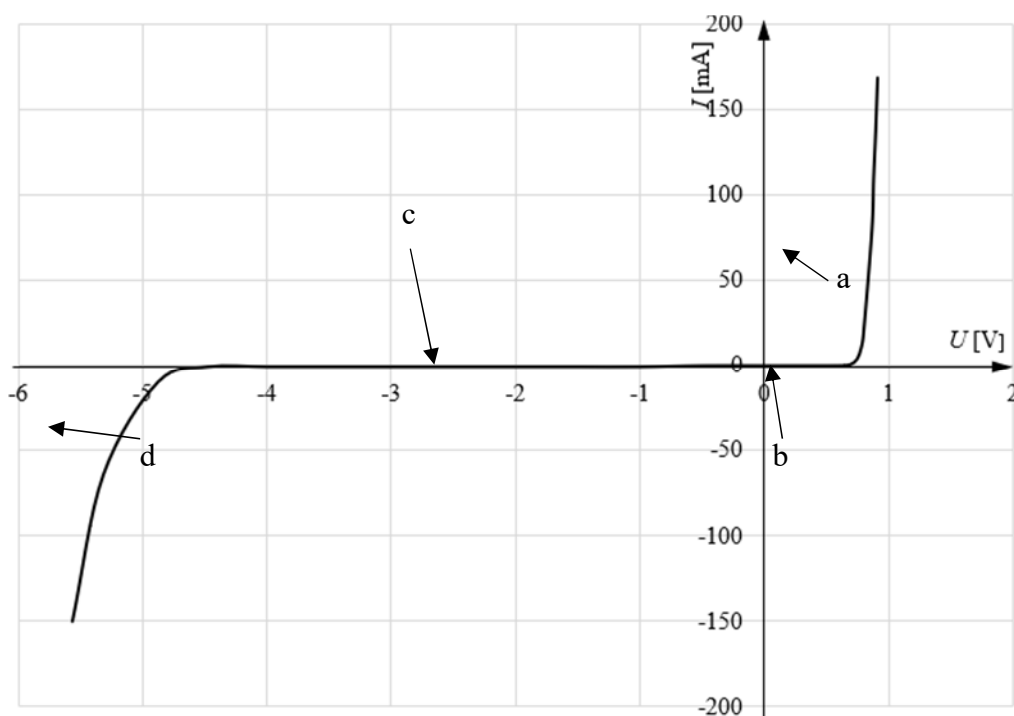


Jeżeli dioda zostanie w obwodzie spolaryzowana w kierunku przewodzenia, to popłynie przez nią duży prąd utworzony przez nośniki większościowe, nazywany prądem przewodzenia. W takim przypadku rezystancja diody jest niewielka. W celu spolaryzowania diody w kierunku przewodzenia łączy się biegun dodatni źródła napięcia z obszarem o przewodnictwie typu p, a biegun ujemny z obszarem o przewodnictwie typu n.

W przypadku diody spolaryzowanej w kierunku zaporowym, pod wpływem napięcia zewnętrznego po obu stronach złącza p-n powstają obszary o zubożonej koncentracji nośników ładunków, czyli obszar izolacyjny nazywany warstwą zaporową. Nie jest on jedną idealnym izolatorem, w związku z czym płynie przez niego niewielki prąd wsteczny (z obszaru p do n) nazywany prądem zaporowym. W takim przypadku rezystancja diody jest duża. Polaryzacja w kierunku zaporowym występuje przy połączeniu bieguna dodatniego źródła napięciowego z obszarem n, a bieguna ujemnego z obszarem p. Jeżeli napięcie polaryzacji wzrośnie nadmiernie, wytrzymałość warstwy zaporowej okaże się za małą i nastąpi jej przebicie i lawinowy wzrost prądu wstecznego, powodując uszkodzenie diody i utratę jej właściwości.

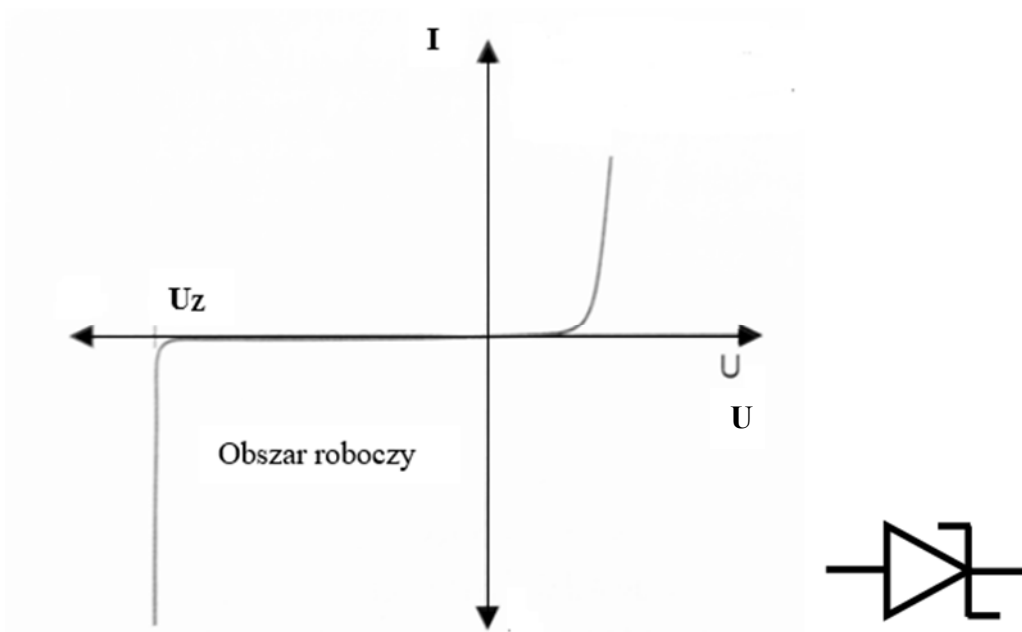
Charakterystyka prądowo-napięciowa diody przedstawiona na rysunku 38 osiada kilka charakterystycznych obszarów.

- Obszar przewodzenia (a) – obejmuje liniowy i gwałtowny wzrost prądu przewodzenia przy niewielkim wzroście napięcia przewodzenia
- Obszar wykładniczego przebiegu (b) – obejmuje początkowy zakres przewodzenia
- Obszar zaporowy (c) – obejmuje charakterystykę przy polaryzacji zaporowej, w całym zakresie prąd zaporowy ma stałą niską wartość
- Obszar przebicia (d) – po przekroczeniu dopuszczalnej maksymalnej wartości napięcia zaporowego (napięcia przebicia) następuje gwałtowny wzrost prądu zaporowego I_R . Następuje przebicie warstwy izolacyjnej



Rys. 38. Charakterystyka prądowo-napięciowa diody

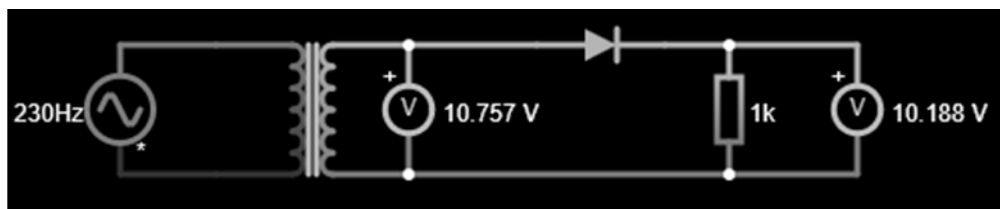
Specyficznym rodzajem diod są diody Zenera. Ich charakterystyka w kierunku przewodzenia jest taka sama jak zwykłych diod, nie jest ona jednak w tym kierunku wykorzystywana. Natomiast w kierunku zaporowym, przy specyficznym dla danego typu diody napięciu, tzw. napięciu Zenera, charakterystyka diody jest stromo opadająca. Tą właściwość wykorzystuje się w niektórych układach zasilających do budowy stabilizatorów napięcia. Charakterystykę prądowo-napięciową diody Zenera oraz jej symbol przedstawia rysunek 39.



Rys. 39. Charakterystyka prądowo-napięciowa diody Zenera oraz jej symbol

Zastosowanie diod w układach prostowniczych zasilaczy

Układy prostownicze, są to układy służące do zamiany prądu przemiennego, na prąd stały albo zbliżony do stałego. Generalnie mogą one być zbudowane na bazie elementów niesterowanych, takich jak diody, albo sterowanych jak na przykład tranzystory. Poniżej zostaną zaprezentowane podstawowe układy zbudowane na bazie diod. Najprostszy prostownik składa się z jednej diody. Na rysunku poniższym (rysunek 40) zastosowano najpierw transformator obniżający wartość napięcia przemiennego, a następnie wpięto do obwodu jedną diodę. Na podłączonym obciążeniu w formie rezystora uzyskano przebieg napięcia prądu w formie pulsacyjnej. Otrzymany przebieg zaprezentowano na rysunku 41.

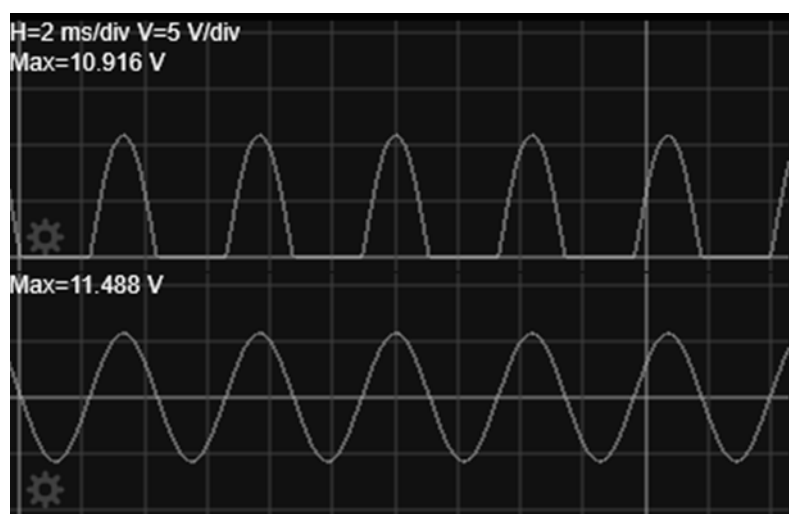


Rys. 40. Schemat prostownika jednopółkowego

Otrzymany przebieg cechują przerwy w zasilaniu o długości połowy okresu prądu przemiennego wejściowego. Takie przerwy są często niepożądane, dlatego stosuje się układy

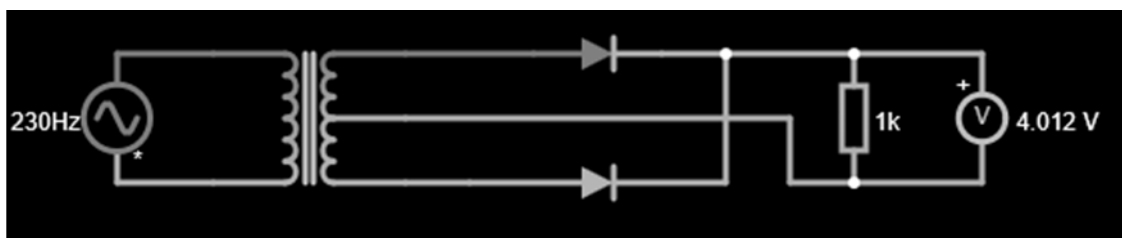


prostownicze dwupółkowe pozwalające efektywnie wykorzystać cały przebieg sinusoidalny.



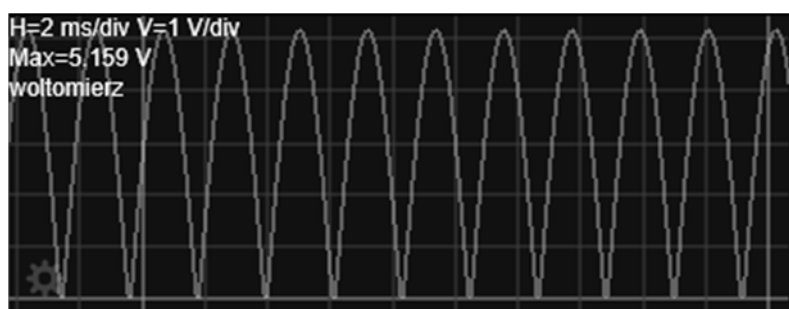
Rys. 41. Przebieg napięcia na wyjściu prostownika jednopołówkowego, oraz wejściowy

Jednym z rozwiązań prostownika dwupółkowego jest prostownik zbudowany na bazie transformatora z dzielonym uzwojeniem wtórnym i dwiema diodami. Jest on zaprezentowany na rysunku 42.



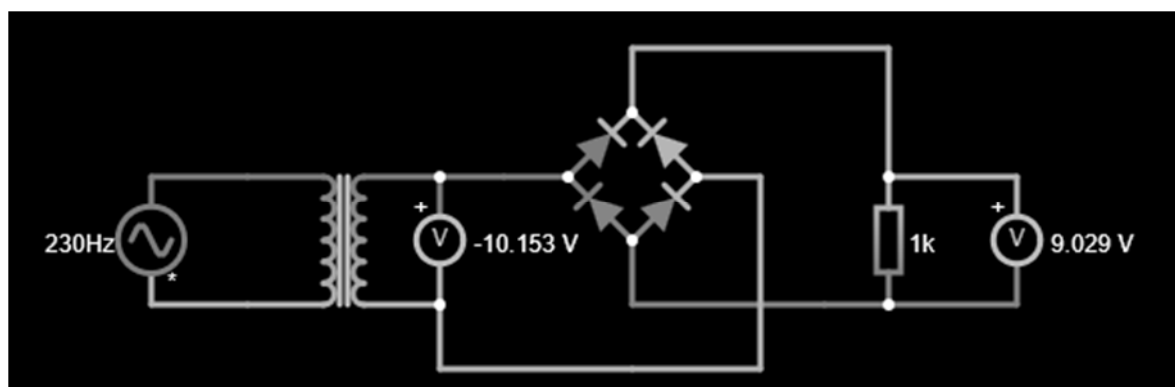
Rys. 42. Prostownik dwupółkowy z transformatorem o dzielonym uzwojeniem

W takim prostowniku wykorzystywany jest transformator z dodatkowym wyprowadzeniem, dzielącym uzwojenie wtórne na dwa o takiej samej liczbie zwojów. Wadą takiego prostownika jest konieczność stosowania odpowiedniego transformatora, ponadto w danej chwili jest używana tylko połowa uzwojenia wtórnego. Powoduje to, że napięcie na wyjściach takiego transformatora jest o połowę niższe niż na wyjściu transformatora o undzielonym uzwojeniu o takiej samej przekładni. Uzyskany przebieg przedstawiono na rysunku 43. Widać na nim, że udało się uzyskać prąd pulsacyjny, w którym efektywnie wykorzystano cały przebieg sinusoidalny prądu przemiennego wejściowego.



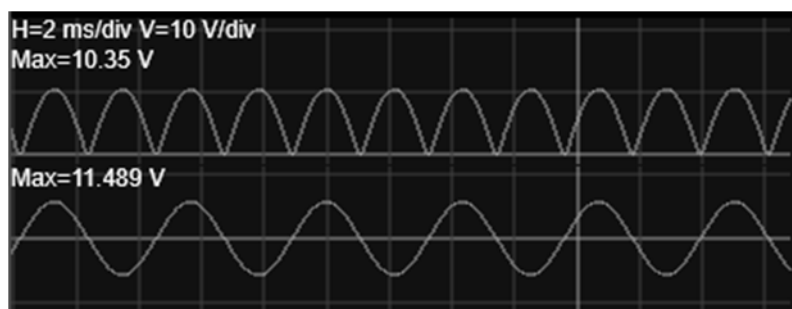
Rys. 43. Przebieg napięcia na wyjściu prostownika dwupołówkowego

Obecnie najbardziej rozpowszechnioną formą prostownika dwupołówkowego jest prostownik zbudowany na bazie 4 diod prostowniczych w układzie tak zwanego mostka Graetza. Prezentuje go rysunek 44.



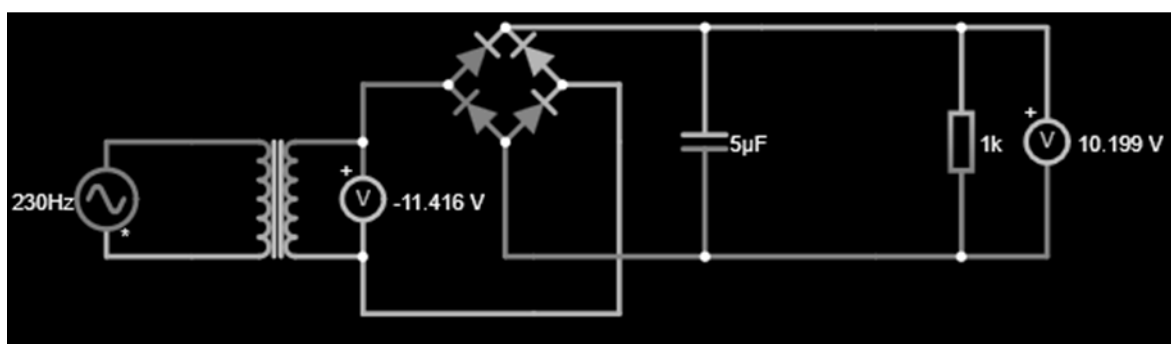
Rys. 44. Prostownik na bazie układu mostka Graetza

Prostownik ten pozwala skutecznie wykorzystać całą sinusoidę, tak jak widać na rysunku 45. Jedyne straty napięcia, jakie występują, to spadki napięcia na dwóch diodach. W danej chwili prąd bowiem płynie przez dwie diody.

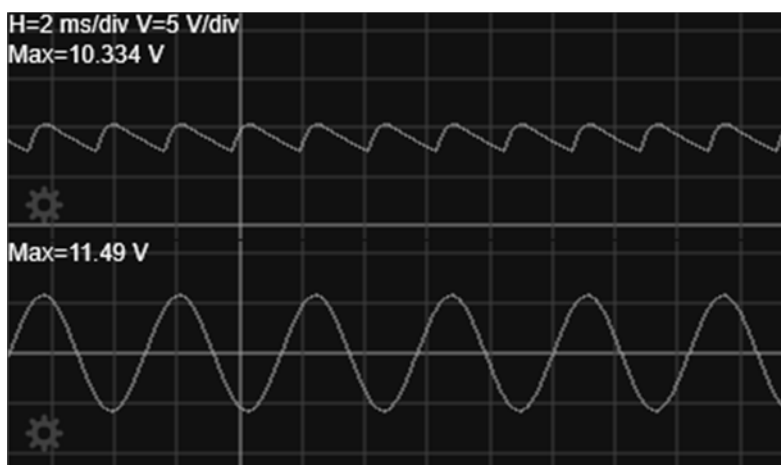


Rys. 45. Napięcie na wyjściu prostownika na bazie mostka Graetza oraz przebieg napięcia wejściowego

Ponieważ na wyjściu prostownika otrzymywany jest przebieg mocno pulsujący, konieczne jest jego wygładzenie. Służą do tego układy filtrujące. Mogą być one budowane na bazie elementów magazynujących energię, takich jak cewki i kondensatory. Przykład zastosowanego układu filtrującego w postaci wpiętego równolegle kondensatora przedstawia rysunek 46. Otrzymany natomiast przebieg przedstawia rysunek 47. Widać na nim znacznie mniejszą pulsację w porównaniu do przebiegu na samym prostowniku z rysunku 45. Ta pulsacja może być jeszcze zmniejszona przez dobranie większego kondensatora.



Rys. 46. Prostownik z układem filtrującym w postaci kondensatora

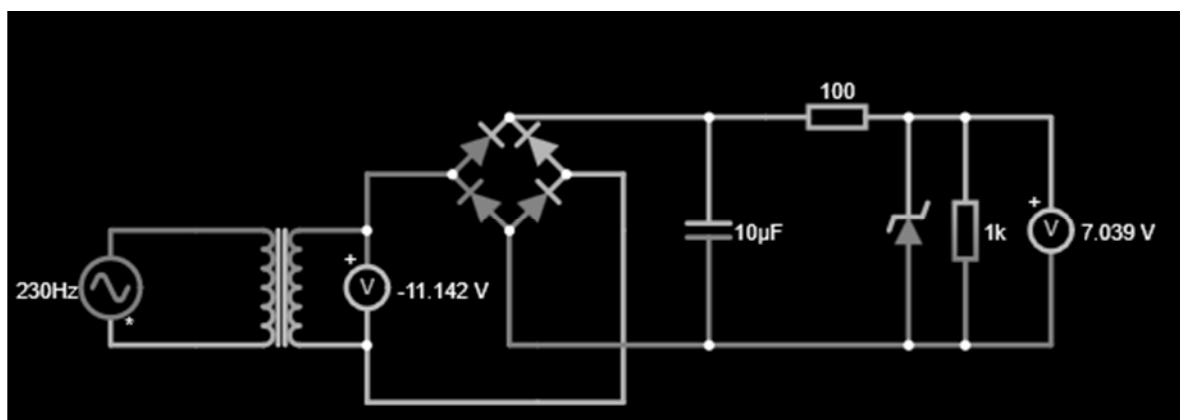


Rys. 47. Przebieg napięcia na wyjściu układu filtrującego oraz wejściowego prądu przemiennego

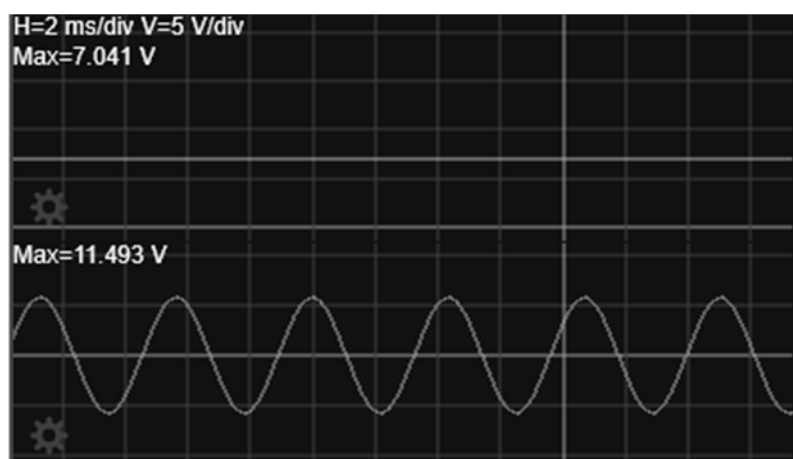
Ostatnim stopniem zasilacza prądu stałego jest stabilizator napięcia, który stabilizuje napięcie wyjściowe na pożądanym poziomie. Do tego celu może być użyta dioda Zenera, z dodatkowym rezystorem. Układ taki prezentuje rysunek 48. Wartość napięcia Zenera diody określa wartość napięcia jaka ma być ustabilizowana. Zadaniem dodatkowego rezystora, jest to by występował na nim spadek nadmiaru napięcia powyżej napięcia Zenera. Otrzymany przebieg wyjściowy prezentuje rysunek 49. Ponieważ charakterystyka diody Zenera nie jest idealnie stroma, dlatego



też wartość napięcia stabilizowanego też nie jest idealna. Obecnie stosuje się często inne stabilizatory na bazie układów scalonych.



Rys. 48. Prostownik z układem stabilizacyjnym napięcia na bazie diody Zenera

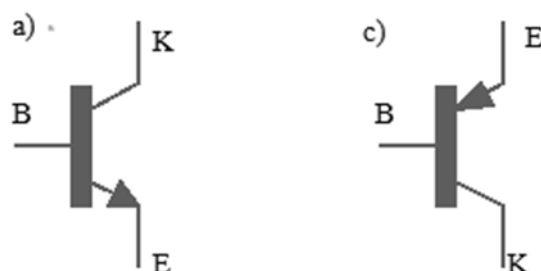


Rys. 49. Napięcie wyjściowe z zasilacza z rysunku 48 oraz napięcie prądu przemiennego na wyjściu transformatora

Tranzystory

Tranzystory są elementami elektronicznymi trójelektrodowymi, w których wykorzystano właściwości złącza p-n. Tranzystory można generalnie podzielić na dwie duże grupy. Są to tranzystory bipolarne oraz tranzystory polowe.

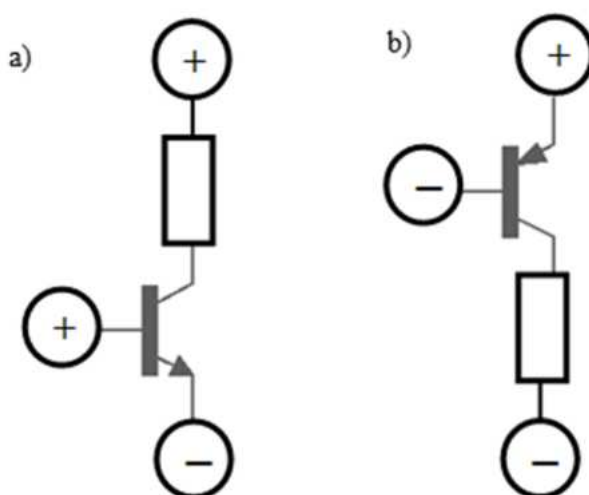
W przypadku tranzystorów bipolarnych zastosowano w nich dwa złącza p-n o wspólnej elektrodzie zwanej bazą. Istnieją dwa rodzaje tranzystorów bipolarnych p-n-p, oraz n-p-n lub w skrócie PNP i NPN. Litera w środku wskazuje typ przewodnictwa warstwy bazy. Na rysunku 50, poniżej przedstawiono symbole obu typów tranzystorów bipolarnych. Poza bazą pozostałe dwie elektrody noszą nazwę kolektora i emitera



Rys. 50. Symbole tranzystorów bipolarnych:
a) NPN, b) PNP, oznaczenie B - baza, E - emiter, K - kolektor

W przypadku tranzystora NPN wytworzenie potencjału ujemnego na emiterze powoduje odepchnięcie znajdujących się w nim elektronów w kierunku bazy. Jeżeli w tym samym czasie na bazie zostanie wytworzony potencjał dodatni, to rozpocznie się przyciąganie elektronów gromadzących się po stronie emitera. Występuje wówczas tak zwana polaryzacja w kierunku przewodzenia. Ponieważ warstwa bazy jest bardzo cienka, to rozpędzone przyciągane elektrony mogą przeskoczyć do obszaru kolektora, który także powinien być spolaryzowany dodatnio. Przepływ elektronów ma oczywiście przeciwny kierunek niż umowny kierunek przepływu prądu. Dlatego też strzałka na symbolu emitera tranzystora NPN jest skierowana na zewnątrz, tak jak umownie płynie prąd.

W przypadku tranzystorów PNP, działają one na tej samej zasadzie, ale nośnikiem ładunków są dodatnie dziury, więc polaryzacje zapewniające przewodzenie muszą być odwrotne. Spolaryzowanie tranzystorów NPN i PNP w kierunku przewodzenia przedstawiono na rys. 51.



Rys. 51. Tranzystory spolaryzowane w kierunku przewodzenia a) NPN, b) PNP



Tranzystory bipolarne, są sterowane za pomocą prądu bazy. Mały prąd przepływający przez bazę powoduje przepływ znacznie większego przez kolektor. Zdolność tranzystora do zwiększania wartości prądu kolektora w stosunku do prądu bazy określa się za pomocą tak zwanego współczynnika wzmocnienia β :

$$\beta = \frac{I_k}{I_b} \quad (109)$$

W tranzystorze NPN prąd wpływa przez dwie elektrody i wypływa przez jedną, z tego powodu:

$$I_e = I_b + I_k \quad (110)$$

gdzie:

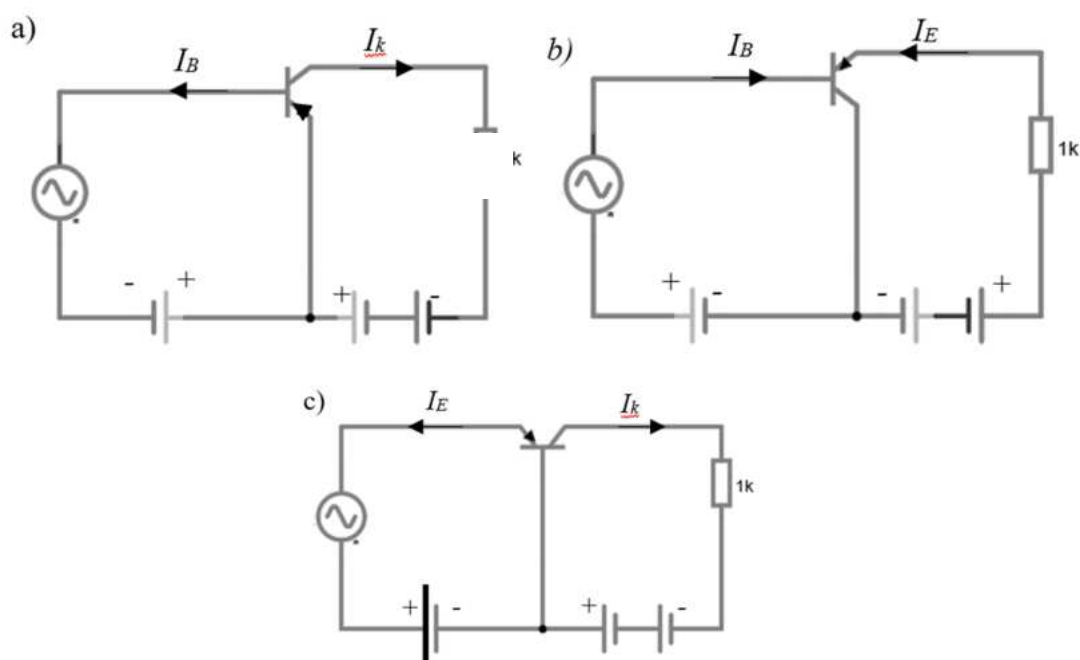
I_e – prąd emitera

I_b – prąd bazy

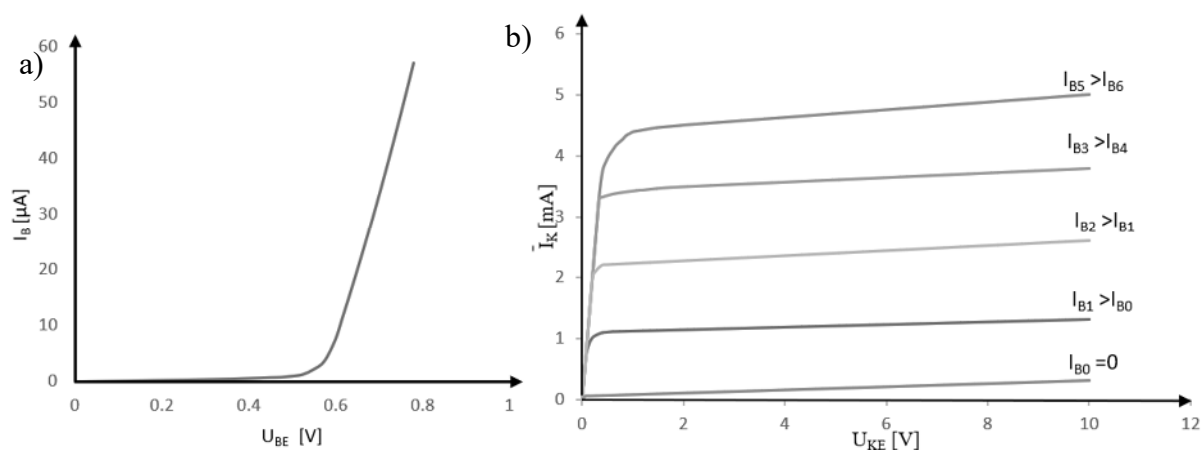
I_k – prąd kolektora

Kiedy różnica potencjałów między bazą i emiterem staje się mniejsza niż, około 0,6 V tranzystor przestaje przewodzić i przechodzi w stan wyłączenia stan ten jest nazywany stanem odcięcia. Natomiast kiedy prąd bazy osiąga taką wartość, że tranzystor nie może już go zwiększyć mówimy, że tranzystor przeszedł w stan nasycenia. Stan pośredni między stanem odcięcia i nasycenia jest nazywany stanem aktywnym. Wówczas współczynnik β zachowuje prawie stałą wartość, a zależność prądu kolektora od prądu bazy jest prawie liniowa.

Na bazie tranzystorów bipolarnych można budować wzmacniacze, czyli układy wzmacniające sygnały elektryczne. Wykorzystując tranzystory można zbudować wzmacniacze w trzech typach układów połączeń. Każdy z tych typów ma jedną elektrodę tranzystora wspólną dla obwodów wejściowego i wyjściowego. Te typy układów są określane jako: ze wspólnym emiterem, wspólnym kolektorem oraz ze wspólną bazą. Układ ze wspólnym emiterem umożliwia uzyskanie dużych wzmocnień napięcia i prądu, Natomiast układ ze wspólnym kolektorem umożliwia uzyskanie dużego wzmocnienia prądowego, natomiast uzyskiwane wzmocnienie napięciowe jest mniejsze od jedności. W przypadku układu ze wspólną bazą nie pozwala na uzyskanie wzmocnienia prądowego, ale pozwala uzyskać bardzo duże wzmocnienie napięciowe. Na rysunku 52 przedstawiono te układy. Natomiast na rysunku 53 przedstawiono przykładowe charakterystyki dla tranzystora w układzie wspólnego emitera.

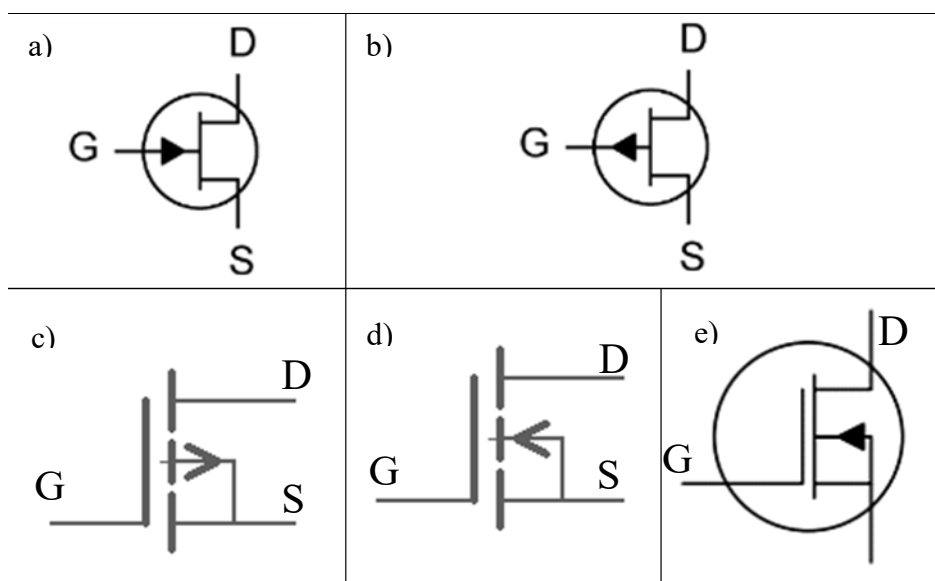


Rys. 52. Układy wzmacniaczy tranzystorowych:
a - z wspólnym emiterem, b - z wspólnym kolektorem, c - z wspólną bazą



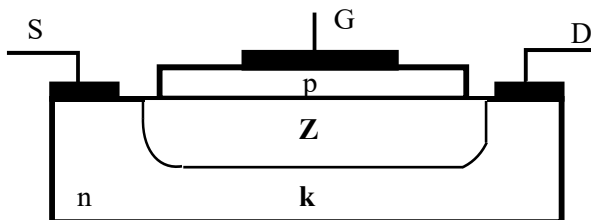
Rys. 53. Przykładowe charakterystyki tranzystora w układzie wspólnego emitera:
a - charakterystyka wejściowa, b - charakterystyka wyjściowa

Innym rodzajem tranzystorów są tranzystory polowe. Występuje kilka rodzajów tranzystorów polowych FET - tranzystory polowy, JFET - tranzystory złączowe, o strukturze metal-izolator-półprzewodnik-MIS, o strukturze metal-tlenek-półprzewodnik-MOS. Przykładowe symbole tranzystorów polowych przedstawia rysunek 54.



Rys. 54. Przykładowe symbole niektórych tranzystorów polowych: a) tranzystor JFET n-kanalowy, b) tranzystor JFET p-kanalowy, c) tranzystor MOSFET wzbożony p-kanalowy, d) tranzystor MOSFET wzbożony n-kanalowy, e) tranzystor MOSFET zubożony n-kanalowy. Wszystkie te podrodzaje mają po trzy elektrody o nazwach bramka - G, dren - D i źródło - S.

Budowę tranzystora JFET przedstawia rysunek 55.



Rys. 55. Schematyczna budowa tranzystora polowego JFET, p - półprzewodnik typu p, n - półprzewodnik typu n, z - obszar zubożony, k - kanał, S - źródło, G - bramka, D - dren

Działanie tranzystorów JEFET opiera się na kontrolowaniu wielkości warstwy zubożonej w nośniki ładunków złącza p-n. Prąd źródło-dren przepływa przez kanał k, którego rozmiar zależy od wielkości warstwy zubożonej - Z. Na jej wielkość można wpływać przez napięcie polaryzujące podłączane do bramki. Sterując napięciem można sterować wielkością strefy zubożonej, a tym samym wielkością prądu źródło-dren. Ponieważ rezystancja bramki jest bardzo duża praktycznie żaden prąd do niej nie wpływa, sterowanie odbywa się samym napięciem. Z tego powodu sam tranzystor polowy pobiera bardzo małą moc. Inne typy tranzystorów polowych różnią się budową wewnętrzną, ale działają na podobnych zasadach. W przypadku tranzystorów JEFET kanał tworzy się w warstwie wzbożonej (p albo n zależnie od podtypu). Jest w niej dostateczna liczba nośników ładunków, aby przewodzić prąd między



źródłem i drenem przy zerowej polaryzacji bramki. Dopiero powiększenie obszaru zubożonego na skutek odpowiedniej polaryzacji bramki może spowodować, że przepływ prądu ustanie. Z tego powodu tranzystory JEFET domyślnie przewodzą, przykładając odpowiednie napięcie do bramki można ograniczyć prąd, albo całkowicie odciąć jego przepływ. W przypadku tranzystorów polowych MOSFET można w ich przypadku wyróżnić tak zwane wzbogacone oraz zubożone. W przypadku zubożonych, są one domyślnie włączone tak jak tranzystory JFET, natomiast tranzystory MOSFET wzbogaconych są domyślnie wyłączone. Oznacza to, że aby prąd mógł popłynąć między źródłem i drenem wymagane jest przyłożenie pewnego potencjału do bramki. Ze względu na minimalny pobór mocy tranzystory polowe znalazły powszechne zastosowanie w układach sterowania mocą oraz w układach scalonych.

Bibliografia

- Chwaleba A., Moeschke B., Płoszajski G. *Elektronika*, Wydawnictwa Szkolne i Pedagogiczne, Warszawa 2008.
- Horowitz P., Hill W., *Sztuka elektroniki*. Wydawnictwa Komunikacji i Łączności, Warszawa 2003.
- Kurdziel R., *Podstawy elektrotechniki*. Wydawnictwa Naukowo-Techniczne, Warszawa 1995.
- Opydo W., *Elektrotechnika i elektronika dla studentów wydziałów nieelektrycznych*. Wydawnictwo Politechniki Poznańskiej, Poznań 2005.
- Pietrzyk W., *Laboratorium z elektrotechniki i elektroniki*, Wyd. Politechniki Lubelskiej, Lublin 1994.
- Platt C., *Elektronika. Od praktyki do teorii. Kolejne eksperymenty*, Helion, Gliwice 2017.
- Platt C., *Encyklopedia elementów elektronicznych*, t. 1, Helion, Gliwice 2021.
- Przeździecki F., *Elektrotechnika i elektronika*, PWN, Warszawa 1982.
- Tietze U., Schenk Ch., *Układy półprzewodnikowe*, Wydawnictwa Naukowo-Techniczne Warszawa 1996.



Kamil Łodygowski

MONITOROWANIE I DIAGNOSTYKA

1. Diagnostyka

Ze względu na duży stopień złożoności maszyn oraz ukierunkowanie na bezpieczeństwo i względy ekonomiczne konstruktorzy, ale głównie użytkownicy tych obiektów zmuszeni są do znajomości ich bieżącego stanu technicznego. Osiągnięcie takich celów możliwe jest na etapie konstruowania przy określeniu właściwych procedur diagnostycznych.

Diagnostyka może dotyczyć ludzi, maszyn, procesów, systemów, środowiska. Pojęcie to łączy kilka innych:

diagnoza – określenie bieżącego stanu technicznego,

geneza – określenie przyczyn zaistnienia obecnego stanu,

prognoza – określenie horyzontu czasowego przyszłej zmiany stanu, wtedy czas pracy maszyny określony jest jej obiektywnym stanem technicznym, a nie statystyką – niezawodnością.

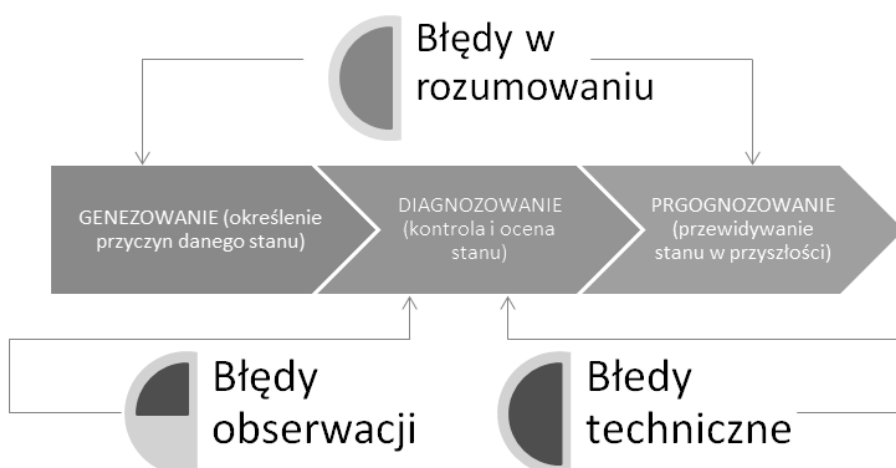
Przeanalizujemy diagnostyczny system eksploatacji (rys. 2). System ten umożliwia zwiększenie efektywności a przede wszystkim poprawę sprawności eksploatacji. W czasie eksploataowania maszyn, ze względu na istnienie szeregu czynników zewnętrznych oraz wewnętrznych w obiekcie technicznym następują zaburzenia stanów równowagi, obrazujące się zmianami – starzeniem, zużywaniem, uszkodzeniami. Sygnałem diagnostycznym jest przebieg zmian mierzalnych wartości pokazujących konkretną informację o stanach badanego obiektu w czasie. Obserwacja sygnałów diagnostycznych ukierunkowanie jest na pozyskiwanie i przetwarzania informacji diagnostycznej oraz klasyfikację stanów maszyny. Sygnał diagnostyczny to zmienna wyjściowa, której parametry musi być jednoznaczny i stabilny. Następnie ważnym etapem jest wnioskowanie diagnostyczne, określenie wartości granicznych oraz wyznaczenie czasu kolejnego diagnozowania.

Podczas procesu diagnozowania należy przede wszystkim ocenić stan obiektu w danej chwili oraz określić niezgodności parametrów regulacyjnych. W przypadku stanu niezdatności należy zlokalizować uszkodzenia i wskazać przewidywany czas naprawy. Ostatnim etapem jest określenie terminu następnego diagnozowania.



Rys. 1. Algorytm postępowania podczas procesu diagnostycznego
Źródło: opracowanie własne

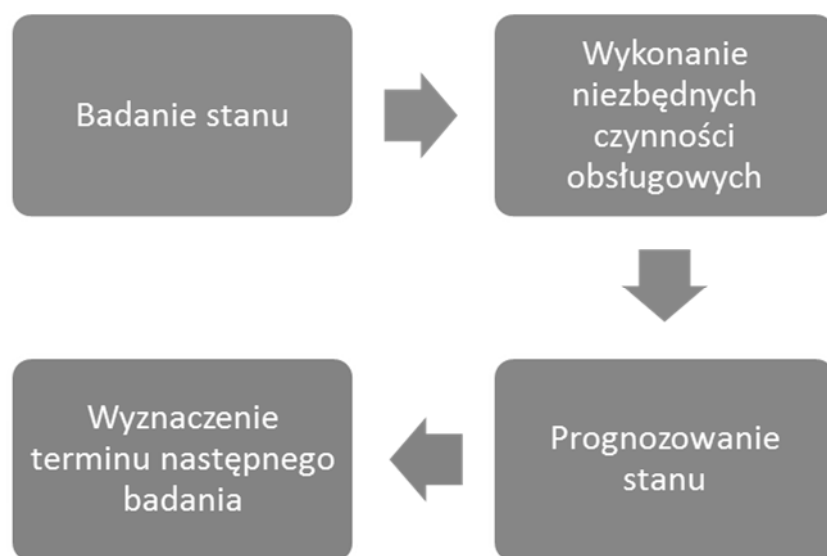
Proces postępowania z obiektem technicznym uzależniony jest od określenia jego stanu. W przypadku, kiedy maszyna jest w stanie zdatnym, postępowanie odbywa się wg algorytmu przedstawionego na rysunku 3. Jeżeli obiekt jest w stanie niezdatności, wtedy procedura wygląda jak na rysunku 4.



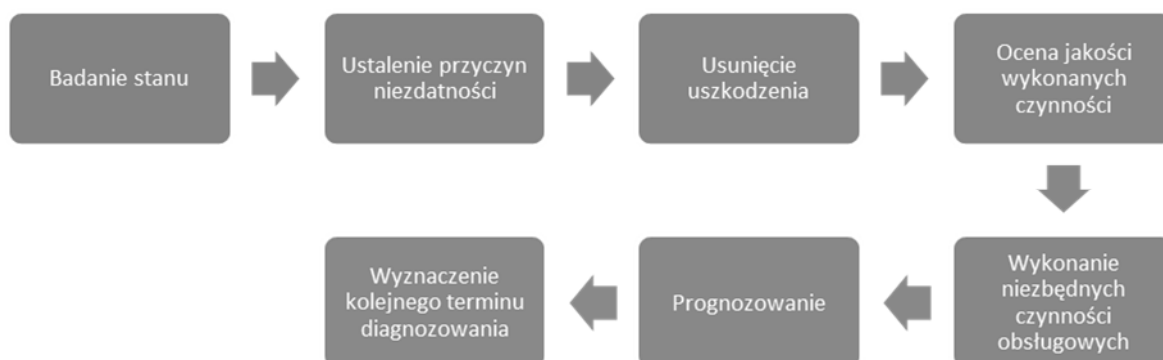
Rys. 2. Diagnostyczny system eksploatacji
Źródło: opracowanie własne



Podstawowym założeniem sterowania procesem eksploatacji maszyn jest to, aby dany obiekt techniczny pozostawał w stanie zdatności. Warunkiem koniecznym utrzymania w tym stanie jest ciągle diagnozowanie i prognozowanie jego stanu. Ustalenie wiarogodnego stanu stanowi podstawę wykonywania koniecznych czynności obsługi. Oceny prawidłowości realizacji czynności obsługowych (wymaganego zakresu, ich jakości) dokonuje się w wyniku stosowania metod diagnostyki technicznej, może pełnić ona również nadrzędną rolę podczas podejmowania decyzji o wycofaniu obiektu technicznego z eksploatacji - likwidacji.



Rys. 3. Proces postępowania diagnostycznego z obiektem technicznym zdatnym



Rys. 4. Proces postępowania diagnostycznego z obiektem technicznym niezdatnym



Stan maszyny określa się za pomocą wielkości fizycznych. Dokonuje się pomiaru wartości wybranych wielkości podczas procesu diagnozowania. Istotne jest prawidłowe określenie wielkości fizycznych mogących być wykorzystanych w pomiarach diagnostycznych, oraz wyznaczenie umiejscowienia czujników pomiarowych. Najważniejszy warunek, jaki musi spełnić wielkość fizyczna, aby można ją było uznać za podstawę do wyznaczania stanu maszyny (procesu), to istnienie zależności między zmianą wartości tej wielkości a zmianą stanu maszyny.

Obiekt techniczny należy traktować jako system, w którym wyróżnia się zmienne:

- stanu (np. luzy, zużycie części mechanicznych),
- wejściowe (np. ilość dostarczonego paliwa, zapotrzebowanie na moc),
- wyjściowe (np. drgania, ilość produktów procesu zużywania, uzyskiwana moc),
- zakłóceń (np. temperatura otoczenia, wilgotność powietrza, zanieczyszczenie paliwa, zanieczyszczenie powietrza)¹.

Symptom stanu jest miarą sygnału, która zmienia się istotnie wraz ze zmianą stanu technicznego obiektu. Tym samym między stanem obiektu w określonym czasie (chwili) a ilustrującą ten stan wartością zmiennej wyjściowej istnieje implikacja. Każdemu konkretnemu stanowi maszyny odpowiada konkretna wartość sygnału pomiarowego. Zbiór zmiennych wejściowych, zwanych także wymuszeniami, określa oddziaływania, którym podlega przedmiot diagnozy podczas badań i oceny jego stanu, warunków pracy.

Zakłócenia w diagnozowaniu powodujące błędy genezy, diagnozy i prognozy obejmują warunki zewnętrzne związane z otoczeniem: temperatura, wilgotność powietrza, stopień zanieczyszczenia atmosfery, których dokładne ustalenie jest niemożliwe. Dodatkowo zakłócenia i tym samym błędy w procesie diagnozowania obejmują warunki techniczne diagnozowania obiektu: prędkość obrotowa, obciążenie, obecność środków smarnych, których dokładne ustalenie nie jest możliwe. Kolejną sprawą są błędy w blokach pomiarowych dopasowujących i przetwarzających oraz w innych urządzeniach diagnostycznych.

¹ Legutko S., Eksploatacja maszyn, Wydawnictwo Politechniki Poznańskiej, Poznań 2007



2. Niezawodność maszyn²

Pojęcie niezawodności może obejmować różne wymagania opisane charakterystykami technicznymi, ekonomicznymi i socjologicznymi obiektów.

Wyróżnia się:

- niezawodność techniczną, która uwzględnia charakterystyki techniczne,
- niezawodność techniczno-ekonomiczną, która uwzględnia charakterystyki techniczne i ekonomiczne,
- niezawodność globalną, która uwzględnia charakterystyki techniczne, ekonomiczne i socjologiczne obiektów.

Niezawodność - bez dodatkowych określeń - jest rozumiana jako niezawodność techniczna. Niezawodność obiektu jest to jego zdolność do spełnienia stawianych mu wymagań. Wielkością charakteryzującą zdolność do spełnienia wymagań może być prawdopodobieństwo spełniania wymagań. Stąd definicja: „niezawodność obiektu jest to prawdopodobieństwo spełnienia przez obiekt stawianych mu wymagań”.

Kiedy wymaganiem jest to, żeby obiekt był zdalny (sprawny) w przedziale $(0, t)$, którego miarą może być czas, ilość wykonanej pracy, liczba wykonanych czynności, długość przebytej drogi itp. wtedy: „niezawodność obiektu jest to prawdopodobieństwo, że obiekt jest zdalny (sprawny) w przedziale $(0, t)$ ” lub: „niezawodność obiektu jest to prawdopodobieństwo, że wartości parametrów określających istotne właściwości obiektu nie przekroczą w ciągu okresu $(0, t)$ dopuszczalnych granic w określonych warunkach eksploatacji obiektu”. W sensie probabilistycznym niezawodność obiektu $R(t)$ w danej chwili t jest prawdopodobieństwem $P(T \geq t)$, że jego trwałość T jest większa od t , tj. $R(t) = P(T \geq t)$. Trwałość T może być wyrażona np. czasem w [s], długością w [km] itp. Z tego wynika, że za każdym razem $R(t)$ jest inna.

Gotowość obiektu naprawialnego - tj. obiektu, któremu przywraca się sprawność, gdy ją utraci - może być definiowana w różny sposób, np.

- Gotowość obiektu jest to prawdopodobieństwo, że obiekt będzie gotowy do pełnienia swych funkcji w chwili t .
- Gotowość obiektu jest to frakcja danego okresu (np. 1 roku), w ciągu którego obiekt jest zdolny do pełnienia swych funkcji lub je pełni.

² Macha E. Niezawodność maszyn; Skrypt Nr 237, Politechnika Opolska.



- Gotowość obiektu jest to frakcja sumy okresów eksploatacji obiektu, w ciągu której obiekt pełni swe funkcje lub jest zdolny do pełnienia swych funkcji.
- Gotowość jest to frakcja całego życia obiektu, w ciągu której obiekt jest zdolny do pełnienia funkcji lub ją pełni.

$$\text{Gotowość: } G = T / (T+Q),$$

gdzie

T – średnia długość okresów sprawności,

Q – średnia długość okresów niesprawności

Do wartości szczególnych niezawodności należą:

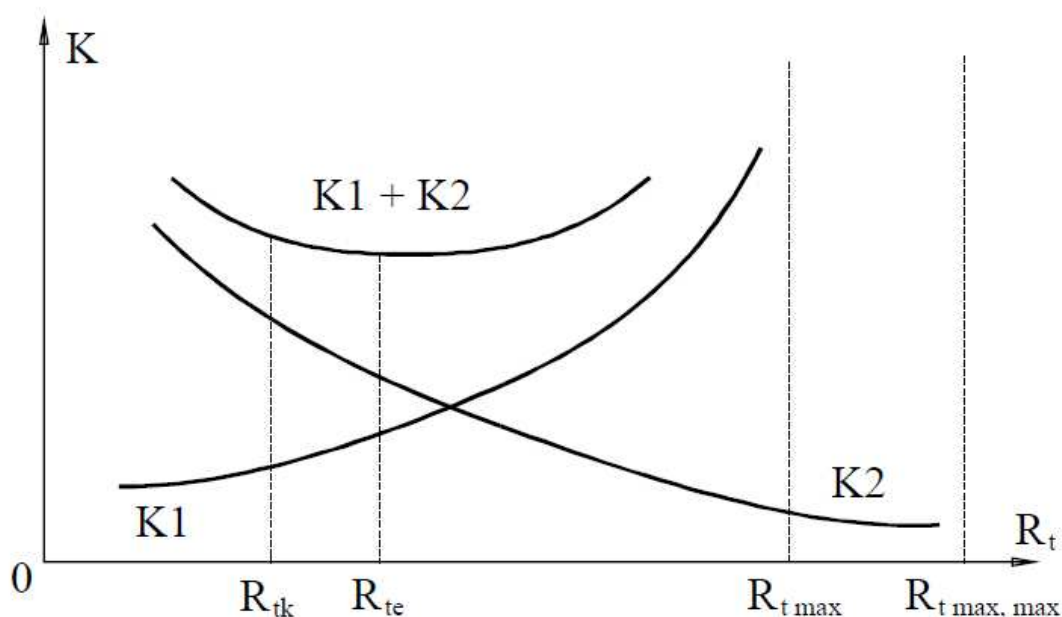
$R_{t \max}$ - wartość maksymalna niezawodności (uzyskiwana lokalnie),

R_{te} - ekonomicznie optymalna wartość niezawodności,

R_{tk} - wartość krytyczna niezawodności (nietolerowana przez użytkowników),

$R_{t \max \max}$ - największa wartość niezawodności uzyskiwana w technice światowej.

Na rys. 5 przedstawiono zależność kosztów od niezawodności.



Rys. 5. Zależność kosztów K od niezawodności R_t ; K_1 – koszty zwiększenia R_t , K_2 – koszty postojów, gwarancji, serwisu itp.

Źródło: Macha E. Niezawodność maszyn; Skrypt Nr 237, Politechnika Opolska.

Z rysunku 5 wynika, że koszty K_1 uzyskania większej niezawodności R_t rosną, natomiast przy dużej niezawodności maleją koszty K_2 postojów, gwarancji, serwisu itp. Istnieje



zatem minimalna suma kosztów $K1 + K2$ przy których uzyskuje się ekonomicznie optymalną wartość niezawodności R_{te} .

Każdy obiekt techniczny ma określoną niezawodność (R), trwałość (T) i gotowość (G). Zależnie od konkretnego zastosowania oraz wymagań (podawanych zwykle w normach, przepisach, umowach handlowych itp.) coraz częściej żąda się od wytwórców podawania wartości liczbowych odpowiednich wskaźników dotyczących niezawodności, trwałości i gotowości. Mając takie dane można racjonalniej podejmować decyzje związane z produkcją wyrobów oraz ich długotrwałą eksploatacją. Ze względu na rodzaj charakterystyki, która jest istotna dla danego obiektu technicznego, obiekty dzieli się na 8 klas (tabela 1).

Tab. 1. Podstawowe klasy obiektów technicznych

R	T	G	Klasy	Oznaczenie klasy
-	-	-	I	typu I
R	-	-	II	typu R
-	T	-	III	typu T
R	T	-	IV	typu RT
-	-	G	V	typu G
R	-	G	VI	typu RG
-	T	G	VII	typu TG
R	T	G	VIII	typu RTG

Źródło: Macha E. Niezawodność maszyn; Skrypt Nr 237, Politechnika Opolska.

- I. Obiekty typu I (ilość), to obiekty, którym nie stawia się wymagań jakościowych związanych z ich niezawodnością, trwałością i gotowością.
- II. Obiekty typu R (niezawodność), od których wymaga się dużej niezawodności. Typowymi przykładami takich obiektów są obiekty specjalnego przeznaczenia (np. samoloty, helikoptery).
- III. Obiekty typu T (trwałość), od których wymaga się przede wszystkim dużej trwałości. Są to drogie i ważne gospodarczo obiekty techniczne, np. budynki, mosty, wiadukty itp.
- IV. Obiekty typu RT (niezawodność i trwałość), dla których podstawowymi wymaganiami są jednocześnie duża niezawodność i duża trwałość. Są to drogie i ważne gospodarczo obiekty o długotrwałej i ciągłej eksploatacji, np. elektrownie (zwłaszcza jądrowe), zapory wodne, statki i inne.



- V. Obiekty typu G (gotowość), od których wymaga się dużej gotowości. Są to głównie pogotowia - medyczne (karetka reanimacyjna), straż pożarna (wóz strażacki) i inne.
- VI. Obiekty typu RG (niezawodność i gotowość), które cechują się zarówno dużą niezawodnością jak i dużą gotowością, np. helikopter pogotowia medycznego, aparaty ratownictwa górniczego i inne.
- VII. Obiekty typu TG (trwałość i gotowość), od których wymaga się głównie dużej trwałości i gotowości, np. statki ratownictwa morskiego i inne.
- VIII. Obiekty typu RTG (niezawodność, trwałość, gotowość). Są to różnego rodzaju obiekty pogotowia, charakteryzujące się dużą trwałością i niezawodnością.

W teorii i inżynierii niezawodności przyjmuje się, że funkcją, która najlepiej charakteryzuje zmiany niezawodności dowolnego obiektu technicznego jest funkcja intensywności uszkodzeń $\lambda(t)$. Z jej przebiegu można wyciągnąć wiele wniosków natury teoretycznej i praktycznej, a także wyznaczyć:

- Funkcję niezawodności

$$R(t) = P(T \geq t) = \exp \left[- \int_0^t \lambda(x) dx \right],$$

- Funkcję zawodności

$$Q(t) = F(t) = P(T < t) = 1 - R(t) = 1 - \exp \left[- \int_0^t \lambda(x) dx \right],$$

- Funkcję gęstości prawdopodobieństwa

$$f(t) = \frac{dF(t)}{dt} = \lambda(t) \exp \left[- \int_0^t \lambda(x) dx \right],$$

- Funkcję wiodącą (skumulowana intensywność uszkodzeń)

$$\Lambda(t) = \int_0^t \lambda(x) dx.$$

Znajomość przebiegu funkcji $\lambda(t)$ umożliwia producentowi i użytkownikowi, podejmowanie ważnych decyzji praktycznych w zakresie:



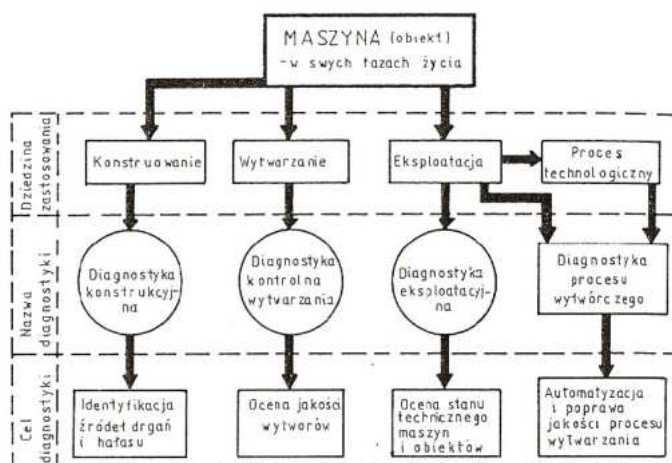
- ustalania niezbędnych okresów starzenia wstępnego produkowanych obiektów,
- ustalenia wielkości i asortymentu części zamiennych,
- planowania optymalnej pracy serwisu technicznego, służb remontowych,
- ustalenia optymalnych okresów wymian profilaktycznych elementów i zespołów,
- ustalania optymalnych okresów eksploatacji obiektów (teoria odnowy),
- innych działań techniczno-ekonomicznych (okres gwarancji).

W wielu przypadkach eksperymentalne przebiegi funkcji $\lambda(t)$ można aproksymować funkcjami analitycznymi (teoretycznymi rozkładami prawdopodobieństwa jak np. rozkładem wykładniczym, beta, równomiernym, Weibulla lub kompozycją tych rozkładów).

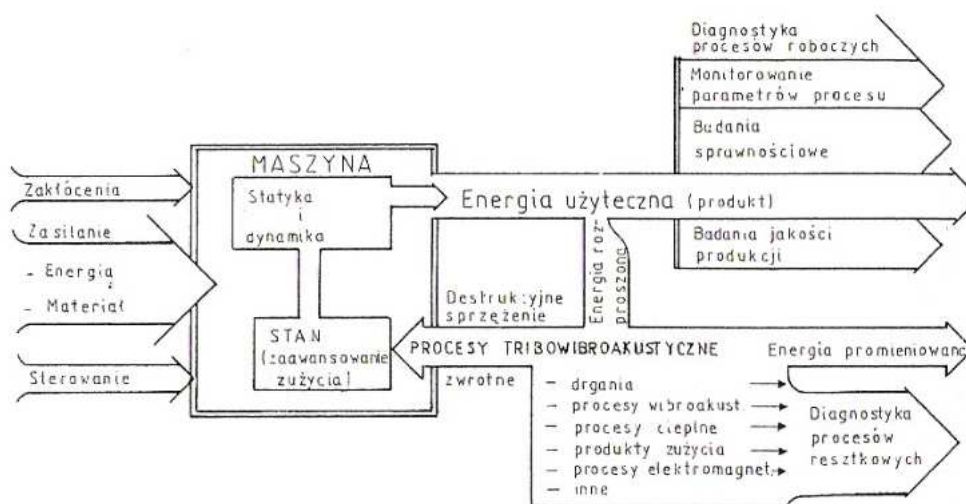
Rzeczywiste przebiegi funkcji $\lambda(t)$ konkretnego obiektu, zależnie od przyjętej strategii eksploatacji, mogą być bardzo różne i mogą być celowo kształtowane.

3. Diagnostyka wibroakustyczna

Diagnostyka wibroakustyczna (WA) oparta jest na obserwacji jednego z procesów resztkowych generowanych podczas ruchu maszyny. Obserwacja procesów WA maszyn daje bezpośrednią informację o zaawansowaniu procesów ich zużycia i degradacji, czyli o stanie technicznym. Informację tę uzyskuje się w sposób nieinwazyjny, co jest niezwykle cenną cechą operacyjną diagnostyki WA. Nieinwazyjność ta dotyczy realizowanego procesu technologicznego, miejsca montażu, a także maszyny jako całości. Dodać należy do tego wielowymiarowość informacyjną procesów WA, co oznacza, że z drgań obserwowanych na obudowie łożyska możemy oceniać zaawansowanie zużycia łożysk, wirnika, osiowość agregatu, jego pewność posadowienia, itp. Generalnie można powiedzieć, że intensywność procesów WA zależy odwrotnie proporcjonalnie od jakości maszyny na każdym etapie jej eksploatacji. Maszyny w ruchu to systemy samogenerujące drgania, hałas, pulsacje medium i inne procesy WA. Są one przyczynami lub efektami zużywania się maszyn.



Rys. 6. Systematyczne ujęcie celów poszczególnych rodzajów diagnostyki WA w przemyśle
Źródło: Engel Z.: Wkład Profesora Czesława Cempela w rozwój wibroakustyki, Diagnostyka 3(47)/2008.



Rys. 7. Maszyna jako system przekształcający energię o różnych możliwościach diagnostyki
Źródło: Engel Z.: Wkład Profesora Czesława Cempela w rozwój wibroakustyki, Diagnostyka 3(47)/2008.

DIAGNOSTYKA SYGNAŁU WIBROAKUSTYCZNEGO:

- Podczas każdej pracy maszyny występują drgania.
- Drgania te są bardzo pojemne, jeśli chodzi o informacje (porównywalne z sygnałami wizualnymi).
- Wydobycie sygnału wibroakustycznego jest niezwykle tanie.
- Można wskazać uszkodzenie maszyny na podstawie sygnału wibroakustycznego.
- Można precyzyjnie wskazać miejsce i obszar zainicjowania uszkodzenia.
- Można oszacować czas pracy maszyny na podstawie sygnałów wibroakustycznych.

Jakie częstotliwości maszyn można zarejestrować?



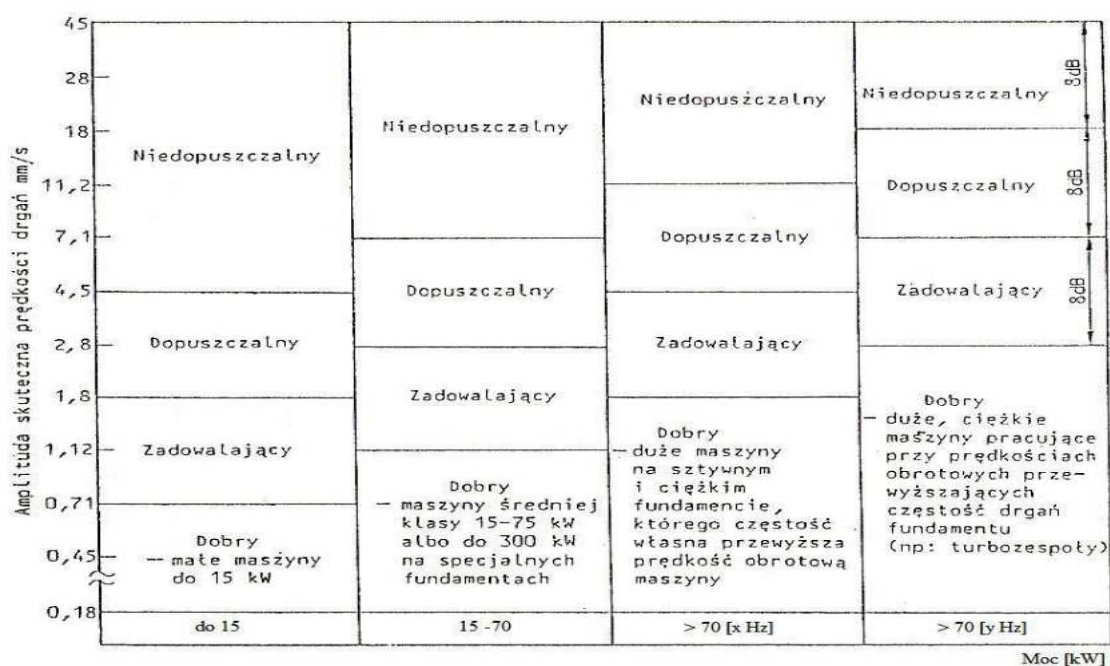
- $\sim 0 \div 10$ Hz
- $\sim 10 \div 1$ kHz $\rightarrow$ większość maszyn drga z tą częstotliwością
- > 1 kHz ($\sim 1 \div 400$ kHz)

Do **niskich wartości** najlepszym parametrem do opisu jest **PRZEMIESZCZENIE**, a dokładniej amplituda drgań (przedział I). Do **średnich częstotliwości** (większa szybkość drgań danego elementu) stosuje się do opisu **PRĘDKOŚĆ** (przedział II). Do **wysokich wartości** drgań najlepiej obrazuje drgania parametr **PRZYSPIESZENIA** (przedział III).

KROTNOŚĆ ZMIANY STANÓW:

x2 ! $\rightarrow$ zmiana stanu

x6 !!! $\rightarrow$ AWARIA



Rys. 8. Klasyfikacja stanu drganiowego czterech grup maszyn (w zależności od mocy w kW)

Źródło: Engel Z.: Wkład Profesora Czesława Cempela w rozwój wibroakustyki, Diagnostyka 3(47)/2008.

4. Badania nieniszczące

Wyróżniamy następujące rodzaje metod badań nieniszczących elementów maszyn:

- badania wizualne (VT), czyli oględziny zewnętrzne;
- badania endoskopowe;
- badania penetracyjne (PT);



- badania magnetyczno-proszkowe (MT);
- badania ultradźwiękowe (UT);
- badania prądami wirowymi (ET);
- radiografia przemysłowa (RT).

Badanie wizualne VT (visual testing) jest metodą wykorzystywaną najczęściej jako badanie wstępne w połączeniu z inną metodą nieniszczącą. Polega ona na bezpośrednim wykryciu i ocenie nieciągłości występujących na powierzchni obiektu. W metodzie tej wykorzystuje się bezpośrednio narząd wzroku, czasami wspomagany prostą optyką (zestaw lusterek, lupa, peryskop, endoskop, wideoskop). W metodzie tej wykrywano są duże nieciągłości powierzchniowe (wklęsnięcia, podtopienia, braki przetopu, pęknięcia kuźnicze, hartownicze, spawalnicze) oraz wady kształtu badanego obiektu (ubytki korozyjne, porowatości, odkształcenia kątowe, pustki). Metoda ta znajduje zastosowanie między innymi do badania części statków, samolotów, pomp, wymienników ciepła, rurociągów, zbiorników turbin, wirników.

Badania endoskopowe polegają na podglądzie wnętrza przy wykorzystaniu aparatów umożliwiających doprowadzenie światła i optyki. Wykorzystujemy do tego celu różnorodne urządzenia, takie jak szkła powiększające, spoinomierze, lusterka kontrolne, mikroskopy czy wideoendoskopy. Badania endoskopowe pozwalają wykryć usterki spowodowane wadami kształtu, odstępstwami wymiarowymi, nieciągłościami powierzchni czy uszkodzeniami eksploatacyjnymi.

Badania penetracyjne są prostą i szybką metodą pozwalającą na wykrywanie nieciągłości powierzchniowych o szerokościach od 10-6m, takich jak pęknięcia zmęczeniowe, pęknięcia szlifierskie, porowatości, rozwarstwienia, wżery, pęknięcia powstałe po kuciu lub po walcowaniu. Metoda ta wykorzystuje zjawisko kapilarności – wnikania (penetracji) cieczy wskazującej (penetrantu) w głąb defektów powierzchni badanej (pęknięć, szczelin, rys, porów). Po oczyszczeniu badanej powierzchni z nadmiaru penetrantu i jej osuszeniu, nanosi się na nią następnie cienką, białą warstwę wywoływacza. Wywoływacz "wyciąga" penetrant z wad i czyni je widzialnymi w formie kolorowych, liniowych lub zaokrąglonych wskazań. Najczęściej do badań penetracyjnych stosuje się penetranty barwne (czerwone) lub fluorescencyjne. W przypadku, gdy zostanie zastosowany penetrant fluorescencyjny, w celu wywołania zjawiska fluorescencji, a tym samym ujawnienia nieciągłości powierzchni badanego materiału, stosuje się lampy o promieniowaniu ultrafioletowym. Barwniki te pod działaniem



promieniowania UV świecą najczęściej kolorem żółto – zielonym i są dobrze widoczne na ciemnym tle. W metodzie tej konieczne jest zaciemnienie stanowiska badawczego w celu zwiększenia wykrywania wad.

Metodę penetracyjną stosuje się w materiałach ferromagnetycznych, nieferromagnetycznych (stal, staliwo, żeliwo, miedź, mosiądz, brąz, wolfram) oraz do badań materiałów niemetalicznych (np. ceramicznych).

Zalety badań penetracyjnych:

- szybki i prosty proces badania,
- możliwość badania różnych materiałów i wyrobów o dowolnych kształtach i wymiarach,
- łatwość wykrywania wad o wielkości od ok. 0,001 mm,
- łatwa ocena wskazań,
- łatwość stosowania w warunkach warsztatowych i terenowych,
- niskie koszty badania,
- możliwość mechanizacji procesu badania,
- duża skuteczność wykrywania wad.

Wady badania penetracyjnego:

- konieczność wstępnego oczyszczenia i odtłuszczenia powierzchni badanej oraz oczyszczenia powierzchni po badaniu,
- wykrywanie tylko wad otwartych,
- wpływ temperatury obiektu na właściwości preparatów,
- starzenie się preparatów,
- duża toksyczność preparatów, a zatem konieczność zapewnienia dobrej wentylacji podczas stosowania w pomieszczeniach zamkniętych.

Badania magnetyczno-proszkowe są jedną z najbardziej czułych, wiarygodnych i wydajnych metod nieniszczących, służących kontroli powierzchni materiałów ferromagnetycznych, czyli wszystkich stali konstrukcyjnych z wyłączeniem stali wysokostopowych (austenitycznych). Metoda ta wykorzystuje oddziaływanie sił strumienia magnetycznego na cząsteczki ferromagnetyczne aplikowane na powierzchni badanego obiektu. W przypadku wystąpienia wady następuje rozproszenie strumienia magnetycznego oraz zmiana układu proszku magnetycznego w tej okolicy. Najlepszą wykrywalność osiąga się w sytuacji, w której kierunek ułożenia wady jest prostopadły do kierunku sił pola magnetycznego. Wraz ze zmniejszeniem



tego kąta, wskazania wady są coraz to słabsze. Metoda magnetyczno-proszkowa charakteryzuje się możliwością wykrywania wąskich i płytkich nieciągłości powierzchniowych i podpowierzchniowych do około 2 mm. Stosowana jest do badania połączeń spawanych, odlewów, wałów korbowych silników spalinowych, przekładni zębatych, lin kolejek i wyciągów górskich.

Zaletą tej metody jest duża szybkość wykonywanego badania oraz natychmiastowy wynik. Materiał badawczy występuje w formie barwnych lub fluorescencyjnych suchych proszków magnetycznych, zawiesin olejowych lub wodnych. Źródłami strumienia magnetycznego są jarzma (w przypadku badań małych lub trudnodostępnych obiektów) lub generatory prądu.

Badania ultradźwiękowe należą do metod "badań objętościowych". Polegają one na wprowadzaniu fal ultradźwiękowych do obiektu, które są odbijane przez nieciągłości, uginane i rozpraszane na krawędziach nieciągłości. Metoda ta pozwala wykrywać pęknięcia, zawałowania, rozwarstwienia, porowatości, nieszczelności na wskroś i inne nieciągłości wewnątrz elementów. Można je również stosować do szacowania zmian mikrostruktury materiału, powstających podczas długotrwałej eksploatacji oraz do pomiaru grubości obiektów. Główne dziedziny zastosowań to połączenia spawane, szyny kolejowe i tramwajowe, łopatki turbin i sprężarek, wyroby ceramiczne (izolatory).

Metoda prądów wirowych jest metodą powierzchniową. Polega na wzbudzaniu zmiennego pola elektromagnetycznego w badanym materiale i odbieraniu reakcji materiału poprzez sondę badawczą i defektoskop prądowirowy. Analiza wartości zmian pola elektromagnetycznego, amplitudy oraz przesunięcia fazowego napięcia i natężenia pozwala na bardzo precyzyjną ocenę stanu badanego materiału, występujących nieciągłości w postaci np. pęknięć, ubytków erozyjnych lub korozyjnych, ocenę ich wielkości oraz głębokości. Metoda ta jest bardzo często wykorzystywana przez przemysł chemiczny, maszynowy, lotniczy, rafineryjny, cukrowniczy, papierniczy, spożywczy, kosmiczny oraz do badania rurek wymienników ciepła w elektrowniach jądrowych i konwencjonalnych.

Metoda radiograficzna umożliwia wykrywanie wewnętrznych oraz powierzchniowych i podpowierzchniowych nieciągłości obiektów. Prowadzenie badań metodą radiograficzną polega na naświetlaniu obiektów promieniowaniem jonizującym tj. promieniowaniem rentgenowskim (X) lub promieniowaniem γ (gamma) otrzymywanym ze sztucznych źródeł izotopowych oraz na rejestracji obrazu prześwietlanego obiektu na kliszy radiograficznej lub w postaci cyfrowej. Metoda ta jest stosowana w kontroli złączy spawanych i zgrzewanych,



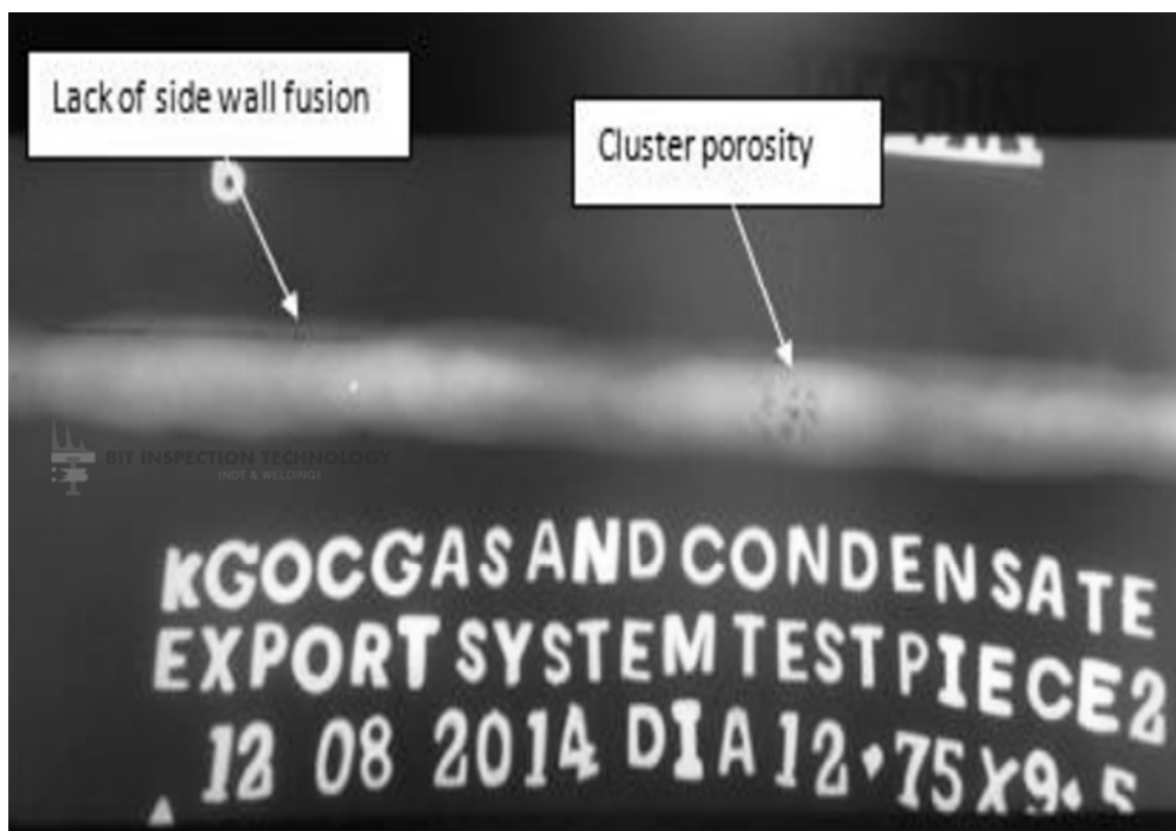
odlewów, odkuwek, rur, wlewków, kęsisk i in. Radiografia ma zastosowanie w badaniach wszystkich metali i ich stopów. Przyjmuje się, że za pomocą techniki radiograficznej wykrywa się różnice grubości wynoszące od 1,5 % do 2 %. Wykrywalność metody przy użyciu promieniowania gamma jest znacznie gorsza niż przy zastosowaniu promieniowania X.



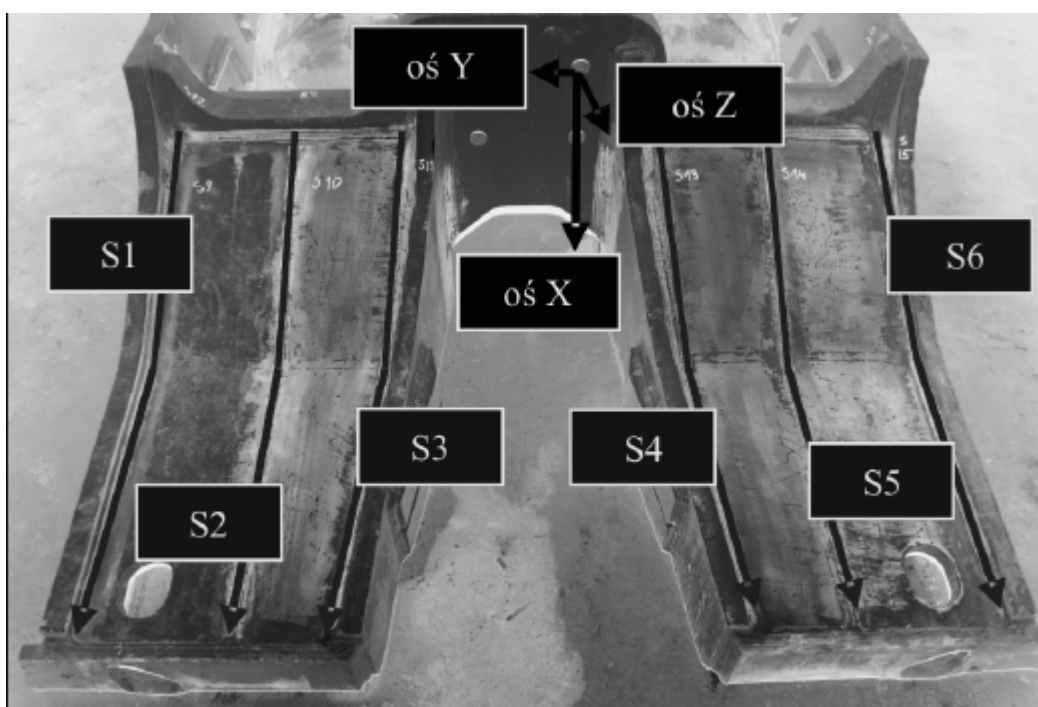
Rys. 9. Badania nieniszczące metodą magnetoproszkową materiałów ferromagnetycznych
Źródło: opracowanie własne



Rys. 10. Badania nieniszczące metodą penetracyjną materiałów stalowych
Źródło: opracowanie własne



Rys. 11. Badania nieniszczące metodą radiograficzną materiałów stalowych



Rys. 12. Miejsca złącza spawanego poddane badaniom nieniszczącym metodą ultradźwiękową.
Źródło: Rakwicz B., Wojtynek R.: *Zastosowanie metody ultradźwiękowej do oceny stanu złączy spawanych spągownicy sekcji ścianowej obudowy zmechanizowanej*. Maszyny Górnicze 3/2013.

Tab. 2. Wyniki badań ultradźwiękowych złączy spawanych

Oznaczenie złącza spawanego	Numer pomiaru	Współrzędne położenia wskazania			Długość wskazania	Wysokość echa względem poziomu odniesienia	Ocena
		X [mm]	Y [mm]	Z [mm]	L [mm]	H [dB]	
spągница sekcji ścianowej obudowy zmechanizowanej po 30 tysiącach cykli wytrzymałościowych							
S2	1	60	0	20	15	10,0	
S5	2	20	0	22	10	7,0	
	3	40	0	22	10	6,5	
spągница sekcji ścianowej obudowy zmechanizowanej po 60 tysiącach cykli wytrzymałościowych							
S2	4	60	0	22	25	11,0	
	5	385	0	24	20	7,0	
S5	6	20	0	23	15	9,0	
	7	40	0	23	15	8,5	
spągница sekcji ścianowej obudowy zmechanizowanej po 100 tysiącach cykli wytrzymałościowych							
S2	8	60	0	22	40	16,0	
	9	295	10	24	10	6,0	
	10	355	10	24	75	11,3	
	11	435	10	23	25	14,0	
	12	770	12	22	150	10,0	



	13	960	10	21	12	8,5	
S5	14	0	10	23	35	14,5	
	15	40	10	23	60	7,3	
	16	115	0	22	115	9,3	
	17	335	10	23	75	13,5	
	18	640	-5	23	20	8,5	

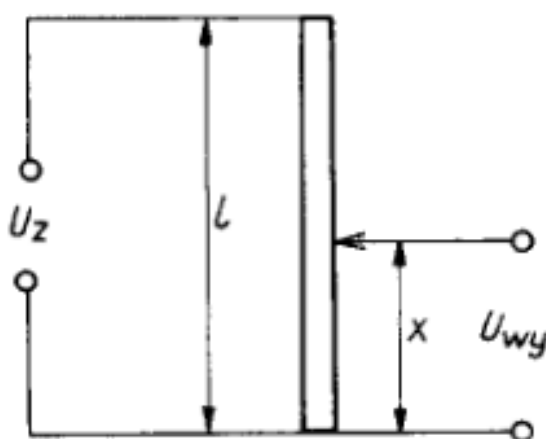
Źródło: Rakwicz B., Wojtynek R.: Zastosowanie metody ultradźwiękowej do oceny stanu złączy spawanych spawnicą sekcji ścianowej obudowy zmechanizowanej. Maszyny Górnicze 3/2013.

5. Czujniki w monitorowaniu maszyn

Czujniki potencjometryczne. Jednym z najbardziej znanych czujników przesunięcia jest czujnik potencjometryczny, którego styk ślizgowy, wykonując ruch prostoliniowy, obrotowy lub śrubowy, przyjmuje położenie odpowiadające mierzonemu przemieszczeniu. Potencjometr włączony w prosty układ elektryczny przetwarza przesunięcie prostoliniowe lub kątowe w zakresie od jednego do kilku obrotów, na napięcie stałe lub przemienne:

$$U_{wy} = \frac{x}{l} U_z$$

Gdzie: U_{wy} – napięcie wyjściowe, x – położenie chwilowe, l – maksymalna wartość położenia, U_z – napięcie zasilające

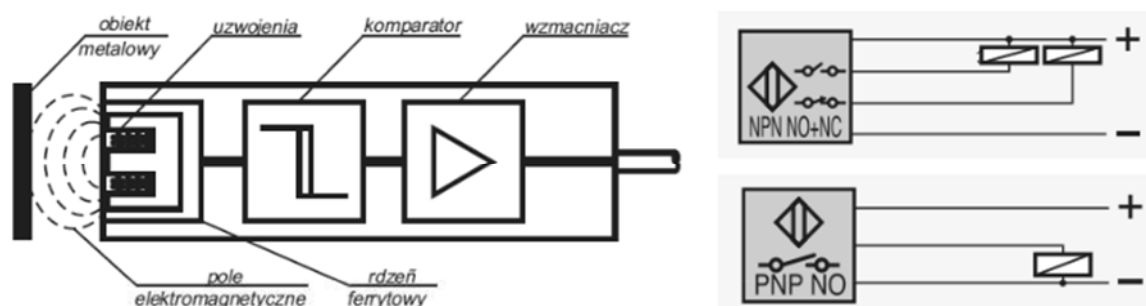


Rys. 13. Idea działania czujnika potencjometrycznego liniowego
Źródło: <http://marcin.kielczewski.pracownik.put.poznan.pl/ZSP07.pdf>



Rys. 14. Zintegrowana pompa paliwa z pływakim ramieniowym wykorzystującym czujnik potencjometryczny do monitorowania stanu paliwa

Czujniki indukcyjne. Zasada działania indukcyjnego czujnika opiera się na zmianie pola elektromagnetycznego, przez przemieszczający się metalowy element przed jego czołem. Wytworzone zmienne pole magnetyczne wzbudza w przewodniku prądy wirowe, jeżeli jest on ferromagnetykiem zostaje namagnesowany. Wymienione dwa zjawiska powodują zmianę drgań – amplitudę lub częstotliwość. Zmiany te są określane przez demodulator i na tej podstawie generowany jest sygnał wyjściowy. Indukowanie prądów wirowych oddziałuje na pole elektromagnetyczne, wytworzone wokół czoła czujnika. Pole wytwarzane jest przez obwód indukcyjny generatora wysokiej częstotliwości. Układ progowy kontroluje amplitudę generowanego napięcia, maleje ona wraz z przybliżaniem metalu. Zmiana stanu wyjścia czujnika następuje po zbliżeniu metalu na odległość włączenia i wyłączenia.

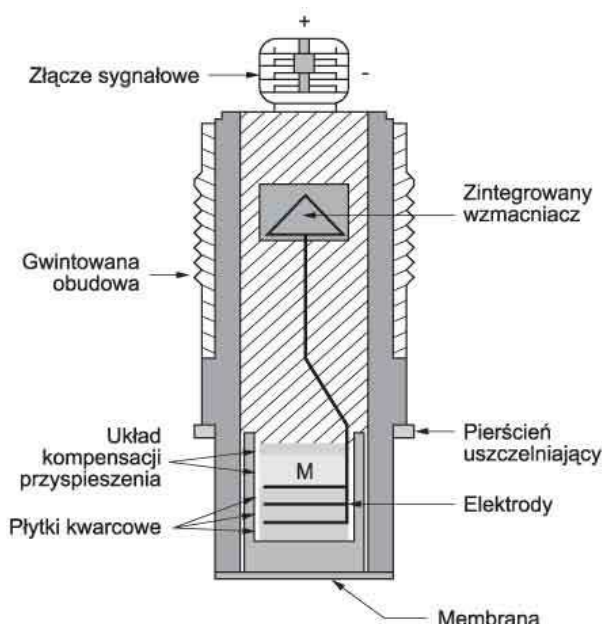


Rys. 15. Idea działania czujnika indukcyjnego

Źródło: <http://marcin.kielczewski.pracownik.put.poznan.pl/ZSP07.pdf>

Czujniki piezoelektryczne. Podstawą działania czujnika jest zjawisko piezoelektryczne, które polega na pojawieniu się ładunków elektrycznych na ściankach kryształu przy jego deformacji sprężystej. Zmiana kierunku odkształcenia powoduje zmianę znaku ładunku. Istnieje też zjawisko odwrotnego przyłożenia napięcia do elektrod przylegających do ścianek kryształu powoduje zmianę jego wymiarów. Usunięcie deformacji daje zanik ładunku.

Czujniki piezoelektryczne posiadają dużą rezystancję wewnętrzną. Napięciowe przetwarzanie sygnału wymaga wbudowania wzmacniacza możliwie blisko czujnika, powiem długie przewody doprowadzające zmieniają rezystancję i pojemność układu, zafałszowując sygnał z sensora.

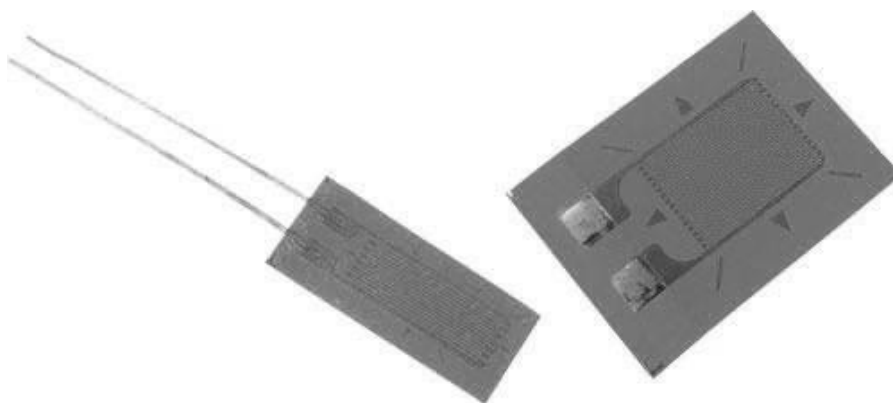


Rys. 16. Budowa i idea działania czujnika piezoelektrycznego

Źródło: <https://elektronikab2b.pl/technika/489-zastosowania-piezoelektrycznych-czujnikow-cisnienia>



Czujniki tensometryczne. Tensometrami nazywamy elementy rezystancyjne w postaci cienkich drutów, siatek, folii, wykonane z metalu lub półprzewodnika. Poddane działaniu sił zewnętrznych ulegają one odkształceniu zarówno w kierunku wzdłużnym, jak i prostopadłym do kierunku działania wypadkowej siły oraz zmieniają wymiary geometryczne i rezystancję.



Rys. 17. Wygląd czujników tensometrycznych

Źródło: <https://www.tranzystor.pl/artykuly-i-schematy/automatyka>

Zasada działania tensometru oporowego opiera się na właściwości fizycznej drutu metalowego, polegającej na zmianie jego rezystancji elektrycznej wraz ze zmianą długości. Drut oporowy (lub folia) naklejany jest z wykorzystaniem kleju na element odkształcający się pod wpływem działających sił lub momentów. Materiał oporowy czujnika ulega identycznym odkształceniom, co element, na którym czujnik został przyklejony.

Bibliografia

Engel Z.: Wkład Profesora Czesława Cempela w rozwój wibroakustyki, Diagnostyka 3(47)/2008.

<http://marcin.kielczewski.pracownik.put.poznan.pl/ZSP07.pdf>.

<https://elektronikab2b.pl/technika/489-zastosowania-piezoelektrycznych-czujnikow-cisnienia>.

<https://www.tranzystor.pl/artykuly-i-schematy/automatyka>.

Legutko S., Eksploatacja maszyn, Wydawnictwo Politechniki Poznańskiej, Poznań 2007.

Macha E. Niezawodność maszyn; Skrypt Nr 237 Politechnika Opolska.



Fundusze Europejskie
dla Wielkopolski

Dofinansowane przez
Unię Europejską



SAMORZĄD
WOJEWÓDZTWA
WIELKOPOLSKIEGO

Rakwicz B., Wojtynek R.: Zastosowanie metody ultradźwiękowej do oceny stanu złączy spawanych spawnicą sekcji ścianowej obudowy zmechanizowanej. *Maszyny Górnicze* 3/2013.

Zientek P.: Metody badań nieniszczących wybranych elementów konstrukcji turbosprężarki małej mocy. „*Maszyny Elektryczne – Zeszyty Problemowe*”, 3 (111)/2016.



Tomasz Kapłon

PODSTAWY AUTOMATYKI I ROBOTYKI**1. Podstawy automatyki**

Automatyka jest dziedziną wiedzy zajmującą się sterowaniem urządzeniami albo procesami bez albo z ograniczonym udziałem człowieka.

Sterowanie jest celowym oddziaływaniem na przebieg procesów technologicznych, ekonomicznych czy społecznych, tak aby uzyskać pożądaną przebieg lub cel.

Systemy binarne**Podstawy algebry Boole'a**

Podstawa teorii układów cyfrowych została podana przez George'a Boole'a jest to tak zwana algebra Boole'a. W jej przypadku liczby mogą przyjmować dwie wartości 0 albo 1. Występują także trzy podstawowe operacje logiczne:

- $+$ - dodawanie logiczne „lub”
- $\cdot$ - mnożenie logiczne „i”
- $\bar{}$ - (daszek nad liczbą) negacja logiczne „not”

Podstawowe aksjomaty można zapisać w postaci:

$$x + 0 = x \quad (1) \qquad x + x = x \quad (2)$$

$$x + 1 = 1 \quad (3) \qquad x + \bar{x} = 1 \quad (4)$$

$$x \cdot 0 = 0 \quad (5) \qquad x \cdot x = x \quad (6)$$

$$x \cdot 1 = x \quad (7) \qquad x \cdot \bar{x} = 0 \quad (8)$$

$$\bar{\bar{x}} = x \quad (9)$$

Natomiast podstawowe prawa to:

- Prawa pochłaniania

$$x_1 + x_1 x_0 = x_1 \quad (10) \qquad x_1(x_1 + x_0) = x_1 \quad (11)$$

$$x_1 + x_1 \bar{x}_0 = x_1 \quad (12) \qquad x_1(x_1 + \bar{x}_0) = x_1 \quad (13)$$

- Reguły sklejania

$$(x_1 + x_0)(x_1 + \bar{x}_0) = x_1 \quad (14) \qquad x_1 + \bar{x}_1 x_0 = x_1 + x_0 \quad (15)$$

$$x_1 x_0 + x_1 \bar{x}_0 = x_1 \quad (16) \qquad \bar{x}_1 + x_1 x_0 = \bar{x}_1 + x_0 \quad (17)$$



- Prawa De Morgana

$$\overline{x_1 + x_2} = \bar{x}_1 \cdot \bar{x}_2 \quad (18)$$

$$\overline{x_1 \cdot x_2} = \bar{x}_1 + \bar{x}_2 \quad (19)$$

Funkcje logiczne

Każdą funkcję logiczną można zdefiniować za pomocą tablicy lub równania. Przykładowo prostą funkcję logicznego dodawania można zapisać jako:

$$Y = x_1 + x_2 \quad (20)$$

Oraz opisać tablicą prawdy (rys. 1):

X	Y	Z
0	0	0
0	1	0
1	0	0
1	1	1

Rys. 1. Tablica prawdy dla funkcji $Y = x_1 + x_2$

W przypadku bardziej złożonych funkcji zawsze można je sprowadzić do postaci

Sumy iloczynów, przykładowo:

$$Y = x_1 x_2 \bar{x}_3 + x_1 x_2 + \bar{x}_1 \bar{x}_2 \quad (21)$$

Albo iloczynu sum, przykładowo:

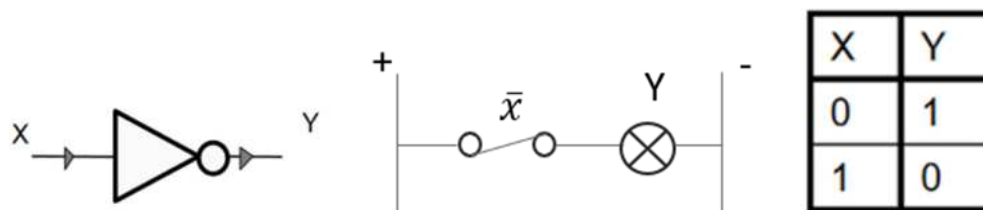
$$Z = (x_1 + \bar{x}_2 + \bar{x}_3) \cdot (\bar{x}_1 + x_2) \cdot (\bar{x}_1 + \bar{x}_2) \quad (22)$$

Układy kombinacyjne – podstawowe funkcje

Podstawowe funkcje logiczne są bardzo często realizowane w technice za pomocą elementów stykowych (najczęściej styków przekaźników i styczników), albo za pomocą elementów bezstykowych - elektronicznych bramek logicznych. Poniżej zostały przedstawione podstawowe funkcje logiczne i ich realizacja na bramkach oraz elementach stykowych.

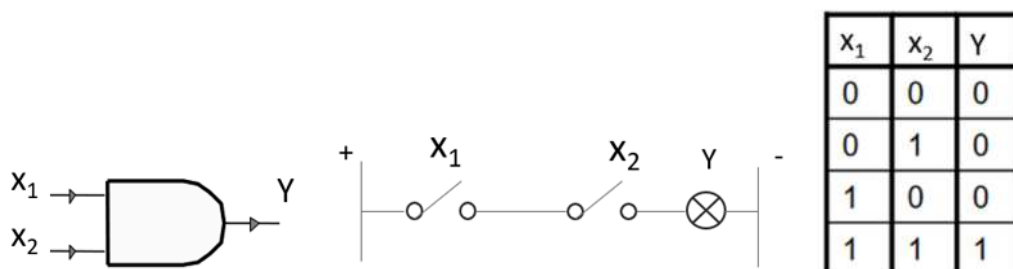


- Negacja - NOT $Y = \bar{x}$



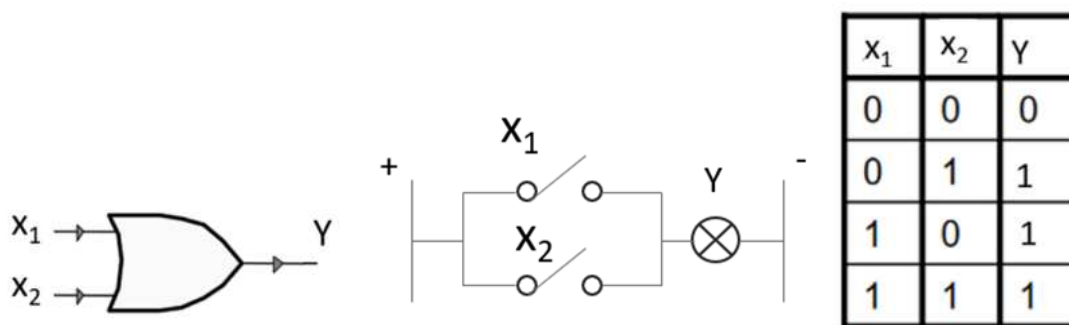
Rys. 2. Realizacja logicznej negacji na bramkach i elementach stykowych

- Iloczyn - AND $Y = x_1 \cdot x_2$



Rys. 3. Realizacja logicznego iloczynu na bramkach i elementach stykowych

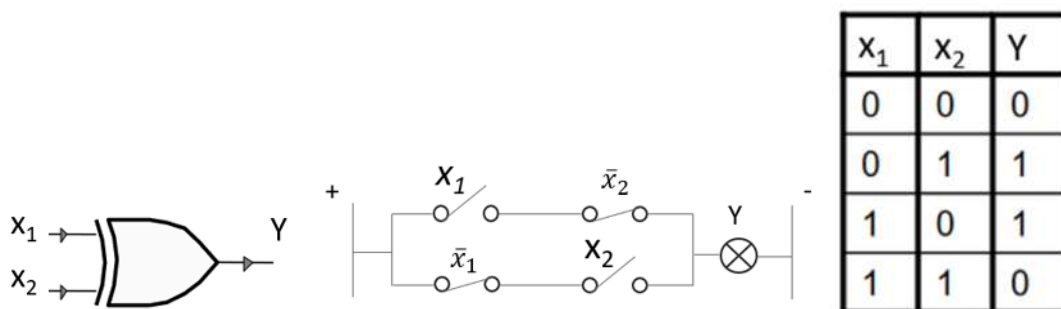
- Suma - alternatywa - OR $Y = x_1 + x_2$



Rys. 4. Realizacja logicznej sumy na bramkach i elementach stykowych

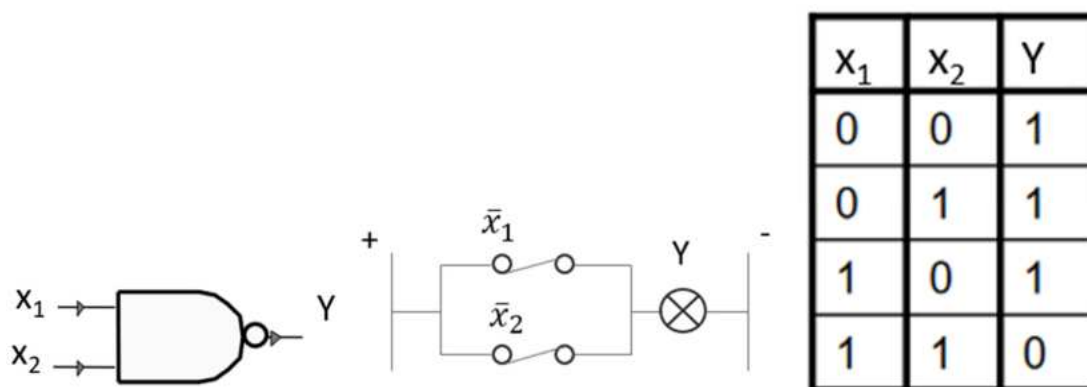


- $XOR\ Y = x_1 \oplus x_2$



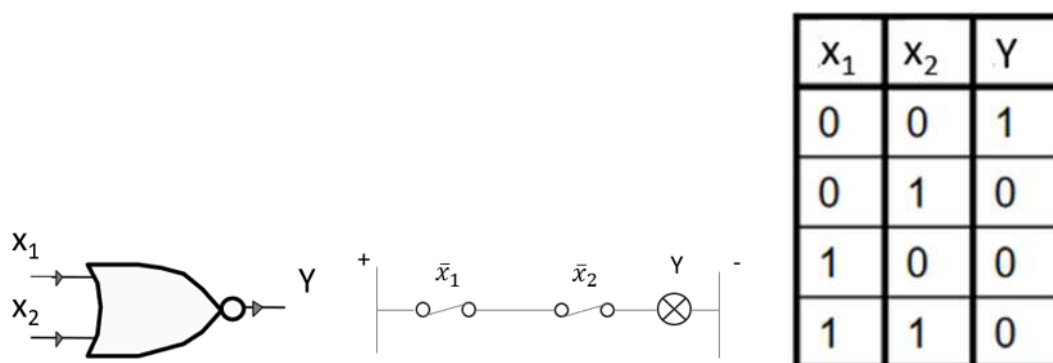
Rys. 5. Realizacja funkcji logicznej XOR na bramkach i elementach stykowych

- $NAND\ Y = \overline{x_1 \cdot x_2} = \overline{x_1} + \overline{x_2}$



Rys. 6. Realizacja zanegowanego iloczynu na bramkach i elementach stykowych

- $NOR\ Y = \overline{x_1 + x_2} = \overline{x_1} \cdot \overline{x_2}$



Rys. 7. Realizacja zanegowanej alternatywy na bramkach i elementach stykowych



- XNOR



x_1	x_2	Y
0	0	1
0	1	0
1	0	0
1	1	1

Rys. 8. Bramka logiczna funkcji XNOR

Opis układów automatyki

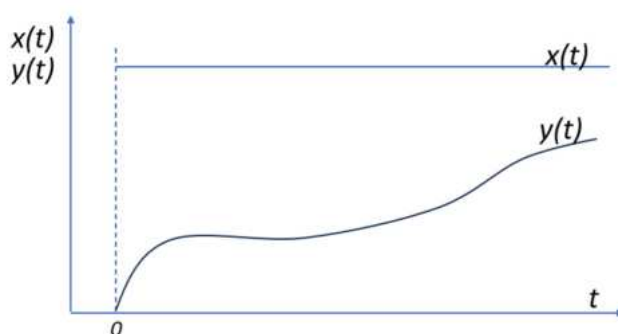
W ramach podstaw automatyki, zajmować będziemy się układami liniowymi stacjonarnymi. Takie układy można opisać za pomocą równań różniczkowych o stałych współczynnikach.

$$a_n \cdot \frac{d^n y(t)}{dt^n} + a_{n-1} \cdot \frac{d^{n-1} y(t)}{dt^{n-1}} + \dots + a_1 \cdot \frac{dy(t)}{dt} + a_0 \cdot y(t) = b_m \cdot \frac{d^m x(t)}{dt^m} + b_{m-1} \cdot \frac{d^{m-1} x(t)}{dt^{m-1}} \dots + b_1 \cdot \frac{dx(t)}{dt} + b_0 \cdot x(t) \quad (23)$$

$a_i \ i = 0, \dots, n$ – stałe współczynniki zależne od struktury i od wartości parametrów układu

$b_j \ j = 0, \dots, m$ – stałe współczynniki zależne od źródła sygnału wejściowego oraz od wartości parametrów układu i jego struktury

Jeżeli znamy powyższe równie oraz zależność sygnału wejściowego $x(t)$, to podstawiając $x(t)$, do powyższego równania można je rozwiązać i otrzymać zależność na sygnał wyjściowy $y(t)$. Dzięki temu możemy uzyskać charakterystykę dynamiczną obiektu. Przykładowy przebieg sygnału $y(t)$ będącego odpowiedzią na sygnał $x(t)$, przedstawia rysunek 9.



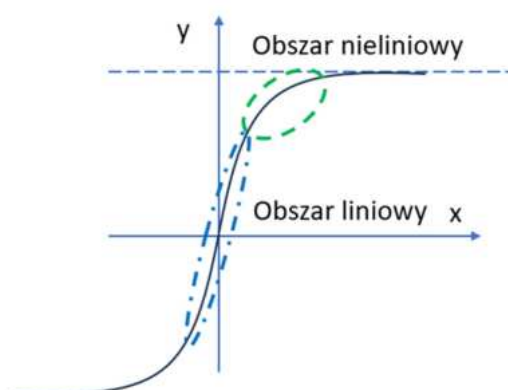
Rys. 9. Charakterystyka dynamiczna obiektu



Jeżeli chcemy wyznaczyć charakterystykę statyczną obiektu a więc w stanach ustalonych gdy $x(t) = \text{const}$ i $y(t) = \text{const}$. To wszystkie pochodne w równaniu należy przyrównać do zera. Wówczas uzyskamy:

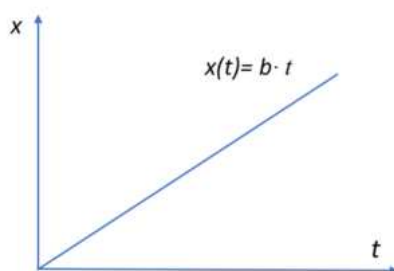
$$a_0 y = b_0 x \quad (24)$$

W takim wypadku otrzymana charakterystyka jest linią prostą, jednakże w rzeczywistych obiektach występuje zawsze pewna nieliniowość na skutek występowania nasycenia, tak jak na rysunku 10.



Rys. 10. Charakterystyka statyczna rzeczywistego obiektu

Rozwiązywanie równań różniczkowych jest w praktyce dość kłopotliwe, nawet przy stosunkowo prostych sygnałach wejściowych. Przykładowy sygnał wejściowy przedstawiony na rysunku 1, można go opisać poniższymi równaniami:



Rys. 11. Przebieg przykładowego sygnału wejściowego

$$x(t) = b \cdot t \quad (25)$$

$$\frac{dx(t)}{dt} = b \quad (26)$$

$$\frac{d^2x(t)}{dt^2} = 0 \quad (27)$$



Wówczas równanie opisujące układ przybierze postać:

$$a_n \cdot \frac{d^n y(t)}{dt^n} + a_{n-1} \cdot \frac{d^{n-1} y(t)}{dt^{n-1}} + \dots + a_1 \cdot \frac{dy(t)}{dt} + a_0 \cdot y(t) = b_1 \cdot b + b_0 \cdot b \cdot t \quad (28)$$

Jego rozwiązanie pozostaje dość kłopotliwe. Z tego powodu do opisu układów w automatyce stosuje się przekształcenie Laplace'a.

Przekształcenie Laplace'a

Przekształcenie Laplace'a definiuje się jako przyporządkowanie danej funkcji czasu $f(t)$ funkcji zmiennej zespolonej $F(s)$, która jest określona wzorem

$$F(s) = \int_0^{\infty} f(t) e^{-st} dt \quad (29)$$

Gdzie s jest zmienną zespoloną, operatorem

Operację przekształcenia Laplace'a, czyli transformatę Laplace'a, zapisujemy w sposób poniższy

$$\mathcal{L}[f(t)] = F(s) \quad (30)$$

Istnieje także odwrotne przekształcenie Laplace'a, przyporządkowującej funkcji w dziedzinie liczby zespolonej s , funkcję w dziedzinie czasu:

$$\mathcal{L}^{-1}[F(s)] = f(t) = \frac{1}{2\pi j} \cdot \int_{C-j\infty}^{C+j\infty} F(s) e^{st} ds \text{ dla } t > 0 \quad (31)$$

Wybrane właściwości transformaty Laplace'a:

$$\mathcal{L}[kf(t)] = k\mathcal{L}[f(t)] = kF(s) \quad (32)$$

$$\mathcal{L}[x(t) + y(t)] = X(s) + Y(s) \quad (33)$$

Przekształcenie Laplace'a pochodnej funkcji

$$\mathcal{L}\left[\frac{dy(t)}{dt}\right] = sY(s) - y(0) \quad (34)$$

$$\mathcal{L}\left[\frac{d^2}{dt^2} f(t)\right] = s^2 F(s) - sf(0) - f'(0) \quad (35)$$

Transformata całki

$$\mathcal{L}\left[\int_0^t f(\tau) d\tau\right] = \frac{F(s)}{s} \quad (36)$$

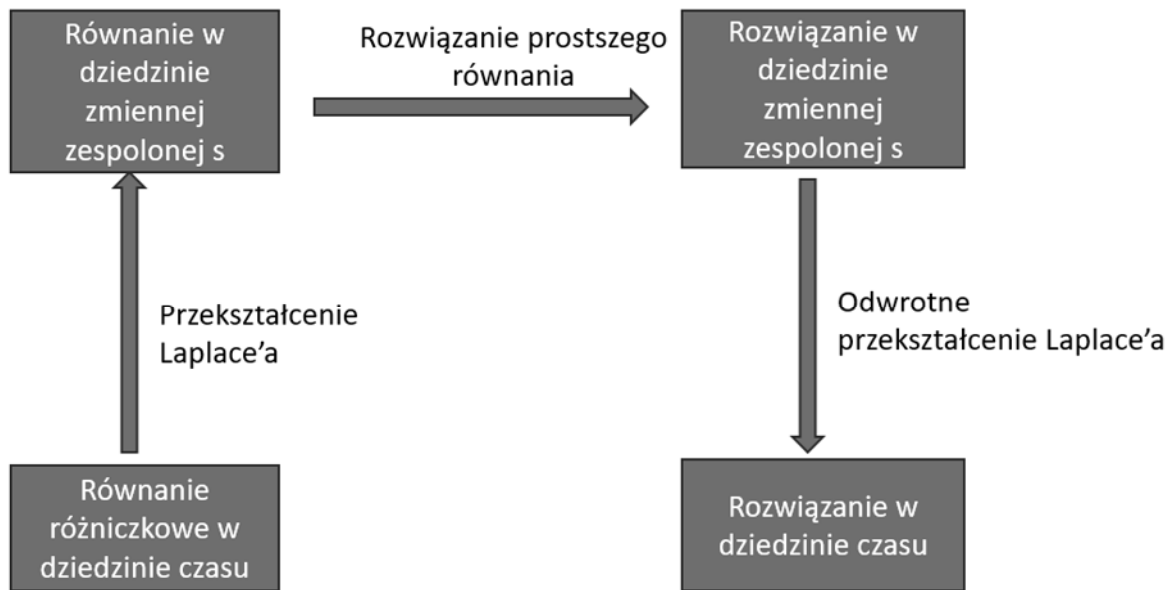
W praktyce dla większości typowych funkcji opracowano gotowe tablice, podobne do tych dla pochodnych i całek. W poniżej tabeli 1 zostało zebranych kilka typowych transformat.



Tablica 1. Przekształcenia Laplace'a niektórych funkcji

Funkcja $f(t)$	Transformata $F(s)$
$\delta(t)$	1
$1(t)$	$\frac{1}{s}$
kt	$\frac{k}{s^2}$
$\frac{kt^2}{2}$	$\frac{k}{s^3}$
$\frac{kt^n}{n!}$	$\frac{k}{s^{n+1}}$
ke^{at}	$\frac{k}{s-a}$
$\frac{k}{a}e^{-\frac{t}{a}}$	$\frac{k}{as+1}$
$\sin \omega t$	$\frac{\omega}{s^2 + \omega^2}$
$\cos \omega t$	$\frac{s}{s^2 + \omega^2}$
$e^{at} \sin \omega t$	$\frac{\omega}{(s-a)^2 + \omega^2}$
$e^{at} \cos \omega t$	$\frac{s-a}{(s-a)^2 + \omega^2}$
$1 - e^{-\frac{t}{a}}$	$\frac{1}{s(as+1)}$
$\frac{1}{a^2}e^{-\frac{t}{a}}$	$\frac{1}{s(as+1)^2}$
$\frac{e^{-\frac{t}{a}} - e^{-\frac{t}{b}}}{a-b}$	$\frac{1}{(as+1)(bs+1)}$
$\frac{1}{a^3}(a-t)e^{-\frac{t}{a}}$	$\frac{s}{(as+1)^2}$

Stosowanie przekształcenia Laplace'a ma swoje uzasadnienie w tym, że równania różniczkowe opisujące dany układ po poddaniu przekształceniu są znacznie prostsze do rozwiązania w tej formie. Z tego powodu, zazwyczaj przy ich rozwiązywaniu postępuje się zgodnie z poniższym schematem (rys. 12).



Rys. 12. Zastosowanie przekształcenia Laplace'a do rozwiązywania równań różniczkowych

Transmitancja operatorowa i widmowa

Równanie wejścia – wyjścia układu liniowego o jednym wejściu i wyjściu można poddać przekształceniu Laplace'a. Wówczas równanie:

$$a_n \cdot \frac{d^n y(t)}{dt^n} + a_{n-1} \cdot \frac{d^{n-1} y(t)}{dt^{n-1}} + \dots + a_1 \cdot \frac{dy(t)}{dt} + a_0 \cdot y(t) = b_m \cdot \frac{d^m x(t)}{dt^m} + b_{m-1} \cdot \frac{d^{m-1} x(t)}{dt^{m-1}} \dots + b_1 \cdot \frac{dx(t)}{dt} + b_0 \cdot x(t) \quad (37)$$

Przy założonych zerowych warunkach początkowych:

$$y(0) = \dot{y}(0) = \dots = y^{n-1} = 0 \quad (38)$$

$$x(0) = \dot{x}(0) = \dots = x^{m-1} = 0 \quad (39)$$

Po wykonaniu przekształcenia można otrzymać równanie w postaci operatorowej:

$$a_n \cdot s^n Y(s) + a_{n-1} \cdot s^{n-1} Y(s) + \dots + a_1 \cdot s Y(s) + a_0 \cdot Y(s) = b_m \cdot s^m X(s) + b_{m-1} \cdot s^{m-1} X(s) \dots + b_1 \cdot s X(s) + b_0 X(s) \quad (40)$$

Stąd:

$$Y(s) = \frac{b_m \cdot s^m + b_{m-1} \cdot s^{m-1} \dots + b_1 \cdot s + b_0 X}{a_n \cdot s^n + a_{n-1} \cdot s^{n-1} + \dots + a_1 \cdot s + a_0} X(s) \quad (41)$$

Albo:

$$Y(s) = G(s) \cdot X(s) \quad (42)$$

Przy czym:

$$G(s) = \frac{b_m \cdot s^m + b_{m-1} \cdot s^{m-1} \dots + b_1 \cdot s + b_0 X}{a_n \cdot s^n + a_{n-1} \cdot s^{n-1} + \dots + a_1 \cdot s + a_0} \quad (43)$$



Z powyższych wzorów wynika że:

$$G(s) = \frac{Y(s)}{X(s)} \quad (44)$$

Powyższą funkcję operatorową $G(s)$, nazywamy **transmitancją operatorową** obiektu liniowego. Jest ona zdefiniowana jako stosunek transformaty Laplace'a sygnału wyjściowego – $Y(s)$, do transformaty Laplace'a sygnału wejściowego – $X(s)$, przy zerowych warunkach początkowych.

Jeżeli w miejsce operatora s , podstawimy liczbę zespoloną $j\omega$ gdzie j - to liczba urojona, a ω to pulsacja (częstotliwość) to otrzymamy:

$$G(j\omega) = \frac{Y(j\omega)}{X(j\omega)} \quad (45)$$

Jest to wielkość zespolona nazywana transmitancją widmową. Jest ona określona przez stosunek sinusoidalnego sygnału wyjściowego zapisanego w postaci zespolonej do sinusoidalnego sygnału wejściowego zapisanego w postaci zespolonej przy zerowych warunkach początkowych. Transmitancję widmową można zapisać w postaci sumy części rzeczywistej i urojonej:

$$G(j\omega) = \text{Re}[G(j\omega)] + j\text{Im}[G(j\omega)] \quad (46)$$

Albo w postaci wykładniczej

$$G(j\omega) = |G(j\omega)|e^{j\varphi(\omega)} \quad (47)$$

Gdzie moduł (amplituda):

$$|G(j\omega)| = \sqrt{\{\text{Re}[G(j\omega)]\}^2 + \{\text{Im}[G(j\omega)]\}^2} \quad (53)$$

Natomiast faza (argument):

$$\varphi(\omega) = \arctg \frac{\text{Im}[G(j\omega)]}{\text{Re}[G(j\omega)]} \quad (48)$$

Amplituda określa dla danej pulsacji stosunek sygnału amplitudy sygnału wyjściowego do amplitudy sygnału wejściowego. Natomiast faza określa przesunięcie fazowe między sinusoidą wejściową i wyjściową.

Podstawowe charakterystyki

W celu zbadania danego elementu, zadaje się mu pewien zdefiniowany sygnał, a następnie rejestruje się jego odpowiedź w czasie. Otrzymuje się w ten sposób charakterystykę dynamiczną w odpowiedzi na zdefiniowany sygnał. Zależnie od typu sygnału charakterystyki te są określane jako skokowa, impulsowa oraz liniowa.



- 1) Charakterystyka skokowa jest odpowiedzią na wymuszenie w postaci skoku jednostkowego $-I(t)$, zdefiniowanego jako:

$$1(t) = \begin{cases} 0 & \text{dla } t < 0 \\ 1 & \text{dla } t \geq 0 \end{cases} \quad (49)$$

- 2) Charakterystyka impulsowa jest odpowiedzią na wymuszenie w postaci delty Diraca $-\delta(t)$, zdefiniowanej jako:

$$\delta(t) = \begin{cases} 0 & \text{dla } t \neq 0 \\ +\infty & \text{dla } t = 0 \end{cases} \quad (50)$$

- 3) Charakterystyka liniowa jest odpowiedzią na wymuszenie liniowe określone jako:

$$x(t) = a \cdot t \quad (51)$$

Znając transmitancję operatorową danego elementu, można wyznaczyć jego charakterystykę mnożąc jego transmitancję operatorową przez transformatę Laplace'a sygnału wymuszającego, a następnie wynik tej operacji poddając odwrotnemu przekształceniu Laplace'a

Podstawowe człony automatyki

W układach automatyki można wyodrębnić kilka grup podstawowych elementów o charakterystycznych cechach. Są one podstawą do budowania układów bardziej złożonych składających się z wielu podstawowych elementów. Poniżej zostaną przedstawione te podstawowe elementy.

Człon proporcjonalny

Równanie wejścia-wyjścia elementu proporcjonalnego można przedstawić w postaci:

$$a_0 y(t) = b_0 u(t) \quad (52)$$

Po podstawieniu $k=b_0/a_0$, przy czym k nazywamy współczynnikiem wzmocnienia, albo po prostu wzmocnieniem:

$$y(t) = ku(t) \quad (53)$$

Transmitancja operatorowa członu operatorowego ma postać:

$$G(s) = k \quad (54)$$

Równania opisujące odpowiedź na podstawowe sygnały wymuszające mają następującą formę:



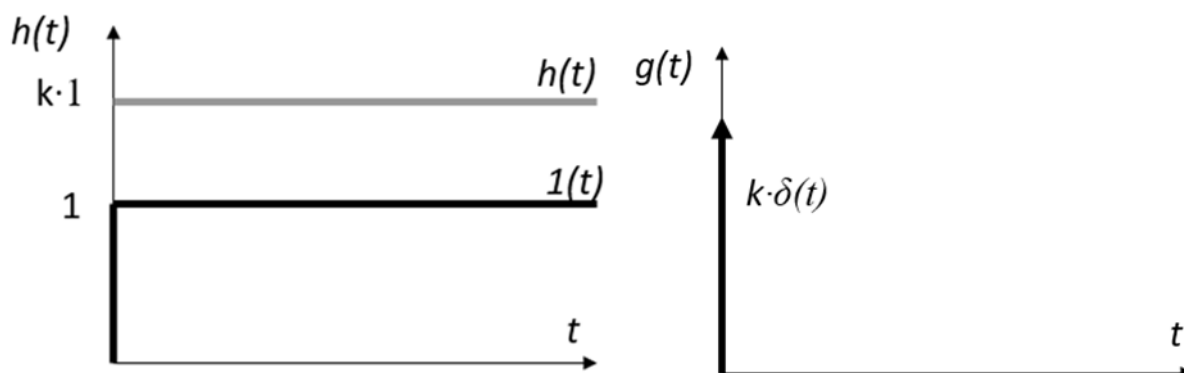
- Odpowiedź impulsowa:

$$g(t) = k\delta(t) \quad (55)$$

- Odpowiedź skokowa:

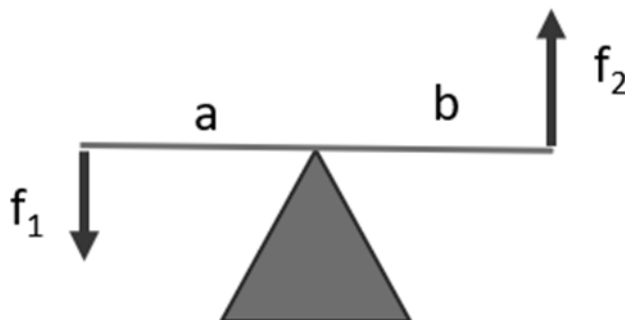
$$h(t) = k \cdot 1(t) \quad (56)$$

Odpowiedzi na te podstawowe sygnały przedstawia rysunek 13.



Rys. 13. Odpowiedzi skokowa $h(t)$ i impulsowa $g(t)$ członu proporcjonalnego.

Przykładem elementu proporcjonalnego może być dźwignia taka jak na rysunku 14 gdzie sygnałem wejściowym jest siła f_1 na końcu ramienia a , natomiast siła f_2 na końcu ramienia b jest sygnałem wyjściowym. Wówczas możemy zależność między siłami opisać jako:



Rys. 14. Dźwignia jako przykład elementu proporcjonalnego

$$f_1(t) \cdot a = f_2(t) \cdot b \quad (57)$$

$$f_2(t) = \frac{b}{a} f_1(t) = k f_1(t) \quad (58)$$

Po wykonaniu transformacji Laplace'a, oraz wiedząc, że transmitancja to stosunek sygnału wyjściowego do wejściowego:

$$G(s) = \frac{F_2(s)}{F_1(s)} = k \quad (59)$$



Człon inercyjny I-go rzędu

Równanie wejścia-wyjścia elementu inercyjnego I-go rzędu można zapisać jako:

$$a_1 \frac{dy}{dt} + a_0 y(t) = b_0 u(t) \quad (60)$$

Po dokonaniu podstawień $T = \frac{a_1}{a_0}$ i $k = \frac{b_0}{a_0}$ gdzie T nazywamy stałą czasową elementu,

a k - wzmocnieniem, można to równanie zapisać:

$$T \frac{dy}{dt} + y(t) = ku(t) \quad (61)$$

Po dokonaniu transformacji Laplace'a dla zerowych warunków początkowych można otrzymać:

$$TsY(s) + Y(s) = kU(s) \quad (62)$$

Wiedząc, że transmitancja to stosunek sygnału wyjściowego- $Y(s)$, do wejściowego- $U(s)$ otrzymujemy:

$$G(s) = \frac{Y(s)}{U(s)} = \frac{k}{Ts+1} \quad (63)$$

Równanie odpowiedzi impulsowej w dziedzinie czasu możemy poznać najpierw mnożąc transmitancję członu inercyjnego przez funkcję opisującą delte Diraca w dziedzinie liczby zespolonej s , a następnie wykonując odwrotne przekształcenie Laplace'a, czyli:

$$g(t) = \mathcal{L}^{-1}[G(s) \cdot 1] = \mathcal{L}^{-1}\left[\frac{k}{Ts+1}\right] = \frac{k}{T} e^{-\frac{t}{T}} \quad (64)$$

W podobny sposób, można wyprowadzić wzór na odpowiedź skokową:

$$h(t) = \mathcal{L}^{-1}\left[G(s) \cdot \frac{1}{s}\right] = \mathcal{L}^{-1}\left[\frac{k}{s(Ts+1)}\right] = k \left(1 - e^{-\frac{t}{T}}\right) \quad (65)$$

Podstawiając do wzoru na odpowiedź skokową $t=T$ otrzymujemy:

$$h(T) = k \left(1 - e^{-\frac{T}{T}}\right) \quad (66)$$

$$h(T) = k(1 - e^{-1}) \quad (67)$$

a więc:

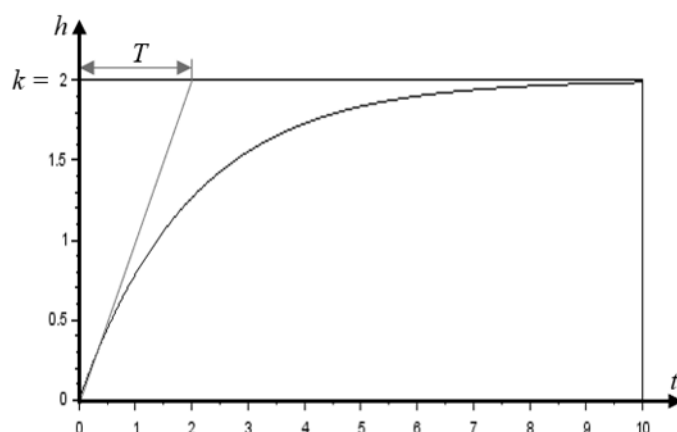
$$h(T) \approx 0,63k \quad (68)$$

Tak więc T , jest to czas po którym odpowiedź elementu inercyjnego I-go rzędu osiągnie wartość $0,63k$. Przy czym k jest wartością maksymalną, do której dąży w nieskończonym czasie.

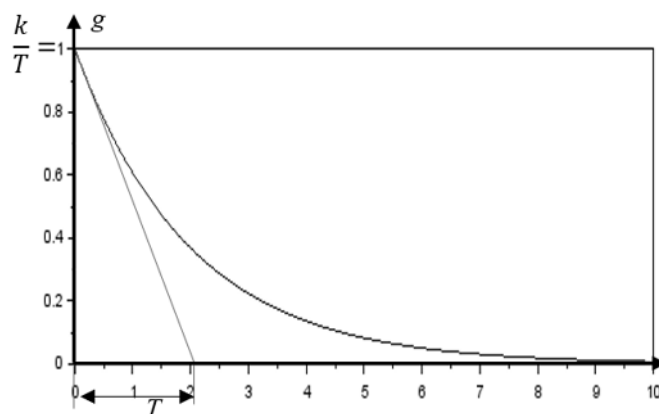
Na poniższych rysunkach 15 i 16 przedstawiono odpowiedzi skokową i impulsową dla członu inercyjnego pierwszego rzędu o transmitancji

$$G(s) = \frac{2}{2s+1} \quad (69)$$

a więc o $k = 2$ i $T = 2$.

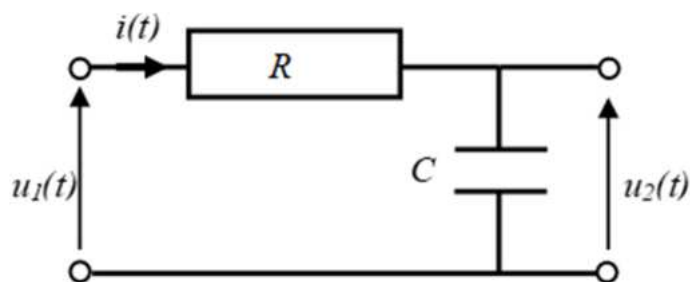


Rys. 14. Odpowiedź skokowa członu inercyjnego I-go rzędu o transmitancji $(s) = \frac{2}{2s+1}$



Rys. 15. Odpowiedź impulsowa członu inercyjnego I-go rzędu o transmitancji $(s) = \frac{2}{2s+1}$

Przykładem członu inercyjnego pierwszego rzędu może być filtr RC, taki jak na rysunku 16, dla którego sygnałem wejściowym jest napięcie $u_1(t)$, a wyjściowym $u_2(t)$,



Rys. 16. Filtr RC jako przykład członu inercyjnego I-go rzędu.

Dla takiego układu można zapisać, przy przyjętych zerowych warunkach początkowych:

$$u_2(t) = u_c(t) = \frac{1}{C} \int_0^T i(\tau) d\tau \quad (70)$$

$$u_1(t) = u_R(t) + u_c(t) = i(t) \cdot t + \frac{1}{C} \int_0^T i(\tau) d\tau \quad (71)$$



Po dokonaniu transformacji Laplace'a obu równań:

$$U_2(s) = \frac{1}{C} \frac{I(s)}{s} \quad (72)$$

$$U_1(s) = I(s)R + \frac{1}{C} \frac{I(s)}{s} \quad (73)$$

Ponieważ transmitancja jest stosunkiem sygnału wyjściowego do wejściowego:

$$G(s) = \frac{U_2(s)}{U_1(s)} = \frac{\frac{1}{C} \frac{I(s)}{s}}{(s)R + \frac{1}{C} \frac{I(s)}{s}} = \frac{\frac{1}{Cs}}{R + \frac{1}{Cs}} = \frac{1}{RCs + 1} \quad (74)$$

Tak więc otrzymuje się transmitancję członu inercyjnego I-go rzędu o $k=1$ i $T=RC$

Człon całkujący idealny

Element całkujący idealny, albo element całkujący bezinercyjny posiada równanie wejścia-wyjścia w postaci:

$$a_1 \frac{dy}{dt} = b_0 u(t) \quad (75)$$

Po dokonaniu podstawienia $k_c = \frac{b_0}{a_1}$ gdzie k_c to współczynnik wzmocnienia i wykorzystując całkę:

$$y(t) = k_c \int_0^t u(\tau) d\tau \quad (76)$$

Następnie wykonując transformację Laplace'a przy zerowych warunkach początkowych, można zapisać

$$Y(s) = \frac{k_c}{s} U(s) \quad (77)$$

Ostatecznie można uzyskać transmitancję operatorową:

$$G(s) = \frac{Y(s)}{U(s)} = \frac{k_c}{s} \quad (78)$$

Równanie odpowiedzi impulsowej ma postać:

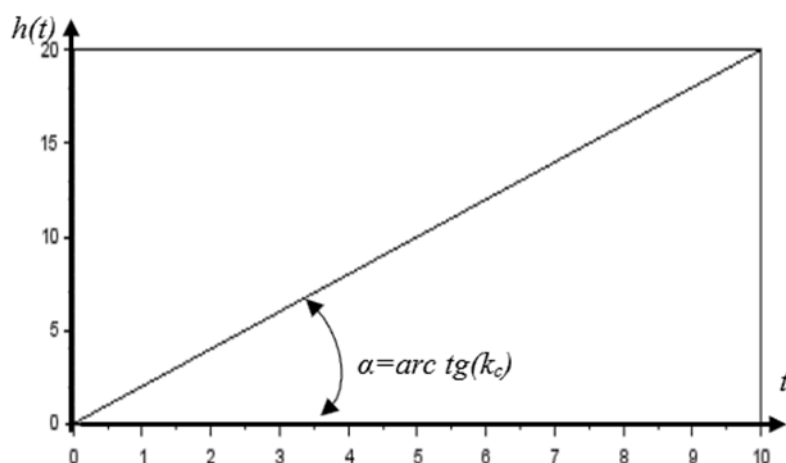
$$g(t) = k_c \cdot 1(t) \quad (79)$$

Natomiast odpowiedzi skokowej:

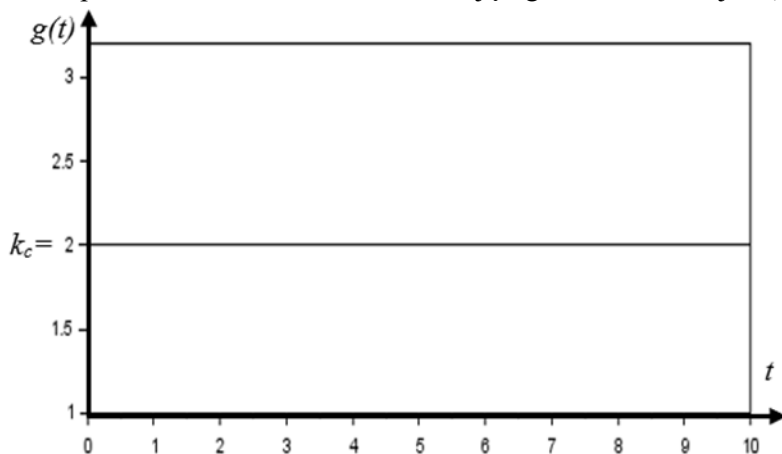
$$h(t) = k_c t \quad (80)$$

Przykładowe odpowiedzi skokowe i impulsowe dla elementu całkującego o transmitancji

$G(s) = \frac{2}{s}$ przedstawiają rysunki 17 i 18.

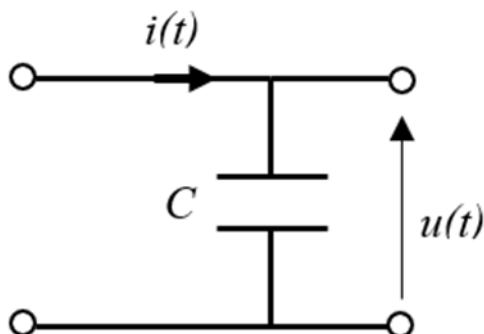


Rys. 17. Odpowiedź skokowa elementu całkującego o transmitancji $G(s) = \frac{2}{s}$



Rys. 18. Odpowiedź impulsowa elementu całkującego o transmitancji $G(s) = \frac{2}{s}$

Przykładem idealnego układu całkującego może być idealny kondensator, taki jak na rysunku 19, dla którego sygnałem wejściowym jest natężenie prądu wpływającego do kondensatora $i(t)$, a sygnałem wyjściowym napięcie na kondensatorze - $u(t)$.



Rys. 19. Idealny kondensator jako człon całkujący idealny



Dla takiego elementu można zapisać:

$$u(t) = \frac{1}{c} \int_0^T i(\tau) d\tau \quad (81)$$

$$U(s) = \frac{1}{c} \frac{I(s)}{s} \quad (82)$$

Ostatecznie transmitancję można zapisać jako:

$$G(s) = \frac{U(s)}{I(s)} = \frac{1}{cs} \quad (83)$$

Człon całkujący rzeczywisty

Element całkujący rzeczywisty albo inaczej całkujący z inercją można opisać równaniem:

$$a_2 \frac{d^2 y}{dt^2} + a_1 \frac{dy}{dt} = b_0 u(t) \quad (84)$$

Po dokonaniu podstawień $T = \frac{a_2}{a_1}$ i $k_c = \frac{b_0}{a_1}$ gdzie T to stała czasowa, a k_c to wzmocnienie, można zapisać

$$T \frac{d^2 y}{dt^2} + \frac{dy}{dt} = k_c u(t) \quad (85)$$

Po wykonaniu przekształcenia Laplace'a otrzymuje się:

$$Ts^2 Y(s) + sY(s) = k_c U(s) \quad (86)$$

Ostatecznie transmitancja przybiera postać:

$$G(s) = \frac{Y(s)}{U(s)} = \frac{k_c}{s(Ts+1)} \quad (87)$$

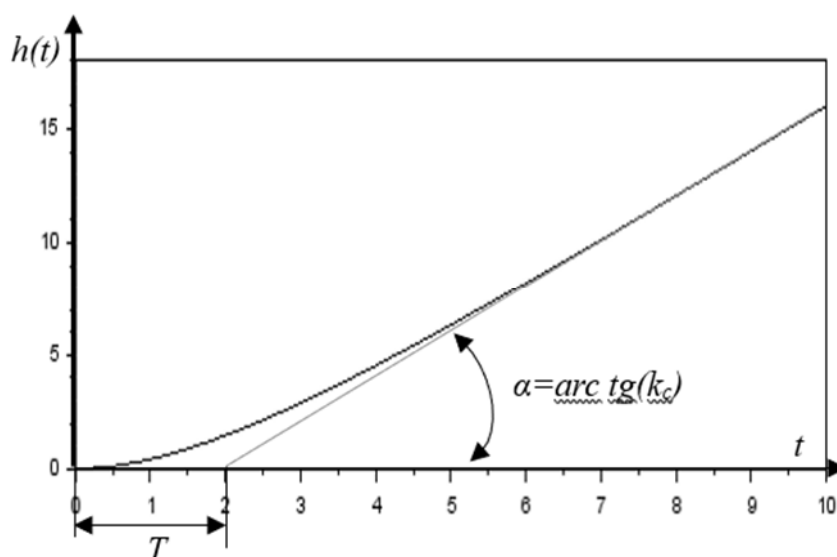
Równanie odpowiedzi impulsowej ma postać:

$$g(t) = k_c \left(1 - e^{-\frac{t}{T}} \right) \quad (88)$$

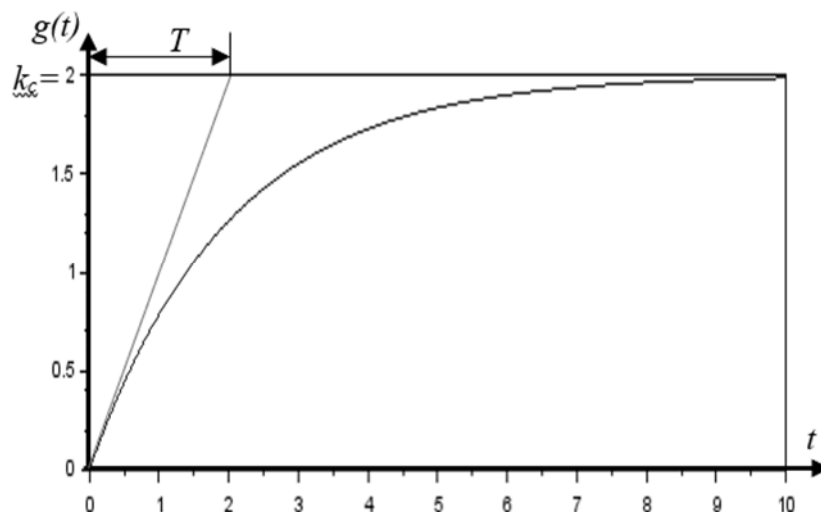
Natomiast odpowiedź skokowa:

$$h(t) = k_c \left[t - T \left(1 - e^{-\frac{t}{T}} \right) \right] \quad (89)$$

Przykładowe odpowiedzi impulsowe i skokowe dla elementu całkującego z inercją o transmitancji $G(s) = \frac{2}{s(2s+1)}$ przedstawiają rysunki 20 i 21.



Rys. 20. Odpowiedź skokowa elementu całkującego rzeczywistego o transmitancji $G(s) = \frac{2}{s(2s+1)}$



Rys. 21. Odpowiedź impulsowa elementu całkującego rzeczywistego o transmitancji $G(s) = \frac{2}{s(2s+1)}$

Człon różniczkujący idealny

Element różniczkujący idealny można opisać równaniem:

$$a_0 y(t) = b_1 \frac{du}{dt} \quad (90)$$

A po dokonaniu podstawienia $k_r = \frac{b_1}{a_0}$, gdzie k_r jest wzmocnieniem można zapisać:

$$y(t) = k_r \frac{du}{dt} \quad (91)$$

Po zastosowaniu przekształcenia Laplace'a:

$$Y(s) = k_r s U(s) \quad (92)$$

Ostatecznie można otrzymać transmitancję:

$$G(s) = k_r s \quad (93)$$

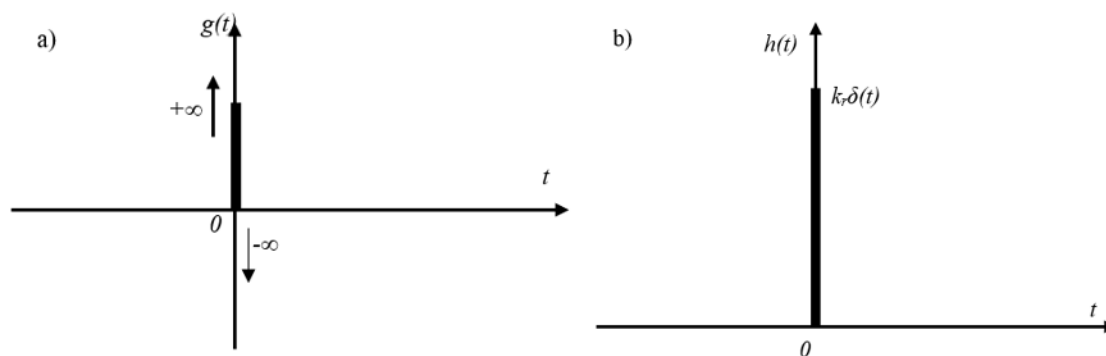
Równanie odpowiedzi impulsowej ma postać:

$$g(t) = k_r \frac{d\delta(t)}{dt} \quad (94)$$

Natomiast odpowiedzi skokowej:

$$h(t) = k_r \delta(t) \quad (95)$$

Rysunki 22a i 22b przedstawiają przykładowe odpowiedzi impulsową i skokową członu różniczkującego idealnego.



Rys. 22. Odpowiedzi elementu różniczkującego idealnego a-impulsowa, b-skokowa

W praktyce niemożliwe jest zbudowanie idealnego członu różniczkującego. Może on jednak znaleźć zastosowanie w uproszczonych modelach symulacyjnych.

Człon różniczkujący rzeczywisty

Element różniczkujący rzeczywisty, albo inaczej różniczkujący z inercją można opisać równaniem:

$$a_1 \frac{dy}{dt} + a_0 y(t) = b_1 \frac{du}{dt} \quad (96)$$

Po dokonaniu podstawień $T = \frac{a_1}{a_0}$ i $k_r = \frac{b_1}{a_0}$

Można otrzymać równanie:

$$T \frac{dy}{dt} + y(t) = k_r \frac{du}{dt} \quad (97)$$

Po zastosowaniu przekształcenia Laplace'a otrzymujemy równanie w formie operatorowej:

$$TsY(s) + Y(s) = k_r \frac{du}{dt} \quad (98)$$

Ostatecznie można z niego wyprowadzić transmitancje:

$$G(s) = \frac{Y(s)}{U(s)} = \frac{k_r s}{Ts + 1} \quad (99)$$



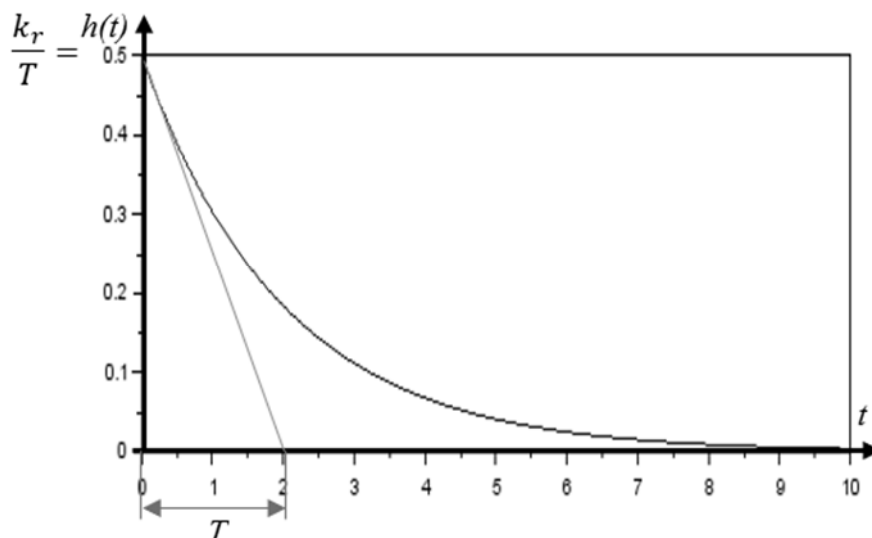
Równanie odpowiedzi impulsowej członu różniczkującego rzeczywistego ma postać:

$$g(t) = \frac{k_r}{T} \delta(t) - \frac{k_r}{T^2} e^{-\frac{t}{T}} \quad (100)$$

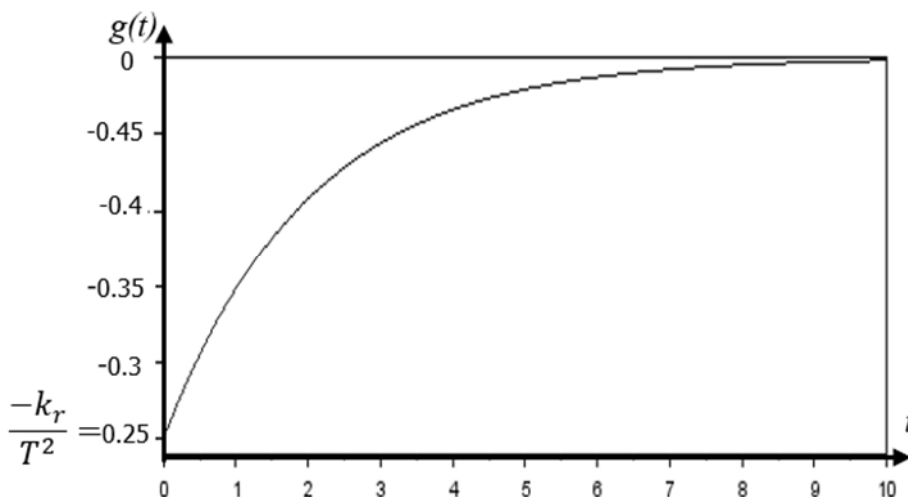
Odpowiedź skokowa natomiast:

$$h(t) = \frac{k_r}{T} e^{-\frac{t}{T}} \quad (101)$$

Przykładowe odpowiedzi skokowe i impulsowe członu różniczkującego rzeczywistego o transmitancji $G(s) = \frac{s}{2s+1}$ przedstawiono na rysunkach 23 i 24.



Rys. 23. Odpowiedź skokowa członu różniczkującego rzeczywistego o transmitancji $G(s) = \frac{s}{2s+1}$



Rys. 24. Odpowiedź impulsowa członu różniczkującego rzeczywistego o transmitancji $G(s) = \frac{s}{2s+1}$

Człon inercyjny II-go rzędu

Człon inercyjny II-go rzędu posiada następujące równanie wejścia-wyjścia:



$$a_2 \frac{d^2 y}{dt^2} + a_1 \frac{dy}{dt} + a_0 y(t) = b_0 u(t) \quad (102)$$

Po dokonaniu transformacji Laplace'a otrzymujemy

$$a_2 s^2 Y(s) + a_1 s Y(s) + a_0 Y(s) = b_0 U(s) \quad (103)$$

Następnie dzielą sygnał wyjściowy przez wejściowy otrzymujemy:

$$G(s) = \frac{Y(s)}{U(s)} = \frac{b_0}{a_2 s^2 + a_1 s + a_0} = \frac{\frac{b_0}{a_0}}{\frac{a_2}{a_0} s^2 + \frac{a_1}{a_0} s + 1} \quad (104)$$

Następnie podstawiając $T_1 = \frac{a_2}{a_0}$, $T_2 = \frac{a_1}{a_0}$, oraz $k = \frac{b_0}{a_0}$ otrzymujemy transmitancję w postaci:

$$G(s) = \frac{k}{T_1 s^2 + T_2 s + 1} \quad (105)$$

Jeżeli równanie charakterystyczne, czyli wyrażenie z mianownika przyrównane do zera posiada pierwiastki w dziedzinie liczb rzeczywistych, to jest to transmitancja członu inercyjnego II-go rzędu.

Oznacza to, że jeżeli dla:

$$T_1 s^2 + T_2 s + 1 = 0 \quad \Delta \geq 0 \quad (106)$$

To można znaleźć pierwiastki s_1 i s_2 . Wówczas mianownik można zapisać w postaci iloczynu podwójnego:

$$G(s) = \frac{k}{(s-s_1)(s-s_2)} = \frac{k}{(-s_1)(\frac{1}{-s_1}s+1)(-s_2)(\frac{1}{-s_2}s+1)} = \frac{\frac{k}{s_1 s_2}}{(T_{11}s+1)(T_{21}s+1)} \quad (107)$$

Ostatecznie otrzymuje się transmitancję o nowych stałych czasowych T_{11} i T_{21} i wzmacnieniu k_1 :

$$G(s) = \frac{k_1}{(T_{11}s+1)(T_{21}s+1)} \quad (108)$$

Jest to postać połączonych szeregowo dwóch członów inercyjnych.

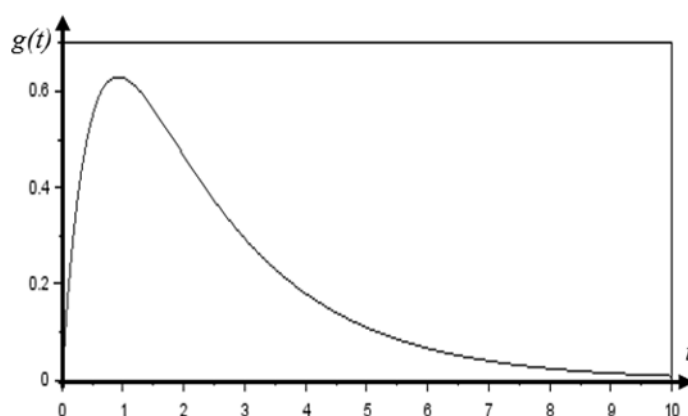
Odpowiedź impulsowa członu inercyjnego drugiego rzędu ma postać:

$$g(t) = k \left[\frac{1}{T_{11}-T_{21}} e^{-\frac{t}{T_{11}}} - \frac{1}{T_{11}-T_{21}} e^{-\frac{t}{T_{21}}} \right] \quad (109)$$

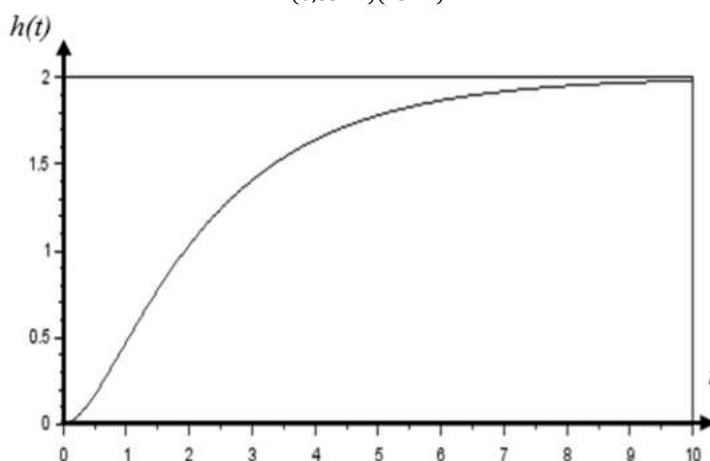
Natomiast odpowiedź skokowa jest opisana równaniem:

$$h(t) = k \left[1 - \frac{T_{11}}{T_{11}-T_{21}} e^{-\frac{t}{T_{11}}} + \frac{T_{21}}{T_{11}-T_{21}} e^{-\frac{t}{T_{21}}} \right] \quad (110)$$

Przykłady odpowiedzi impulsowych i skokowych dla elementu inercyjnego II-go rzędu o transmitancji $G(s) = \frac{2}{(0,5s+1)(2s+1)}$ przedstawiają rysunki 25 i 26. Warto zauważyć, że przebieg odpowiedzi skokowej elementu inercyjnego II-go rzędu posiada dwa punkty przegięcia, podczas, kiedy element I-go rzędu tylko jeden.



Rys. 25. Odpowiedź impulsowa elementu inercyjnego II-go rzędu o transmitancji $G(s) = \frac{2}{(0,5s+1)(2s+1)}$



Rys. 26. Odpowiedź skokowa elementu inercyjnego II-go rzędu o transmitancji $G(s) = \frac{2}{(0,5s+1)(2s+1)}$

Człon oscylacyjny

Człon oscylacyjny jest opisany takim samym równaniem wejścia-wyjścia jak człon inercyjny II-go rzędu:

$$a_2 \frac{d^2 y}{dt^2} + a_1 \frac{dy}{dt} + a_0 y(t) = b_0 u(t) \quad (111)$$

Jednakże w jego wypadku równanie charakterystyczne transmitancji:

$$G(s) = \frac{k}{T_1 s^2 + T_2 s + 1} \quad (112)$$

Nie posiada pierwiastków w dziedzinie liczb rzeczywistych czyli dla:

$$T_1 s^2 + T_2 s + 1 = 0 \quad \Delta < 0 \quad (113)$$

W takim wypadku wygodniej jest się posługiwać innym zapisem transmitancji

$$G(s) = \frac{k \omega_n^2}{s^2 + 2\zeta \omega_n s + \omega_n^2} \quad (114)$$



Gdzie:

k - współczynnik wzmocnienia członu

ω_n - pulsacja drgań własnych nietłumionych

ζ - zredukowany współczynnik tłumienia

W przypadku odpowiedzi skokowej tego członu odpowiedź ma kształt oscylacji tłumionych o stałym okresie drgań. Pulsacja drgań tłumionych – ω , ma wartość:

$$\omega = \omega_n \sqrt{1 - \zeta^2} \quad (115)$$

A okres tych drgań wynosi $T = \frac{2\pi}{\omega}$. Gdy $\zeta = 1$ odpowiedź jest bezoscylacyjna, natomiast im ζ jest mniejsze, tym większe są oscylacje, oraz ich okres jest krótszy. Skrajnym przypadkiem, jest nierzeczywisty wariant gdy $\zeta = 0$. Wystąpiłby wówczas stan braku tłumienia. W układach mechanicznych byłyby to układy bez tarcia, a w elektrycznych układy bez rezystancji. Przypadek dla $-1 < \zeta < 0$, oznaczałby, że oscylacje rosłyby do nieskończoności bez żadnego czynnika wzmacniającego je. Jest to przypadek możliwy tylko matematycznie.

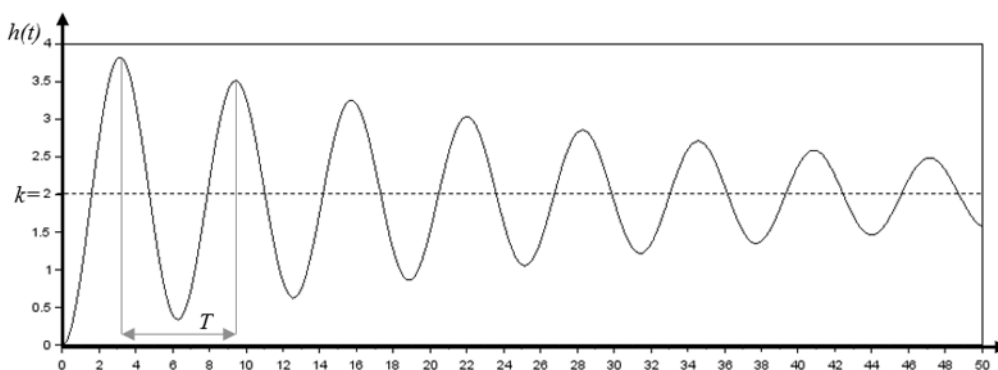
Odpowiedź impulsowa może zostać opisana wzorem:

$$h(t) = \left[1 - e^{-\zeta \omega_n t} \left(\cos(\omega_n \sqrt{1 - \zeta^2} t) + \frac{\zeta}{\sqrt{1 - \zeta^2}} \sin(\omega_n \sqrt{1 - \zeta^2} t) \right) \right] \quad (116)$$

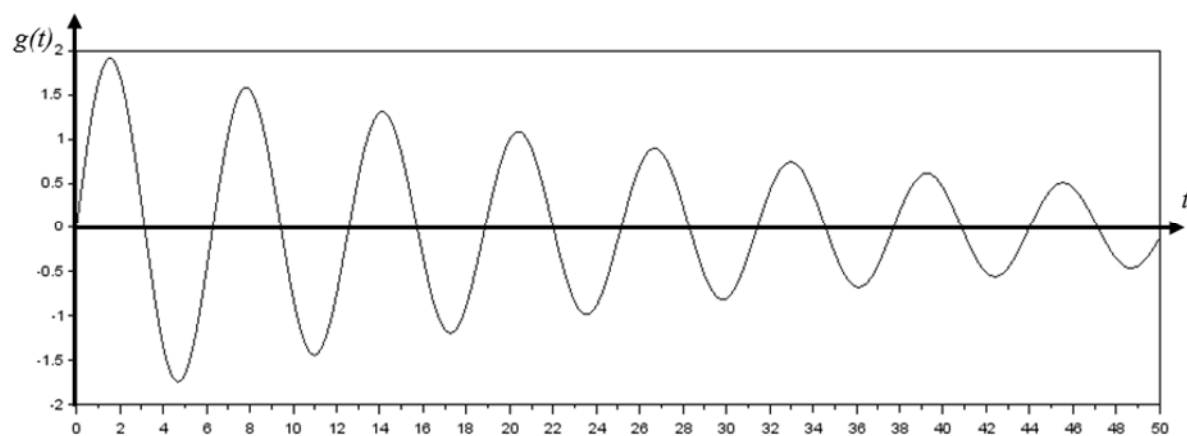
Natomiast odpowiedź impulsowa:

$$g(t) = \frac{k \omega_n e^{-\zeta \omega_n t}}{\sqrt{1 - \zeta^2}} \sin(\omega_n \sqrt{1 - \zeta^2} t) \quad (117)$$

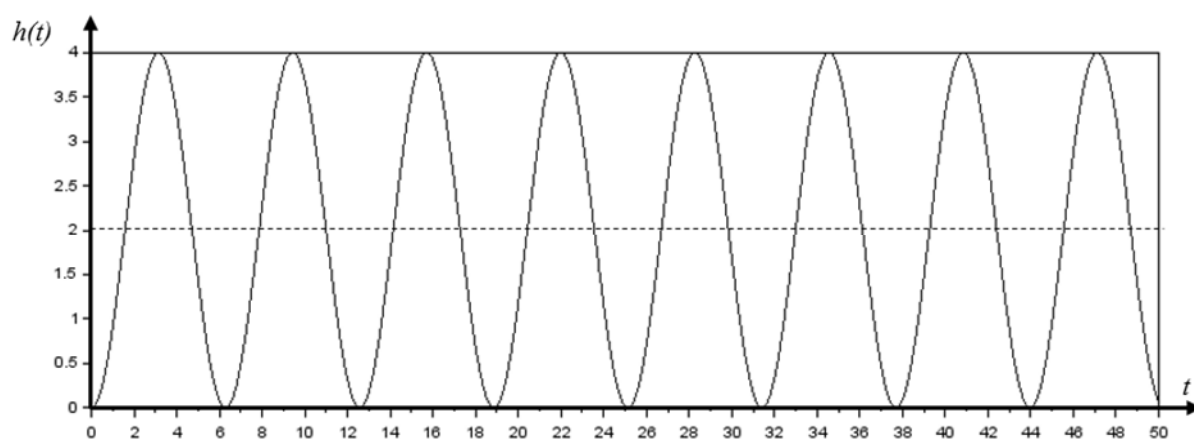
Przykład odpowiedzi skokowej i impulsowej dla elementu oscylacyjnego o $\zeta = 0.03$, $\omega_n = 1$, i $k = 2$, czyli o transmitancji $G(s) = \frac{2}{s^2 + 2 \cdot 0.03 \cdot 1s + 1^2}$ przedstawiają rysunki 27 i 28. Natomiast odpowiedzi dla członu o współczynniku tłumienia $\zeta = 0$ rysunki 29 i 30.



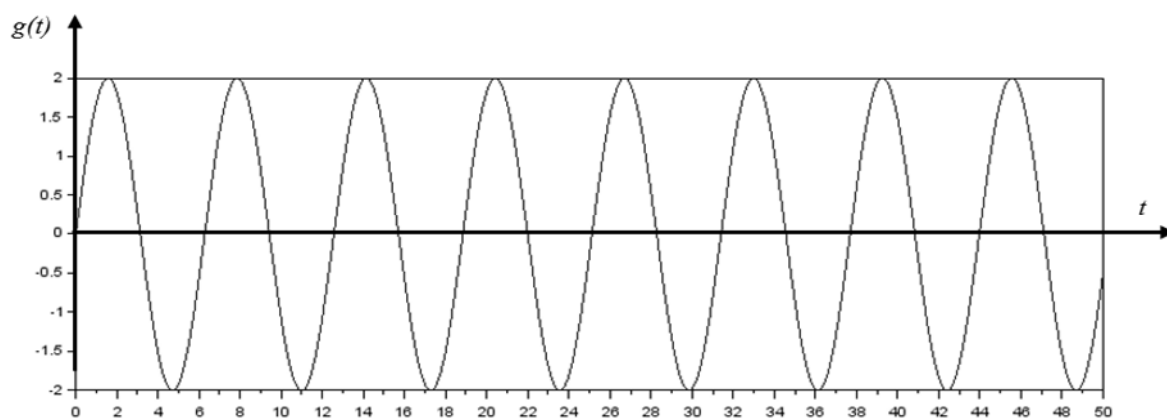
Rys. 27. Odpowiedź skokowa elementu oscylacyjnego o transmitancji $G(s) = \frac{2}{s^2 + 2 \cdot 0.03 \cdot 1s + 1^2}$



Rys. 28. Odpowiedź impulsowa elementu oscylacyjnego o transmitancji $G(s) = \frac{2}{s^2 + 2 \cdot 0.03 \cdot 1s + 1^2}$



Rys. 29. Odpowiedź skokowa elementu oscylacyjnego o transmitancji o współczynniku tłumienia $\zeta=0$



Rys. 30. Odpowiedź impulsowa elementu oscylacyjnego o transmitancji o współczynniku tłumienia $\zeta=0$



Człon opóźniający

Elementem opóźniającym jest taki element w przypadku, którego sygnał wyjściowy $y(t)$, pojawia się po pewnym czasie - T_0 nazywanym czasem opóźnienia, od pojawienia się sygnału wejściowego $u(t)$ na wejściu. Jego równanie wejścia-wyjścia można zapisać w postaci:

$$y(t) = ku(t - T_0) \quad (118)$$

Po dokonaniu przekształcenia Laplace'a:

$$Y(s) = kU(s)e^{-sT_0} \quad (119)$$

Jego transmitancja ma natomiast postać:

$$G(s) = ke^{-sT_0} \quad (120)$$

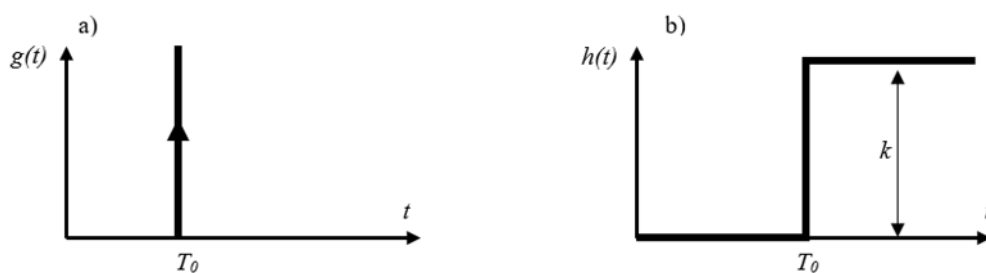
Odpowiedź impulsowa może być opisana równaniem:

$$g(t) = k\delta(t - T_0) \quad (121)$$

Natomiast odpowiedź skokowa:

$$h(t) = k \cdot 1(t - T_0) \quad (122)$$

Na rysunku 31 przedstawiono odpowiedzi impulsową i skokową członu opóźniającego. Typowym przykładem członu opóźniającego jest taśmociąg o długości l poruszający się ze stałą prędkością v . W jego przypadku sygnałem wejściowym będzie pojawienie się materiału na jego początku, a sygnałem wyjściowym pojawienie się materiału na jego końcu po czasie równym $T_0 = l/v$

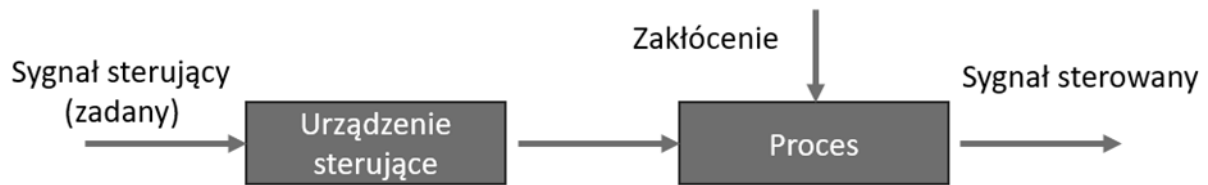


Rys. 31. Odpowiedzi członu opóźniającego a - impulsowa, b - skokowa

Otwarty i zamknięty układ sterowania

Otwarty układ sterowania

Otwarty układ sterowania jest to układ, w którym sygnał sterujący oddziałuje na proces poprzez urządzenie sterujące, bez udziału sygnału na wyjściu obiektu. Sygnał wyjściowy nie wpływa na działanie urządzenia sterującego. W układzie tym może występować sygnał wejściowy niepodlegający sterowaniu – jest to zakłócenie (rys. 32).

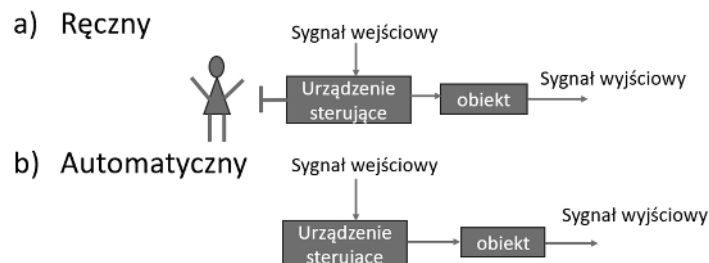


Rys. 32. Otwarty układ sterowania

Występowanie zakłóceń jest wadą każdego takiego systemu. Źródłami zakłóceń mogą być np.:

- niedokładność urządzeń
- zużycie urządzenia
- luzy
- wpływ temperatury, wilgotności
- zmiany napięcia zasilającego.

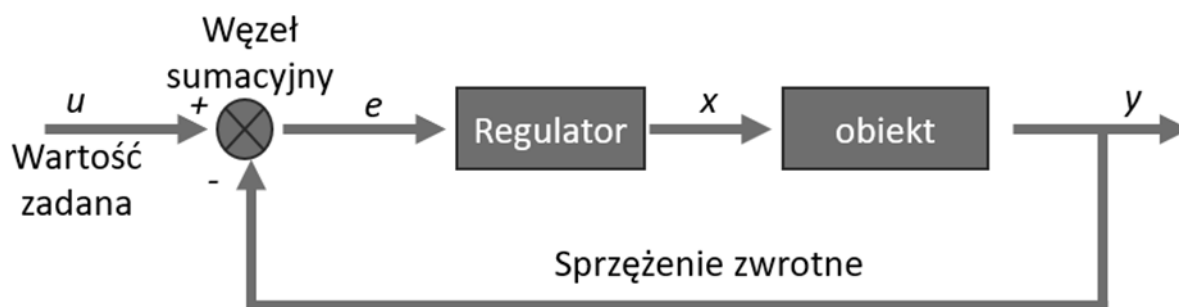
Układ otwarty może być stosowany skutecznie tylko wtedy, gdy wprowadza się modyfikacje zwiększające skuteczność urządzenia, jednakże w zamian wprowadzające ograniczenia. Przez ograniczenia mechaniczne zwykle zmniejsza się jednakże uniwersalność danego systemu. Otwarte układy sterowania można podzielić na sterowane automatycznie i ręcznie (rys. 33). W przypadku sterowanych ręcznie to człowiek pełni rolę urządzenia zadającego.



Rys. 33. Otwarty układ sterowania ze sterowaniem ręcznym i automatycznym

Zamknięty układ sterowania

Zamknięty układ sterowania jest to układ sterowania, w którym sygnał wyjściowy jest podawany jest na wejście za pomocą ujemnego sprzężenia zwrotnego. Eliminuje przez to wady układu otwartego, poprawia dokładność, niezawodność, kompensuje zakłócenia (rys. 34).



Rys. 34. Zamknięty układ sterowania.

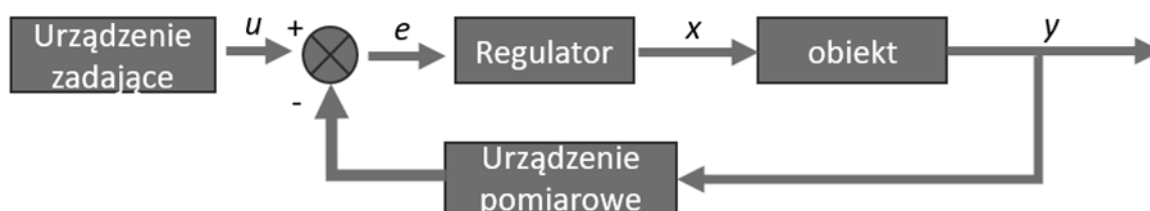
Istotą zamkniętego układu sterowania jest sprzężenie zwrotne, czyli przekazywania sygnału z wyjścia na wejście:

- 1) Sygnał wyjściowy y jest przekazywany do węzła sumacyjnego ze znakiem minus (ujemne sprzężenie zwrotne)
- 2) W węźle sumacyjnym następuje wyznaczenie różnicy – e – nazywanej odchyłką, uchybem, błędem regulacji, pomiędzy wartością zadana – u , a sygnałem wyjściowym – y

$$e = u - y \quad (123)$$

- 3) Wartość uchybu – e jest podawane na regulator
- 4) Zadaniem regulatora jest natomiast sprowadzenie wartości e do zera

Ponieważ najczęściej nie można podać sygnału wyjściowego y bezpośrednio na wejście trzeba dokonywać jego pomiaru za pomocą urządzenia pomiarowego (rys. 35).



Rys. 35. Zamknięty układ sterowania z urządzeniem pomiarowym.

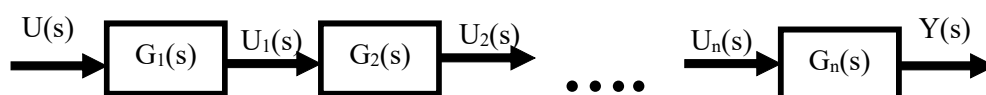
Schematy blokowe

Większość systemów składa się z kombinacji wielu podstawowych członów automatyki. Można je przedstawić za pomocą schematów blokowych. W przypadku tych bardziej skomplikowanych systemów, aby je analizować trzeba wyznaczyć ich transmitancje zastępcze. Określić ją można na podstawie znajomości transmitancji poszczególnych elementów składowych i sposobu ich połączenia (sposobu przepływu sygnałów między nimi).



Połączeniem szeregowym, jest przypadek, kiedy sygnał wychodzący z jednego członu wchodzi na wejście następnego. Przykład takiego połączenia przedstawia rysunek 36. Transmittancja zastępcza jest wówczas iloczynem transmittancji składowych.

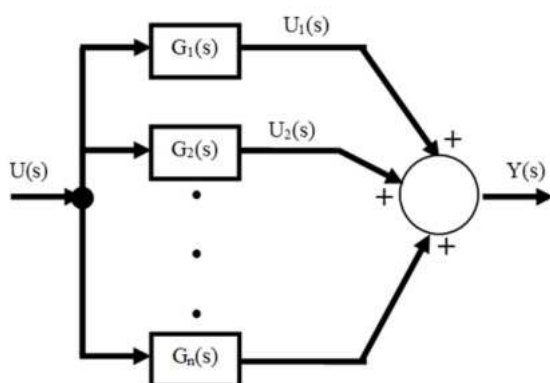
$$G(s) = G_1(s) \cdot G_2(s) \cdot \dots \cdot G_n(s) \quad (124)$$



Rys. 36. Połączenie szeregowe transmittancji

Połączeniem równoległym, jest przypadek, kiedy sygnał pochodzący z jednego źródła przesyłany równoległymi trasami trafia do różnych członów, a sygnały wyjściowe z tych członów są sumowane formując ostateczny sygnał wyjściowy. Przykład takiego połączenia przedstawia rysunek 37. Transmittancja zastępcza jest wówczas sumą transmittancji składowych.

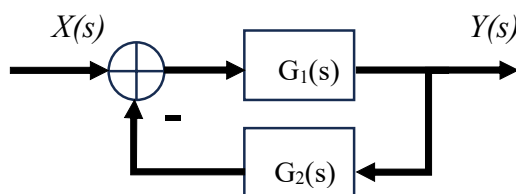
$$G(s) = G_1(s) + G_2(s) + \dots + G_n(s) \quad (125)$$



Rys. 37. Połączenie równoległe transmittancji

Sprzężenie zwrotne

W zamkniętych układach sterowania występuje sprzężenie zwrotne, które w sposób istotny modyfikuje transmittancję układu. Przykład układu ze sprzężeniem zwrotnym przedstawia poniższy rysunek 38.



Rys. 38. Schemat układu z ujemnym sprzężeniem zwrotnym



Transmitancję zastępczą przedstawionego układu można wyznaczyć w sposób przedstawiony poniżej.

$$Y(s) = G_1(s)[X(s) - G_2(s)Y(s)] \quad (126)$$

$$Y(s) + G_1(s)G_2(s)Y(s) = G_1(s)X(s) \quad (127)$$

$$Y(s)[1 + G_1(s)G_2(s)] = G_1(s)X(s) \quad (134)$$

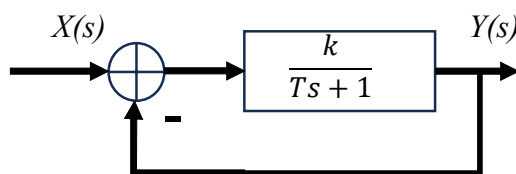
$$G_z(s) = \frac{Y(s)}{X(s)} = \frac{G_1(s)}{1 + G_1(s)G_2(s)} \quad (128)$$

W przypadku zastosowania układu z dodatnim sprzężeniem zwrotnym otrzymalibyśmy transmitancję zastępczą w postaci:

$$G_z(s) = \frac{Y(s)}{X(s)} = \frac{G_1(s)}{1 - G_1(s)G_2(s)} \quad (129)$$

Taki układ nie byłby jednak stabilny, z tego powodu nie stosuje się dodatniego sprzężenia zwrotnego w automatycznych układach regulacji. Znajduje jednak zastosowanie w generatorach.

Rozpatrując poniższy układ składający się z elementu inercyjnego z sprzężeniem zwrotnym (rys. 39):



Rys. 39. Element inercyjny z sprzężeniem zwrotnym

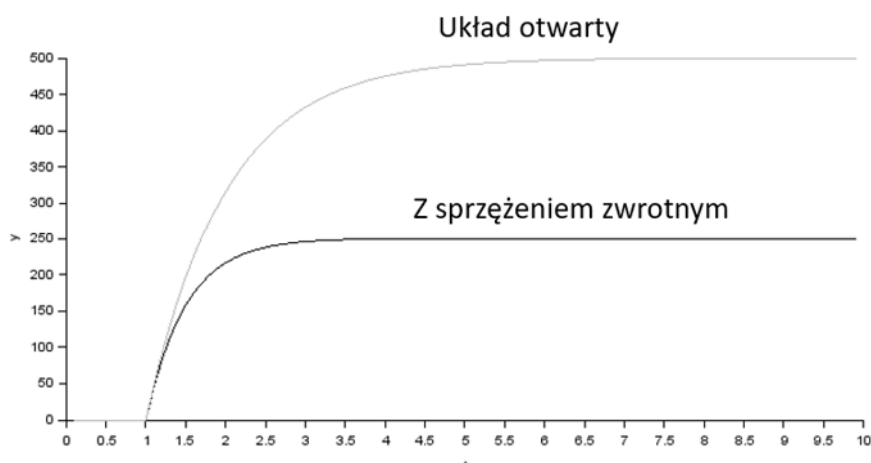
$$G(s) = \frac{G_0}{1 + G_0 G_1} = \frac{G_0}{1 + G_0 \cdot 1} = \frac{\frac{K}{Ts+1}}{1 + \frac{K}{Ts+1}} = \frac{\frac{K}{Ts+1}}{\frac{Ts+1+K}{Ts+1}} = \frac{\frac{K}{1+K}}{\frac{T}{1+K}S+1} = \frac{K_z}{T_z S+1} \quad (130)$$

Gdzie:

$$K_z = \frac{K}{1+K} \quad (131)$$

$$T_z = \frac{T}{1+K} \quad (132)$$

Otrzymaliśmy człon inercyjny, którego wzmacnienie i stała czasowa są mniejsze od wyjściowych. Uzyskaliśmy więc skrócenie stałej czasowej kosztem wzmacnienia. Rysunek 40 pokazuje, jak wyglądałyby odpowiedzi skokowe takiego układu działającego jako układ otwarty i z sprzężeniem zwrotnym.

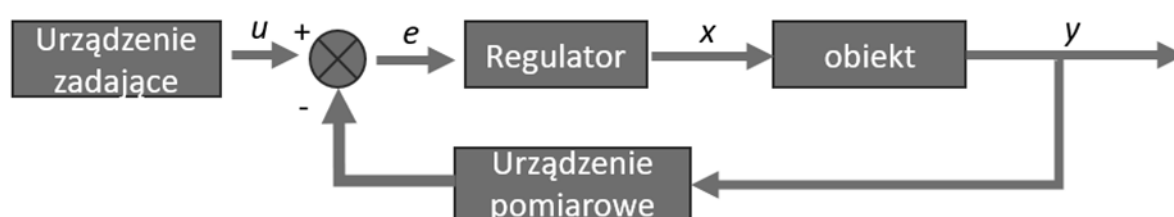


Rys. 40. Przykładowe odpowiedzi skokowe elementu inercyjnego w układzie otwartym i z sprzężeniem zwrotnym

Aby poprawić pracę układu ze sprzężeniem zwrotnym należy dodać do niego dodatkowy układ regulacji nazywany regulatorem. Jednym z najpowszechniej stosowanych regulatorów jest regulator PID.

Regulator PID

Regulator jest elementem układu automatyki, który po wystąpieniu odchylenia sygnału regulowanego od sygnału zadanego powoduje, poprzez oddziaływanie na obiekt regulacji, powrót przebiegu sygnału regulowanego do przebiegu zgodnego z sygnałem zadanym. Sygnałem wejściowym regulatora jest sygnał błędu regulacji (uchyb), natomiast sygnałem wyjściowym – sygnał sterujący (sterowanie), oddziałujący na obiekt regulacji. Schemat zamkniętego układu sterowania z regulatorem przedstawia rysunek 41.



Rys. 41. Zamknięty układ sterowania z regulatorem

Jednym z najpowszechniej obecnie stosowanych regulatorów jest regulator PID, czyli regulator proporcjonalno-całkujaco-różniczkujący. Jego skrócona nazwa pochodzi od angielskich nazw członów P - proportional, I - integrative, D - derivative.



Rozróżnia się 5 form regulatora PID:

- Tylko proporcjonalny – P
- Tylko całkujący – I
- Proporcjonalno-różniczkujący – PD
- Proporcjonalno-całkujący – PI
- Proporcjonalno-całkująco-różniczkujący - PID

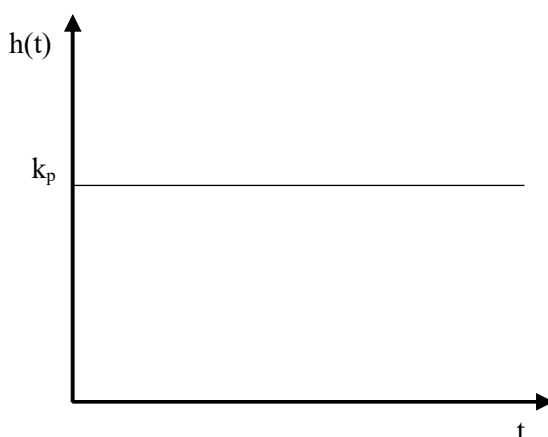
Regulator typu P generuje na swoim wyjściu sygnał proporcjonalny do błędu regulacji. Przy jego zastosowaniu uchyb zawsze będzie niezerowy. Jego transmitancja jest taka sama jak członu proporcjonalnego, jego parametrem jest stała wzmocnienia k_p .

$$G_r(s) = k_p \quad (133)$$

Sygnał wyjściowy z regulatora (sterujący dla obiektu sterowanego) ma postać:

$$x(t) = k_p e(t) \quad (134)$$

Jego odpowiedź na wymuszenie skokowe jest taka sama jak dla członu proporcjonalnego (rys. 42).



Rys. 42. Odpowiedź regulatora proporcjonalnego na skok jednostkowy

W przypadku regulatora typu PI, człon całkujący działa jak sumator – im dłużej uchyb regulacji jest różny od zera tym większy sygnał na wyjściu członu I. Dzięki temu w przypadku regulatora PI możliwe jest zredukowanie uchybu do zera. Postać transmitancji tego regulatora ma poniższą postać:

$$G_r(s) = k_p \left(1 + \frac{1}{T_i s} \right) \quad (135)$$

Gdzie:

k_p – stała wzmocnienia

T_i – stała czasowa całkowania



Czasami opisując regulator PI, zamiast stałej czasowej T_i dokonuje się podstawienia $k_i = \frac{k_p}{T_i}$

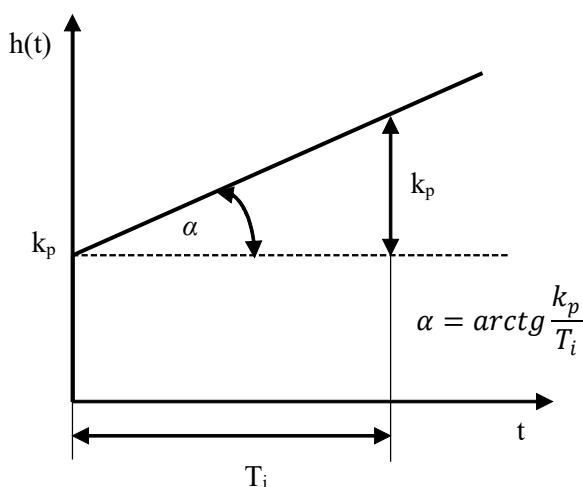
Przebieg czasowy sygnału wyjściowego z regulatora ma postać:

$$x(t) = k_p \left[e(t) + \frac{1}{T_i} \int_0^t e(\tau) d\tau \right] \quad (136)$$

Odpowiedź skokowa dla skoku o wartości 1 ma postać:

$$h(t) = k_p + \frac{k_p}{T_i} t \quad (137)$$

Została ona też zilustrowana na rysunku 43.



Rys. 43. Odpowiedź regulatora PI na skok jednostkowy

Regulator PD posiada człon różniczkujący D, którego wielkość sygnału wyjściowy zmienia się w zależności od szybkości zmian uchybu. Poprawia to dynamikę regulatora Im ta zmiana jest szybsza tym jego reakcja jest większa. Postać transmitancji tego regulatora ma poniższą postać:

$$G_r(s) = k_p(1 + T_d s) \quad (138)$$

Gdzie:

k_p – stała wzmocnienia

T_d – stała czasowa różniczkowania

Czasami opisując regulator PD, zamiast stałej czasowej T_d dokonuje się podstawienia $k_d = k_p T_d$

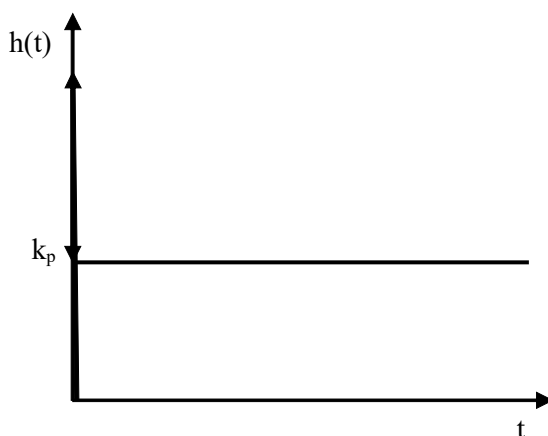
Odpowiedź na sygnał wejściowy uchybu w czasie ma postać:

$$x(t) = k_p \left[e(t) + T_d \frac{de}{dt} \right] \quad (139)$$



Odpowiedź skokowa (rys. 44) ma postać:

$$h(t) = k_p[1(t) + T_d\delta(t)] \quad (140)$$



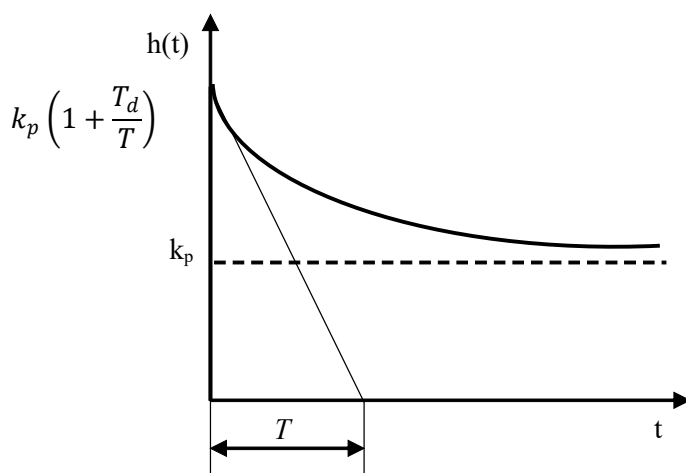
Rys. 44. Odpowiedź regulatora PD na skok jednostkowy

Ponieważ w fizycznym świecie nie można zrealizować idealnego członu rzeczywistego definiuje się także regulator PD z inercją określona przez stałą czasową T . Wówczas jego transmitancja ma postać:

$$G_r(s) = k_p \left(1 + \frac{T_d s}{T s + 1} \right) \quad (141)$$

Odpowiedź skokowa ma natomiast postać (rys. 45):

$$h(t) = k_p \left[1(t) + \frac{T_d}{T} e^{-\frac{t}{T}} \right] \quad (142)$$



Rys. 45. Odpowiedź regulatora PD z inercją na skok jednostkowy



Regulator PID posiada wszystkie trzy wymienione powyżej człony. Jego transmitancja ma postać:

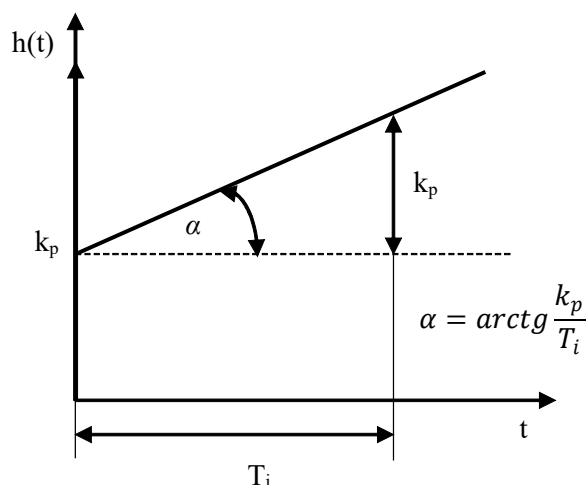
$$G_r(s) = k_p \left(1 + \frac{1}{T_i s} + T_d s \right) \quad (143)$$

Odpowiedź na sygnał wejściowy uchybu w czasie ma postać:

$$x(t) = k_p \left[e(t) + \frac{1}{T_i} \int_0^t e(\tau) dt + T_d \frac{de}{dt} \right] \quad (144)$$

Natomiast odpowiedź skokowa (rys. 46):

$$h(t) = k_p \left[1(t) + \frac{1}{T_i} t + T_d \delta(t) \right] \quad (145)$$



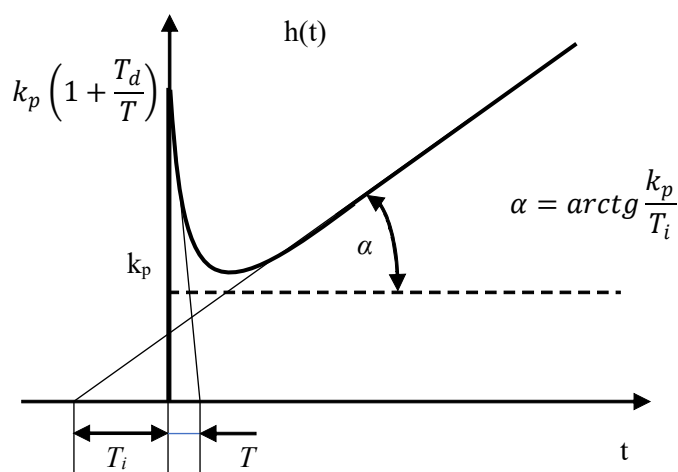
Rys. 46. Odpowiedź regulatora PID na skok jednostkowy

Określa się także tak zwany rzeczywisty regulator PID, w którym człon różniczkujący posiada inercję. Wówczas transmitancja ma postać:

$$G_r(s) = k_p \left(1 + \frac{1}{T_i s} + \frac{T_d s}{Ts+1} \right) \quad (146)$$

A odpowiedź skokowa (rys. 47):

$$h(t) = k_p \left[1(t) + \frac{1}{T_i} t + \frac{T_d}{T} e^{-\frac{t}{T}} \right] \quad (147)$$



Rys. 47. Odpowiedź rzeczywistego regulatora PID na skok jednostkowy

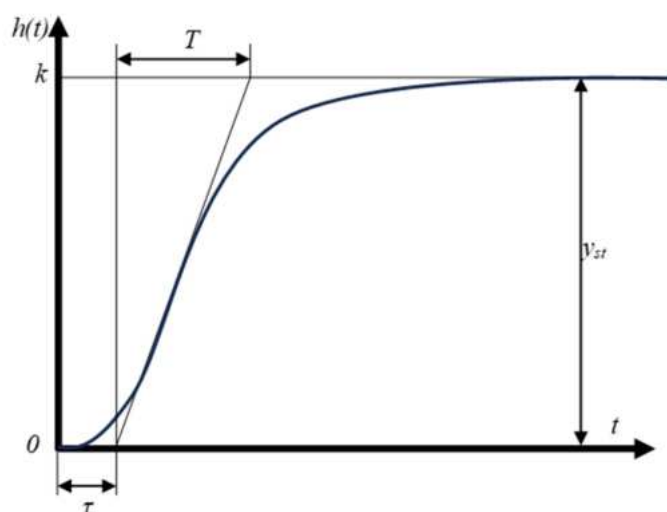
Strojenie regulatora PID polega na doborze parametrów regulatora PID, a więc wzmocnienia k_p , i stałych czasowych T_i i T_d . Jedną z podstawowych metod jest metoda Zigglera-Nicholsa. Wyróżnia się przy tym jej dwa warianty.

I metoda:

- Zakłada że obiekt opisany jest członem inercyjnym (bezoscylacyjnym) z opóźnieniem z dominującą stałą czasową T oraz opóźnieniem transportowym τ
- Taki obiekt przybliżony jest transmitancją:

$$G(s) = \frac{k}{Ts+1} \cdot e^{-\tau s} \quad (148)$$

- Ustawia się regulator na strojenie w nominalnym punkcie pracy przy małym wzmocnieniu, oraz wyłączonych członach I oraz D.
- Następnie rejestruje się odpowiedź skokową obiektu
- Na zarejestrowanej odpowiedzi znajdują się punkt przegięcia krzywej i kreśli styczną przechodzącą przez ten punkt oraz odczytuje się parametry T , τ oraz wartość y_{st} - wartość ustaloną sygnału odpowiedzi (rys. 48).
- Następnie wyznacza się parametr $R = \frac{y_{st}}{T}$
- Na koniec z tabeli (tab. 2) podanej przez Zigglera-Nicholsa odczytuje się odpowiednie parametry regulatora.



Rys. 48. Wyznaczanie parametrów do strojenia regulatora PID metodą Zigglera-Nicholsa

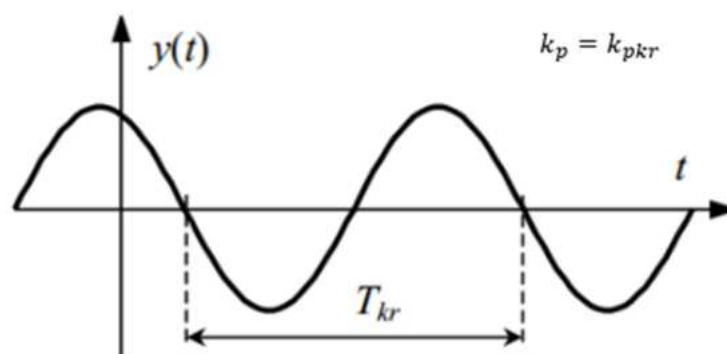
Tablica 2. Nastawy regulatora PID zgodnie z metodą Zigglera-Nicholsa

P $G(s) = k_p$	$k_p = \frac{1}{R\tau}$		
PI $G(s) = k_p(1 + \frac{1}{T_I s})$	$k_p = \frac{0,9}{R\tau}$	$T_I = 3,3\tau$	
PID $G(s) = k_p(1 + \frac{1}{T_I s} + T_D s)$	$k_p = \frac{1,2}{R\tau}$	$T_I = 2\tau$	$T_D = 0,5\tau$

II metoda

Jest stosowana, gdy układ jest oscylacyjny albo nie może pracować bez regulatora.

- Ustawiamy tylko działanie członu k_p
- Nastawiamy wzmocnienie k_p na małym poziomie
- Zwiększamy k_p rejestrując odpowiedź aż do wystąpienia niegasnących oscylacji, czyli o stałej amplitudzie
- Należy odczytać wartość wzmocnienia k_{pkr} , zwanego wzmocnieniem krytycznym, przy którym wystąpiły drgania, oraz okres tych drgań T_{kr} , zwany okresem krytycznym (rys. 49).



Rys. 49. Wyznaczanie parametrów dla regulatora PID drugą metodą Zigglera-Nicholsa

Zgodnie z regułami Zieglera-Nicholsa należy dobrać następujące parametry:

- Dla regulatora PID $k_p=0,6k_{pkr}$ $T_i=0,5T_{kr}$, $T_d=0,12T_{kr}$
- Dla regulatora PI $k_p=0,45k_{pkr}$ $T_i=0,85T_{kr}$,
- Dla regulatora PI $k_p=0,5k_{pkr}$

Jakość regulacji

Jakość układów regulacji ocenia się za pomocą wskaźników jakości przebiegu wielkości regulowanej. Wskaźniki te można podzielić na:

- Wskaźniki dokładności statycznej układu
- Wskaźniki szybkości działania (szybkości reakcji układu na wymuszenie)
- Wskaźniki całkowite
- Wskaźniki stabilności

Rozpatrując dokładność statyczną należy zauważyć, że jednym z podstawowych wymagań, które muszą spełnić układy regulacji jest utrzymywanie odchylenia sygnału regulowanego od sygnału zadanego w dopuszczalnych granicach. Odchylenie to nazywamy uchybem ustalonym – e_u .

$$e_u = \lim_{t \rightarrow \infty} e(t) \quad (149)$$

Badane są uchyby:

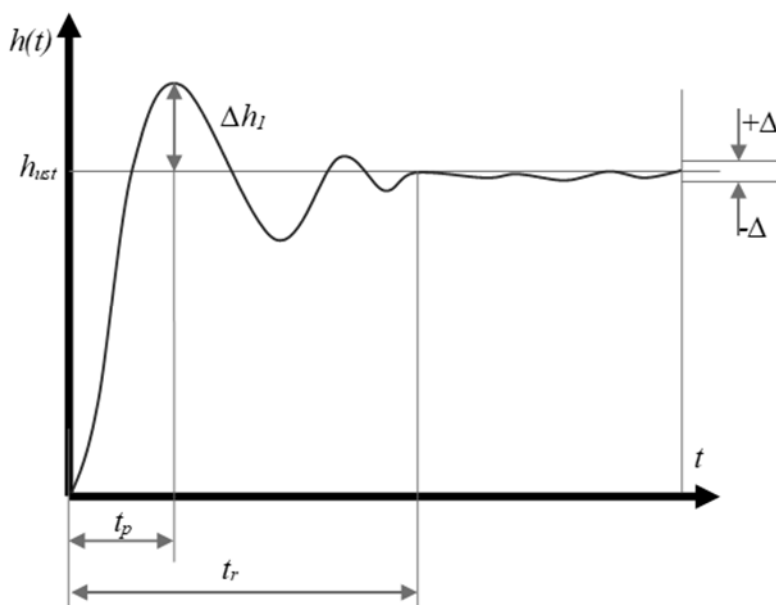
- Od skoku
- Od sygnału liniowo narastającego
- Od parabolicznego



Jakość dynamiczną bada się na podstawie odpowiedzi układu na skok jednostkowy. Parametrami, które się bada, są:

- Czas regulacji t_r - czas od wymuszenia sygnału do osiągnięcia przez sygnał wyjściowy asymptoty, albo do momentu, gdy błąd zaczyna mieścić się w sposób trwały poniżej $\Delta=0,01$ błędu maksymalnego
- Czas do osiągnięcia wartości maksymalnej t_p
- Maksymalne odchylenie dynamiczne Δh_1
- Przeregulowanie $\kappa = \frac{\Delta h_1}{h_{ust}}$

Przedstawione parametry zostały zilustrowane na rysunku 50.



Rys. 50. Wybrane parametry dynamicznej jakości regulacji

Wyróżniamy dwa całkowite kryteria jakości regulacji:

- Wskaźnik I_1 , reprezentujący pole zawarte pomiędzy przebiegiem błędu $e(t)$, stosowany dla aperiodycznego przebiegu uchybu.

$$I_1 = \int_0^{\infty} e(t) dt \quad (150)$$

- Wskaźnik I_2 reprezentujący pole zawarte pomiędzy kwadratem błędu $e^2(t)$ i asymptotą, stosowany dla oscylacyjnego przebiegu uchybu $e(t)$.

$$I_2 = \int_0^{\infty} e^2(t) dt \quad (151)$$



Kryteria stabilności

Jako układ niestabilny określa się w automatyce układ który raz wytrącony z równowagi nigdy do niej nie wraca, a sygnał wyjściowy oddala się od położenia początkowego, dążąc do nieskończoności.

Układem stabilnym nazywamy natomiast układ, który po ustaniu sygnału zakłócającego powraca do stanu równowagi. Sygnał wyjściowy w takim układzie nie dąży do nieskończoności (mieści się w określonych granicach)

Szczególnym przypadkiem układu stabilnego jest układ stabilny asymptotycznie - powraca on do takiego samego stanu równowagi, jaki był przed wystąpieniem zakłócenia.

Do oceny stabilności układów może posłużyć kilka kryteriów. Poniżej zaprezentowano wybrane:

- Kryterium odpowiedzi skokowej

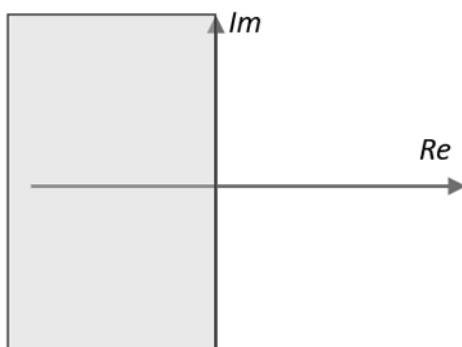
Układ zamknięty jest stabilny, gdy w odpowiedzi na sygnał skokowy osiąga stan ustalony w czasie dążącym do nieskończoności

- Kryterium biegunów

Wszystkie pierwiastki równania charakterystycznego $M(s)$ leżą w lewej półpłaszczyźnie płaszczyzny zmiennej zespolonej s (mają ujemne części rzeczywiste). Dla poniższej transmitancji układu, równaniem charakterystycznym jest funkcja $M(s)$ przyrównana do 0.

$$G(s) = \frac{Y(s)}{X(s)} = \frac{b_m \cdot s^m + b_{m-1} \cdot s^{m-1} + \dots + b_1 \cdot s + b_0}{a_n \cdot s^n + a_{n-1} \cdot s^{n-1} + \dots + a_1 \cdot s + a_0} = \frac{L(s)}{M(s)} \quad (152)$$

Jeżeli obliczone dla niej pierwiastki znajdują się w obszarze zaznaczonym poniżej to układ jest stabilny (rys. 51).



Rys. 51. Obszar, w którym muszą znaleźć się pierwiastki równania charakterystycznego, aby układ był stabilny



- Analityczne kryterium – kryterium Hurwitza

Ma ono dwa warunki:

- 1) Wszystkie współczynniki równania charakterystycznego muszą być dodatnie
- 2) Podwyznaczniki macierzy Hurwitza Δ_i od $i=2$ do $i=n-1$ są większe od 0

Równanie charakterystyczne $M(s) = 0$ ma postać:

$$a_n s^n + a_{n-1} s^{n-1} + \dots + a_1 s + a_0 = 0 \quad (153)$$

Natomiast utworzona na jego podstawie macierz Hurwitza i jej wyznacznik ma postać:

$$\Delta_n = \begin{vmatrix} a_{n-1} & a_n & 0 & 0 & 0 & 0 & \dots & 0 & 0 & 0 \\ a_{n-3} & a_{n-2} & a_{n-1} & a_n & 0 & 0 & \dots & 0 & 0 & 0 \\ a_{n-5} & a_{n-4} & a_{n-3} & a_{n-2} & a_{n-1} & a_n & \dots & 0 & 0 & 0 \\ \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots & \dots \\ 0 & 0 & 0 & 0 & 0 & 0 & \dots & a_1 & a_3 & a_4 \\ 0 & 0 & 0 & 0 & 0 & 0 & \dots & a_0 & a_1 & a_2 \\ 0 & 0 & 0 & 0 & 0 & 0 & \dots & 0 & 0 & a_0 \end{vmatrix} \quad (154)$$

Macierz ma wymiar $n \times n$ gdzie n jest stopniem wielomianu równania charakterystycznego. Wypełnianie macierzy należy zacząć od początku pierwszego wiersza współczynnikiem o numerze $n-1$, następnie w następnej pozycji należy wpisać współczynnik o numerze n . Pozostałe miejsca w wierszu należy uzupełnić zerami, gdyż nie ma już współczynników o wyższym numerze. Drugi wiersz należy rozpocząć od współczynnika o numerze $n-3$, a trzeci od $n-5$ i z następnymi postępować analogicznie. Wolne miejsca w wierszach należy uzupełnić zerami. Poniżej zaprezentowano przykład zastosowania kryteriów Hurwitza:

Równanie charakterystyczne:

$$5s^3 + 3s^2 + 1s + 8 = 0 \quad (155)$$

Równanie to posiada następujące współczynniki:

$$a_3=5, a_2=3, a_1=1, a_0=8$$

Zgodnie z I kryterium:

$$5>0, 3>0, 1>0, 8>0$$

Zgodnie z II kryterium

$$\Delta_3 = \begin{vmatrix} 3 & 5 & 0 \\ 8 & 1 & 3 \\ 0 & 0 & 8 \end{vmatrix} = 3 \cdot 1 \cdot 8 - 8 \cdot 5 \cdot 8 = 24 - 320 = -296 \quad (156)$$

Wyznacznik jest ujemny, więc układ jest niestabilny.



2. Podstawy robotyki

Robotyka jest dziedziną nauki i techniki zajmującą się zagadnieniami z obszarów mechanik, sterowania, projektowania, zastosowania i eksploatacji robotów. Jej przedmiotem są również zastosowania robotów w badaniach naukowych, szeroko rozumianej technice, medycynie i innych sferach działalności człowieka.

Definicja robota przemysłowego

Zgodnie z PN-EN ISO 8373 2.6 – **robot przemysłowy** (pot. Robot) jest to automatycznie sterowany, przeprogramowalny, uniwersalny manipulator, programowalny w trzech lub więcej osiach, który może być stacjonarny lub mobilny, przeznaczony do przemysłowej automatyzacji

- Przeprogramowalność - oznacza, że zaprogramowane ruchy lub funkcje pomocnicze robota mogą być zmieniane bez zmiany struktury mechanicznej lub układu sterowania
- Uniwersalność – oznacza, że może być on adaptowany do różnych zastosowań po fizycznej zmianie (zmiana struktury mechanicznej lub układu sterowania)
- Manipulator przemysłowy - Urządzenie przemysłowe przeznaczone do wspomagania albo częściowego lub całkowitego zastępowania człowieka przy wykonywaniu czynności manipulacyjnych w przemysłowym procesie produkcyjnym, sterowane ręcznie lub automatycznie za pomocą własnego układu sterującego stałoprogramowego lub zewnętrznego układu sterującego

Budowa robotów przemysłowych

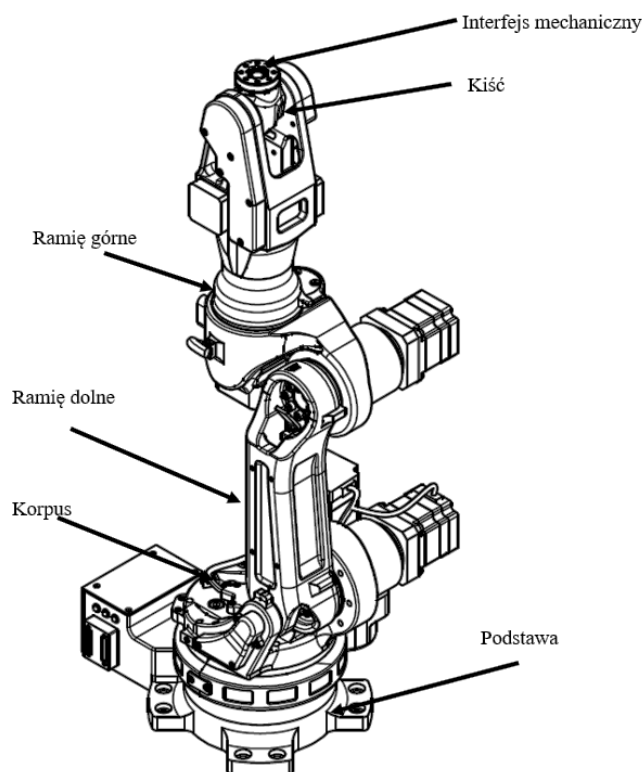
Robot przemysłowy składa się z trzech zasadniczych części:

- Manipulatora
- Układu sterowania
- Układu zasilania

Przeważnie układy sterowania jak i zasilania znajdują się w osobnej szafie stojącej w pobliżu manipulatora. Zdarzają się jednakże rozwiązania, w których są zintegrowane z manipulatorem i umieszczone w jego podstawie. Należy pamiętać, że bez układu sterownia, robot byłby tylko konstrukcją wieloprzegubową. Budowę samego manipulatora przedstawiono na rysunku 52. Przedstawiony robot składa się głównie z ramion połączonych przegubami obrotowymi. Ostatnie ramię robota zakończone jest kłosem z interfejsem mechanicznym. Interfejs mechaniczny służy do podłączenia chwytaka, albo narzędzia roboczego, przykładowo



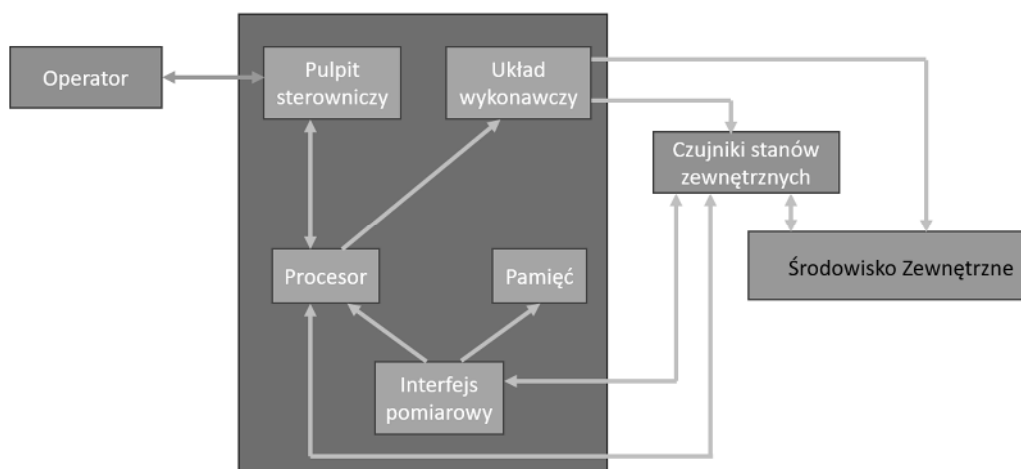
zgrzewadła. Możliwość zamontowania różnego oprzyrządowania w interfejsie mechanicznym jest jednym z czynników, dzięki któremu roboty mogą być dostosowywane do różnych zadań.



Rys. 52 Budowa manipulatora

Ruch poszczególnych ramion robota jest zapewniany przez napędy, którymi mogą być silniki elektryczne albo siłowniki pneumatyczne albo hydrauliczne.

Układ sterowania zapewnia możliwość programowania robota, wykonywania programu, wymiany danych pomiędzy operatorem a robotem, a także zbierania danych o środowisku zewnętrznym jak i o stanie samego robota. Tak więc składa się z elektroniki w skład której wchodzi procesor przetwarzający program, specjalizowane układy służące do wyznaczania trasy ruchu, urządzenia stanowiące interfejs między robotem i operatorem oraz szeroko rozbudowane układy pomiarowe. W ramach układów pomiarowych znajdują się czujniki zbierające dane o stanie robota, a więc prędkościach i położeniach poszczególnych ramion, jak również mierzące pobierane prądy oraz powstające siły i momenty. Czujnikami zbierającymi dane zewnętrzne mogą być różnego rodzaju czujniki odległości, detektory kolizji, detektory obecności albo systemy wizyjne. Zależności między różnymi zespołami robota przedstawia rysunek 53.



Rys. 53. Zależności zespołów robota

Układ zasilania zapewnia dostarczanie energii do zespołów wykonawczych i sterujących tak by te mogły wykonywać swoje zadania.

Klasyfikacja robotów

Podziału robotów można dokonać ze względu na wiele kryteriów. Poniżej przedstawiono kilka podziałów. Poniżej zaprezentowano podział względu na specjalizację:

- Uniwersalne – przeznaczony do różnych operacji technologicznych i manipulacyjnych
- Specjalizowane – przeznaczone do wykonywania operacji technologicznych jednego rodzaju lub do operacji manipulacji przy współpracy z jednym rodzajem wyposażenia technologicznego
- Specjalne – przeznaczony do wykonywania jednej operacji technologicznej lub czynności manipulacyjnej z jedną odmianą wyposażenia technologicznego

Dokonyuje się też podziału ze względu na rodzaj napędu robotów na pneumatyczne, hydrauliczne, elektryczny albo będące kombinacją wcześniej wymienionych. Wśród obecnie stosowanych robotów przeważają te o napędzie elektrycznym. Innym podział robotów polega na ich podziale ze względu na ich sposób przemieszczania się. Dzieli się tu roboty na stacjonarne i mobilne. Te stacjonarne mogą być następnie podzielone ze względu na sposób instalowania. Wyróżnić tutaj można:

- Podłogowe, których podstawa albo korpus są przystosowane do bezpośredniego montażu na podłodze
- Portalowe, które są mocowane do ścian oraz sufitów



- Konsolowe przeznaczone do montażu na urządzeniach, które obsługują. Przykładowo montowane bezpośrednio na obrabiarkach CNC odbierające wyroby gotowe i załadowujące materiał, albo montowane na wtryskarkach odbierające wypraski.

Przestrzeń robocza robota

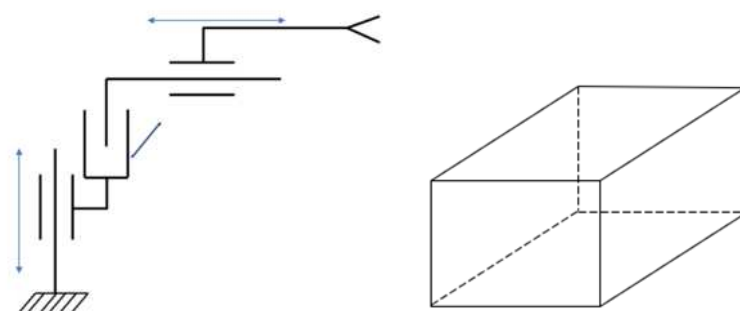
Przestrzenią roboczą robota przemysłowego zgodnie z PN-EN ISO 8373 4.8.4 nazywamy przestrzeń, która może być osiągnięta przez punkt odniesienia kiści z dodaniem zasięgu obrotu lub przemieszczenia każdego połączenia kiści. Kształt przestrzeni roboczej zależy od przeznaczenia robota przemysłowego i jest związany z wyborem jego struktury kinematycznej. Z tego względu przemieszczenie kiści robota może być realizowane w następujących układach:

- Kartezjańskim, czyli prostokątnym (3 pary postępowe PPP)
- Cylindrycznym (1 para obrotowa i 2 pary postępowe OPP)
- Sferycznym (3 pary obrotowe lub 2 pary obrotowe i 1 para postępową – OOO/OOP)
- Typu SCARA (1 para postępową i 2 pary obrotowe – POO/OOP)
- Przegubowym (antropomorficznym) (3 pary obrotowe – OOO)

Układy kinematyczne robotów

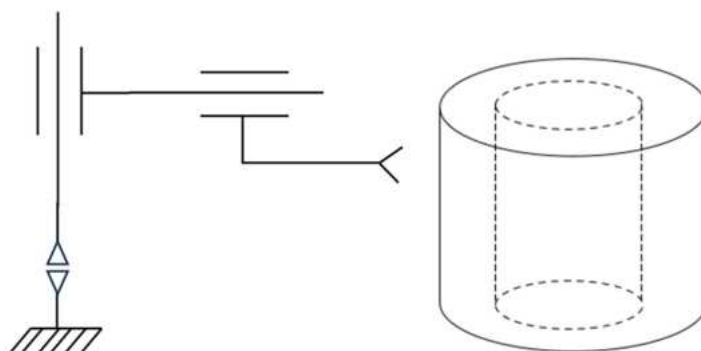
Układy kinematyczne robotów mogą być w postaci otwartego łańcuch kinematycznego jak i zamkniętego. Poniżej zaprezentowano typowe układy kinematyczne robotów o otwartym łańcuchu kinematycznym o trzech stopniach swobody. Dla oznaczenia pary kinematycznej użyto oznaczenie o-obrotowa, p-przesuwna prostoliniowa. Układ kinematyczny w dużej mierze wpływa na możliwości robota i kształt jego przestrzeni roboczej, a więc przestrzeni, w której może się znaleźć jego końcówka robocza.

Robot kartezjański jest robotem, w którym występują tylko postępowe pary kinematyczne. Jego przestrzeń robocza ma kształt prostopadłościenny. Cechą tego robota jest proste programowanie przemieszczeń końcówki roboczej. Rysunek 54 przedstawia schematycznie układ kinematyczny robota kartezjańskiego o trzech stopniach swobody.



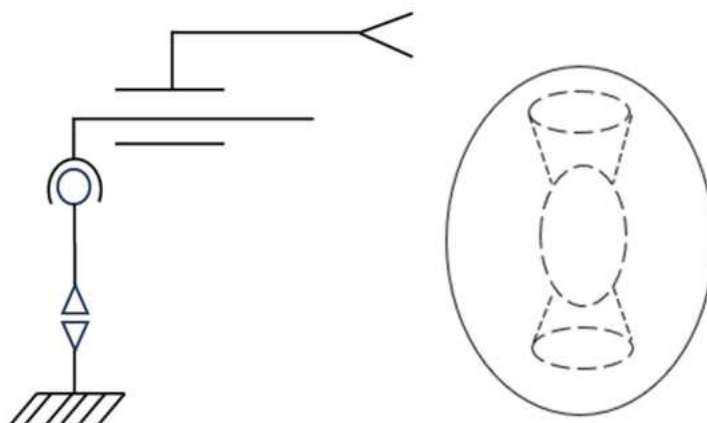
Rys. 54. Robot kartezyjski i kształt jego przestrzeni roboczej

Robot cylindryczny o trzech stopniach swobody posiada jedną parę kinematyczną obrotową i dwie przesuwne. Kształt jego przestrzeni roboczej ma formę wydrążonego cylindra. Roboty te dobrze sprawdzają się w operacjach podnieś i przenieś. Rysunek 55 przedstawia schematycznie układ kinematyczny robota cylindrycznego



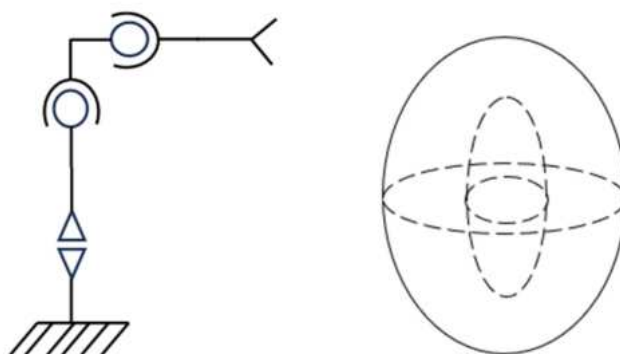
Rys. 55. Robot cylindryczny i kształt jego przestrzeni roboczej

Roboty o układzie kinematycznym sferycznym rozpoczęły masową robotyzację zakładów przemysłowych. Pojawiły się już w latach 60-tych XX w. w fabrykach samochodów w USA. Współcześnie ten układ kinematyczny jest rzadko spotykany. Roboty sferyczne zostały wyparte przez roboty antropomorficzne (przegubowe), które są w stanie wykonywać bardziej skomplikowane ruchy i dotrzeć z końcówką roboczą do miejsc nieosiągalnych dla robotów sferycznych. Rysunek 56 przedstawia układ kinematyczny robota sferycznego.



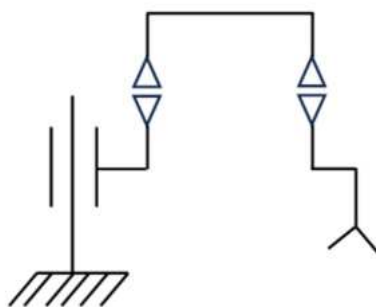
Rys. 56. Robot sferyczny i kształt jego przestrzeni roboczej

Roboty przegubowe są obecnie najbardziej rozpowszechnionymi robotami. Posiadają pary kinematyczne tylko obrotowe. Rysunek 57 przedstawia układ kinematyczny robota przegubowego.



Rys. 57. Robot przegubowy - antropomorficzny i kształt jego przestrzeni roboczej

Roboty typu SCARA posiadają przy trzech stopniach swobody jedną parę kinematyczna przesuwna i dwie obrotowe podobnie jak roboty cylindryczne. W tym przypadku jednak para kinematyczna jest zawsze albo pierwsza albo ostatnia w łańcuchu, a ponadto osie par obrotowych są do siebie równoległe. Układ taki jest korzystny do operacji podnieś i przenieś, umożliwiając osiągnięcie dużych prędkości przez te roboty. Rysunek 58 przedstawia układ kinematyczny robota typu SCARA.



Rys. 58. Robot SCARA

W przypadku robotyki istotną kwestią są tak zwane zagadnienia kinematyki proste i odwrotne. Zagadnienie proste kinematyki polega na określeniu na podstawie struktury kinematycznej robota, wymiarów oraz parametrów nastawnych położenia końcówki roboczej robot. Polega ono, więc na wyliczeniu położenia końcówki roboczej przy znajomości w danej chwili położenia kąтового poszczególnych przegubów i przesunięć elementów o ruchu prostoliniowym. W przypadku tego zagadnienia istnieje zawsze tylko jedno rozwiązanie.

W przypadku zagadnienia odwrotnego na mając dane położenie docelowe końcówki robota należy ustalić, jakie położenie kątowe muszą przyjąć poszczególne przeguby oraz jakie przemieszczenia poszczególne zespoły o ruchu liniowym. W tym przypadku rozwiązanie nie zawsze jest możliwe, ponieważ zadany punkt może znajdować się poza strefą docelową. Ponadto mogą wystąpić osobliwości, może wystąpić wiele możliwych rozwiązań dla tego samego punktu. Powoduje to zagrożenie, że przy zadanym programie, w którym robot ma się przemieścić do danego punktu, może on przeskakiwać między dwoma rozwiązaniami, co może doprowadzić do nieprzewidywanych ruchów ramion robota. Jest to sytuacja, której należy unikać.

Układy współrzędnych robotów

W przypadku robotów przemysłowych przy ich programowaniu w celu ustalenia współrzędnych docelowych posługujemy się kilkoma układami współrzędnych. Poniżej wymieniono trzy najważniejsze.

- **Globalny układ współrzędnych**- jego osie są oznaczane jako: X_0 , Y_0 , Z_0 . Początek tego układu współrzędnych jest określany przez użytkownika, zgodnie z jego wymaganiami. Kierunek osi Z_0 jest zgodny, a zwrot $+Z_0$ jest przeciwny do wektora

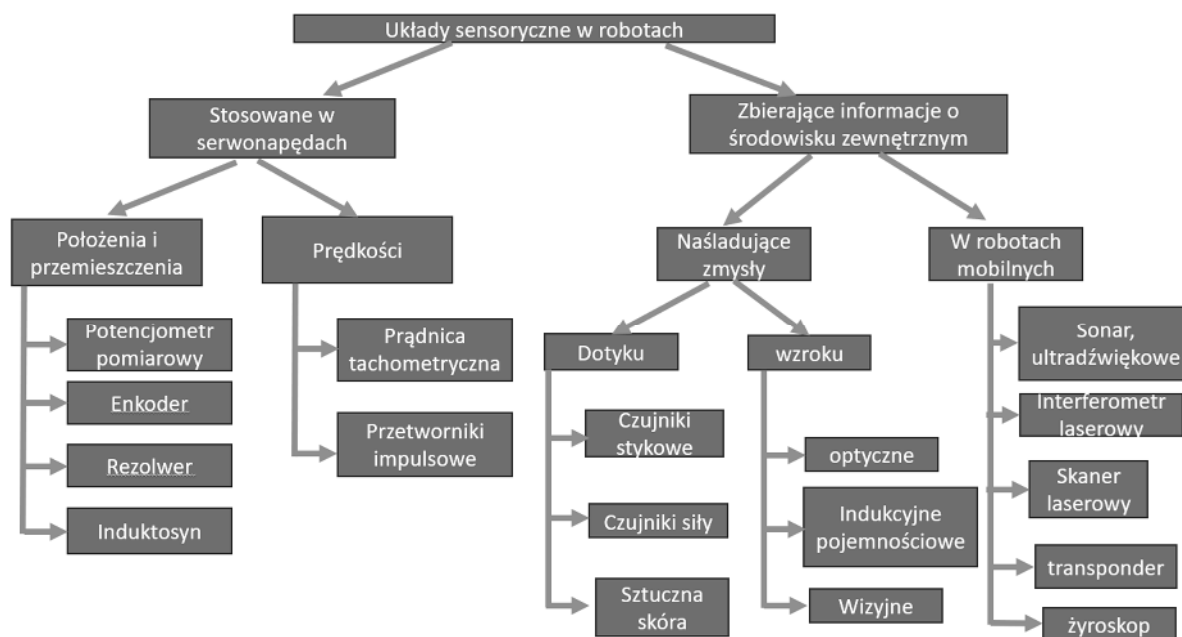


przyspieszenia ziemskiego. Natomiast oś $+X_0$ jest określana dowolnie przez użytkownika.

- **Podstawowy układ współrzędnych (regionalny):** X_n, Y_n, Z_n (gdzie n to kolejna oś robota). Początek podstawowego układu współrzędnych (czyli osie X_1, Y_1, Z_1) powinien być określony przez producenta robota. Kierunek osi $+Z_1$ jest prostopadły do podstawy, natomiast jej zwrot $+Z_1$ jest skierowany od powierzchni zamocowania robota. Zwrot osi $+X_1$ jest skierowany od początku układu współrzędnych i przechodzi przez rzut środka przestrzeni roboczej. Jeżeli konfiguracja robota uniemożliwia stosowanie tej zasady to zwrot $+X_1$ powinien być określony przez producenta.
- **Układ współrzędnych interfejsu mechanicznego robota (lokalny):** X_m, Y_m, Z_m (m -liczba osi robota $n+1$). Początkiem układu współrzędnych jest środek interfejsu mechanicznego robota. Zwrot $+Z_m$ jest skierowany od interfejsu mechanicznego do elementu roboczego. Oś $+X_m$ określa przecięcie płaszczyzny interfejsu mechanicznego robota z płaszczyzną X_1, Z_1 (albo płaszczyzną równoległą do X_1, Z_1).

Czujniki w robotach

W celu zbierania informacji o stanie robota, jak również o stanie środowiska zewnętrznego roboty są wyposażane w liczne czujniki. Podział tych czujników przedstawiono na rysunku 59.



Rys. 59. Podział czujników stosowanych w robotach



Programowanie robotów

Zadaniem robotów przemysłowych jest wykonywanie pewnego pożądanego cyklu pracy. Celem programowania robota jest nauczenie go tego cyklu. Tworzony program składa się w znacznej części z opisu drogi, jaką ma się przemieszczać końcówka robocza robota. Opisywane są punkty w przestrzeni roboczej, między, którymi ma się przemieszczać oraz z jakimi prędkościami. Dodatkowo w programie znajdują się instrukcje nie dotyczące opisu ruchów robota. Mogą one dotyczyć sposobu reakcji na sygnały od czujników, informacje, kiedy włączać lub wyłączać efektor zamontowany na końcu robota, albo jak obsługiwać urządzenia zewnętrzne.

Generalnie programowanie robotów można podzielić na programowanie on-line, czyli na programowanie przy stanowisku pracy robota, oraz off-line, czyli poza stanowiskiem pracy robota.

Programowanie on-line można podzielić natomiast na programowanie ręczne i programowanie przez nauczanie. W przypadku programowania ręcznego mamy do czynienia z sytuacją, w której program obejmuje kolejność i rodzaj wykonywanych ruchów, ale ich wartości wynikają z ustawionych ograniczników powiązanych konstrukcyjnie z robotem. Taka forma programowania jest odpowiednia dla prostych robotów, których układ sterowania opiera się na sterownikach PLC.

Natomiast przez nauczanie jest obecnie najpowszechniejszą metodą programowania robotów. W tym przypadku osoba programująca robota musi przemieszczać manipulator w sposób ręczny lub mechaniczny wzdłuż pożądanego toru, w celu zapisania tego toru, albo jego fragmentów do pamięci robota, w celu jego późniejszego odtworzenia. Wyróżnia się w tym przypadku tak zwane programowanie dyskretne oraz programowanie ciągłe.

Przy programowaniu dyskretnym robot jest sterowany za pomocą sterownika ręcznego (tzw. teach pedant), tak by pokonywał szereg punktów w przestrzeni. Każdy z tych punktów jest zapisywany w pamięci, aby później robot na podstawie ich znajomości mógł odtworzyć całą trajektorię. Sterowanie robotem, w którym ten przemieszcza się między punktami jest określane jako PTP (point-to-point).

Programowanie ciągłe jest stosowane w przypadku, kiedy wymagane są płynne ruchy ramion robota wzdłuż skomplikowanych krzywych. W tym wypadku podczas nauki osoba programująca prowadzi końcówkę robota – jeżeli robot jest mały, albo końcówkę urządzenia symulującego końcówkę ramienia robota. W tym wypadku zapamiętywane są setki tysięcy



punktów tworzących skomplikowaną trajektorię, między którymi później program interpoluje drogę. Ten sposób sterowania określa się jako CP (continuous-path).

W przypadku programowania off-line odbywa się ono poza stanowiskiem robota, a więc robot może wykonywać nadal swój stary program, podczas gdy jest już tworzony nowy. W celu programowania stosuje się języki programowania i programy przeznaczone dla danych typów robotów kompatybilnych z rozwiązaniami danego producenta robotów.

Układy sterowania robotów i manipulatorów

Układy sterowania robotów można generalnie podzielić na układy sterowania teleoperatorów, układy sterowania sekwencyjne oraz układy sterowania numeryczne klasy CNC (Computer Numerical Control).

W przypadku układów sterowania teleoperatorów operator zadaje ruchy robotowi poprzez urządzenie zwane zadajnikiem, a układ sterujący przetwarza ruchy zadajnika oraz jego położenie, prędkości i przyspieszenia na sygnał sterujący ramionami robota. Tak więc w tym wypadku człowiek stanowi element układu sterowania.

Układy sekwencyjne bazują obecnie na sterownikach PLC. Programy wykonywane przez te układy bazują głównie na przechodzeniu między kolejnymi stanami, przy czym warunki przejścia są opisywane funkcjami logicznymi.

Układy komputerowe sterowania numerycznego są oparte głównie na układach mikroprocesorowych, są stosowane w przypadku robotów mogących wykonywać skomplikowane płynne ruchy. Struktura sprzętowa takiego układu sterowania może składać się z:

- Zasilacza
- Jednostki centralnej
- Układu odpowiedzialnego za poruszanie robotem
- Obwodów wejścia/wyjścia
- Interfejsu użytkownika
- Programatora ręcznego

Sama jednostka centralna składa się z procesora centralnego, pamięci RAM sterownika PLC, dysku twardego oraz modułu komunikacyjnego. Procesor centralny pełni rolę komputera centralnego przetwarzającego główny program. Rolą sterownika PLC jest sterowanie i nadzór nad urządzeniami wejścia/wyjścia. Monitoruje także zmienne wejściowe i procesowe.



Głównymi składnikami układu odpowiedzialnego za poruszanie robotem są procesor interpolatora, oraz sterowniki serwonapędów poszczególnych osi. Interpolator przyjmuje od procesora centralnego współrzędne docelowego położenia przegubów robota oraz dane określające rodzaj trajektorii i prędkość ruchu. Po przeliczeniu tych danych według algorytmów interpolujących, generuje on sygnały reprezentujące wartości przemieszczeń dla każdej osi osobno i przesyła je do poszczególnych sterowników serwonapędów. Interpolator jest specjalizowanym procesorem do szybkich obliczeń matematycznych jednakże w układach mniej skomplikowanych interpolacja może być wykonywana przez procesor główny

Chwytyki robotów

Niezbędnym wyposażeniem robotów przemysłowych są efektory, do których zalicza się urządzenia chwytające (chwytyki) oraz narzędzia wykonujące w procesie produkcyjnym zadania technologiczne.

Zadaniami urządzeń chwytających są:

- Pobranie (uchwycenie) obiektu manipulacji w położeniu początkowym
- Trzymanie obiektu (przedmiotu) w trakcie trwania czynności manipulacyjnych tzn. oddziaływanie na przedmiot siłą zapobiegającą zmianie jego położenia względem chwytaka
- Poprawiania (ewentualnego) orientacji manipulowanego przedmiotu
- Uwolnienia obiektu manipulacji w miejscu docelowym

Rozróżnia się następujące rodzaje chwytania przedmiotu:

- Chwytanie przez obejmowanie (kształtowe)
- Chwytanie cierne (siłowe)
- Chwytanie przez przyssanie
- Chwytanie magnetyczne

Podczas wstępnego doboru chwytaka należy określić najbardziej niezawodny sposób chwytania danego obiektu, przy czym głównym wyróżnikiem jest jego kształt. Ponieważ chwytyki mechaniczne mogą realizować chwytanie kształtowe oraz cierne występują ograniczenia związane z wymiarami manipulowanych obiektów oraz początkowym usytuowaniem manipulowanych przedmiotów, uniemożliwiającym dostęp do ich powierzchni.

Chwytyki podciśnieniowe mogą być stosowane tylko w przypadku obiektów mających powierzchnię o określonej gładkości i płaskości, które nie są zabrudzone (opiółki, kurz itp.). Ponadto ich wadą jest zużywanie się przyssawki w trakcie eksploatacji. Ponadto musi istnieć



możliwość wytworzenia podciśnienia, co może być trudne w przypadku blachy perforowanej. W przypadku chwytaka podciśnieniowego ze względu na konieczność wytworzenia podciśnienia sam proces uchwycenia trwa także dłużej niż w przypadku chwytaka mechanicznego. Ponieważ przyssawki są wytwarzane z różnego rodzaju gum, istotnym ograniczeniem w stosowaniu tych chwytaków jest temperatura manipulowanych obiektów. Tworzywa sztuczne nie są bowiem odporne na wysoką temperaturę.

Chwytyki elektromagnetyczne można natomiast stosować do obiektów ferromagnetycznych. Magnetyzm szczątkowy przyciąga drobiny metalowe, przez co zmniejsza się siłą chwytu i utrudnione jest uwolnienie obiektu. Chwytyki magnetyczne nie można także stosować, jeżeli technologia procesu nie dopuszcza magnesowania manipulowanych przedmiotów.

W tabeli nr 3 zebrano orientacyjne wskazania odnośnie możliwości zastosowania danego chwytaka w zależności od kształtu obiektu manipulacji.

Tablica 3. Możliwość stosowania danych chwytaków w zależności od kształtu obiektu manipulacji

Chwytyki Obiekt manipulacji	Mechaniczny	Podciśnieniowy	Elektromagnetyczny
Wąłki, tuleje	Tak	Nie	Warunkowo (tylko płaskie krążki)
Płytki	Warunkowo	Tak	Tak
Arkusze blach, płyty	Nie	Tak	Tak
Prostopadłościany	Tak	Warunkowo	Tak
Obiekty o złożonych kształtach	Tak (specjalne konstrukcje)	Nie	Tak (z wieloma magnesami)

Bibliografia

- Dębowski A., Automatyka Podstawy teorii, Wydawnictwo Naukowe PWN, WNT, 2018.
- Kasprzyk J., Hajda J., Programowanie sterowników PLC, Wydawnictwo Pracowni Komputerowej Jacka Skalmierskiego, 1998.
- Kostro J., Elementy, urządzenia i układy automatyzacji, Wydawnictwa Szkolne i Pedagogiczne, 1993.
- Mikulczyński T., Samsonowicz Z., Więclawek R., Automatyzacja procesów produkcyjnych, wydawnictwo: Mikulczyński Tadeusz, Samsonowicz Zdzisław, Więclawek Rafał, 2015.



Fundusze Europejskie
dla Wielkopolski

Dofinansowane przez
Unię Europejską



SAMORZĄD
WOJEWÓDZTWA
WIELKOPOLSKIEGO

Szelerski M. W., Automatyka przemysłowa w praktyce, Projektowanie, modernizacja i naprawa, Wydawca: KaBe, 2016.

Szelerski M. W., Robotyka przemysłowa Teoria, budowa, eksploatacja, Wydawnictwo KaBe, 2019.

Morecki A., Knapczyk J., Podstawy robotyki teoria i elementy manipulatorów i robotów, WNT, 2001.

Kaczmarek W., Panasiuk J., Robotyzacja procesów produkcyjnych, Wydawnictwo Naukowe PWN, 2017.

Rumatowski K., Podstawy regulacji automatycznej, Wydawnictwo Politechniki Poznańskiej, 2008.

Puławczewski J., Szacki K., Manitus A., Zasady automatyki, Wydawnictwo naukowo-techniczne 1974.

Honczarenko J., Roboty przemysłowe Budowa i zastosowanie, Wydawnictwo Naukowo-Techniczne, 2010.



Andrzej Milecki

PROJEKTOWANIE LINII ZAUTOMATYZOWANYCH

Wprowadzenie

Automatyzacja masowych procesów produkcyjnych została po raz pierwszy zastosowana na początku lat 50. XX wieku w firmie Ford. Wprowadzono w niej ruchomą linię montażową, która znacznie zwiększyła szybkość i wydajność produkcji. Ta rewolucyjna innowacja polegała także na tym, że każdy pracownik wykonywał wielokrotnie to samo zadanie. Takie podejście usprawniło produkcję i umożliwiło zwiększenie wydajności. W latach 1920-1921 obniżono cenę samochodu Forda do 355 dolarów, a wielkość produkcji osiągnęła 1 250 000 sztuk.

W latach 60. XX w. elektronika osiągnęła na tyle wysoki poziom rozwoju, że przestała być przedmiotem zainteresowania wyłącznie twórców radioodbiorników, telewizorów i komputerów budowanych na bazie układów cyfrowych TTL średniej skali integracji, a stała się bardzo interesująca dla konstruktorów i wytwórców urządzeń przemysłowych. To wtedy rozwinęły się np. serwonapędy prądu stałego a za nimi obrabiarki sterowane numerycznie, roboty oraz zautomatyzowane i skomputeryzowane linie produkcyjne. Układy i elementy pomiarowe bezpośrednio współpracowały z urządzeniami elektronicznymi. Elektronika zaczęła pojawiać się w samochodach i maszynach roboczych, a coraz większa rzesza hobbystów zaczęła budować różnorodne drobne urządzenia elektroniczne, typu zamki, sterowniki oświetlenia, zabawki itp. To wtedy także pojawiły się poważne konstrukcje, w których zastosowano komputery do sterowania maszynami.

Najważniejsze postępy w automatyzacji zostały dokonane dzięki rozwojowi, w zakresie elektroniki i automatyzacji. Ich wprowadzenie do budowanych urządzeń stało się najczęściej stosowaną metodą prowadzącą do poprawy ich właściwości i funkcjonalności. W drugiej połowie lat 70. XX w. zaczęto produkować mikroprocesory, a na początku lat 80. Mikrokontrolery, które były miniaturowymi sterownikami. Dzięki miniaturyzacji możliwe stało się ich stosowanie w niemalże dowolnym urządzeniu. Oznaczało początek używania bardziej zaawansowanych technik automatyzacji w przemyśle.

W wytwarzanych już od połowy lat 70. XX w. urządzeniach powszechna zaczęła być koegzystencja mechaniki z elektroniką. Wtedy to powstała nowa technika projektowania, nazywana mechatroniką. Równolegle rozwijana była automatyzacja całych linii



produkcyjnych, na których coraz częściej zaczęły pojawiać się roboty przemysłowe. Automatyzacja linii produkcyjnych i to warunek rozwoju oraz utrzymania się na rynku każdej z firm. Zautomatyzowane linie produkcyjne cechują się wysoką powtarzalnością i dokładnością. Bez tego trudno jest uzyskać najwyższą jakość oferowanych produktów. Czynności powtarzalne, monotonne można zautomatyzować za pomocą robotów, które wykonują je o wiele szybciej i dokładniej niż człowiek. Skróceniu ulega czas, niezbędny do wykonania określonych działań i wytworzenia produktu.

Rozwój techniki mikroprocesorowej doprowadził do nowych możliwości przetwarzania i przepływu informacji w sterowanych urządzeniach. Zalety mikroprocesora, takie jak miniaturowe rozmiary, duże moce obliczeniowe, łatwość programowania oraz niska cena pozwoliły na łatwe ich wbudowywanie do urządzeń. Mikrokontrolery są używane w niemalże każdym urządzeniu, np. w robocie, podajniku, podnośniku oraz w maszynach: tokarkach, wiertarkach, frezarkach, prasach, krawędziarkach itp. Dzisiaj prawie każde zaawansowane technicznie urządzenie, ma przynajmniej prosty moduł elektroniczny z małym mikrokontrolerem wykonującym pomiary oraz sterującym silnikami i innymi elementami wykonawczymi.

Prawdziwą rewolucją w automatyce było upowszechnienie się komputerów klasy PC. Po odpowiednim dostosowaniu do warunków przemysłowych, zaczęły one służyć zarówno do nadzorowania układów sterowania, jak i do sterowania zaawansowanymi stacjami i całymi liniami produkcyjnymi. Wraz z ich wprowadzeniem automatyka uzyskała nowe możliwości, a zakres automatyzacji poszerzył się o zupełnie nowe obszary zastosowań, związane z kompleksową automatyzacją produkcji oraz z elastycznym zarządzaniem i sterowaniem przedsiębiorstwem. Dzisiaj trudno jest spotkać nowoczesną linię produkcją bez sterowania komputerowego.

1. Podstawy automatyzacji urządzeń

Automatyzacja w praktyce sprowadza się do zastosowania urządzeń pomiarowych i napędów oraz sterowników, dzięki którym urządzenia albo procesy produkcyjne będą działały samoczynnie, bez albo z ograniczonym udziałem ludzi – pracowników. W celu automatyzacji stosowane są następujące elementy i/albo urządzenia:

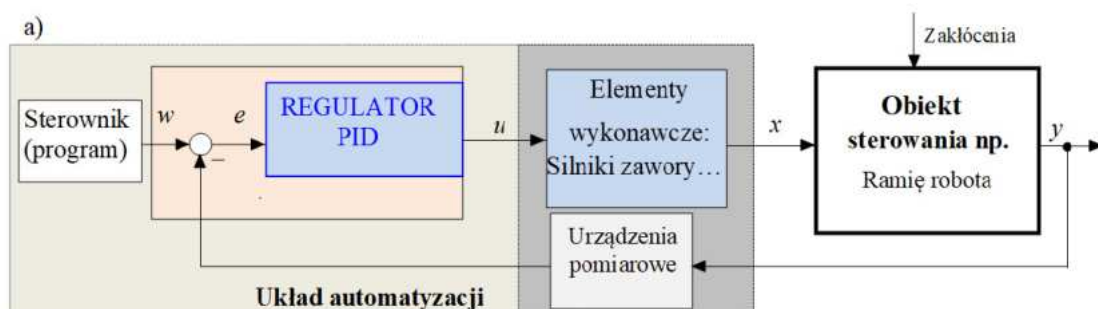
- 1) pomiarowe (czujniki, przetworniki oraz zespoły pomiarowe)



- 2) wykonawcze (silniki, zawory, zasuw, elektromagnesy, siłowniki, zespoły napędowe, pompy, podajniki, regulatory)
- 3) części centralnej i sterującej (panele operacyjne, przełączniki, nastawniki, wskaźniki, wyświetlacze)
- 4) sterowniki (PLC, regulatory, komputery przemysłowe – IC, moduły komunikacji).

Pierwsze z nich przeznaczone są do mierzenia wielkości wyjściowych z poszczególnych, sterowanych podzespołów. Najczęściej mierzone są wielkościami mechanicznymi takimi jak: siła, przyspieszenie, prędkość, położenie, ciśnienie, temperatura itp. Do tego stosowane są różne czujniki i urządzenia pomiarowe, które generują na wyjściach sygnały elektryczne, czyli napięcie U albo natężenie prądu I a niekiedy impulsy. Definicja mówi, że **sygnał to dowolna wielkość fizyczna będąca funkcją (zmienną) w czasie**. Ze względu na rodzaj sygnału, w automatyzacji wyróżnia się sygnały elektryczne:

- **analogowe** (ciągłe) – wartość sygnału określona jest przez wartości wielkości fizycznej i może przyjmować dowolną wartość z określonego przydziału, zwykle: napięciowe $-10V$ do $+10V$, 0 do $10V$, 0 do $5V$ albo prądowe 0 do 20 mA , 4 do 20 mA)
- **binarne** (przełączające, dwustanowe) – sygnał może przyjmować tylko jedną z dwóch wartości: 0 lub 1 , kodowaną odpowiednio jako: $0V$ albo $24VDC$.

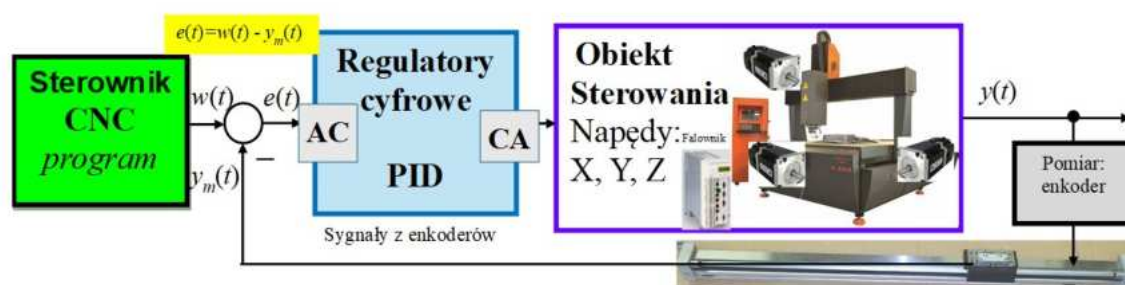


Rys. 1. Zamknięty układ automatyzacji schemat blokowy: w – sygnał zadany, e – uchyb regulacji, x – sygnał sterujący, y – sygnał wyjściowy.

W zależności od tego, jaki rodzaj sygnałów jest wykorzystywany w procesie sterowania, mówimy o sterowaniu analogowym (ciągłym) albo o sterowaniu binarnym. W pierwszym przypadku powszechnie stosuje się układy sterowania ze sprzężeniem zwrotnym, polegającym na pomiarze rzeczywistej wartości sygnału y na wyjściu sterowanego obiektu i przekazaniu go na wejście, w celu porównania z wartością sygnału zadanego w (rys. 1). Porównanie to polega na wyznaczeniu różnicy między tymi sygnałami. Uzyskany wynik $e = w - y$ nazywa się błędem



albo uchybem regulacji. Wykorzystuje się go do skorygowania oddziaływania regulatora sterującego obiektem, w taki sposób, aby zminimalizować ten uchyb, a w ostateczności doprowadzić go do zera. Taki układ nazywa się zamkniętym układem regulacji z ujemnym sprzężeniem zwrotnym (rys. 1). Sygnały wyjściowe są mierzone za pomocą różnego rodzaju czujników i urządzeń pomiarowych. Od ponad 100 lat, do regulacji stosowane są głównie regulatory proporcjonalno-całkująco-różniczkujące (ang. Proportional-Integral-Derivative – PID). Wielkości fizyczne, które stanowią wyjście automatyzowanego urządzenia są przedmiotem oddziaływania układu regulacyjnego. Nazywamy je wielkościami bądź sygnałami regulowanymi y , np.: siła, ciśnienie, przyspieszenie, prędkość, przemieszczenie, temperatura. Oddziaływanie regulatora na przebieg procesu odbywa się za pośrednictwem elementów wykonawczych. Wartości zadane są najczęściej ustalane przez systemu sterowania nadrzędnego, w którym zapisany jest program działania urządzenia a nawet całej linii produkcyjnej. W układach regulacji ręcznej, rolę regulatora pełni człowiek.



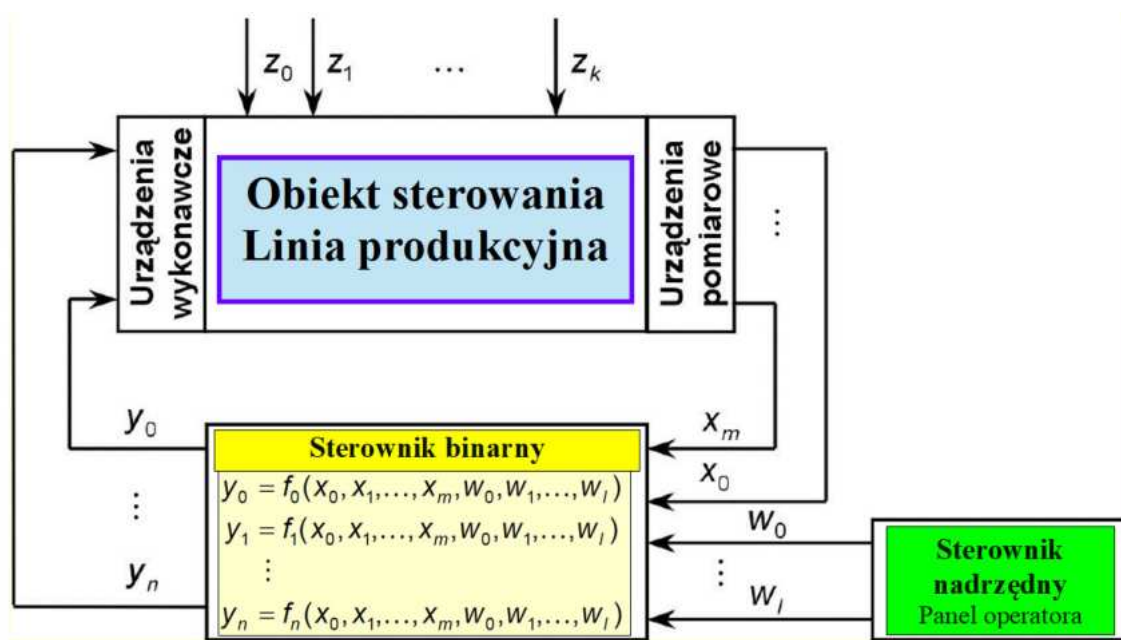
Rys. 2. Zamknięty układ automatyzacji obrabiarki CNC

Na rysunku 2 przedstawiono schemat układu sterowania obrabiarki CNC. W samym sterowniku CNC jest zapisany program obróbki tj. kolejne pozycje na trajektorii i parametry ruchu narzędzia itp. Na tej podstawie sterownik CNC, generuje sygnały zadane do poszczególnych sterowników elektronicznych silników obrabiarki. Przesuwają one elementy obrabiarki, tzn. najczęściej uchwyt z narzędziem obróbkowym np. wiertło, nóż tokarski albo frez. Przesuwania w poszczególnych osiach są mierzone za pomocą elementów pomiarowych, których sygnały wyjściowe są podawane jako sprzężenie zwrotne na wejście układu – do węzła odejmującego. Na wejściach i wyjściach regulatora zastosowane są przetworniki Analogowo-Cyfrowe (AC) oraz Cyfrowo-Analogowe (CA), które zamieniają odpowiednio sygnały ciągłe w postaci napięcia na liczby oraz liczby na napięcie.



W praktyce przemysłowej, większość procesów na liniach produkcyjnych, dotyczy jednak wykonywania operacji typu włącz/wyłącz. W związku z tym, w układzie sterowania występują sygnały binarne „0” albo „1”, kodowane napięciem 24VDC (poziom wysoki tzn. element jest obecny albo włącz) oraz 0V (poziom niski – brak elementu albo wyłącz). Automatyzacji podlega wielowejściowy obiekt, składający się z szeregu urządzeń tj. napędów, transporterów, podajników, zespołów pomiarowych, robotów i chwytaków, obrabiarek, automatów montażowych, pakujących itp. W każdym z tych elementów zastosowane są różne czujniki i układy pomiarowe generujące na wyjściu sygnały binarne 0/1 (jest/nie ma) oraz elementy wykonawcze typu silniki, elektromagnesy, elektrozawory itp.

Do sterowania układów binarnych stosowane są sterowniki przemysłowe, realizujące funkcje boolowskie, określone macierzą $F=[f_0(x_1, x_2, \dots, x_m, w_1, \dots), f_1(\dots), \dots, f_n(\dots)]$. Zaprogramowany **sterownik binarny** (ang. Programmable Logic Controller – PLC) odpowiedzialny jest za generowanie zadań sterowania, w postaci sygnałów binarnych. W każdej chwili czasowej (kroku), do sterownika podawany jest wektor wejść $X=[x_0, x_1, \dots, x_m]$ oraz wektor sygnałów zadanych $W=[w_1, w_2, \dots, w_l]$, a sterownik wysyła do urządzeń wykonawczych wektor, określający stany wyjść $Y=[y_1, y_2, \dots, y_n]$. Schemat układu sterowania binarnego pokazano na rys. 3.



Rys. 3. Schemat blokowy automatyzacji binarnej linii produkcyjnej: W – wektor sygnałów zadanych, X – wektor sygnałów mierzonych, F – wektor funkcji sterowania, Y – wektor sygnałów wyjściowych, Z – wektor zakłóceń (wszystkie sygnały są binarne 0/1)



Obecnie najczęściej stosowanymi formami automatyzacji są:

- sztywna (Fixed Automation) – stosowana do produkcji dużej ilości danego lub bardzo podobnych do siebie produktów. Zakłada się, że ich produkcja będzie trwała przynajmniej kilka lat i nie będą wprowadzane do jej automatyzacji żadne zmiany. Obecnie, ta forma automatyzacji w zasadzie nie jest stosowana.
- programowalna (Programmable Automation) – zaprojektowana do miejsc, gdzie produkty wytwarza się w dużych seriach. Działanie linii produkcyjnej może być dostosowywane poprzez zmianę programów sterowania, do aktualnych potrzeb firmy.
- procesowa (Process Automation) – stosowana w przemysłach, gdzie produkcja odbywa się w sposób ciągły m.in. w rafineriach, elektrowniach, czy zakładach przetwórstwa chemicznego. Koncentruje się na monitorowaniu i kontrolowaniu zmiennych procesowych, takich jak temperatura, ciśnienie, przepływ, skład chemiczny itp.
- elastyczna (Flexible Automation) – zaawansowana forma automatyzacji umożliwiająca szybkie dopasowanie produkcji do aktualnie wytwarzanych elementów bez żadnych przestojów. Przykładem tego typu rozwiązania są np. systemy produkcji just-in-time w produkcji samochodów.
- zintegrowana (Integrated Automation) – polega na integracji wszystkich systemów przedsiębiorstwa od zarządzania, przez projektowanie, produkcję, aż po dystrybucję. Jej przykładem są systemy zarządzania wszystkimi zasobami przedsiębiorstwa typu ERP (Enterprise Resource Planning) albo systemy sterowania cyklem życia produktu (PLM).

W automatyzacji sztywnej algorytm sterowania zostaje zapisany w sposób trwały w postaci połączeń elektrycznych poszczególnych elementów elektrycznych i w przypadku konieczności zmiany programu, należy zmienić te połączenia oraz zastosować nowe elementy w układ sterowania. Przykładem może być układ przekaźnikowy albo scalony typu, płytka z elementami elektronicznymi i/albo elementy elektromechaniczne. Zmienne-programowe układy sterowania tzn. automatyzacja programowalna, procesowa i zintegrowana, to układy sterujące, których algorytm sterowania został zapisany w sposób nietrwały w pamięci urządzenia sterującego i w przypadku konieczności zmiany programu należy tylko zmienić zawartość pamięci układu.

System automatyki to najbardziej rozbudowane rozwiązanie techniczne w zakresie automatyzacji. Składa się z wielu elementów, układów i urządzeń automatyki połączonych ze



sobą w jedną całość – system. Realizuje on kompleksowo problem sterowania całym zakładem albo jego złożonym fragmentem, np. linią do produkcji karoserii albo stacją obróbczą itp.

2. Wybrane elementy stosowane w automatyzacji

Wprowadzanie automatyzacji do sterowania linią produkcyjną wymaga zastosowania na niej różnorodnych elementów automatyzacji. Na rynku dostępna jest obecnie cała ich gama i konieczna jest ich znajomość, aby właściwie je dobrać i zaprojektować układ automatyzacji linii produkcyjnej. Dzięki osiągnięciom mikroelektroniki powstały także różnorodne elementy wykonawcze. Elementy automatyki to człony spełniające w układzie automatyki bądź w urządzeniu automatyzowanym proste funkcje, takie jak: zmiana postaci sygnałów, ich wzmocnienie oraz porównanie. Elementem jest więc czujnik wykrywający albo pomiarowy, wzmacniacz sygnału, elektrozawór, silnik, prądnica, grzałka itp., a także dowolny regulator albo sterownik.

Człony spełniające funkcje bardziej złożone nazywane są urządzeniami automatyki, np. urządzenia pomiarowe (zawierające czujniki, przetworniki, wzmacniacze, filtry itp.), urządzenia napędowe zawierające: sterownik-falownik, czujnik pomiaru położenia, silnik, przekładnie, prowadnice itp. Układ automatyki powstaje z połączenia elementów i urządzeń w pewien zespół, wykonujący określone zadanie. Możemy przykładowo rozróżnić układ sterowania robotem, obrabiarką itp. Z kolei system automatyki to kompletny zestaw elementów, urządzeń i układów, czyli zestaw aparatury kontrolno-pomiarowej, które zastosowano do automatyzacji przynajmniej np. stacji obróbczej. Zestaw taki, aby mógł być nazwany systemem, musi być na tyle duży i kompletny, aby umożliwiał sterowanie liniami produkcyjnymi albo wręcz całymi fabrykami. Systemy do automatyzacji linii przemysłowych wymagają zastosowania czujników do wykrywania obiektów, pomiaru parametrów procesu oraz szeregu różnych napędów elektrycznych, pneumatycznych i hydraulicznych. Do sterowania stosowane są różne sterowniki, połączone w sieć. Odpowiednio dobrane mogą sprostać wielu wyzwaniom stając się niezawodnym rozwiązaniem w wielu aplikacjach.

Czujniki binarne

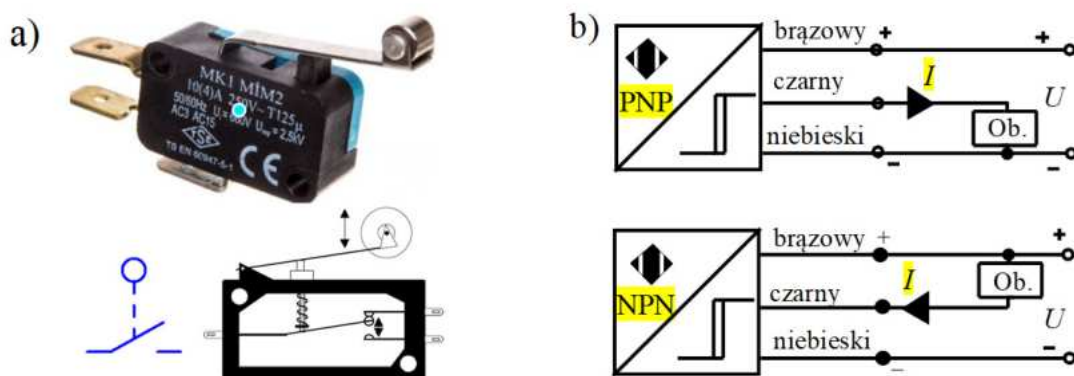
Czujniki przemysłowe to elementy układu automatyki, które przeznaczone są do wykrywania i rejestrowania sygnałów z urządzeń na liniach produkcyjnych. W binarnych układach automatyzacji stosowane są różnego rodzaju czujniki wykrywające obecność lub brak, jakiegoś



wybranego elementu np. obecność butelki, nakrętki, śruby itp., bądź parametru np. przekroczenie ciśnienia 8 Pa albo temperatury 50°C. Stosowane są czujniki stykowe (włączniki, wyłączniki) w postaci tzw. krańcówek oraz bezkontaktowe czujniki zbliżeniowe. Najbardziej popularne są czujniki: indukcyjnościowe, pojemnościowe, fotoelektryczne albo ultradźwiękowe. Wykorzystują one różne zasady działania i mogą wykrywać elementy, wykonane z różnych materiałów, ale ich wspólną cechą jest zwieranie i rozwieranie połączenia elektrycznego na wyjściu czujnika, za pomocą styków albo tranzystorów.

Najstarsze, ale też najtańsze, są czujniki mechaniczne – stykowe, które do zmiany stanu wymagają bezpośredniego przełączenia zestyków przez wykrywany obiekt (rys. 4a). Są stosowane do wykrywania położenia lub obecności obiektów w zadanych położeniach, najczęściej na krańcach ruchu. Stosowane są w nich elementy, które poruszają się w reakcji na kontakt z obiektem. Przykładem może być włącznik albo przycisk aktywowany ręcznie albo przez kontakt mechaniczny np. z wózkiem. Dzięki swojej prostocie są niezawodne, choć ich wadą jest iskrzenie, zużywanie się styków, ich drgania itp.

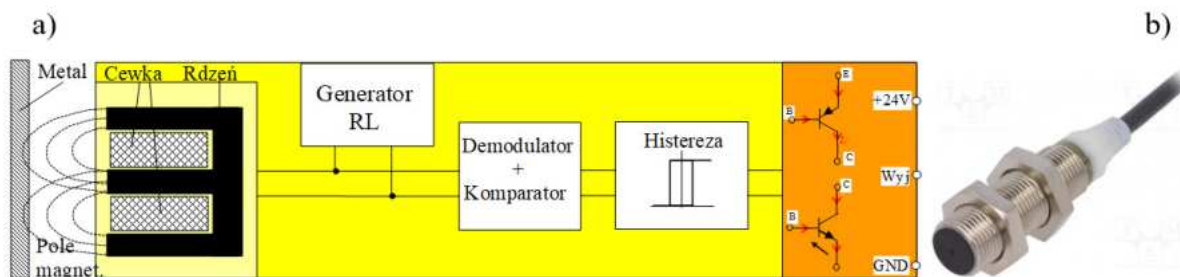
Dwustanowe wyjście wszystkich bezstykowych, czujników umożliwia bezpośrednią współpracę z przekaźnikami i programowanymi sterownikami logicznymi (PLC). W stopniach wyjściowych tych czujników stosowane są tranzystory typu NPN albo PNP. Każde z tych dwóch typów wyjść jest zwykle zabezpieczone przez nieprawidłowym połączeniem oraz przed przepięciem za pomocą diod. Czujniki są zwykle zasilane napięciem stałym z zakresu od 10 do 30 V. Stosowane w nich tranzystory wyjściowe mogą być obciążone prądem równym maksymalnie 0,2 A. Schemat i sposób połączenia tych czujników pokazano na rys. 4b. Istotny jest kierunek przepływu prądu wyjściowego: w stanie „1” w czujniku PNP prąd wypływa z czujnika, a w czujniku PNP wpływa.



Rys. 4. Schematy czujników: a) krańcówki aktywowane przez nacisk (zdjęcie, symbol, budowa),
b) bezkontaktowy czujnik indukcyjnościowy z wyjściami tranzystorowymi typu: PNP i PNP



Czujniki indukcyjne są elementami automatyki wykrywającymi wprowadzenie metalu w strefę czułości/działania czujnika. W układach automatyki zastępują one stykowe wyłączniki krańcowe, łączniki drogowe itp. Służą do określenia położenia ruchomych części maszyn i urządzeń. Stosuje się je tam, gdzie są wymagane: duża niezawodność, częstotliwość przełączania, precyzja i powtarzalność działania oraz możliwość w pracy w trudnych warunkach środowiskowych.



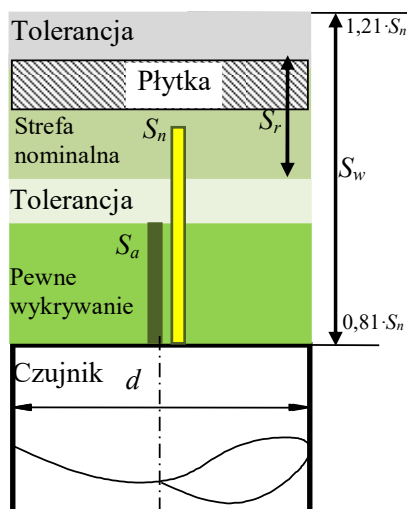
Rys. 5. Budowa bezkontaktowego czujnika indukcyjnego: a) schemat, b) widok

Częścią wykrywającą obiekt jest cewka nawinięta na rdzeniu ferrytowym, połączona z generatorem (rys. 5). Generuje ona w otoczeniu czoła czujnika zmienne pole magnetyczne, które w metalu wytwarza prądy wirowe, co z kolei powoduje „obciążenie” układu oscylatora i w efekcie spadek amplitudy oscylacji. Zmiany te są wykrywane przez komparator i przy pewnej, charakterystycznej dla danego typu czujnika odległości obiektu metalowego od jego czoła, następuje skokowa zmiana stanu (napięcia) na wyjściu komparatora. Strefa działania to podstawowy parametr czujnika bezkontaktowego indukcyjnego (rys. 6). Wynika ona z typu oraz gabarytów czujnika. W katalogach wyrobu i na tabliczkach znamionowych jest podawana najczęściej strefa nominalna S_n . Wynosi ona zwykle od ok. 0,5 do 1,0 średnicy czujnika. Strefa rzeczywista S_r (uwzględnia fabryczną tolerancję wykonania wyrobu), spełnia warunek:

$$0,9 S_n < S_r < 1,1 S_n. \quad (1)$$

Pomiar strefy rzeczywistej w warunkach fabrycznych polega na zbliżeniu do powierzchni czołowej w osi czujnika kwadratowej płytki ze stali o grubości 1 mm i o boku równym średnicy czujnika. Najistotniejsza jest jednak strefa robocza S_w , która gwarantuje działanie czujnika w pełnym zakresie temperatury i napięcia zasilającego dla długiego okresu eksploatacji. Jest ona kreślona nierównościami:

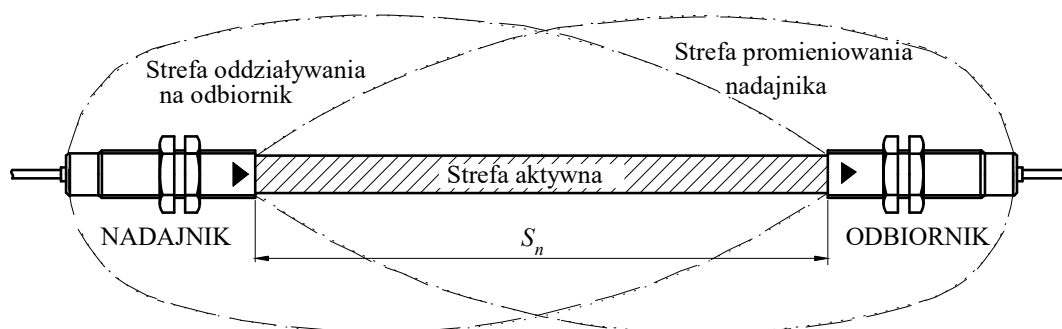
$$0,81 S_n < S_w < 1,21 S_n. \quad (2)$$



Rys. 6. Strefa działania czujnika

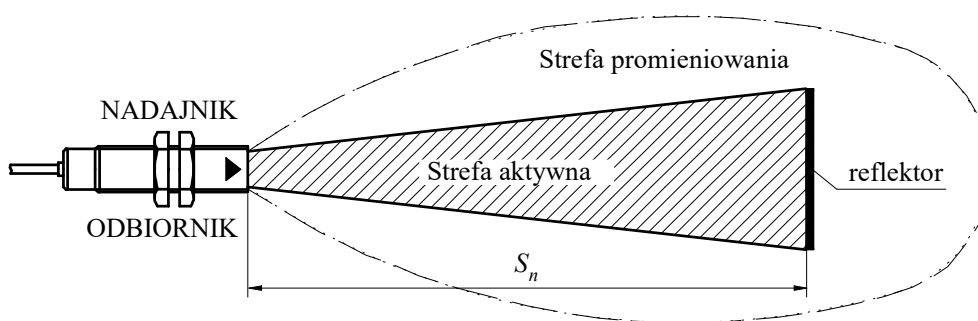
W realnych warunkach eksploatacji strefa działania może różnić się aż o 21% od wartości nominalnej. Strefa działania, tzn. odległość obiektu od czoła czujnika pojemnościowego, przy której następuje przełączenie wyjścia czujnika, zależy także od rodzaju materiału wykrywanego. W katalogach podaje się następujące współczynniki korekcyjne strefy działania: metale – 1,0; szkło – 0,5; PCW – 0,6; woda – 1,0; olej – 0,1. Maksymalna częstotliwość przełączania tych czujników wynosi od ok. 0,5 do 2,5 kHz.

W przemyśle najczęściej stosowane są czujniki indukcyjnościowe, ponieważ większość wykrywanych elementów wykonanych jest ze stali. Do wykrywania materiałów niemetalowych stosowane są czujniki pojemnościowe albo optyczne. W tych pierwszych, zamiast cewki indukcyjnej jest stosowany kondensator. Zbliżenie obiektu do czoła czujnika powoduje zmianę pojemności kondensatora, którego okładki znajdują się w „czole” czujnika. Zakres zmian pojemności zależy od przewodności wykrywanego materiału, jego stałej dielektrycznej, jak też jego masy, powierzchni i odległości od elektrody. Zmiana pojemności kondensatora, spowodowana zbliżeniem jakiegoś materiału/elementu powoduje wzrost napięcia generowanego przez oscylator RC w czujniku. Sygnał ten steruje wyjściem. Czujniki pojemnościowe to elementy automatyki reagujące na wprowadzenie w ich strefę czułości metali, szkła, tworzyw, drewna, materiałów sypkich, cieczy i różnych substancji. Ich częstotliwość pracy jest znacznie mniejsza od częstotliwości czujników indukcyjnych i wynosi zwykle do ok. 50 Hz.



Rys. 7. Jednowiązkowa bariera świetlna

Czujniki optoelektroniczne wykrywają obiekty, które znajdują się na drodze przebiegu wiązki światła. Wykorzystuje się je itp. do kontroli położenia ruchomych części maszyn, do identyfikacji obiektów znajdujących się w zasięgu działania czujników, itp. przesuwających się na taśmach transportowych, do określania poziomu cieczy i materiałów sypkich itp. Zaletą czujników optoelektronicznych jest duży zasięg działania uzyskiwany dla małych obudów czujników. Najczęściej czujniki te wykorzystują modulowane światło z zakresu bliskiej podczerwieni. Zaletą takiego rozwiązania jest niewrażliwość czujników na widzialne światło z otoczenia. Dodatkowo dzięki wzajemnej synchronizacji nadajnika i odbiornika, gwarantowana jest duża odporność na zakłócenia i możliwość pracy w warunkach zanieczyszczenia powietrza i przy zabrudzeniu układu optycznego czujnika. Wytworzony w nadajniku impuls świetlny, nawet osłabiony rozproszeniem, dociera do odbiornika lub nie, co jest sygnalizowane stanem 0 albo 1 na wyjściu. Zaletą czujników optycznych jest duża strefa działania, która może sięgać nawet kilkudziesięciu metrów. Zanieczyszczenie powietrza i zabrudzenie układu optycznego skraca tę strefę. Jednym z najstarszych czujników optycznych jest jednowiązkowa bariera świetlna, w której nadajnik i odbiornik są umieszczone w oddzielnych obudowach (rys. 7). Stosuje się także bariery świetlne, w których występuje wiązka (ściana) światła, przebiegająca od nadajników do odbiorników, które znajdują się naprzeciw siebie w skrajnych punktach zasięgu. Przesłonięcie wiązki promieni świetlnych przez obiekt powoduje przerwanie transmisji fali świetlnej i przełączenie obwodu wyjściowego czujnika.



Rys. 8. Czujnik refleksyjny

W czujnikach optoelektronicznych refleksyjnych i odbiciowych nadajnik i odbiornik są umieszczone we wspólnej obudowie (rys. 8). Są one skierowane w końcowy punkt zasięgu, w którym jest umieszczony specjalny reflektor odblaskowy. Od tego reflektora odbija się wysyłana przez nadajnik wiązka promieni świetlnych. Jej przesłonięcie przez obiekt powoduje przerwanie transmisji i przełączenie obwodu wyjściowego czujnika z „0” na „1”. Czujniki odbiciowe, wykrywają światło odbite nie od reflektora a od przedmiotu, którego obecność wykrywają. Podobnie jak reflektor, przedmiot wykrywany musi mieć zdolność odbijania światła.

Czujniki ultradźwiękowe wykorzystują fale ultradźwiękowe (o częstotliwości powyżej 20 kHz), do wykrywania elementów albo do pomiaru ich odległości od czujnika. Jeśli w strefie działania czujnika jest jakiś obiekt, to wysyłane przez nadajnik impulsy dźwiękowe są od niego odbijane i wracają do odbiornika. Mierzony jest czas upływający od wysyłania impulsu ultradźwiękowego do powrotu echa odbitego od przeszkody. Odległość jest obliczana na podstawie tego czasu i prędkości dźwięku. Czujniki ultradźwiękowe pozwalają na bezdotykowe wykrycie obiektów metalowych i niemetalowych. Są wykorzystywane głównie do pomiaru odległości obiektów od czujnika oraz do detekcji przedmiotu lub poziomu cieczy.

3. Czujniki analogowe

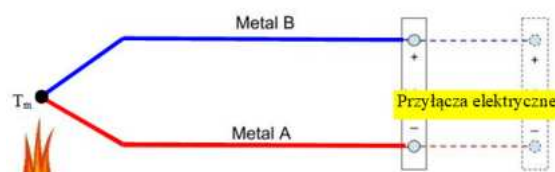
Osobną grupę elementów automatyki stanowią czujników pomiarowe. Zwykle muszą one współpracować z odpowiednim do nich układem elektronicznym, który formuje elektryczny sygnał wyjściowy, w taki sposób, aby jego wartość była proporcjonalna do wartości mierzonej. Czujniki te mierzą różne wielkości fizyczne, pozwalając na monitorowanie parametrów takich jak położenie, prędkość, siła, naprężenie ciśnienie, temperatura. Przetwarzają one wielkość fizyczną na sygnał elektryczny.



Przykładem czujnika analogowego jest termopara. Temperatura jest jedną z wielkości fizycznych i procesowych, która musi być mierzona w wielu procesach technologicznych i przesyłowych. W automatyce, celu dokonania pomiaru temperatury stosuje się różnego rodzaju przetworniki z wyjściem elektrycznym. Bardzo często do pomiaru temperatury wykorzystywane są termopary (termoogniwa). Są to czujniki temperatury wykorzystujące zjawisko Seebecka (rys. 9). Polega ono na powstawaniu zależnej od temperatury siły elektromotorycznej na styku dwóch różnych metali. Termopara składa się z dwóch różnych metali (drucików), spojenych na jednym końcu (punkt pomiaru).

Najczęściej stosowane są termopary typu:

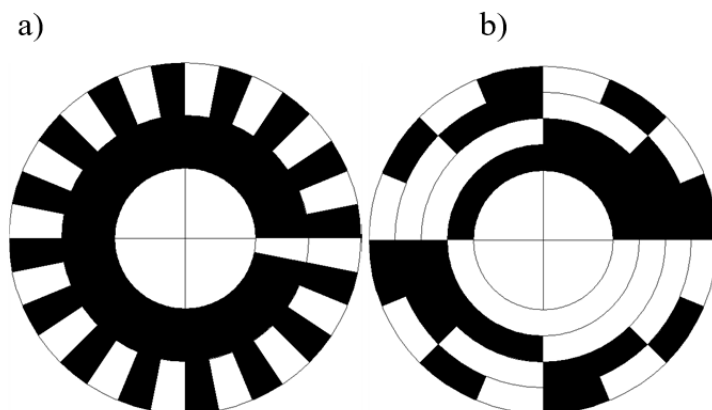
- T: miedź (Cu) — konstantan (Cu+Ni),
- J: żelazo (Fe) — konstantan (Cu+Ni),
- K: chromel (Ni+Cr) — alumel (Ni+Al),
- N: (Ni+Cr+Si) — (Ni+Si),
- E: chromel (Ni+Cr) — konstantan (Cu+Ni).



Rys. 9. Termopara

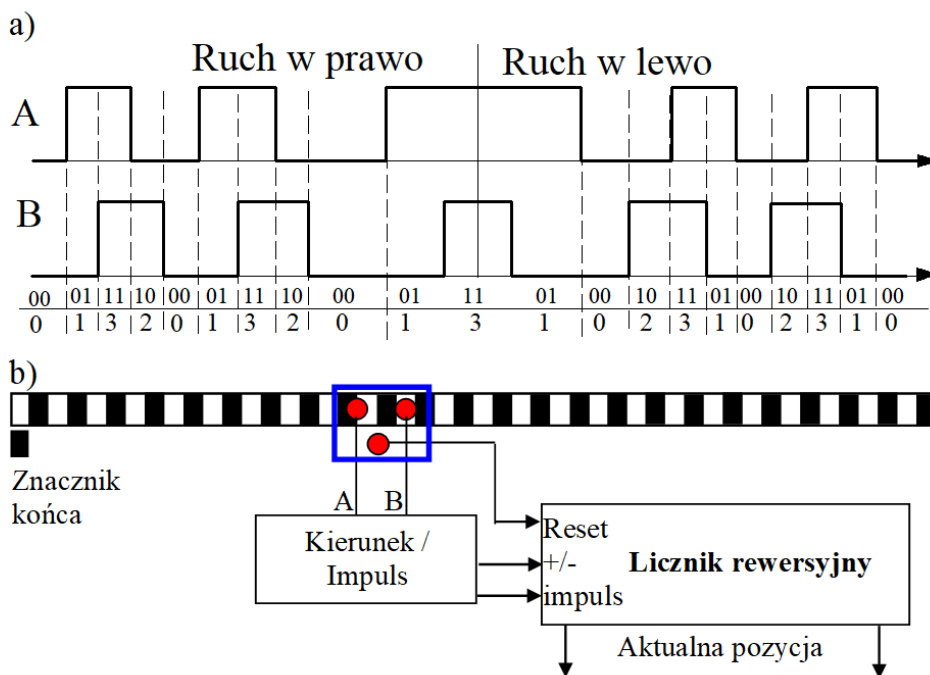
Przykładowo, termopara J może mierzyć temperatury od -210 do $+1200$ °C, generując napięcie od $-8,1$ do $69,5$ mV.

Termorezystory to czujniki, których rezystancja zmienia się wraz ze zmianą temperatury. Wielkość zmiany rezystancji określana jest przez temperaturowy współczynnik rezystancji. Termorezystory platynowe mogą mierzyć temperaturę w zakresie -200 °C ÷ $+850$ °C. W pomiarach przemysłowych najczęściej wykorzystuje się czujniki platynowe Pt100 o rezystancji 100 Ohm w temperaturze 0 °C. W automatyce wykorzystuje się również termorezystory niklowe typu: Ni100, Ni1000, które mogą być wykorzystywane do pomiaru temperatury w zakresie -60 °C ÷ $+180$ °C. Termorezystory miedziane zastosowanie znajdują głównie w chłodnictwie. Ich zakres pomiarowy to 0 °C ÷ $+150$ °C. Drugim sposobem pomiaru temperatur są pomiary bezkontaktowe, przy pomocy pirometrów oraz kamer termowizyjnych. Pirometr mierzy głównie temperaturę punktu lub małej powierzchni elementu. Kamera termowizyjna rejestruje rozkład temperatur na powierzchni, tj. kilkaset tysięcy punktów w przestrzeni obejmowanej przez jej obiektyw. Pozwala na obrazowanie rozkładu przestrzennego pola temperatury na powierzchni obiektu.



Rys. 10. Tarcze enkoderów: a) inkrementalnego, b) absolutnego 4-ro bitowego

Do pomiaru pozycji i przemieszczenia bardzo powszechnie stosowanymi w systemach automatyki są enkodery (rys. 10). Budowane są wersje obrotowe i liniowe. Dzieli się je na przyrostowe (inkrementalne) i absolutne (kodowe). W trakcie ruchu, enkodery inkrementalne generują impulsy, które są zliczane przez specjalne liczniki rewersyjne tj. dodające-odejmujące. Enkoder pokazany na rysunku 10a generuje 16b impulsów na obrót. Enkodery absolutne generują liczby, określające aktualne położenie. Pokazany na rysunku, generuje na jeden obrót 16 liczb 4-ro bitowych.



Rys. 11. Enkoder inkrementalny: a) generowane impulsy, b) liniał podłączony do układu pomiarowego z licznikiem rewersyjnym



Enkodery liniowe wyposażone są w liniał, będący wzorcem długości (rys. 11a). Ma on naniesione w równych odległościach, na przemian ciemne i jasne pola albo szczeliny, przez które przechodzi wiązka nadawana do odbiornika. Wzdłuż liniału przesuwa się czytnik z fotodiodą, rozpoznający jasne i ciemne pola albo szczeliny. Od liczby szczelin zależy rozdzielczość urządzenia. Na tej podstawie odbiornik wysyła sygnały w kształcie impulsów 0 albo 1. Enkodery przyrostowe, inaczej inkrementalne, generują impuls co określone przesunięcie np. co $10\ \mu\text{m}$ i dlatego muszą współpracować ze specjalnym licznikiem rewersyjnym, czyli liczącym w przód (dodającym) i w tył (odejmującym). Enkodery te wyposażone są w dwa czujniki generujące sygnały A i B przesunięte względem siebie (rys. 11b) o $+90^\circ$ albo o -90° . Na podstawie tego przesunięcia układ licznika wykrywa czy ruch odbywa się w lewo i wtedy dodaje impulsy, czy też w prawo i wtedy odejmuje impulsy (rys. 11c). Zmiana stanu licznika dokonywana jest przy zmianie dowolnego zbocza dla obu sygnałów A i B. Po włączeniu zasilania konieczne jest przesunięcie głowicy pomiarowej czujnika do punktu odniesienia, tj. krańcowego, co zeruje licznik. Stanowi to istotną niedogodność enkodera inkrementalnego.

Rozwiązaniem konieczności zerowania układu pomiarowego po włączeniu zasilania jest zastosowanie enkodera absolutnego, w którym każdemu położeniu przyporządkowana jest konkretna wartość liczbowa, określająca aktualną pozycję (rys. 10). Dzięki temu enkoder absolutny podaje zaraz po włączeniu zasilania, na swoim wyjściu aktualną pozycję. Zanik napięcia nie powoduje utraty informacji o aktualnym położeniu i nie ma konieczności ponownego ustalania punktu zerowego, względem którego będą określone przemieszczenia. Enkodery wykonywane są w wersji obrotowej i liniowej. Te pierwsze mogą być jednoobrotowe albo wieloobrotowe.

Enkodery absolutne są droższe od inkrementalnych o takich samych parametrach. Na rynku można także spotkać Enkodery optyczne i magnetyczne. Enkodery magnetyczne obrotowe mają pierścień i czujnik, który mierzy przemagnesowania, natomiast magnetyczne enkodery liniowe z elementu pomiarowego poruszającego się nad taśmą magnetyczną. Są one nazywane liniałami magnetycznymi. Enkodery liniowe są stosowane w maszynach CNC, centrach obróbkowych i obrabiarkach elektroiskrowych. W najdokładniejszych stosuje się interferencyjne układy pomiarowe, w których prążki (wzorce długości) uzyskuje się w wyniku interferencji przenikających się wiązek lasera. Dzięki interpolacji można uzyskać rozdzielczości nawet do $0,01\ \mu\text{m}$.



W tabeli 1 zestawiono różne, najważniejsze czujniki.

Tab. 1. Klasyfikacja czujników

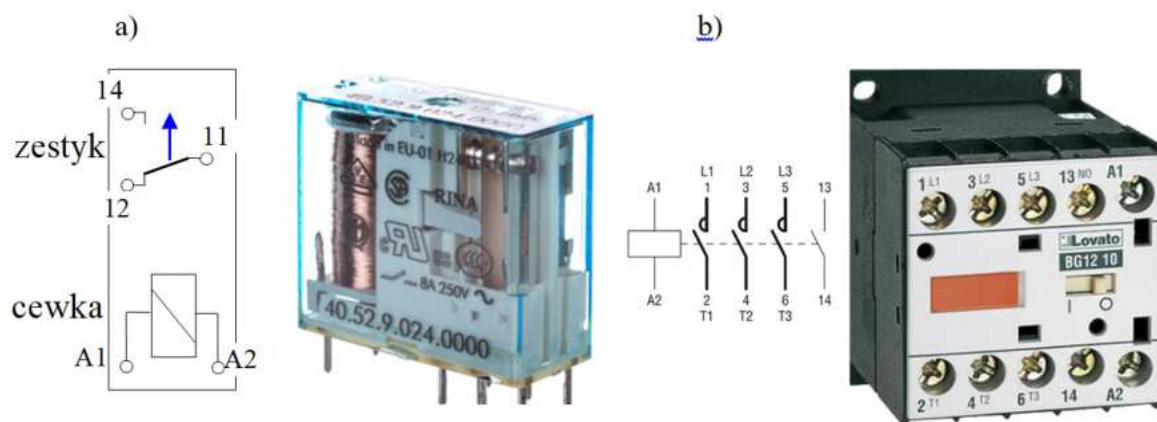
Wielkość mierzona	Typ czujnika
Położenie kątowe	Potencjometryczne, indukcyjne: transformatory położenia kątowego, selsyny
Przemieszczenie liniowe	Potencjometryczne, indukcyjne różnicowe LVDT, magnetostrykcyjne
	Enkodery absolutne optyczne
	Enkodery inkrementalne optyczne i magnetyczne
Prędkość i przyspieszenie	Prądnice: tachometryczne, synchroniczne, indukcyjne
	Akcelerometry
Siła	Tensometry rezystancyjne i półprzewodnikowe, piezoelektryczne
Dotyk	Stykowe: mikroprzełączniki, krańcówki
	Pojemnościowe, rezystancyjne
Zbliżenie, lokalizacja i orientacja	Ultradźwiękowe
	Indukcyjne
	Optyczne diodowe i laserowe: bariera, odbiciowe, refleksyjne
Identyfikacja obiektu	Systemy wizyjne
	Paskowe, RFID
Temperatura	Termopary, termorezystory
	Termistory półprzewodnikowe, pirometry, diody PIR

4. Napędy i elementy wykonawcze

Przełączniki i styczniki są elementami elektromechanicznymi, wykonawczymi stosowanymi powszechnie w systemach automatyki do włączania i wyłączania urządzeń. Maksymalna częstotliwość ich przełączania mieści się w zakresie 0,5÷5 Hz. Bardzo małe przełączniki mogą być przełączane sygnałem o częstotliwości nawet do 70 Hz, jednak charakteryzują się niewielką trwałością, ograniczeniami w stosowaniu oraz ceną, co powoduje, że są one rzadko stosowane. Sygnałem wejściowym przełącznika jest napięcie i natężenie prądu elektrycznego podawanego na cewkę, a sygnałem wyjściowym – stan styków (zamknięte lub otwarte). Wyróżniamy przełączniki, których cewki są zasilane napięciem stałym albo przemiennym. W katalogach podaje się następujące podstawowe parametry przełączników: napięcie znamionowe cewki (np. 6; 12; 24VDC albo 24VAC, 230VAC), moc znamionową cewki, liczbę i rodzaj styków, moc łączeniową przy obciążeniu rezystancyjnym, maksymalną częstotliwość przełączania oraz trwałość łączeniową i mechaniczną. Bardzo istotnym parametrem jest właśnie trwałość przełącznika elektromechanicznego. Minimalna trwałość mechanizmu wynosi od 1 do 100 mln cykli, a trwałość styków zależy od rodzaju obwodu, który jest sterowany przełącznikiem. Można ją zwiększyć tylko przez prawidłowe dobranie typu przełącznika do rodzaju obciążenia. Większość rzeczywistych obciążeń ma charakter



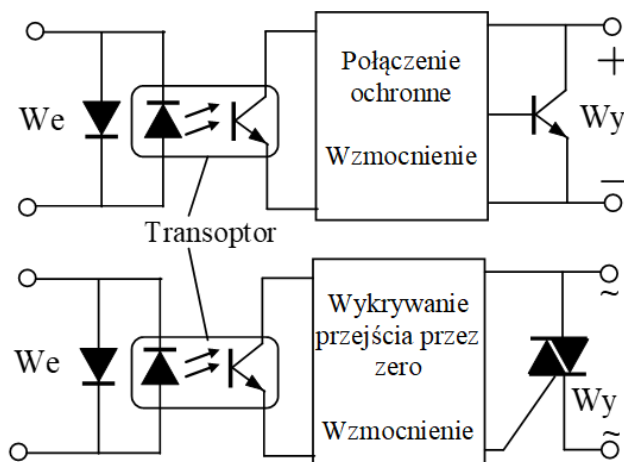
pojemnościowy, co oznacza, że prąd pobierany w pierwszej chwili po zamknięciu obwodu jest znacznie większy niż kilka chwil później.



Rys. 12. Symbol i widok: a) przekaźnika, b) stycznika 3-fazowego

Jeśli przez styki płynie natężenie prądu bliskie maksymalnemu, to dla zapewnienia wysokiej żywotności, napięcie przełączane nie powinno przekraczać $10\div 20\%$ wartości maksymalnej. Na trwałości styków oraz ich prądowo-napięciową charakterystykę duży wpływ mają materiały wykorzystywane na pokrycie pól kontaktowych. Renomowani producenci produkują przekaźniki ze stykami wykonanymi z materiałów, które są przystosowane do określonych warunków fizycznych oraz rodzaju zasilanego odbiornika. Rodzaj kategorii użytkowania stycznika lub przekaźnika jest ważnym parametrem, a prawidłowy jego dobór sprzyja długiej i bezawaryjnej pracy.

Styczniki to łączniki mechaniczne, które działają na tej samej zasadzie co przekaźniki elektromagnetyczne. Są też podobne pod względem budowy, ale znacznie większe. Różnica dotycząca typu łączonych za ich pomocą obwodów określa ich zastosowanie. Styczniki służą do włączania i wyłączania obwodów głównych elektrycznych dużej mocy, typu grzałki, silniki i elektromagnesy, a przekaźniki do przełączania obwodów pomocniczych. Styczniki mają zwykle trzy zestyki główne dużej mocy do zasilania obwodów trójfazowych np. $3 \times 600\text{VAC}/50\text{A}$ oraz zestyki pomocnicze (np. do podtrzymania). Ich cewki są najczęściej przystosowane do zasilania napięciem 230 VAC oraz 24VDC . Częstotliwość przełączania styczników jest niewielka - mogą załączać obwody tylko $0,5\div 2$ razy na sekundę.



Rys. 13. Przekazniki półprzewodnikowe: tranzystorowe i triakowe

Na rynku dostępne są także przekaźniki półprzewodnikowe (ang. solid state relays – SSR). Stanowią one elektroniczny odpowiednik przekaźników elektromechanicznych. Zamiast cewki stosowane w nich są transoptory, składające się z diody świecącej i fototranzystora (rys. 13). Zapewniają one izolację między układami o różnicy potencjałów nawet do 4 kV, chroniąc dodatkowo wrażliwe elementy elektroniczne. Na wyjściu jest używany tranzystor dużej mocy, tyrystor albo triak. Tym samym przełączanie poszczególnych obwodów realizowane jest bez udziału części ruchomych, czyli bez ruchomych styków. Przekazniki wykonują sterowanie bezstykowe lub elektroniczne. Wszystkie przekaźniki i styczniki zapewniają oddzielenie elektryczne obwodów automatyki od obwodów dużej mocy, zapewniając tzw. separację galwaniczną. Przekazniki SSR charakteryzują się następującymi cechami:

- dużą częstotliwością przełączania, dochodzącą do kilku kHz,
- bezstykową, bezgłośnie oraz iskrobezpieczną pracę,
- bardzo dużą trwałością i niezawodnością,
- zdolnością do załączania dużej mocy (prąd od 90 do 150 A i napięcie do 1000 V),
- sterowaniem chwili włączenia i wyłączenia np. przy przechodzeniu sinusoidy przez zero.

Do napędzania urządzeń montowanych na automatyzowanych liniach produkcyjnych najczęściej stosuje się silniki elektryczne. Do wykonywania prostych ruchów liniowych od jednej pozycji krańcowej do drugiej najczęściej stosowane są siłowniki pneumatyczne, sterowane zaworami elektropneumatycznymi. Silniki wirujące dzieli się na dwie zasadnicze grupy: prądu stałego i prądu przemiennego. Silnik prądu stałego jest silnikiem elektrycznym



zasilanym prądem stałym i służy do zamiany energii elektrycznej na energię mechaniczną. Jako maszyna elektryczna prądu stałego może pracować jako silnik lub prądnica. W tym drugim przypadku wirnik napędzany jest energią mechaniczną dostarczoną z zewnątrz, a na zaciskach uzwojenia twornika odbierana jest wytworzona energia elektryczna. Większość silników prądu stałego to silniki komutatorowe, to znaczy takie, w których uzwojenie twornika zasilane jest prądem poprzez komutator. Jednak istnieje wiele silników prądu stałego, które nie posiadają komutatora lub też komutacja przebiega na drodze elektronicznej.

Maszyny prądu przemiennego buduje się jako trójfazowe i jednofazowe. Ze względu na stosunkowo prostą ich budowę maszyny prądu przemiennego znalazły powszechne zastosowanie jako silniki napędzające różnego rodzaju urządzenia mechaniczne oraz jako generatory wielkiej mocy. Maszyny prądu stałego mają wprawdzie budowę bardziej skomplikowaną od maszyn prądu przemiennego, lecz równocześnie odznaczają się lepszymi właściwościami regulacyjnymi. Dlatego też są stosowane w bardziej skomplikowanych, wymagających np. regulacji prędkości obrotowej, układach napędowych.

5. Silniki trójfazowe

Obecnie najczęściej stosowanymi w maszynach i urządzeniach technologicznych są silniki trójfazowe indukcyjne. Na ich stojanie nawinięte są trzy albo sześć uzwojeń tworzących odpowiednio, jedną albo dwie pary biegunów. Są one zasilane napięciem trójfazowym, które jest źródłem wirującego pola magnetycznego o prędkości synchronicznej n_s :

$$n_s = \frac{60f}{p} [\text{obr/min}], \quad (3)$$

gdzie: f – częstotliwość napięcia zasilającego, p – liczba par biegunów.

Dla częstotliwości sieci 50 Hz i dla $p=1$ prędkość wirowania pola n_s jest równa 3000 obr/min, a dla $p=2$, n_s wynosi 1500 obr/min. Wirnik silnika asynchronicznego może być wykonany w postaci:

- szeregu prętów zwartych po obu końcach (tzw. silniki klatkowe),
- szeregu uzwojeń, których końce są wyprowadzone do pierścieni umieszczonych na wale i także zwartych (tzw. silniki pierścieniowe).

Wirujące, czyli zmienne pole magnetyczne powoduje indukowanie siły elektromotorycznej (SEM) w prętach albo uzwojeniach wirnika. Wartość tej siły jest proporcjonalna do różnicy prędkości wirującego pola i wirnika. W rezultacie przez zwarte uzwojenia wirnika płynie prąd wytwarzający pole magnetyczne. Wirujące pole magnetyczne



stojana przyciąga i wprawia w ruch obrotowy wirnik, który generuje moment obrotowy na wale wyjściowym. Wirnik obraca się z prędkością n_w , z prędkością mniejszą od prędkości wirowania pola stojana. W rezultacie występuje tzw. poślizg, czyli opóźnianie się wirnika w stosunku do pola stojana, określony wzorem:

$$s = \frac{n_s - n_w}{n_s} 100\% . \quad (4)$$

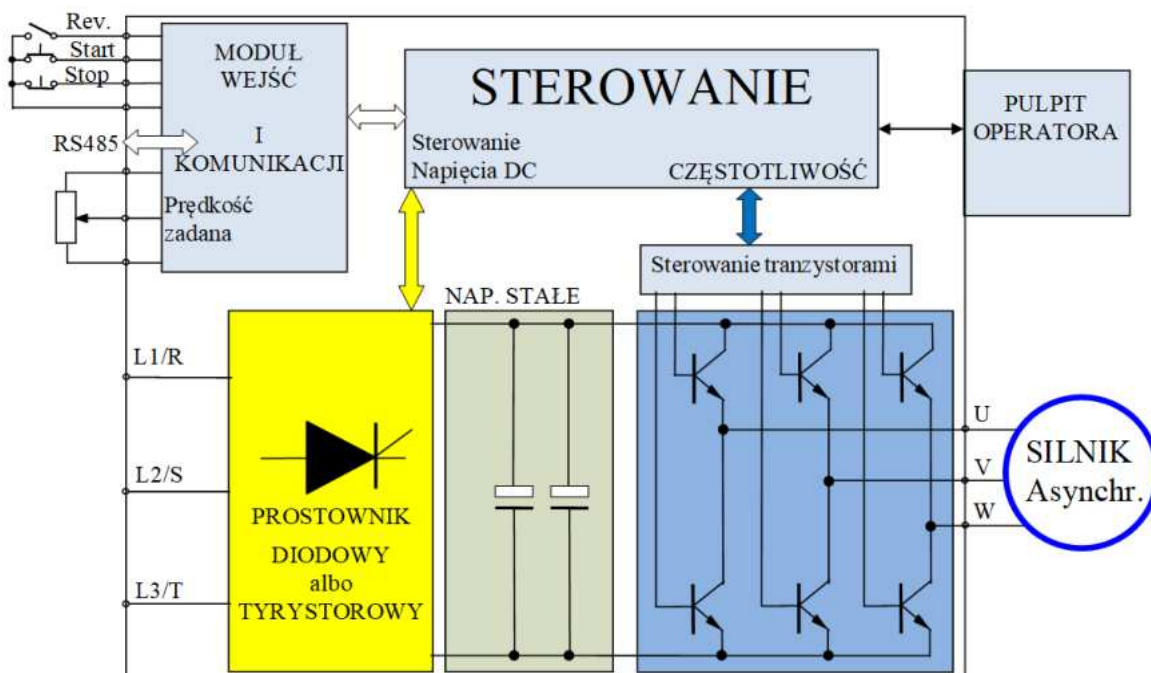
W trakcie rozruchu, tj. przy zatrzymanym wirniku silnik rozwija maksymalny moment zwany rozruchowym M_r . W trakcie pracy poślizg zmienia się od bliskiego 1% do ok. 15%. Zwiększenie momentu obciążenia powyżej wartości M_k prowadzi do zatrzymania wirnika.

Prędkość obrotową wirnika silnika prądu przemiennego wyraża wzór (3). Wynika z niego, że prędkość obrotową można regulować przez zmianę prędkości wirującego pola, tzn. przez zmianę częstotliwości napięcia zasilającego f . Dlatego aby wytworzyć w stojanie pole magnetyczne wirujące o różnej częstotliwości, stosuje się przemienniki częstotliwości, nazywane też falownikami. Są to zasilacze zmiennoprądowe z regulowaną bezstopniowo częstotliwością napięcia zasilającego silnik. Zmiana częstotliwości f napięcia zasilającego powoduje nie tylko zmianę prędkości silnika, lecz również zmianę momentu elektromagnetycznego rozwijanego przez silnik. Zmniejszanie częstotliwości (poniżej 50 Hz) powoduje wzrost momentu elektromagnetycznego, a zwiększenie częstotliwości prowadzi do zmniejszenia momentu elektromagnetycznego. Dlatego, dla poprawnej pracy silnika, niezależnie od aktualnej częstotliwości zasilania, konieczne jest spełnienie warunku: $U/f = \text{const}$. Przemiennik częstotliwości, nazywany też falownikiem, zawiera 4 bloki funkcjonalne:

- układ sterowania i komunikacji,
- prostownik trójfazowy, najczęściej regulowany,
- pośredni obwód filtracji napięcia stałego z baterią kondensatorów,
- blok tranzystorowy przekształcający prąd stały w trójfazowy prąd przemienny o częstotliwości wynikającej z zadanej prędkości obrotowej.

Schemat blokowy falownika został pokazany na rys. 14. Napięcie przemienne 3 x 400VAC o częstotliwości 50Hz z sieci zasilającej jest prostowane do napięcia stałego oraz filtrowane (wygładzane) za pomocą kondensatorów. Może być sterowany albo nie. Zaletą prostownika sterowanego jest możliwość regulowania napięcia DC przez tyrystory. Stopień pośredni w przetwornicy, to bateria kondensatorów. Można ją traktować jako magazyn energii, z którego układ sterowanych tranzystorów kształtuje trójfazowe napięcia o regulowanej częstotliwości i wartości. Tranzystory pracują jak klucze, tzn. otwierają lub zamykają przepływ prądu ze

źródła napięcia stałego do odpowiedniej fazy silnika. Układ sterowania otrzymuje i obsługuje sygnały komunikacyjne z otoczenia przetwornicy oraz steruje pracą inwertera mocy. Sygnały te mogą pochodzić z zewnętrznych urządzeń sterujących bądź z panelu operatora.



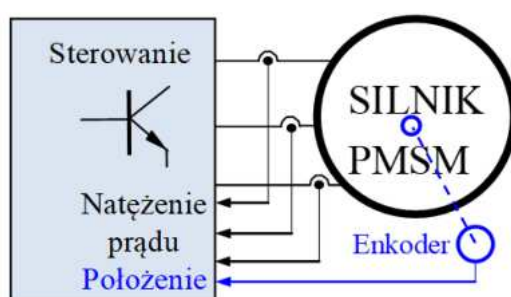
Rys. 14. Schemat blokowy przemiennika częstotliwości – falownika

Spośród silników synchronicznych w urządzeniach technologicznych praktyczne zastosowanie w układach napędowych znalazły silniki synchroniczne z magnesami trwałymi umieszczonymi na wirniku. Na ich stojanie nawinięte jest trójfazowe uzwojenie. Takie silniki nazywane są też Permanent Magnet Synchronous Motor – PMSM. Są stosowane głównie w napędach posuwowych obrabiarek i w robotach. Silniki elektryczne z magnesami trwałymi umieszczonymi na wirniku mają najwyższą sprawność energetyczną ze wszystkich znanych i stosowanych rodzajów maszyn elektrycznych. Stojan posiada trójfazowe uzwojenie, wytwarzające kołowe pole wirujące. Rozruch i sterowanie w warunkach zmian obciążenia musi być kontrolowane elektronicznie. Silniki te mogą pracować, w zależności od sposobu zasilania i sterowania, jako silniki:

- synchroniczne PMSM,
- bezszczotkowe prądu stałego z komutatorem elektronicznym sterowane skokowo.



W obu przypadkach właściwości napędu są inne, decyduje o tym rozkład pola magnetycznego w szczeliny silnika oraz sposób zasilania i sterowania. Silniki PMSM powinny być sterowane napięciami sinusoidalnymi. Do tego konieczne jest stosowanie dokładnego pomiaru położenia wirnika za pomocą np. enkodera. Informacja o położeniu wirnika dostarczana jest do sterownika, który odpowiednio steruje zespołem mocy falownika. Dodatkowo konieczny jest pomiar natężenia prądu w poszczególnych uzwojeniach. Silnik synchroniczny z magnesami trwałymi sterowany jest z falownika metodą PWM, kluczami tranzystorowymi, na podstawie pomiaru prądów poszczególnych faz silnika oraz sygnału z enkodera (rys. 15).



Rys. 15. Podłączenie silnika PMSM

Zależnie od sposobu i rodzaju dostarczanej energii wyróżniamy urządzenia napędowe:

- pneumatyczne (brak zagrożenia pożarowego, odporność na wpływ płynów i związków agresywnych, łatwość w eksploatacji, niezawodność w działaniu, nie wymagają stosowania zabezpieczeń przed przeciążeniem),
- elektryczne (nośnikiem jest sygnał elektryczny – najczęściej napięcie lub prąd stały, bardzo bogata i rozbudowana dostępność urządzeń pomiarowych i sterujących, łatwość wykonania operacji matematycznych, przesyłanie sygnałów na dowolne odległości),
- hydrauliczne (możliwość uzyskania bardzo dużych sił przy dobrej dynamice, duża trwałość, bardzo dobry stosunek energii do masy).

Przeważająca liczba zadań sterowania urządzeniami w praktyce przemysłowej sprowadza się do załączania i wyłączania wybranych elementów napędowych, grzałek, elektromagnesów i lamp. Nazywa się to sterowaniem dwustanowym, przełączającym lub binarnym. Projektowanie takiego sterowania sprowadza się do zbudowania schematu urządzenia sterującego z wykorzystaniem przełączników i przekaźników albo bramek logicznych.



Schemat połączeń tych elementów stanowi algorytm sterowania, którego techniczna realizacja polega na wprowadzeniu tego schematu do pamięci uniwersalnego sterownika. Jest on zdolny do wykonywania algorytmu symulując niejako wprowadzony schemat. Mówimy wtedy o sterowniku programowanym pamięciowo. Program to ciąg kolejno następujących i logicznie powiązanych ze sobą instrukcji do przetwarzania sygnałów i danych przez urządzenie sterujące - sterownik.

6. Podstawy sterowników przemysłowych typu PLC

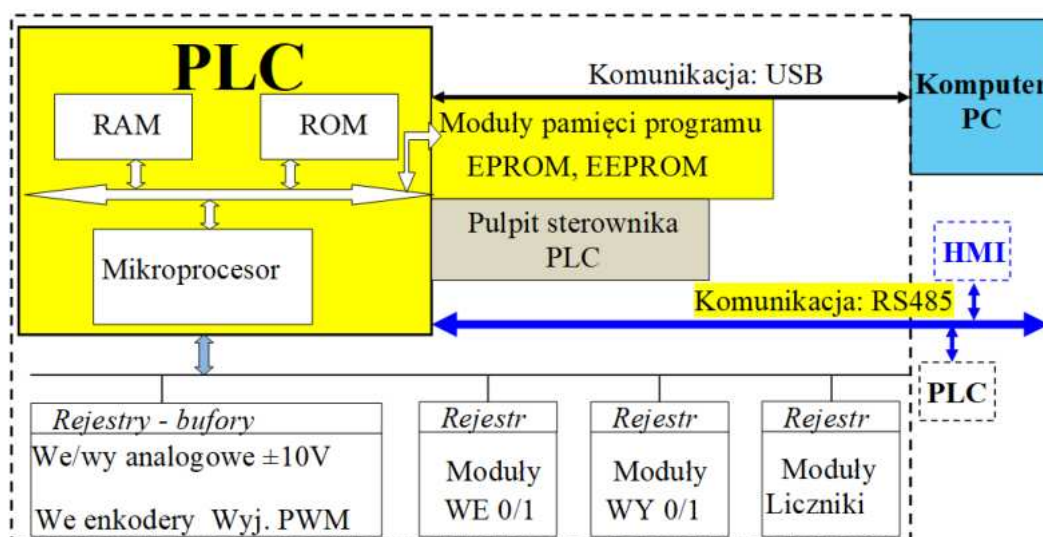
Do dyspozycji automatyków jest dzisiaj bardzo różnorodna aparatura pomiarowa tj. czujniki i systemy pomiarowe, silniki i kompletne napędy, regulatory, sterowniki, panele do wprowadzania i wyświetlania informacji oraz systemy komunikacji.

Sterowniki programowalne (ang. Programmable Logic Controllers – PLC), nazywane także binarnymi, są obecnie powszechnie stosowane w przemyśle. Na podstawie sygnałów analogowych i cyfrowych o stanie urządzenia, otrzymywanych z czujników oraz urządzeń pomiarowych, sterowniki PLC przetwarzają dane zgodnie z zapisanym w pamięci programem i generują sygnały sterujące. Pierwsze sterowniki PLC zbudowano w końcu lat 60. XX w. Ich główną zaletą jest możliwość szybkiej zmiany wykonywanego algorytmu przez programowanie pamięci sterownika. Charakteryzują się one dużą niezawodnością w warunkach przemysłowych, przy małych wymiarach oraz łatwością utrzymania w ruchu produkcyjnym, dzięki możliwości napraw przez wymianę całych modułów. Na początku lat 80. XX w. pojawiły się inteligentne moduły we/wy, mające własne procesory, które mogły realizować złożone funkcje obliczeniowe. Od lat 90. XX w. wykorzystuje się komputery PC do programowania sterowników PLC, co znacznie rozszerzyło możliwości programowe i komunikacyjne sterowników, pracujących obecnie w sieciach.

Na rysunku 16 pokazano schemat blokowy typowego sterownika PLC. Jego głównym elementem jest mikroprocesorowa jednostka centralna (CPU), która wykonuje obliczenia logiczne i arytmetyczne oraz zarządza sterownikiem. Jest wyposażona w pamięć typu ROM (ang. Read Only Memory) przeznaczoną do odczytu. Są w niej zapisane informacje o konfiguracji sterownika. W pamięci programowanej (EEPROM) zapisywany jest program sterujący, napisany przez użytkownika. W pamięci RAM (ang. Random Access Memory) są przechowywane dane odczytywane z wejść oraz wyniki wykonywanych operacji, w tym te określające stany wyjść. Pamięć RAM jest zwykle wykorzystywana do budowy elementów



czasowych i licznikowych. Do podłączenia czujników binarnych, sterowniki PLC są wyposażone w moduły wejść 0/1. Podobnie do podłączania elementów wykonawczych, które można tylko włączyć albo wyłączyć, stosowane są przekaźnikowe moduły wyjściowe. Poza tym, sterowniki wyposażone są także w liczniki, które mogą służyć do zliczania impulsów podawanych na wejścia albo do odmierzania czasu (Timery).



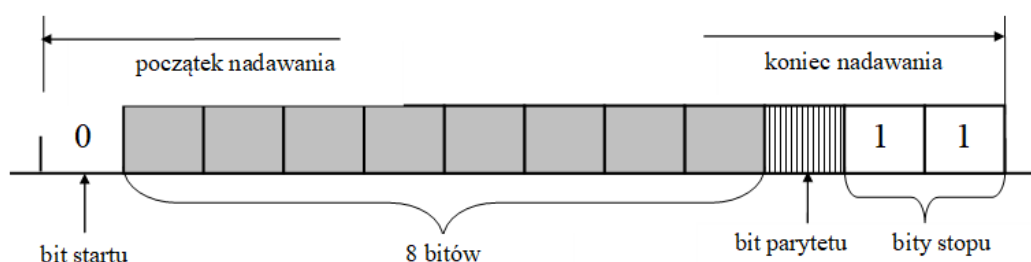
Rys. 16. Schemat blokowy sterownika przemysłowego PLC

Sterowniki PLC były początkowo przeznaczone tylko do wykonywania jednobitowych operacji przełączających i dlatego posiadały tylko wejścia binarne (0/1) na napięcie 0V/24VDC oraz wyjścia binarne – przekaźnikowe, mogące włączać elementy o napięciu max. 250V i natężeniu prądu do 1A. Obecnie są dostępne także wyjścia tranzystorowe (24VDC) bądź triakowe (230VAC). Współczesne sterowniki PLC są także wyposażone w moduły wejść i wyjść analogowych, napięciowych (–10V... +10V). Do sterowników można także dołączyć enkodery inkrementalne – do wejść licznikowych. Moduły wyjściowe umożliwiają sterowanie silnikami krokowymi oraz układami mocy z modulacją szerokości impulsów (PWM) itd. Realizacja zadań sterowania rozbudowanymi zestawami maszyn wymaga połączenia sterowników PLC w sieć np. RS485. Opcjonalnie do sterowników można też podłączyć monitory i panele operatorskie (ang. Human-Machine Interface – HMI), wykorzystywane do wizualizacji procesu i sterowania nim w trybie bezpośrednim (ang. *on-line*). Pozwalają one na wybór odpowiednich opcji sterowania, wprowadzanie nastaw regulatorów, kontrolę i korektę ewentualnych błędów sterowania. Mają one duże możliwości i dlatego pozwalają (w zależności od użytego oprogramowania) zarówno na wizualizację i kontrolę procesu, jak i na edycję,



wprowadzanie i testowanie samego programu sterującego. Komputer PC podłączany jest do sterownika na czas programowania.

Ponieważ te sterowniki są obecnie budowane na bazie komputerów, które przetwarzają tylko sygnały cyfrowe – liczby binarne, to napięcie elektryczne zostaje przetworzone na liczby w przetwornikach analogowo-cyfrowych (AC) na liczby. Przetwarzanie przebiega w dwóch fazach: próbkowanie i kwantowanie. W pierwszej napięcie wejściowe zostaje zapamiętane, a w drugiej przetworzone na liczbę zapisaną w kodzie binarnym. Przetwornika AC zamienia sygnał analogowy (napięcie) na ciąg liczb L1, L2, L3 ... podawanych na wyjście co określony czas. Jego podstawowe parametry to: zakres napięć wejściowych np. 0 do 1V, 0 do 2,5V; liczba bitów na wyjściu np. 8, 12, 16 oraz czas przetwarzania (wyzwalania) T , zwykle od 1 μ s do 1 ms. Przetwornik 8-mio bitowy generuje na wyjściu liczby z zakresu od 0 do 255 a 12-to bitowy od 0 do 4095. Przetwornik cyfrowo-analogowy (CA), przetwarza w kierunku odwrotnym tzn. zamienia liczby binarne na odpowiadające im napięcia wyjściowe. Jednak sygnał analogowy w postaci napięcia może być przesyłany tylko na niewielkie odległości, bowiem jest podatny za zakłócenia. Na większe odległości można przysyłać sygnał prądowy o zakresie prądów 4 mA 20 mA, który jest znacznie bardziej odporny na zakłócenia. Jednak najlepszy do transmisji na duże odległości jest sygnał cyfrowy, pozwalający na przesyłanie dowolnych informacji i danych.



Rys. 17. Przebieg transmisji szeregowej słowa

W układach do komunikacji sieciowej kolejne liczby są przesyłane za pomocą transmisji szeregowej. Jest to rodzaj cyfrowej transmisji danych, w którym dane są przesyłane jednym przewodem (albo jedną parą). Poszczególne bity słowa (liczb) są przesyłane kolejno, bit po bicie. Stosowane są następujące interfejsy: RS-232, RS-422A, RS-485, I2C, SPI, USB, Ethernet, USB. W praktyce stosowana jest też łączność bezprzewodowa. Używane są 3 główne



standardy: WiFi (60%), ZigBee (20%) oraz Bluetooth (20%). Dane są poprzedzone bitem startu (stan logiczny: 0), który jest sygnałem do rozpoczęcia transmisji (rys. 17). Po nim są przesyłane bity danych, od najmłodszego do najstarszego, potem bit parytetu (parzystości), a po nim następuje odstęp, nazywany bitem albo bitami stopu przed następnym znakiem (stan logiczny: 1). Stop może on być dłuższy bo nie ma ograniczenia czasowego.

Na rysunku 18 pokazano podłączenie różnych elementów do sterownika PLC. Ze względu na rodzaj wyróżnić można sygnały:

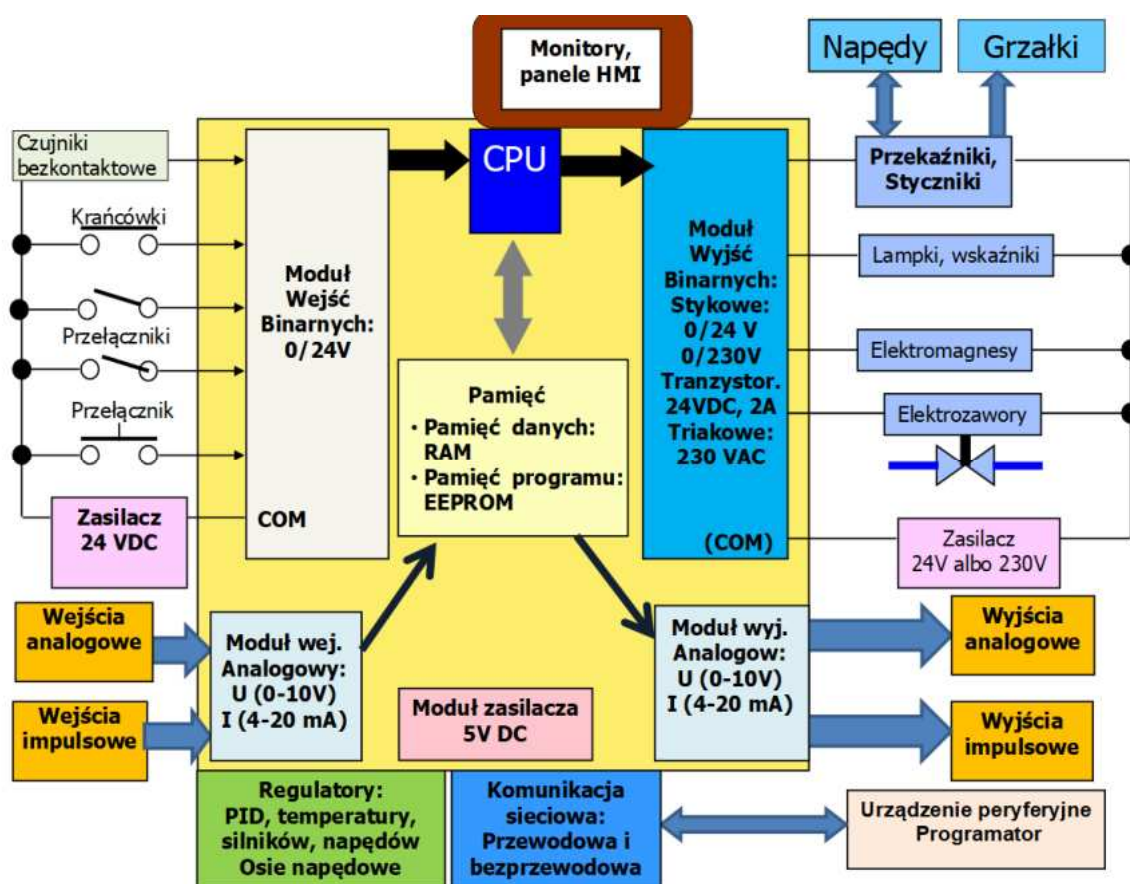
- binarne, przełączające, w których sygnał może przyjmować tylko jedną z dwóch wartości 0 lub 1, kodowaną napięciowo 0V/24VDC albo 0/230VAC i 0V/3x400VAC,
- analogowe albo ciągłe, w których sygnał może przyjmować dowolną wartość z określonego przydziału napięcia: -10V do +10V, 0 do 10V, 0 do 5V, albo natężenia prądu 0 do 20 mA albo 4 do 20 mA.

Na wejścia sterownika PLC można podawać sygnały:

- binarne, które dostarczają informacje: on/off (24VDC/0VDC) z takich elementów jak:
 - wyłączniki krańcowe stykowe,
 - czujniki indukcyjne, pojemnościowe, optyczne itp.
 - przełączniki graniczne poziomu, ciśnienia itp.
- pomiarowe (analogowe), które dostarczają informacje ciągłe o aktualnych:
 - położeniu, prędkości, sile,
 - ciśnieniu, naprężeniu,
 - temperaturze, napięciu, natężeniu
- impulsowe w postaci upływu czasu T oraz o częstotliwości f określających:
 - zmiany położenia z enkoderów inkrementalnych
 - czas jaki upływa od wysłania impulsu do jego powrotu.

Wyjścia sterownika mogą być następujące:

- binarne typu on/off (0VDC/24VDC):
 - do włączania albo wyłączenia różnych elementów i urządzeń,
 - zwolnienie (enable), hamuj, kierunek L/P
- Analogowe do określenia zadanych (np. -10 do +10V):
 - generowanej siły, prędkości, położenia
 - ciśnienia, natężenia przepływu, temperatury, napięcia itp.
- Impulsowe, czasowe do wykonywania kroków napędu albo jego prędkości (PWM).



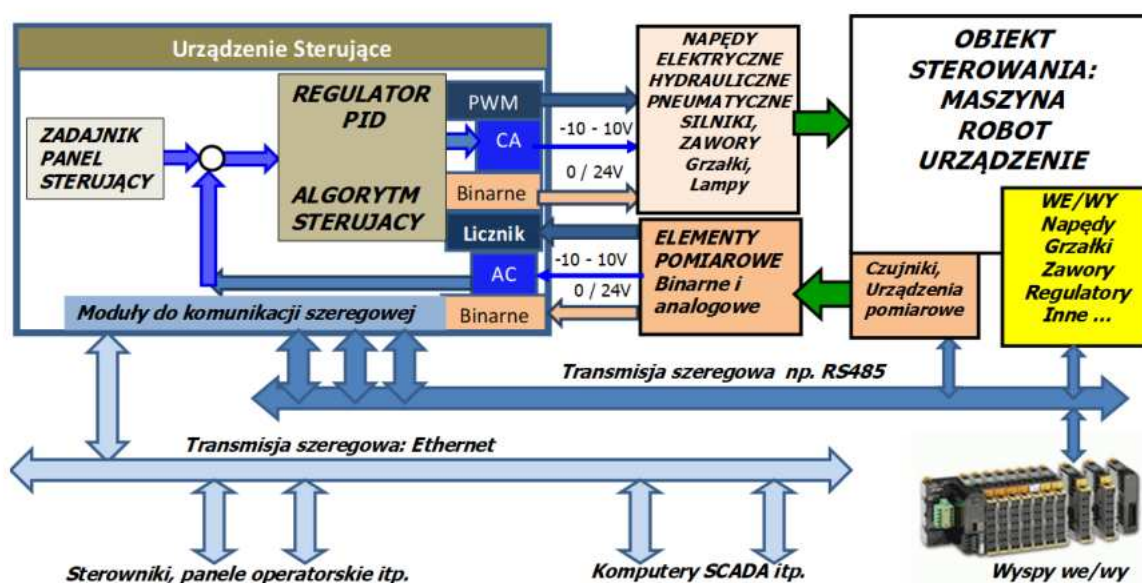
Rys. 18. Podłączenie elementów do sterownika przemysłowego PLC

Na rysunku 18 pokazano podłączenie różnych elementów do wejść binarnych (przełączniki stykowe, krańcówki, czujniki bezkontaktowe itp.) oraz do wyjść binarnych (lampki, przekazniki i styczniki, elektromagnesy oraz elektrozawory). Przekazniki i styczniki przeznaczone są do włączania i wyłączania elementów dużej mocy typu: grzałki, silniki i napędy). W sterowniku występują także moduły wejść i wyjść analogowych, do których można dołączać elementy generujące bądź sterowane, za pomocą sygnałów elektrycznych. Nośnikiem informacji jest napięcie elektryczne o zakresie np. 0V – 10 V albo natężenie prądu 4 mA – 20 mA. Do wejść mogą być podłączone urządzenia pomiarowe, np. potencjometry, linały transformatorowe, wagi tensometryczne natomiast do wyjść można podłączyć np. sterowniki napędów elektrohydraulicznych albo elektrycznych. Enkodery inkrementalne można podłączyć do wejść impulsowych. Do wyjść impulsowych można podłączyć np. silnik krokowy.

Współcześnie produkowane sterowniki PLC wyposażone są także odrębne moduły różnych regulatorów np. PID albo regulatory dwupołożeniowe, przeznaczone do regulacji temperatury



oraz regulatory silników i napędów. Dostępne są również moduły do regulacji i pozycjonowania od 1 do 4 napędów, tzw. regulatory kilku osi np. obrabiarek albo robotów. Sterowniki PLC mogą również komunikować się z różnymi urządzeniami pomiarowymi oraz napędowymi a także z innymi sterownikami PLC oraz komputerami PC, a tym samym z całym Internetem, za pośrednictwem łączności przewodowej bądź bezprzewodowej. Komunikacja ta jest prowadzona w sposób szeregowy, tzn. poszczególne słowa są przesyłane bit po bicie.



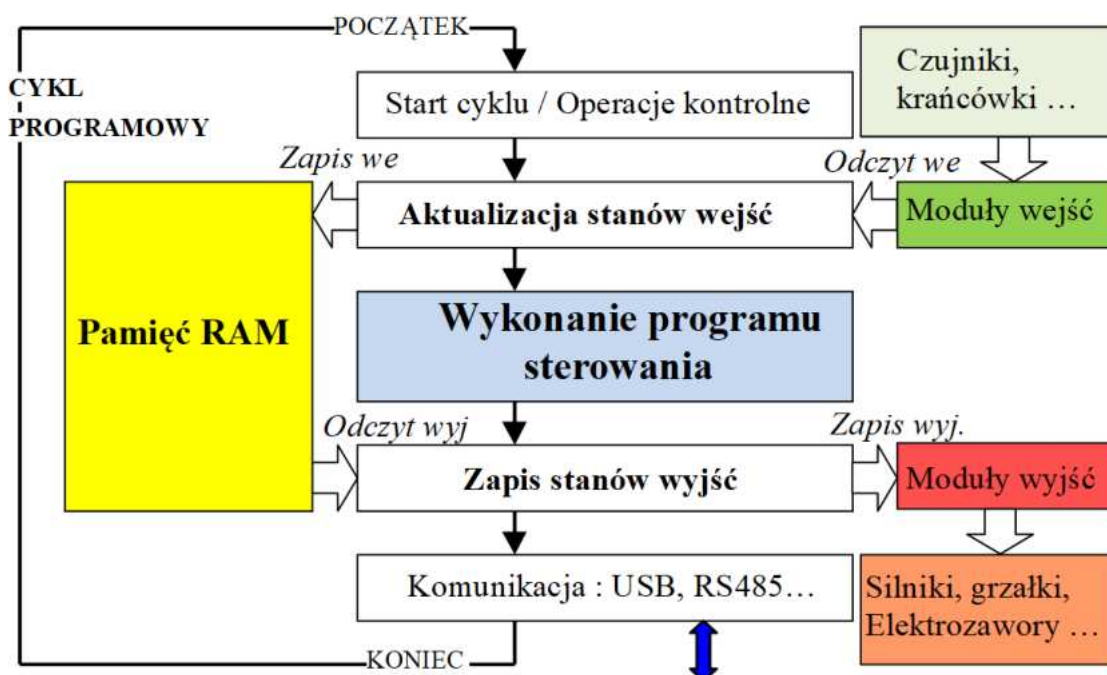
Rys. 19. Komunikacja elementów i układów automatyki

Do komunikacji systemu sterowania wykorzystywane są już od kilkunastu lat również magistrale szeregowe. Wśród nich, bardzo popularny jest standard RS485. Składa się on z różnicowego nadajnika, dwuprzewodowego toru transmisyjnego i różnicowego odbiornika. Przesyłanie różnicowe chroni transmisję danych przed zakłóceniami zewnętrznymi od np. silników i przełączników. RS485 jest często stosowanym interfejsem w sieciach przemysłowych, bo charakteryzuje się dużą szybkością transmisji, nawet na znaczne odległości. W ostatnich latach wzrasta znaczenie połączenia typu przemysłowy Ethernet. W odróżnieniu od innych technologii komunikacyjnych umożliwia funkcjonowanie w jednej sieci wielu różnych protokołów. Wśród najważniejszych można wymienić: EtherCAT, EtherNet/IP, Modbus TCP, Powerlink, Profinet i SERCOS. Przemysłowy Ethernet pozwala na połączenie pomiędzy standardową siecią biurową i siecią przemysłową. Jako medium transmisyjne stosuje się często miedzianą skrętkę ekranowaną. Umożliwia przesyłanie danych



klasy Fast Ethernet, czyli z prędkością 100Mb/sekundę z pełnym duplexem. Maksymalna długość przewodu połączeniowego może wynosić 100m. Coraz ważniejszym medium transmisyjny stają się światłowody, które pozwalają na maksymalny przesyłanie na odległość 14 km. Podłączenie różnych elementów automatyki do sterowania pokazano na rys. 19.

Kolejne rozkazy sterownika PLC są przetwarzane w sposób cykliczny – po zakończeniu poprzedniego rozkazu jest pobierany następny. W zasadzie nie wykorzystuje się skoków. Gdy zostanie przetworzony ostatni rozkaz z pamięci programu, to następuje zakończenie tzw. cyklu programu, po czym program jest przetwarzany ponownie od początku. Takie działanie jest określane pracą szeregowo-cykliczną, w zamkniętej, niekończącej się pętli (rys. 20). Sterownik wykonuje zapisany w pamięci program w tzw. cyklu programowym. Na jego początku są wykonywane operacje kontrolne. Należą do nich wszystkie zadania przygotowujące sterownik do kolejnego cyklu niezbędne do jego prawidłowej pracy. Odbywa się wtedy aktualizacja tablicy błędów oraz opóźnienie programowe czasu trwania cyklu (tylko jeśli sterownik jest uruchomiony w trybie stałej długości cyklu). Następnie sterownik przechodzi do wykonywania właściwego programu. Zasadniczo rozkazy są wykonywane jeden po drugim, niezależnie od wyniku poprzedniej operacji. Tylko w wyjątkowych sytuacjach dopuszczalne jest wykonywanie skoku do innych fragmentów programu. Ponieważ czas wykonania pojedynczej instrukcji jest bardzo krótki (od 0,1 do 10 μ s), przeto z punktu widzenia procesu wydaje się, że rozkazy wykonywane w rzeczywistości kolejno, są wykonywane równocześnie. Program kończy instrukcja END. Po zakończeniu całej sekwencji następuje faza komunikacji. W jej ramach sterownik nawiązuje komunikację z dołączonymi urządzeniami i wymienia informacje, np. przekazuje dane o stanie we/wy, zmienia parametry itp. Po zakończeniu tej fazy cykl jest wykonywany od początku. Jednym z ważniejszych parametrów sterownika jest czas cyklu sterownika, składający się z czasu potrzebnego do pełnego obiegu programu złożonego z tysiąca statystycznych instrukcji. Wynosi on od 0,1 do 10 ms.



Rys. 20. Cykl programowy sterownika PLC

Odczyt wejść odbywa się jednorazowo na początku cyklu. Stan wejść zostaje zapisany w pamięci sterownika. W następnym kroku sterownik wykonuje program, w trakcie którego pobiera informacje o stanach wejść z pamięci i modyfikuje stan wyjść w pamięci. Po zakończeniu wykonywania programu następuje przepisanie stanów wyjść z pamięci do rzeczywistych modułów wyjściowych. Dzięki takiemu podejściu następuje skrócenie czasu cyklu i brak błędów wynikających ze zmiany sygnałów na wejściach i wyjściach podczas trwania jednego cyklu.

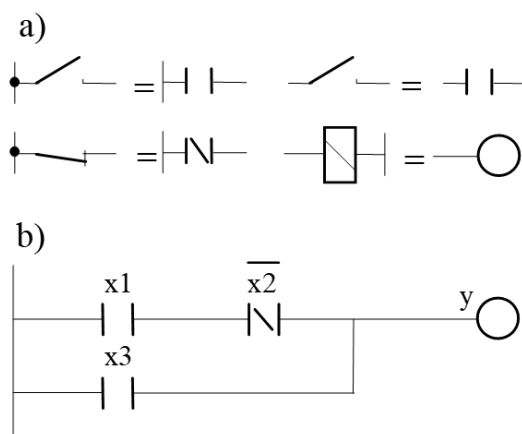
7. Programowanie sterowników PLC

Na wstępie należy zaznaczyć, że każda rodzina sterowników danego producenta ma własną listę rozkazów. Stosuje się następujące języki programowania sterowników binarnych:

- logiki albo schematów drabinkowych (LD),
- bloków funkcyjnych (FBD),
- schematów sekwencyjnych (Sequential Function Chart – SFC)
- listy instrukcji (IL),
- strukturalny, typu język C, Pascal itp.



Język logiki drabinkowej jest językiem graficznym, który stworzono pierwotnie, aby zastąpić układy przekaźnikowe. Z pewnym uproszczeniem można uznać, że logika drabinkowa polega na narysowaniu schematu sterowania przekaźnikowego na ekranie programatora. Wynika to z faktu, iż sterowniki zastąpiły używane wcześniej panele przekaźnikowe. Język ten jest łatwo przyswajany i zrozumiały dla personelu obsługującego. Opiera się na standaryzowanych symbolach graficznych, odpowiadających elementom stykowym, czyli różnym zestynom i cewkom przekaźników. Przy programowaniu, schemat stykowo-przekaźnikowy należy obrócić o 90°. Możliwe jest dodatkowo używanie funkcji: arytmetycznych i logicznych oraz porównań i relacji, a także bloków funkcyjnych, takich jak np.: przerzutniki, czasomierze, liczniki, regulatory PID i inne. Język schematów drabinkowych jest językiem powszechnie używanym i uwzględnionym w normach IEC1.



Rys. 21. Symbole graficzne stosowane w języku LD (a) oraz przykładowy program (b)

Na rysunku 21a pokazano zestyk i cewkę przekaźnika oraz odpowiadające im symbole graficzne stosowane w języku LD. Są one rysowane poziomo – tak jak stopnie drabiny, stąd też potoczna nazwa języka LD jako schemat drabinkowy. Podstawowym elementem obwodu jest styk, który może być zwierny (normalnie otwarty) lub rozwierny (normalnie zamknięty). Przykładowy program napisany w języku LD pokazano na rys. 21b. Poszczególne szczeble schematu (linie poziome) są nazywane obwodami i składają się z połączonych ze sobą pojedynczych elementów wejściowych (zestyków). Na końcu każdej linii musi występować element wykonawczy (wyjście binarne), np. cewka, przerzutnik, licznik, pamięć itp. Obwód jest ograniczony z lewej strony szyną zasilania tj. 24VDC. Zasadą działania języka LD jest przepływ wirtualnego prądu od linii zasilania przez poszczególne modele elementów do elementy wyjściowego. Element wyjściowy zostaje ustawiony w odpowiedni stan (1 albo 0)

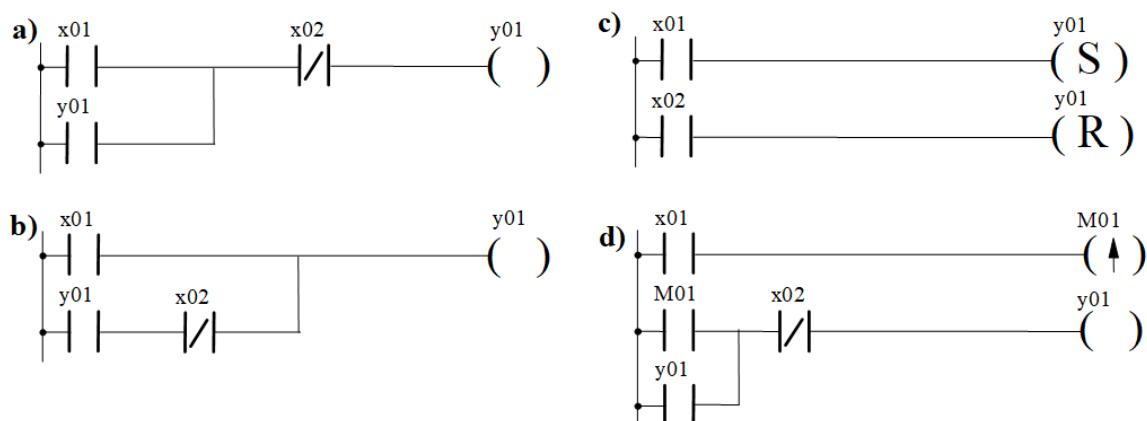


w zależności od tego, czy ten wirtualny prąd odpowiednio dopływa do elementu wykonawczego, czy też nie dopływa. Program jest wykonywany krok po kroku (obwód po obwodzie) aż do ostatniej linii albo do instrukcji END. Wyjątkiem od tej reguły są tzw. skoki – mogą być warunkowe postaci, jeśli warunek, to skok do etykieta, lub bezwarunkowe postaci – skok do etykieta. Jednak skoki w programach występują bardzo rzadko. Przy programowaniu algorytmów sterowania konieczne jest wykonywanie różnorodnych operacji uzależnionych od czasu, typu opóźnienia. Mogą one być programowane z wykorzystaniem elementów czasowych. Ważne są również moduły licznikowe, które pozwalają na zliczanie różnorodnych, występujących w procesie stanów.

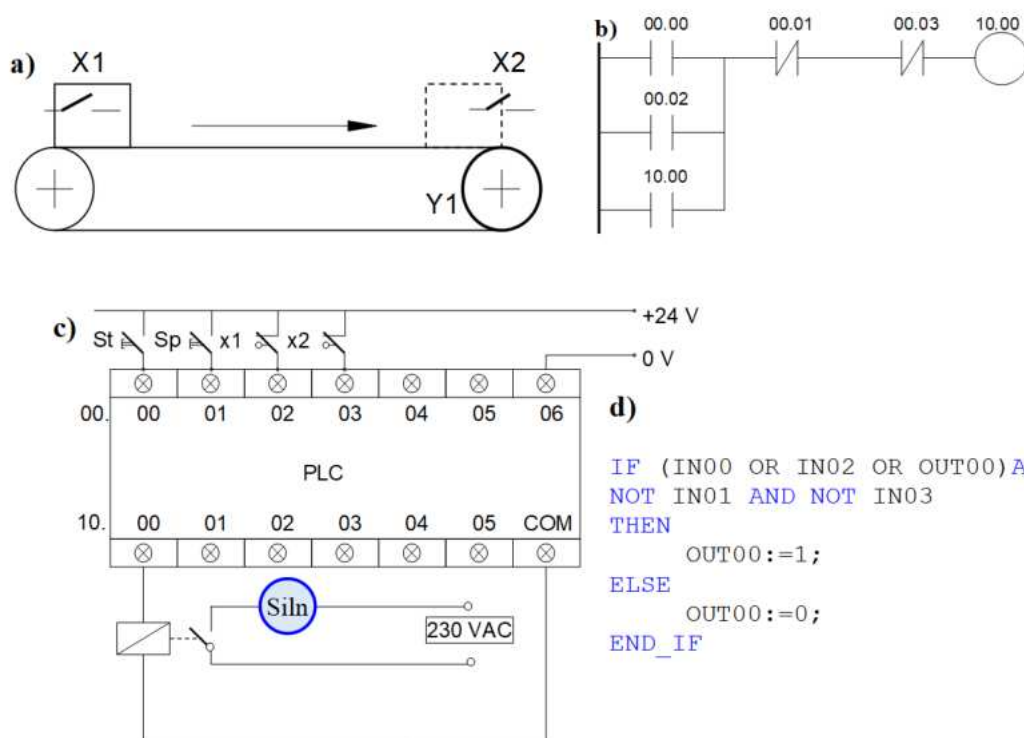
Tab. 2. Styki i cewki stosowane w języku LD

Lp.	Funkcja	Opis
1	— —	Styk normalnie otwarty (zwierny)
2	— /—	Styk normalnie zamknięty (rozwierny)
3	—()—	Cewka
4	—(/)—	Cewka negująca
5	—(S)—	Cewka ustawiająca
6	—(R)—	Cewka kasująca
7	—(↑)—	Cewka impulsowa reagująca na zbocze narastające
8	—(↓)—	Cewka impulsowa reagująca na zbocze opadające
9	—(M)—	Cewka z zapamiętaniem stanu

W tabeli 2 zamieszczono symbole najczęściej stosowanych w języku drabinkowym elementów. Elementy nr 5 i 6 to wejścia przerzutnika R-S, przy czym elementy (S) oznacza Set – ustaw/zapamiętaj, a element (R) oznacza Rset, czyli kasuj/zeruj. Elementy ze strzałkami tj. nr 7 i 8, ustawiają na wyjściu stan „1” tylko przez jeden cykl obiegu programu odpowiednio, gdy: w poprzednim cyklu było „0” a w obecnym jest „1” albo gdy: w poprzednim cyklu było „1”, a w obecnym jest „0”. Element 9 to pamięć – Memory (M).



Rys. 22. Układy pamięci: a) z przewagą kasowania, b) z przewagą ustawiania, c) z wykorzystaniem elementów S i R, d) aktywowanej zbroczem narastającym



Rys. 23. Sterowanie taśmociągiem: a) rysunek taśmociągu, b) program w języku LD, c) schemat podłączenia do sterownika elementów, d) program w języku strukturalnym

Na rysunku 22a zamieszczono program realizujący funkcję pamięci z podtrzymaniem, opisaną poniższym równaniem:

$$y01 = (x01 + y01) \cdot \overline{x02}$$

Na rysunku 22a zamieszczono program realizujący funkcję pamięci z podtrzymaniem, opisaną równaniem:

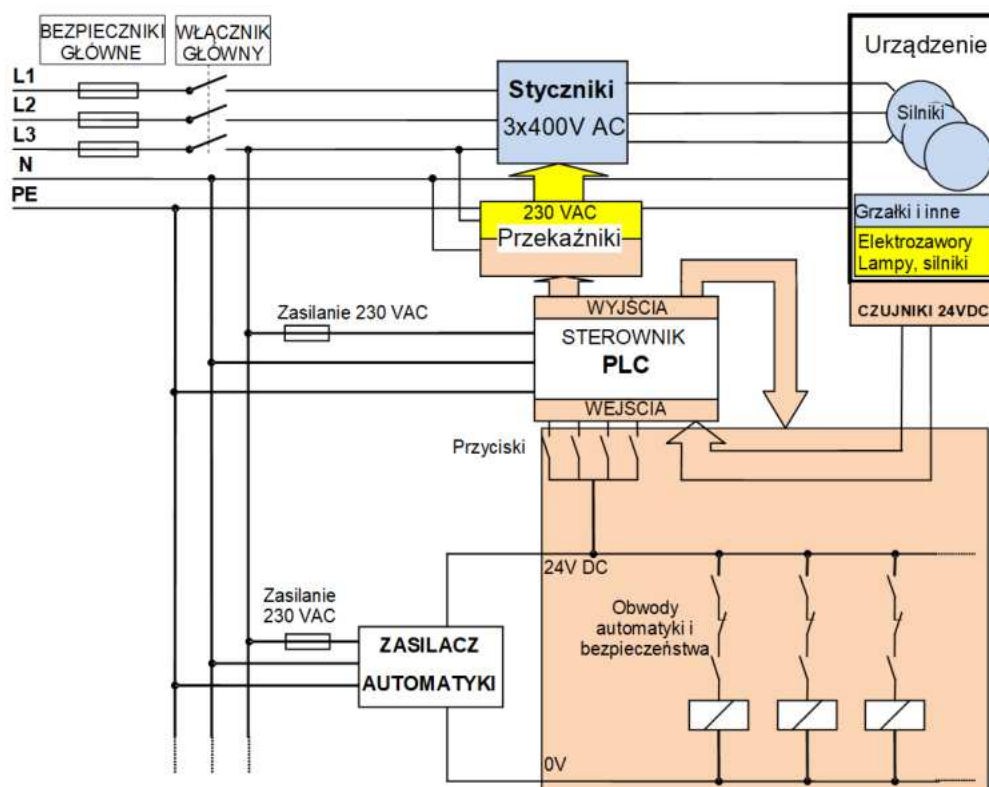
$$y01 = x01 + (y01 \cdot \overline{x02})$$



Na rysunku 23c przedstawiono program, do uruchomienia silnika taśmociągu, po pojawieniu się przedmiotu na jego początku, co jest sygnalizowane czujnikiem x1 (obecność przedmiotu = zwarcie, czyli „1”) lub wciśnięcie startu przełącznikiem St (zwarcie). Taśmociąg powinien zatrzymać się po osiągnięciu pozycji końcowej przez przemieszczany przedmiot, sygnalizowane czujnikiem x2 (obecność przedmiotu = zwarcie) lub naciśnięciem stopu awaryjnego przyciskiem Sp. Narysuj schemat połączeń do sterownika.

8. Zasilanie i połączenia urządzeń zautomatyzowanych

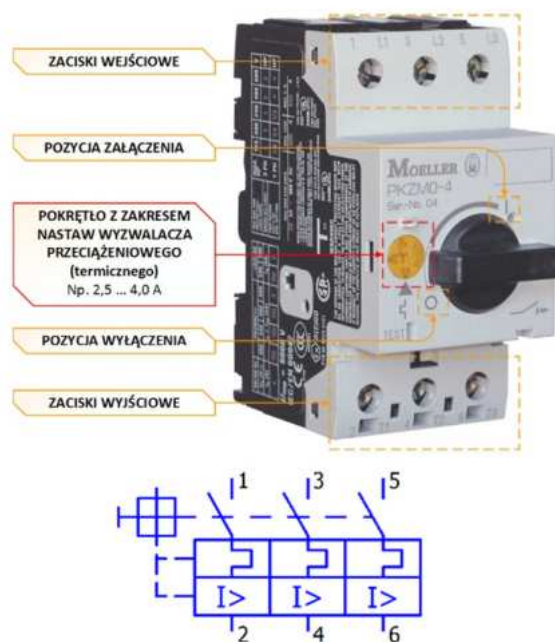
Sterowniki przemysłowe, różne czujniki oraz elementy wykonawcze małej mocy są zwykle zasilane napięciem 24 VDC. Takie napięcie jest zarówno bezpieczne dla człowieka, jak i odpowiednio wysokie dla zapewnienia ochrony układów automatyki przed zakłóceniami. Pozwala na włączanie i wyłączanie elementów o mocy do ok. 100 W, takich jak elektrozawory, przekaźniki, małe silniki, diody sygnalizacyjne itp. Przekaźniki są najczęściej wykorzystywane do włączania urządzeń o mocach do 2-3 kW, zasilanych napięciem jednofazowym 230 VAC. Do włączania i wyłączania elementów o większej mocy stosowane są zwykle trójfazowe styczniki. Schemat układu zasilania urządzenia zautomatyzowanego pokazano na rys. 24.



Rys. 24. Układ zasilania automatyk



Zabezpieczenie silników zarówno 1- jak i 3-fazowych jest niezbędne we wszystkich układach w automatyce przemysłowej. Systemy ochronne, mają za zadanie wykrywanie i reagowanie na potencjalne problemy, zanim doprowadzą one do uszkodzenia układu. Zabezpieczenia monitorują funkcjonowanie silników 3-fazowych, śledząc takie parametry jak natężenia prądu, napięcia i temperaturę. Zastosowanie zabezpieczeń dla silników 3-fazowych minimalizują ryzyko przeciążeń i skoków napięcia, które mogą grozić uszkodzeniem silnika. Ponadto, zabezpieczenia te zwiększają bezpieczeństwo procesów przemysłowych, chroniąc personel i maszyny przed potencjalnymi niebezpieczeństwami. Ich istotną funkcją jest także poprawa wydajności pracy silników, co przekłada się na zwiększenie wydajności całego systemu przemysłowego.



Rys. 25. Bezpiecznik termiczny trójfazowy i nadprądowy do silników 3-fazowych oraz jego symbol elektryczny

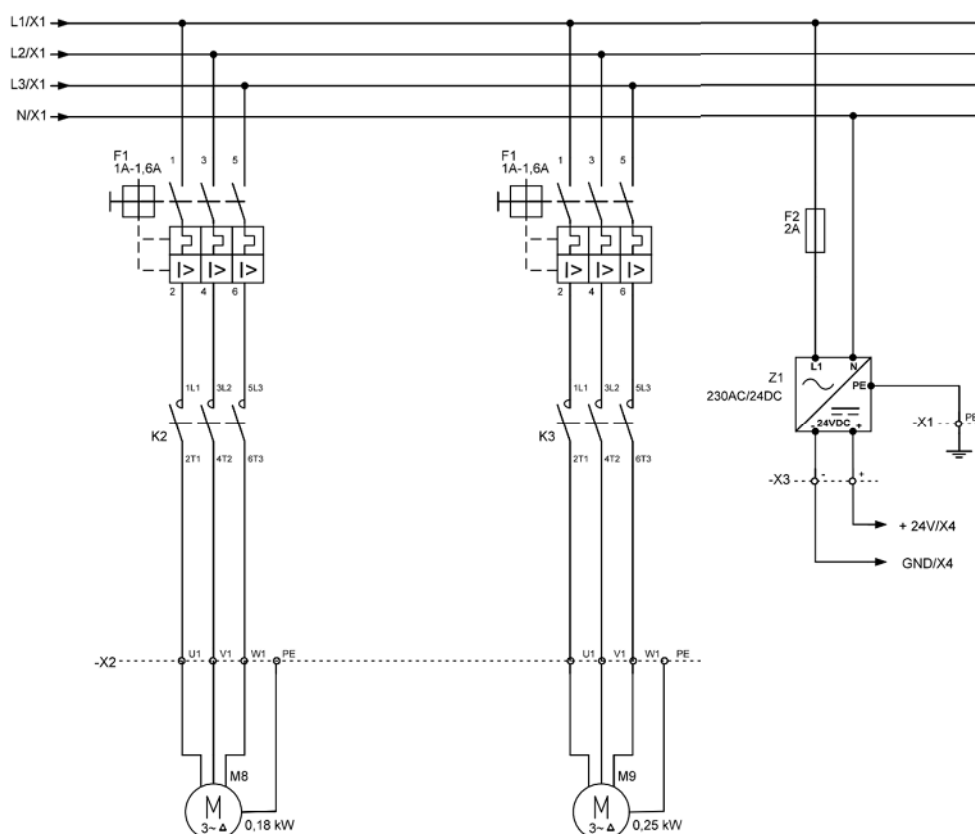
Niezabezpieczony silnik narażony jest na różne uszkodzenia powodowane przeciążeniami, zwarciami czy brakiem fazy. Stosowanymi urządzeniami do zabezpieczenia silnika trójfazowego są:

- wyłącznik różnicowo-prądowy – jest stosowany, gdy silnik zasilany jest bezpośrednio prądem z sieci tzn. nie ma falownika; zapewnia ochronę przeciwporażeniową,
- detektor zaniku fazy – zapobiega przeciążeniom spowodowanym przez zanik jednej z faz podczas pracy urządzenia,



- wyłączniki silnikowe (rys. 26) – wyłączają i załączają silniki; wyposażone są w wyłączacz termiczny i elektromagnetyczny. Wyłączają silnik w ciągu kilku milisekund, gdy dojdzie do przekroczenia jego prądu znamionowego. Umieszcza się je przed stycznikami, aby chroniły cały układ,
- czujnik zaniku fazy – łączy się ze stycznikiem albo z wyłącznikiem silnikowym; zabezpiecza silnik przed uszkodzeniem, spowodowanym w przypadku działania tylko dwóch faz. Ma za zadanie rozłączać obwód, gdy braknie jednej z faz,
- wyłączacz zanikowy – podłącza się go do wyłącznika silnika. Zapobiega samoczynnemu uruchomieniu się silnika tuż po ponownym włączeniu zasilania, co mogłoby stworzyć realne zagrożenie dla zdrowia lub życia osoby obsługującej maszynę. W przypadku zaniku napięcia wyłączy go i nie dopuści do jego automatycznego załączenia,
- czujnik asymetrii napięć zasilania – stosuje się go, aby wyłączyć zasilanie, gdy występuje różnica w ich wartościach napięć lub pojawi się niewłaściwe przesunięcie kątów.

Zastosowanie omówionych urządzeń pomaga skutecznie chronić silnik przed spaleniem czy innym uszkodzeniem na wypadek zaistnienia nieprawidłowości w sieci zasilającej lub instalacji elektrycznej. Oczywiście, aby sprawnie działały i należycie spełniały swoją funkcję, należy je dobrać do mocy silnika i parametrów układu.



Rys. 26. Podłączenie silników 3-fazowych przez wyłączniki silnikowe oraz zasilacza do sieci

9. Podstawy doboru elementów i projektowania linii zautomatyzowanych

Przed przystąpieniem do zaprojektowania automatyzacji linii produkcyjnej należy dokonać rozpoznania jej budowy działania itd. Proces projektowania powinien być realizowany w kilku krokach:

1. Rozpoznać ogólną budowę i działanie zautomatyzowanych poszczególnych urządzeń i całej linii produkcyjnej, w szczególności zadań automatyzacji:
 - ustalić zasady (funkcje), według których mają działać poszczególne urządzenia zautomatyzowane
 - zaproponować tabele prawdy, grafy reguły włączania i wyłączania poszczególnych urządzeń
 - określić reguły sterowania całą linią produkcyjną.
2. Rozpoznać ogólnie zadania urządzeń/linii produkcyjnej i ich parametry (wymagane) tj.:
 - wydajność, szybkość (np. ruchu)
 - dokładność wykonania
 - wymagania dotyczące niezawodności



- określić, co (jakie wielkości fizyczne) trzeba będzie mierzyć i z jaką dokładnością i precyzją
 - zdefiniować, jakie trzeba będzie zastosować silniki napędy w zakresie: wymaganych mocy średniej (znamionowej) i maksymalnej (chwilowej), momentów i sił, prędkości, zakresu ruchu, dokładności regulacji siły, prędkości i przyspieszenia
 - określić wstępne wymagania dotyczące zastosowania sterowników w zakresie:
 - jakimi elementami i urządzeniami trzeba sterować
 - które zadania mają charakter sterowania ciągłego, a które binarnego
 - jakie algorytmy będzie można zastosować.
3. Dobrać metody pomiarów oraz czujniki i urządzenia pomiarowe
- określić gdzie, tj. w których miejscach zastosować czujniki binarne
 - określić wymagania stawiane tym czujnikom (co wykrywać, z jakiej odległości, jaka będzie maksymalna częstotliwość itp.)
 - dobrać czujniki obecności (jest/nie ma - binarne: 0/1)
 - określić jakie wielkości fizyczne mierzyć (przemieszczenie, prędkość, siła, temperatura ...) i z jaką dokładnością oraz z jakim czasem pomiaru
 - dobrać jakie urządzenia pomiarowe zastosować, np. enkodery, liniały pomiarowe, wagi tensometryczne, czujniki laserowe, profilometry, termopary itp.
 - określić układy połączeń wybranych czujników i urządzeń ze sterownikami.
4. Dobrać zespoły napędowe
- jakie, ile i gdzie zastosować silniki, napędy itp.
 - określić ich wymagane parametry, a w szczególności czy powinna być stosowana regulacja np. prędkości i pozycjonowania
 - wybrać rodzaj napędów: elektryczne, hydrauliczne, pneumatyczne
 - ustalić wymagane: moc, prędkości, siły, momenty, dokładności
 - dobrać napędy i ich układy załączania i regulacji
 - wybrać sposób komunikacji napędów ze sterownikami
5. Dobrać sterowniki
- określić liczbę wejść i wyjść binarnych oraz analogowych
 - określić inne wymagane moduły np. regulatory, przetworniki itp.



- określić wymagania dotyczące sterowników w zakresie szybkości działania
 - wybrać sieć komunikacyjną
 - dobrać sterowniki napędów, a w szczególności sposób ich komunikacji z PLC
6. Rozpoznać wymagania dotyczące: minimalizacji zużycia energii, BHP, diagnostyki itp.

Na zautomatyzowanych liniach przemysłowych wyróżniamy

– Sterowniki (PLC, PAC, μ C)



– Czujniki



– napędy elektryczne hydrauliczne, pneumatyczne



Rys. 27. Podział elementów stosowane na zautomatyzowanych liniach produkcyjnych



Rys. 28. Rozmieszczenie elementów na zautomatyzowanej linii produkcyjnej

Sterowniki PLC, kontrolery PAC oraz panele operatorskie i komputery panelowe to istotne elementy nowoczesnej automatyki przemysłowej. Ich rola w Przemśle 4.0 stale rośnie, a użytkownicy oczekują coraz większej wydajności, elastyczności i zaawansowanych funkcji komunikacyjnych. PLC i PAC pozostają filarem automatyki dzięki niezawodności i łatwej integracji, podczas gdy HMI odgrywają ważną rolę w wizualizacji i sterowaniu procesami. Technologie te znajdują zastosowanie w przemyśle, energetyce, transporcie i logistyce, a ich rozwój napędza cyfryzacja i analiza danych w czasie rzeczywistym. Ważnymi trendami na tym rynku będą też: upowszechnianie się systemów sterowania opartych na PLC wykorzystujących algorytmy sztucznej inteligencji i postęp w zakresie ochrony sterowników programowalnych przed nieautoryzowanym dostępem i złośliwym oprogramowaniem w reakcji na rosnące zagrożenie cyberatakami. Najszybciej rozwijać się będzie segment modułowych PLC, cenionych za skalowalność oraz elastyczność – średnio o 11% rocznie.

Sterowniki przemysłowe można podzielić na grupy: nanosterowniki (do 32 we/wy), mikrosterowniki (do 128 we/wy), PLC średniej wielkości (do 1024 we/wy), PLC duże (> 1024 we/wy).

Oprócz tego podział obejmuje:

- Sterowniki kompaktowe/modułowe
- Sterowniki do systemów rozproszonych/do pracy wieloprocesorowej
- Wersje z redundancją/typu fail-safe



- Wersje z web serwerami/z OPC
Wersje zintegrowane z HMI oraz z interfejsami GSM
- Kontrolery automatyki (PAC)
- Systemy embedded
- Sterowniki oparte na technologiach PC (soft PLC, PC-based)
- Sterowniki dedykowane (np. mikroprocesorowe).

Najważniejszym kryterium doboru sterownika PLC jest to, aby można było podłączyć do niego wszystkie elementy, konieczne do zrealizowania zadania automatyzacji. Chodzi o odpowiednie liczby wejść i wyjść, tak aby każdy element bądź urządzenie mogły być do PLC podłączone. Ponadto dobrany sterownik powinien mieć możliwość rozszerzenia swoich możliwości, w tym zwiększenia liczby wejść/wyjść, aby umożliwić ewentualną rozbudowę.

Przykład 1. Sterowanie automatyczną linią do segregacji oraz napełniania i zamykania opakowań

Przydzielenie adresów dla poszczególnych elementów wejściowych zapisano w tabeli 3.

Tabela 3. Adresy elementów wejściowych w sterowniku PLC

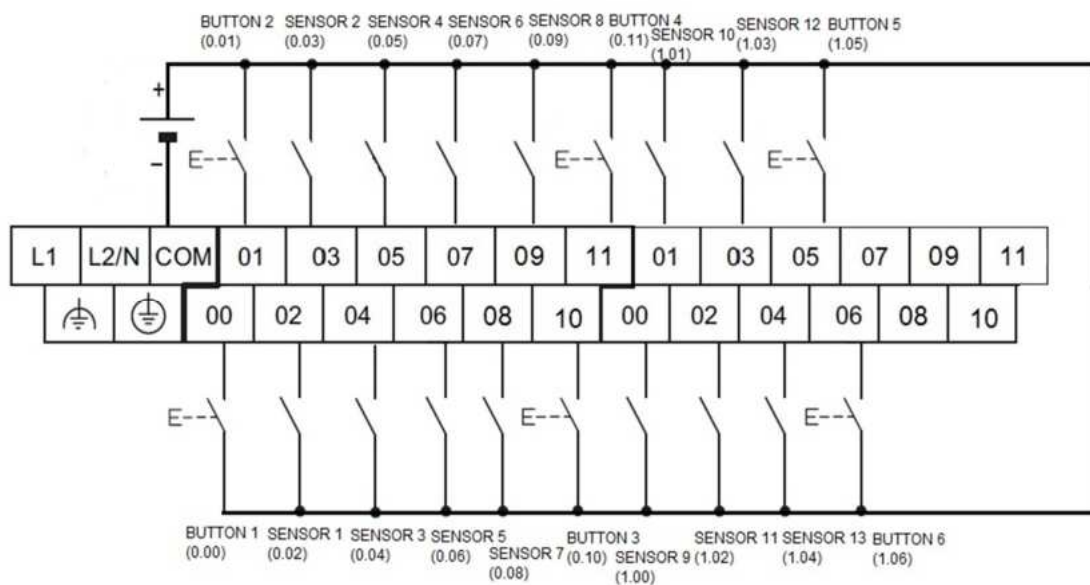
URZĄDZENIE	STYK	ADRES
Przycisk ręczny START 1	BUTTON 1	0.00
Przycisk ręczny STOP 1	BUTTON 2	0.01
Czujnik indukcyjny 1	SENSOR 1	0.02
Czujnik indukcyjny 2	SENSOR 2	0.03
Czujnik fotooptyczny	SENSOR 3	0.04
Czujnik indukcyjny 3	SENSOR 4	0.05
Czujnik indukcyjny 4	SENSOR 5	0.06
Czujnik indukcyjny 5	SENSOR 6	0.07
Czujnik indukcyjny 6	SENSOR 7	0.08
Czujnik indukcyjny 7	SENSOR 8	0.09
Przycisk ręczny START 2	BUTTON 3	0.10
Przycisk ręczny STOP 2	BUTTON 4	0.11
Czujnik indukcyjny 8	SENSOR 9	1.00
Czujnik indukcyjny 9	SENSOR 10	1.01
Czujnik indukcyjny 10	SENSOR 11	1.02
Czujnik indukcyjny 11	SENSOR 12	1.03
Czujnik indukcyjny 12	SENSOR 13	1.04
Przycisk ręczny START 3	BUTTON 5	1.05
Przycisk ręczny STOP 3	BUTTON 6	1.06

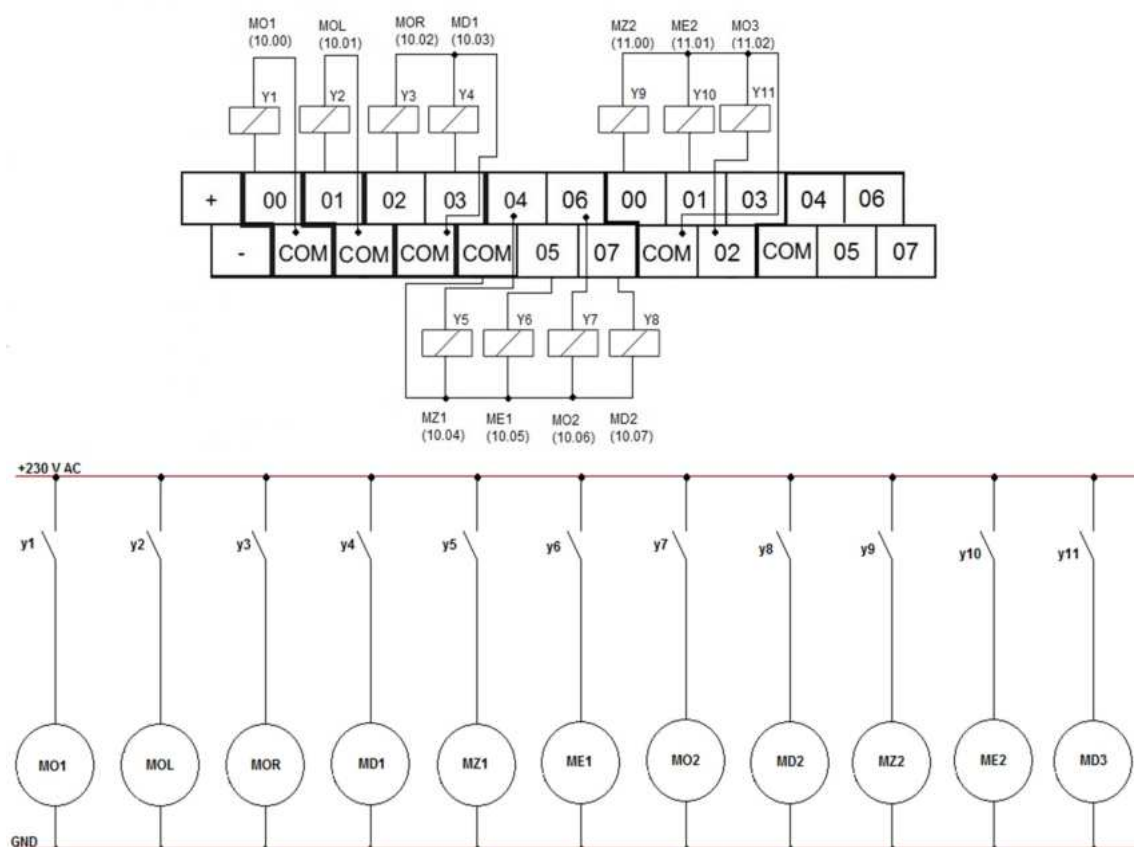


Przydzielenie adresów dla poszczególnych urządzeń wyjściowych (tab.4).

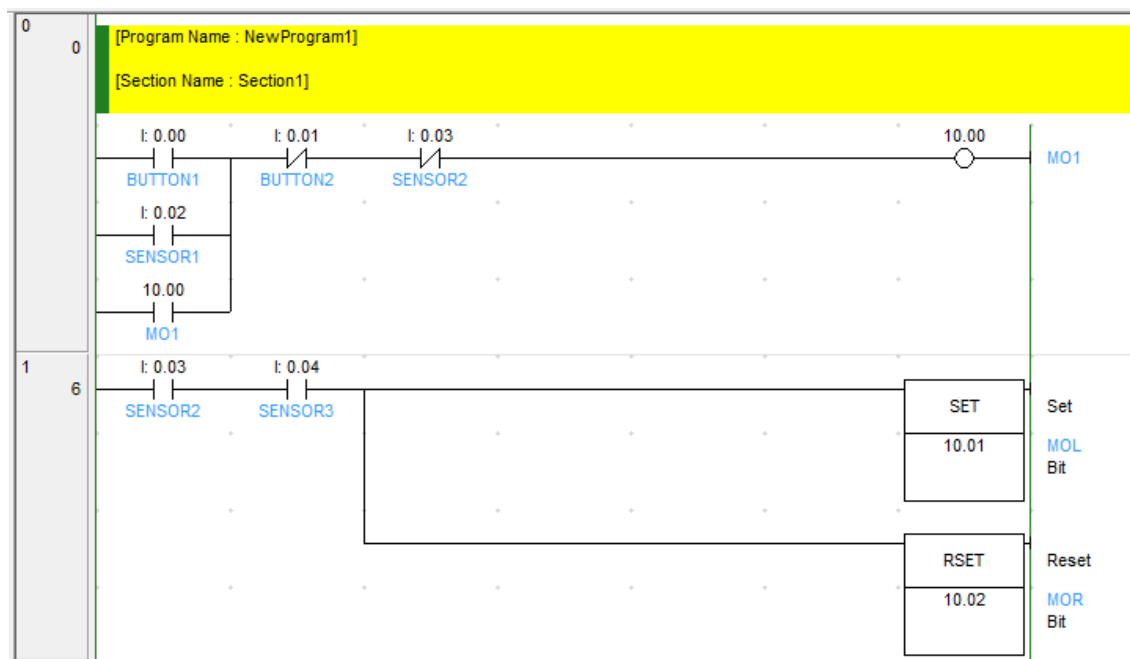
Tabela 4. Adresy urządzeń wyjściowych

URZĄDZENIE	STYK	ADRES
Silnik taśmociągu 1	MO1	10.00
Obroty lewe silnika mechanizmu segregującego	MOL	10.01
Obroty prawe silnika mechanizmu segregującego	MOR	10.02
Napęd dozownika 1	MD1	10.03
Napęd maszyny zamykającej 1	MZ1	10.04
Napęd maszyny etykietującej 1	ME1	10.05
Silnik taśmociągu 2	MO2	10.06
Napęd dozownika 2	MD2	10.07
Napęd maszyny zamykającej 2	MZ2	11.00
Napęd maszyny etykietującej 2	ME2	11.01
Silnik taśmociągu 3	MO3	11.02





Rys. 29. Podłączenie elementów do wejść i wyjść sterownika OMRON CP1L



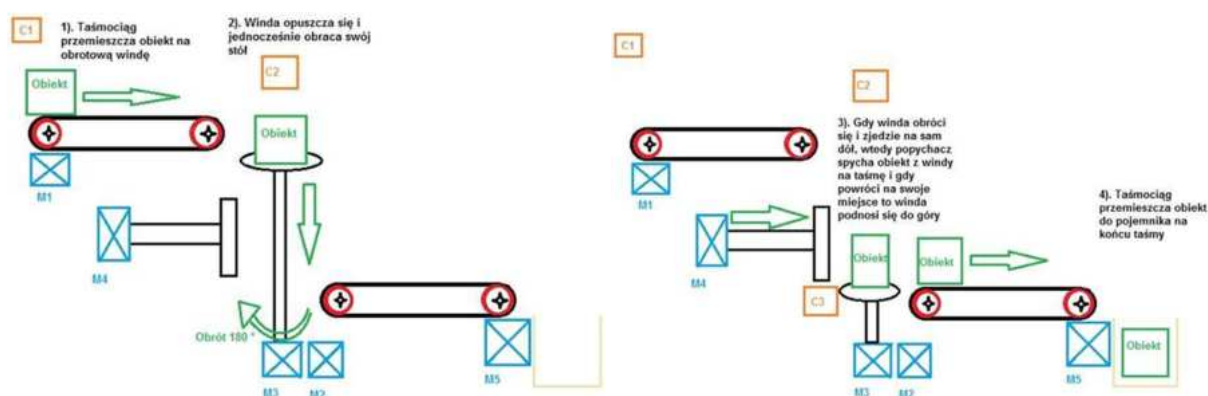
Rys. 30. Przykładowa linia programu na sterownik PLC



Przykład 2. Sterowanie dwoma taśmociągami i windą

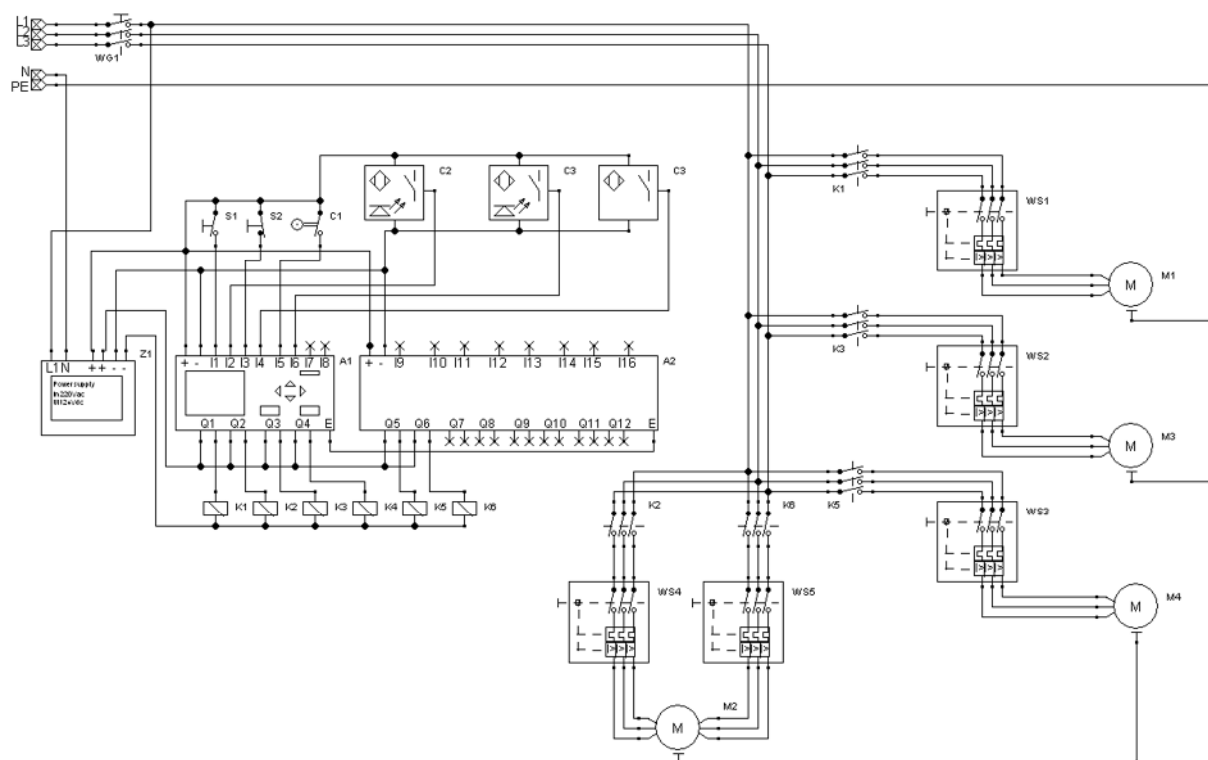
Zadanie projektowe polega na sterowaniu:

1. przesunięciem elementu na taśmociąg 1,
2. opuszczeniu elementu przez windę,
3. przesunięciu elementu na taśmociąg 2
4. przesunięciu elementu na koniec taśmociągu 2.



Zasada działania:

1. Gdy czujnik C1 wykryje element na taśmie nr 1, załączony zostaje silnik M1, który napędza taśmociąg, który przesuwa obiekt na koniec i zrzuca go na windę.
2. Gdy czujnik C2 wykryje element na windzie, uruchamia silniki M3 oraz M2, które powodują ruch windy w dół, przy jednoczesnym obrocie obiektu na windzie.
3. Gdy czujnik C3 wykryje element, załączony zostaje silnik M4, który popycha obiekt z windy na taśmociąg nr 2 i włącza się silnik M5 powodujący ruch taśmociągu 2 i wrzucenie obiektu do pojemnika.



Rys. 31. Schemat połączeń do zasilania (zastosowano sterownik LOGO)

Bibliografia

- Kasprzyk J., Hajda J., Programowanie sterowników PLC, Wydawnictwo Pracowni Komputerowej Jacka Skalmierskiego, 1998.
- Kosmol J., Automatyzacja obrabiarek i obróbki skrawaniem, WNT, Warszawa 1995
- Kost G., Łebkowski P., Węsierski Ł., Automatyzacja i robotyzacja procesów produkcyjnych, PWE, 2013.
- Kostro J., Elementy, urządzenia i układy automatyzacji, Wydawnictwa Szkolne i Pedagogiczne, 1993.
- Mikulczyński T., Samsonowicz Z., Więclawek R., Automatyzacja procesów produkcyjnych, 2015, wydawnictwo: Mikulczyński Tadeusz, Samsonowicz Zdzisław, Więclawek Rafał.
- Szelerski M. W., Automatyka przemysłowa w praktyce, Projektowanie, modernizacja i naprawa, 2016, Wydawca: KaBe.



Robert Rogaczewski

PLANOWANIE I STEROWANIE PRODUKCJĄ

Wstęp

Planowanie i sterowanie produkcją to fundamenty zarządzania procesami w każdym przedsiębiorstwie produkcyjnym. Współczesne firmy, funkcjonujące w dynamicznym i konkurencyjnym środowisku, muszą sprostać rosnącym wymaganiom rynku, jednocześnie dążąc do optymalizacji swoich procesów produkcyjnych. Skuteczne planowanie pozwala przewidzieć i zaplanować niezbędne zasoby, czas oraz metodologie, które umożliwią realizację założonych celów produkcyjnych. Z kolei sterowanie produkcją jest procesem bieżącego monitorowania oraz modyfikowania działań w odpowiedzi na zmieniające się warunki w trakcie realizacji planu produkcji. Celem tych działań jest zapewnienie płynności, efektywności oraz jakości procesów wytwórczych, a także szybkie reagowanie na pojawiające się problemy czy zmiany w zapotrzebowaniu.

Współczesne przedsiębiorstwa produkcyjne muszą zmierzyć się z licznymi wyzwaniami, takimi jak zmieniające się preferencje konsumentów, sezonowość popytu, zmieniające się ceny surowców czy rosnąca globalizacja łańcuchów dostaw. W odpowiedzi na te wyzwania, skuteczne planowanie i sterowanie produkcją stają się kluczowe. Aby procesy produkcyjne były optymalne, muszą być one nie tylko dobrze zaplanowane, ale także elastyczne i dostosowane do zmieniających się warunków.

Planowanie produkcji obejmuje szereg działań związanych z prognozowaniem zapotrzebowania na produkty, określeniem zasobów potrzebnych do ich wytworzenia, a także harmonogramowaniem poszczególnych etapów produkcji. Dobrze zaplanowany proces produkcyjny pozwala na efektywne wykorzystanie dostępnych zasobów, zminimalizowanie kosztów operacyjnych i zapewnienie ciągłości produkcji. Natomiast sterowanie produkcją jest niezbędne do monitorowania realizacji planów produkcyjnych oraz szybkiego reagowania na pojawiające się problemy, takie jak przestoje maszyn, braki w materiałach, zmiany w harmonogramach czy opóźnienia.

Aby procesy planowania i sterowania produkcją były skuteczne, przedsiębiorstwa coraz częściej sięgają po nowoczesne technologie i narzędzia informatyczne, takie jak systemy ERP, MRP, a także narzędzia do zarządzania jakością i optymalizacji procesów produkcyjnych.



Dzięki tym systemom możliwe jest dokładniejsze prognozowanie zapotrzebowania, precyzyjne planowanie zasobów oraz bieżące monitorowanie postępów produkcji. Integracja tych narzędzi w codziennym zarządzaniu produkcją umożliwia szybsze podejmowanie decyzji oraz zwiększa elastyczność w reagowaniu na zmiany rynkowe.

Podsumowując, planowanie i sterowanie produkcją są kluczowymi procesami, które mają bezpośredni wpływ na efektywność i rentowność przedsiębiorstwa produkcyjnego. Zrozumienie tych procesów oraz umiejętność ich skutecznego zarządzania pozwala na optymalizację zasobów, redukcję kosztów, poprawę jakości produktów oraz zwiększenie satysfakcji klientów. Długoterminowy sukces przedsiębiorstwa zależy od jego zdolności do efektywnego planowania i zarządzania produkcją, a także dostosowywania się do zmieniających się warunków rynkowych.

1. Planowanie produkcji

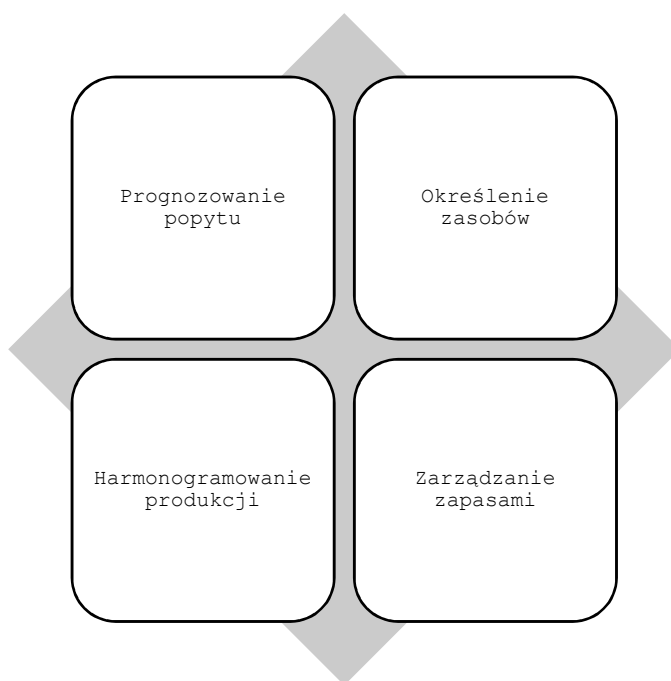
Planowanie produkcji to proces opracowywania szczegółowego harmonogramu i strategii działań, który ma na celu zapewnienie optymalnej organizacji wszystkich zasobów potrzebnych do wytworzenia określonych produktów. Jego głównym celem jest przewidywanie i organizowanie wszystkich aspektów związanych z produkcją, w tym zapotrzebowania na materiały, ludzkie zasoby, czas pracy maszyn oraz inne zasoby niezbędne do realizacji zleceń. Planowanie produkcji stanowi kluczowy element w zarządzaniu procesem produkcyjnym, mając na celu zminimalizowanie kosztów, eliminację nadmiernych zapasów, a także poprawę efektywności wytwórczej. Dobre **planowanie produkcji** pozwala przedsiębiorstwom nie tylko na optymalizację wykorzystania zasobów, ale także na elastyczną reakcję na zmieniające się warunki rynkowe, zapotrzebowanie czy zmieniające się wymagania jakościowe (Heizer, Render, 2014).

Planowanie produkcji obejmuje następujące kluczowe elementy (rys. 1):

- a) prognozowanie popytu (przewidywanie zapotrzebowania na surowce, materiały czy półprodukty, z których zostaną wykonane wyroby końcowe oraz na produkty gotowe, które potencjalni klienci skłonni są zakupić),
- b) określenie zasobów (określenie, jakie zasoby są niezbędne do realizacji planu produkcji),



- c) harmonogramowanie produkcji (ustalenie szczegółowego harmonogramu procesów produkcyjnych, który uwzględnia dostępność zasobów i czas potrzebny do wytworzenia produktów),
- d) zarządzanie zapasami (zapewnienie, aby surowce, materiały i półprodukty były dostępne na czas i w odpowiedniej ilości, aby wyeliminować ryzyko powstania przestojów w produkcji).



Rys. 1. Kluczowe obszary planowania produkcji
Źródło: opracowanie własne

Planowanie produkcji należy rozpatrywać na kilku płaszczyznach. Pierwszym takim aspektem jest bez wątpienia analizowanie planowania produkcji jako procesu organizacyjnego. Według Filipiaka (2003) planowanie produkcji należy rozpatrywać jako proces organizowania i koordynowania działań związanych z produkcją towarów, który obejmuje określenie, jaki produkty mają być wytwarzane, w jakiej ilości oraz w jakim czasie, uwzględniając dostępność zasobów (surowców, maszyn czy siły roboczej) oraz zapotrzebowanie rynkowe. Planowanie produkcji należy również postrzegać jako podejście strategiczne, które ma na celu określenie, jak najlepiej wykorzystać dostępne zasoby w celu efektywnego wytwarzania produktów, z uwzględnieniem zarówno krótkoterminowych, jak i długoterminowych celów organizacji (Sierżant, 2000). Dość istotnym aspektem jest rozpatrywanie planowania produkcji jako



procesu ustalania harmonogramów, który obejmuje informacje, kiedy i w jakiej ilości produkty są wytwarzane, aby spełnić oczekiwania klienta i jego zapotrzebowanie.

Planowanie produkcji jest ściśle związane z systemami zarządzania produkcją, takimi jak MRP (Material Requirements Planning) czy ERP (Enterprise Resource Planning), które automatyzują i usprawniają procesy planowania i zarządzania zasobami w firmach.

Planowanie potrzeb materiałowych

System **MRP³** to system planowania potrzeb materiałowych, pozwalający na określenie zapotrzebowania na dane komponenty. MRP oblicza zapotrzebowanie na materiały w oparciu o harmonogram produkcji, zapewniając, że odpowiednie materiały będą dostępne w odpowiednich ilościach w odpowiednim czasie, aby umożliwić nieprzerwaną produkcję. MRP to system planowania zapotrzebowania materiałowego, które przekształca główny harmonogram produkcji (MPS) w szczegółowe zapotrzebowanie na materiały, komponenty i półprodukty niezbędne do produkcji w ustalonych terminach. Niewątpliwie system umożliwia efektywne zarządzanie zapasami i terminami dostaw, co wspiera płynność procesów produkcyjnych (Volmann i in., 2012).

Planowanie potrzeb materiałowych obejmuje każdy element wyrobu finalnego, w każdej fazie procesu produkcji i określa harmonogram zapotrzebowania materiałowego, wynikający z asortymentu, wielkości i terminu wykonania partii produkcyjnej. Zgodnie z metodą MRP obliczane są terminy dostawy materiałów i elementów koniecznych do wytworzenia wyrobu gotowego, zgodnie z harmonogramem produkcji. W wyniku planowania potrzeb materiałowych opracowywany jest harmonogram dostaw, będący podstawą planowania zapotrzebowania materiałowego. Metoda MRP:

- określa harmonogram zapotrzebowania materiałowego wynikający z asortymentu, wielkości i terminu wykonania partii produkcyjnej,
- korzysta z obliczonych terminów dostaw surowców potrzebnych do wykonania produktu gotowego zgodnie z harmonogramem produkcji,
- obejmuje każdy element wyrobu finalnego.

Planowanie potrzeb materiałowych obejmuje planowanie wielkości i terminów potrzeb poszczególnych materiałów wchodzących w skład wyrobu gotowego. Wielkość potrzeb materiałowych – potrzeb netto (P_N) są określane na podstawie:

³ Planowanie potrzeb materiałowych.



- potrzeb brutto wyrobów (P_B) – liczby wyrobów gotowych określanych przez wielkość partii produkcyjnej w głównym harmonogramie produkcji,
- struktury wyrobu gotowego – określającej zapotrzebowanie na wszystkie materiały i elementy składowe wchodzące w skład wyrobu,
- aktualnego stanu zapasów dysponowanych poszczególnych materiałów (S_{ZAP}) - zapasów możliwych do wykorzystania po odliczeniu, np. zapasów zarezerwowanych do innych potrzeb i wymaganych zapasów końcowych – bezpieczeństwa oraz złożonych zamówień (S_{ZAM}), pozostających w trakcie realizacji.

Podstawą tego systemu jest ustalenie potrzeb materiałowych brutto i netto, wzór natomiast kształtuje się następująco:

$$P_N = P_B - (S_{ZAP} + S_{ZAM})^4$$

Omawiany system bazuje na trzech podstawowych płaszczyznach, tj.:

- głównym harmonogramie produkcji,
- zbiorze struktury wyrobu⁵,
- głównym zbiorze zapasów.

Metoda MRP obejmuje każdy element wyrobu końcowego oraz określa zapotrzebowanie, które wynika z asortymentu, wielkości i terminu wykonania określonej partii zamówionego asortymentu. Wskazać należy również na zalety stosowania systemu MRP. Należy zaliczyć do nich bez wątpienia:

- możliwość obniżenia stanów zapasów oraz poprawy poziomu obsługi odbiorców,
- możliwość tworzenia realnych harmonogramów produkcji,
- możliwość śledzenia zleceń na składniki z zestawienia materiałowego BOM i analizowania wpływu ich opóźnień na realizację MPS,
- możliwość szybkiego informowania dostawców o planowanych potrzebach.
- dążenie do utrzymania zapasu bezpieczeństwa na rozsądnym poziomie,
- korygowanie czynności związanych z zamawianiem materiałów we wszystkich miejscach systemu,

⁴ P_N – potrzeby netto, P_B – potrzeby brutto.

⁵ Strukturę materiałową wyrobu przedstawia się za pomocą listy materiałowej i drzewa struktury. Lista materiałowa to lista wszystkich zespołów, podzespołów i części, które wchodzi w skład wyrobu. Informuje ona również o licznie poszczególnych elementach wchodzących w skład wyrobu finalnego. Drzewo struktury przedstawia graficznie, w jaki sposób elementy wymienione w liście materiałowej są łączone ze sobą w celu otrzymania gotowego wyrobu. Informuje również o planowanych czasach realizacji (CR) wszystkich elementów. Tworzy się wówczas wielopoziomowy obraz struktury wyrobu.



- przydatność w produkcji w partiach lub produkcji przerywanej, a także przy procesach montażu.

Z kolei do wad rozwiązań opartych na MRP należą:

- wdrożenie rozwiązań wymaga zastosowania szybkich komputerów, a kiedy system już funkcjonuje, wprowadzenie do niego zmian może być utrudnione,
- koszty zamówień jak i koszty transportu mogą rosnąć w miarę jak firma obniża poziom zapasów i dąży do stworzenia bardziej skoordynowanego systemu, w którym zamawia się mniejsze ilości produktów dostarczanych wtedy, kiedy są one potrzebne.

Reasumując – główne cele systemu MRP to:

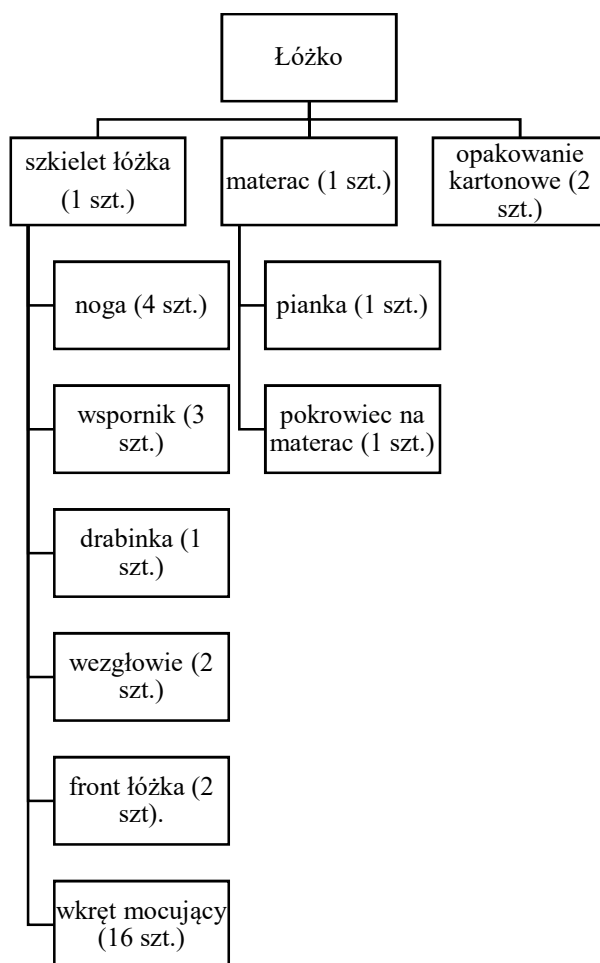
- zapewnienie wystarczającej ilości materiałów, części i produktów na potrzeby zaplanowanej produkcji i dostaw do klienta,
- utrzymanie możliwie najniższego poziomu zapasów,
- planowanie działań produkcyjnych, harmonogramu dostaw i zakupów.

Przykład 1⁶

System MRP jest kluczowy dla planowania potrzeb materiałowych i głównego harmonogramu produkcji. Warto przy tej okazji prześledzić przykład funkcjonowania tego systemu w przedsiębiorstwie branży drzewnej. Jest to stolarnia, która zajmuje się produkcją różnego rodzaju mebli. System planowania potrzeb materiałowych w tym przedsiębiorstwie opiera się na trzech filarach, tj. głównym harmonogramie produkcji, strukturze wyrobu oraz zapasach. Połączenie wyżej wymienionych elementów i uzupełnienie ich o listę materiałową, czy też planowanie „wstecz” pozwoliło na uzyskanie sprawnie funkcjonującego systemu.

Jednym z przykładowych produktów analizowanego przedsiębiorstwa jest łóżko, którego strukturę wyrobu przedstawiono poniżej:

⁶ Na podstawie: G. Karpus, Zbiór zadań z logistyki, WSiP, Warszawa, 2017.



Rys. 2. Struktura wyrobu „łóżka”

Źródło: opracowanie własne

Wraz z wpłynięciem zamówienia do stolarni (na potrzeby niniejszego przykładu założono, iż będzie to 300 łóżek, zapotrzebowanie wystąpi w dn. 10.07.2021 r.) dokonywana jest analiza potrzeb materiałowych. Rozpatrywane są zazwyczaj następujące kwestie: (zestawienie potrzeb materiałowych brutto i netto niezbędne do realizacji, struktura wyrobu (rys. 2), stan zapasów magazynowych, produkcji w toku oraz cen materiałów (tab. 1).

Tabela 1. Stan zapasów magazynowych, produkcji w toku oraz cen materiałów

Lp.	Zapas	Stan magazynowy	Cena jednostkowa netto zakupu
1	Łóżko – wyrób gotowy	15 szt.	
2	Szkielet łóżka	20 szt.	
3	Materac	35 szt.	
4	Opakowanie kartonowe	120 szt.	2,35 zł/szt.
5	Noga	88 szt.	32,00 zł/szt.
6	Wspornik	56 szt.	51,50 zł/szt.



7	Front łóżka	36 szt.	157,00 zł/szt.
8	Drabinka	52 szt.	143,00 zł/szt.
9	Wezgłowie	48 szt.	120,00 zł/szt.
10	Pianka	30 szt.	415,00 zł/szt.
11	Pokrowiec na materac	300 szt.	57,25 zł/szt.
12	Wkręty (opak. 100 szt.)	210 opak.	27,50 zł/opak.

Źródło: opracowanie własne

Ilościowo-wartościowe zestawienie potrzeb materiałowych będzie kształtowało się zatem na następującym poziomie:

Tabela 2. Zestawienie ilościowo-wartościowe potrzeb materiałowych

Wyszczególnienie	Zapotrzebowanie brutto (ilość)	Zapotrzebowanie brutto (wartość)	Zapasy dysponowany (ilość)	Zapotrzebowanie netto (ilość)	Zapotrzebowanie netto (wartość)
Łóżko	300 szt.		15 szt.	285 szt.	
Opakowanie kartonowe	570 szt.	1339,50 zł	120 szt.	450 szt.	1057,50 zł.
Szkielet łóżka	285 szt.		20 szt.	265 szt.	
Materac	285 szt.		35 szt.	250 szt.	
Pianka	250 szt.	103750 zł.	30 szt.	220 szt.	91300 zł.
Pokrowiec na materac	250 szt.	14312,50 zł.	300 szt.	0 szt.	0 zł.
Noga	1060 szt.	33920 zł.	88 szt.	972 szt.	31104 zł.
Wspornik	795 szt.	40942,50 zł.	56 szt.	739 szt.	38058,50 zł.
Front łóżka	530 szt.	83210 zł.	36 szt.	494 szt.	77558 zł.
Drabinka	265 szt.	37895 zł.	52 szt.	213 szt.	30459 zł.
Wezgłowie	530 szt.	63600 zł.	48 szt.	482 szt.	57840 zł.
Wkręt	43 opak.	1182,5 zł.	210 opak.	0 opak.	0 zł.
RAZEM		418047 zł.			327377 zł.

Źródło: opracowanie własne

Należy również określić, jaki jest czas realizacji bądź dostawy poszczególnych elementów produktu finalnego. Znajomość cyklu dostawy (T_{CD}) lub też źródła zakupu jest kluczowe dla planowania potrzeb materiałowych. Na potrzeby niniejszego przykładu czas realizacji przedstawiono w tabeli 3.



Tabela 3. Czas realizacji poszczególnych elementów struktury wyrobu

Lp.	Element	T _{CD}	Źródło
1	Łóżko – wyrób gotowy	2 dni	Produkcja - montaż
2	Szkielet łóżka	1 dzień	Montaż
3	Materac	1 dzień	Montaż
4	Opakowanie kartonowe	2 dni	Zakup
5	Noga	2 dni	Produkcja
6	Wspornik	2 dni	Produkcja
7	Front łóżka	2 dni	Produkcja
8	Drabinka	3 dni	Zakup
9	Wezgłowie	3 dni	Zakup
10	Pianka	4 dni	Zakup
11	Pokrowiec na materac	4 dni	Zakup
12	Wkręty (opak. 100 szt.)	2 dni	Zakup

Źródło: opracowanie własne

Planowanie potrzeb materiałowych dla ww. zlecenia będzie kształtowało się następująco:



Fundusze Europejskie
dla Wielkopolski

Dofinansowane przez
Unię Europejską



SAMORZĄD
WOJEWÓDZTWA
WIELKOPOLSKIEGO

Dni (lipiec)	1	2	3	4	5	6	7	8	9	10
Łóżko – potrzeby brutto										300
Zapas	15	15	15	15	15	15	15	15	15	15
Potrzeby netto										285
T _{ZAM} – zlecenie produkcyjne								285		
Szkielet łóżka – potrzeby brutto								285		
Zapas	20	20	20	20	20	20	20	20		
Potrzeby netto								265		
T _{ZAM} – zlecenie produkcyjne							265			
Materac – potrzeby brutto								285		
Zapas	35	35	35	35	35	35	35	35		
Potrzeby netto								250		
T _{ZAM} – zlecenie produkcyjne							250			
Opakowanie kart. – potrzeby brutto								570		
Zapas	120	120	120	120	120	120	120	120		
Potrzeby netto								450		
T _{ZAM} – zakup						450				
Noga – potrzeby brutto							1060			
Zapas	88	88	88	88	88	88	88			
Potrzeby netto							972			
T _{ZAM} – zlecenie produkcyjne					972					
Wspornik – potrzeby brutto							795			
Zapas	56	56	56	56	56	56	56			
Potrzeby netto							739			
T _{ZAM} – zlecenie produkcyjne					739					
Front łóżka – potrzeby brutto							530			



Fundusze Europejskie
dla Wielkopolski

Dofinansowane przez
Unię Europejską



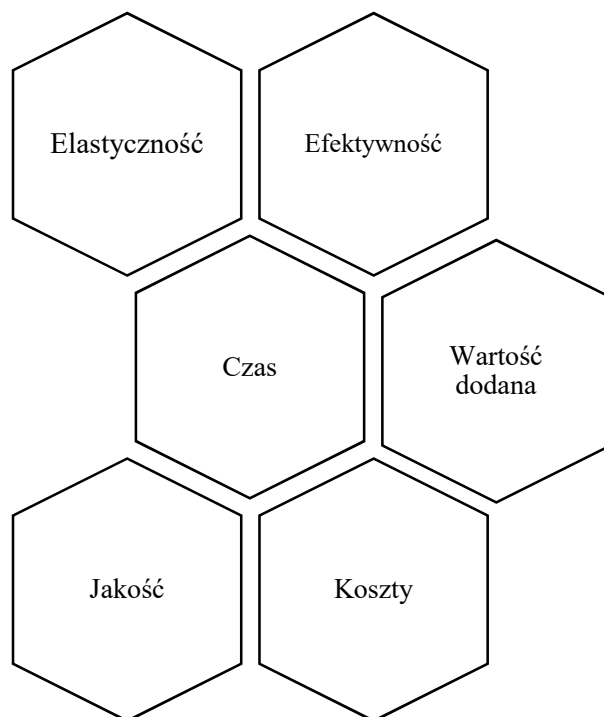
SAMORZĄD
WOJEWÓDZTWA
WIELKOPOLSKIEGO

Zapas	36	36	36	36	36	36	36			
Potrzeby netto							494			
T _{ZAM} – zlecenie produkcyjne					494					
Drabinka – potrzeby brutto							265			
Zapas	52	52	52	52	52	52	52			
Potrzeby netto							213			
T _{ZAM} – zakup				213						
Wezglowie – potrzeby brutto							530			
Zapas	48	48	48	48	48	48	48			
Potrzeby netto							482			
T _{ZAM} – zakup				482						
Pianka – potrzeby brutto							250			
Zapas	30	30	30	30	30	30	30			
Potrzeby netto							220			
T _{ZAM} – zakup			220							
Pokrowiec na mat. – potrzeby brutto							250			
Zapas	300	300	300	300	300	300	300	50	50	50
Potrzeby netto							0			
T _{ZAM} – zakup			0							
Wkręty – potrzeby brutto							43			
Zapas	210	210	210	210	210	210	210	167	167	167
Potrzeby netto							0			
T _{ZAM} – zakup					0					



Harmonogramowanie produkcji

Harmonogramowanie to rozkład zadań (operacji) dla ludzi i maszyn. Obejmuje ono przydział zadań do stanowisk produkcyjnych, ustalenie kolejności wykonywania zadań oraz wprowadzanie zmian do opisu. Harmonogramowanie produkcji to proces planowania, w którym ustala się szczegółowy plan realizacji poszczególnych etapów produkcji. Celem harmonogramowania produkcji jest zapewnienie, że proces produkcji przebiega płynnie, zgodnie z ustalonymi terminami, przy minimalizacji przestojów, optymalnym wykorzystaniu zasobów oraz zachowaniu wysokiej jakości wytwarzanych produktów. Harmonogramowanie pozwala na efektywne zarządzanie czasem, ludźmi, maszynami i materiałami, tak aby produkcja odbywała się zgodnie z planem i spełniała wymagania rynku (Jacobs, Chase, 2011). Należy również wskazać na ważniejsze cele harmonogramowania produkcji (rys. 3)



Rys. 3. Kluczowe cele harmonogramowania produkcji
Źródło: opracowanie własne

Harmonogram to chronologiczny porządek czynności, wykonywanych przez różne jednostki organizacyjne. Wskazuje w jakiej kolejności i w ciągu jakiego czasu muszą być wykonane poszczególne zadania w celu realizacji zlecenia produkcyjnego. Możliwe jest określenie:

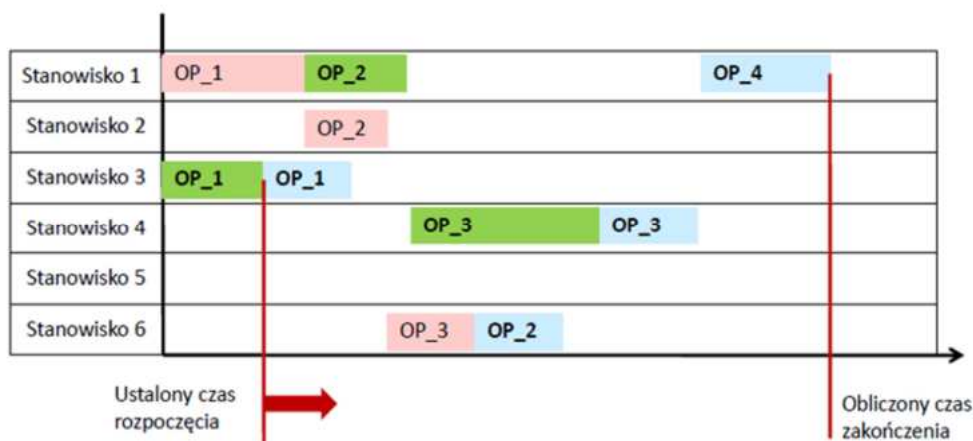
- kiedy poszczególne zadania muszą się rozpocząć i zakończyć,
- w jakim stopniu obciążone są zasoby produkcyjne (pracownicy, środki produkcji)



- tworzy czasowe powiązania wszystkich stanowisk biorących udział w realizacji zlecenia produkcyjnego.

Elementami harmonogramu produkcji są przede wszystkim zadania produkcyjne (poszczególne etapy produkcji, takie jak przygotowanie produkcji, produkcja komponentów, montaż, kontrola jakości czy pakowanie), czas rozpoczęcia i zakończenia (każdemu zadaniu przypisany jest czas rozpoczęcia oraz zakończenia, z kolei czas realizacji poszczególnych etapów może być wyznaczany na podstawie doświadczenia, norm czasowych, bądź danych historycznych), zasoby (określenie zasobów (maszyny, urządzenia, surowce, ludzie) przypisanych do każdego zadania, dzięki czemu można zaplanować dostępność zasobów i zapobiec ich przeciążeniu), priorytety (w przypadku dużej liczby zadań, harmonogram może zawierać priorytety, które wskazują na zadania o wyższej ważności), zależność między zadaniami (określenie, które zadania muszą być wykonane przed innymi; często wykorzystywane w przypadku złożonych procesów produkcyjnych, gdzie jedna czynność musi zakończyć się przed rozpoczęciem kolejnej) oraz przestój (określenie potencjalnych przestojów oraz zaplanowanie działań minimalizujących ich wpływ na cały proces produkcyjny).

Wyróżnić należy dwie podstawowe zasady harmonogramowania. Pierwsza dotyczy **harmonogramowania w przód** (rys. 4). Oznacza to, iż zakładając termin rozpoczęcia realizacji zlecenia, analizuje się jego przebieg, otrzymując najwcześniejsze możliwe terminy rozpoczęcia poszczególnych zadań oraz datę zakończenia całego zlecenia.

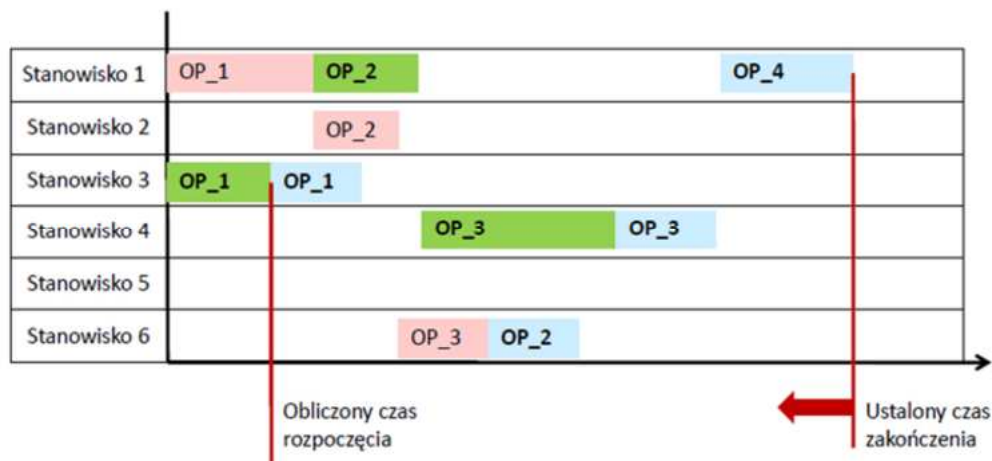


Rys. 4 Harmonogramowanie w przód
Źródło: materiały PP

Kolejnym rodzajem harmonogramowania jest **harmonogramowanie wstecz** (rys. 5). Oznacza to, iż zakładając termin zakończenia realizacji zlecenia, analizuje się jego przebieg (od końca),



otrzymując najpóźniejsze dopuszczalne terminy rozpoczęcia poszczególnych zadań oraz datę rozpoczęcia całego zlecenia, która umożliwi jego zakończenie w założonym terminie.



Rys. 5 Harmonogramowanie wstecz
Źródło: materiały PP

Inną metodą harmonogramowania produkcji jest **reguła priorytetu**. Jest ona jedną z kluczowych metod ustalania kolejności realizacji zadań produkcyjnych. Celem tej reguły jest określenie, które zadania powinny być wykonywane w pierwszej kolejności, aby osiągnąć optymalną efektywność procesu produkcyjnego. Reguły priorytetu pomagają zorganizować pracę, zminimalizować czas oczekiwania, zredukować koszty i poprawić przepustowość produkcji. Reguły priorytetu mogą być stosowane wszędzie tam, gdzie tworzą się kolejki przed stanowiskami produkcyjnymi i należy zdecydować, który z elementów poddać produkcji w pierwszej kolejności. Najczęściej stosowane **reguły priorytetu**:

- 1) najkrótszy czas realizacji (Shortest Processing Time, SPT) - zadania o najkrótszym czasie realizacji są wykonywane w pierwszej kolejności,
- 2) najdłuższy czas realizacji (Longest Processing Time, LPT) - zadania o najdłuższym czasie realizacji mają wyższy priorytet i są realizowane w pierwszej kolejności,
- 3) pierwsze przyszło, pierwsze obsłużone (First Come, First Served, FCFS) - zadania są realizowane w kolejności, w jakiej zostały zgłoszone do produkcji
- 4) najmniejsza liczba operacji pozostałych (Least Work Remaining, LWR) - zadania, które wymagają najmniejszej liczby operacji do wykonania, mają wyższy priorytet,
- 5) priorytet na podstawie terminu realizacji (Earliest Due Date, EDD) - zadania o najwcześniejszym terminie realizacji mają wyższy priorytet.



- 6) priorytet na podstawie wag (Weighted Shortest Processing Time, WSPT) - zastosowanie wag dla każdego zadania i realizacja zleceń, które mają najwyższy stosunek wagi do czasu realizacji.

Przykład 2

Należy zaplanować produkcję trzech zleceń w zakładzie produkcyjnym, przy wykorzystaniu dwóch maszyn. Zlecenia różnią się czasem realizacji oraz terminami dostawy. Twoim zadaniem jest stworzenie harmonogramu produkcji, który uwzględni zasadę Earliest Due Date (EDD), czyli realizację zleceń w kolejności ich terminów dostawy.

Dane:

1. Zlecenie 1 (Z1): Czas produkcji 4 godziny, termin dostawy za 2 dni.
2. Zlecenie 2 (Z2): Czas produkcji 3 godziny, termin dostawy za 1 dzień.
3. Zlecenie 3 (Z3): Czas produkcji 5 godzin, termin dostawy za 3 dni.

Zasoby:

- Maszyna 1: Może produkować wszystkie komponenty.
- Maszyna 2: Może produkować tylko komponenty Z2.

Polecenie:

1. Priorytet: Zastosuj regułę Earliest Due Date (EDD), czyli zlecenia mają być realizowane w kolejności ich terminów dostawy.
2. Harmonogram: Przygotuj harmonogram produkcji, uwzględniając dostępność maszyn i czas produkcji dla każdego zlecenia. Przypisz każde zlecenie do odpowiedniej maszyny i oblicz czas rozpoczęcia oraz zakończenia produkcji.

Po zaplanowaniu, harmonogram produkcji będzie wyglądał następująco:

Tabela 5. Harmonogram produkcji

Zlecenie	Czas produkcji	Termin dostawy	Maszyna	Czas rozpoczęcia	Czas zakończenia
Z2	3 godziny	Dzień 1	Maszyna 2	00:00	03:00
Z1	4 godziny	Dzień 2	Maszyna 1	03:00	07:00
Z3	5 godzin	Dzień 3	Maszyna 1	07:00	12:00

Źródło: opracowanie własne



Wnioski:

- Zlecenie Z2 zostało zrealizowane pierwsze, ponieważ ma najbliższy termin dostawy.
- Zlecenie Z1 następnie, a Zlecenie Z3 jako ostatnie, zgodnie z ich terminami dostawy.

Koniec przykładu 2

Zarządzanie zapasami

Zarządzanie zapasami to proces planowania, kontrolowania i monitorowania zasobów materialnych w firmie w celu zapewnienia ich dostępności w odpowiednich ilościach, w odpowiednim czasie, przy minimalizacji kosztów. Właściwe zarządzanie zapasami jest kluczowe dla utrzymania płynności produkcji i sprzedaży, jednocześnie ograniczając niepotrzebne koszty związane z przechowywaniem nadmiarowych zapasów.

Celem zarządzania zapasami są zatem:

- a) zabezpieczenie ciągłości produkcji i sprzedaży,
- b) minimalizacja kosztów przechowania,
- c) ograniczenie ryzyka braków,
- d) optymalizacja zamówień.

Jedną z metod zarządzania zapasami jest **ekonomiczna wielkość zamówienia** określana mianem formuły Wilsona. Jest ona jedną z kluczowych formuł matematycznych, określających, jaką wielkość zapasów powinno się zamawiać, aby zminimalizować łączne koszty składania zamówień oraz koszty utrzymywania zapasów. Warunki, jakie trzeba spełnić, aby stosowanie ww. mechanizmu było zasadne, to przede wszystkim to, że zużycie zapasu winno być równomierne w czasie, natomiast łączne zapotrzebowanie znane. Pozyskiwanie surowca ma miejsce w określonych interwałach czasowych (Szołtysek, 2009). Aby obliczyć ekonomiczną wielkość zamówienia należy skorzystać z następującego wzoru:

$$EOQ = Q = \sqrt{\frac{2K_D \cdot D}{K_S \cdot C}}$$

gdzie:

K_D – koszt pojedynczego zamówienia,

K_S – jednostkowy koszt składowania,

D – roczne zapotrzebowanie,

C – cena jednostkowa



Przykład 3

Ustal ekonomiczną wielkość zamówienia EWZ paneli dotykowych do sterowania odbiornikami do przedsiębiorstwa na podstawie formuły Wilsona. Miesięczna prognoza sprzedaży urządzeń wynosi 120 szt. Koszt realizacji dostawy kształtuje się na poziomie 50 zł, a koszt miesięcznego magazynowania jednego urządzenia to 20,00 zł.

Rozwiązanie

$$EWD = \sqrt{\frac{2 \cdot 120 \text{ szt.} \cdot 50 \text{ zł}}{20 \text{ zł/mc} \cdot 1 \text{ mc}}} = 24,49 \sim 24 \text{ szt.}$$

Ekonomiczna wielkość dostaw dla analizowanego przypadku wynosi 24 sztuki.

Koniec przykładu 3

Innymi metodami zarządzania zapasami są **metoda minimalnego poziomu zamówienia** (w tej metodzie zapasy są uzupełniane, gdy osiągną określony poziom, zwany "punktem zamówienia") oraz **metoda min-max** (metoda ta zakłada, że ustala się minimalny i maksymalny poziom zapasów dla każdego indeksu magazynowego, natomiast gdy zapasy spadną poniżej poziomu minimalnego, składane jest zamówienie na uzupełnienie zapasów do poziomu maksymalnego).

Przykład 4

Firma produkująca akcesoria komputerowe musi zarządzać zapasami gotowych produktów w magazynie. Celem jest ustalenie minimalnego i maksymalnego poziomu zapasów dla najpopularniejszego produktu – klawiatury komputerowej. Firma chce utrzymać odpowiedni poziom zapasów, aby zaspokoić popyt, ale jednocześnie uniknąć nadmiaru zapasów, który wiąże się z wysokimi kosztami magazynowania.

Dane:

1. Średnia dzienna sprzedaż: 100 sztuk klawiatury komputerowej.
2. Czas realizacji zamówienia (lead time): 5 dni.
3. Zapas bezpieczeństwa: 2 dni (poziom zapasu, który zapewnia firmie ochronę przed nieoczekiwanym wzrostem popytu lub opóźnieniem dostaw).



4. Minimalny poziom zapasów: określa się na podstawie średniej dziennej sprzedaży oraz zapasu bezpieczeństwa.
5. Maksymalny poziom zapasów: określa się na podstawie minimalnego poziomu zapasów oraz ustalonego zapasu na pokrycie ewentualnych wzrostów popytu.

Zadanie:

1. Oblicz minimalny poziom zapasów klawiatury komputerowej, który powinien być utrzymywany w magazynie, aby firma mogła sprostać zwykłemu popytowi i zapobiec brakowi zapasów.
2. Oblicz maksymalny poziom zapasów, biorąc pod uwagę, że firma chce utrzymać dodatkowy zapas na 5 dni sprzedaży.
3. Określ, kiedy należy złożyć nowe zamówienie, aby nie przekroczyć maksymalnego poziomu zapasów.

Rozwiązanie

1. Minimalny poziom zapasów: Minimalny poziom zapasów obliczamy na podstawie średniego dziennego zapotrzebowania oraz zapasu bezpieczeństwa

$$\text{Min} = (\text{średnia dzienna sprzedaż} \times \text{czas realizacji zamówienia}) + (\text{średnia dzienna sprzedaż} \times \text{zapas bezpieczeństwa})$$

$$\text{Min} = (100 \times 5) + (100 \times 2) = 700 \text{ sztuk}$$

2. Maksymalny poziom zapasów: Maksymalny poziom zapasów ustalamy jako sumę minimalnego poziomu zapasów oraz zapasu na pokrycie dodatkowych 5 dni sprzedaży:

$$\text{Max} = \text{Min} + (\text{średnia dzienna sprzedaż} \times 5)$$

$$\text{Max} = 700 + (100 \times 5) = 700 + 500 = 1200 \text{ sztuk}$$

3. Złożenie zamówienia: Nowe zamówienie należy złożyć, gdy zapasy spadną poniżej poziomu minimalnego (700 sztuk), aby zapewnić odpowiedni zapas do momentu dostawy.

Podsumowanie:

- minimalny poziom zapasów: 700 sztuk.
- maksymalny poziom zapasów: 1200 sztuk.
- złożenie zamówienia: gdy zapasy spadną poniżej 700 sztuk.

Dzięki metodzie min-max firma może kontrolować swoje zapasy i zapewnić ciągłość produkcji oraz sprzedaży, jednocześnie unikając wysokich kosztów związanych z nadmiernym przechowywaniem zapasów.

Koniec przykładu 4



2. Sterowanie produkcją

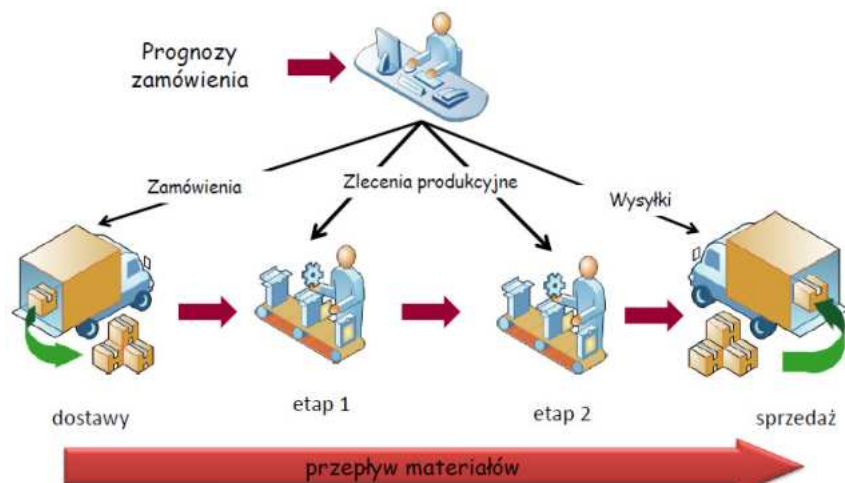
Sterowanie produkcją to proces zarządzania i nadzorowania wszystkich działań związanych z produkcją, mający na celu osiągnięcie maksymalnej efektywności i optymalnego wykorzystania zasobów. Obejmuje planowanie, monitorowanie oraz kontrolowanie procesów produkcyjnych w firmie, aby zapewnić, że produkty są wytwarzane na czas, zgodnie z wymaganiami jakościowymi i przy minimalizacji kosztów. Sterowanie produkcją jest nieodłącznym elementem zarządzania produkcją, a jego celem jest synchronizacja wszystkich czynników produkcyjnych, takich jak maszyny, ludzie, materiały oraz czas, w celu realizacji założonych planów produkcyjnych (Charytonow, Nowakowski, 2014). Sterowanie przepływem produkcji to funkcja kierowania i regulacji przepływu materiałów obejmująca cały cykl wytwarzania począwszy od wydania materiałów do produkcji, aż do przygotowania dostawy wyrobów finalnych. Sterowanie produkcją ma za zadanie realizację kilku podstawowych celów:

- 1) zapewnienie ciągłości produkcji (działania produkcyjne muszą odbywać się bez zakłóceń, a w przypadku wystąpienia awarii należy natychmiast podjąć odpowiednie kroki w celu ich eliminacji),
- 2) optymalne wykorzystanie zasobów (maszyny, urządzenia, surowce i pracownicy muszą być wykorzystywani w sposób efektywny, minimalizując przestoje i niepotrzebne koszty),
- 3) minimalizacja kosztów (w procesie sterowania produkcją szczególną uwagę zwraca się na minimalizację kosztów związanych z produkcją, takich jak koszty pracy, surowców, energii, transportu i magazynowania),
- 4) zgodność z harmonogramem (sterowanie produkcją ma zapewnić, aby produkcja była realizowana zgodnie z wcześniej zaplanowanym harmonogramem, a produkty były dostarczane na czas).
- 5) kontrola jakości (zapewnienie wysokiej jakości produktów jest kluczowe, w związku z czym proces produkcji musi być monitorowany i kontrolowany w taki sposób, aby spełniał wymagania jakościowe).

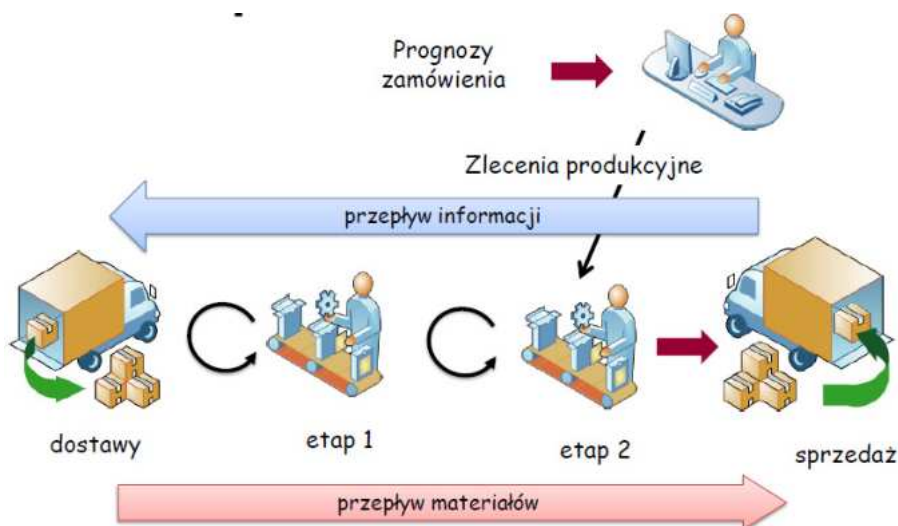
Metody sterowania przepływem produkcji

Wśród metod sterowania produkcją należy wyróżnić sterowanie produkcją w systemie push i pull. W systemie **push** (rys. 6) produkcja jest planowana na podstawie przewidywań

dotyczących popytu. Produkcja jest uruchamiana z wyprzedzeniem, zanim pojawi się zapotrzebowanie na konkretne produkty. W przypadku systemu **pull** (rys. 7) produkcja jest uruchamiana w odpowiedzi na konkretne zapotrzebowanie, tj. produkcja odbywa się na podstawie zleceń klientów. Typowym przykładem takiego systemu jest metoda **Just In Time (JIT)**.



Rys. 6. Metoda sterowania przepływem produkcji – metoda push
Źródło: materiały PP



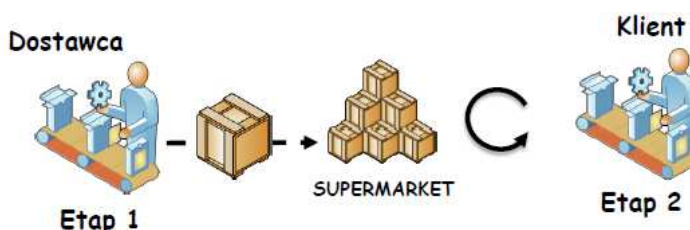
Rys. 7. Metoda sterowania przepływem produkcji – metoda pull
Źródło: materiały PP

W tradycyjnym „pchaniu” dział sterowania produkcją generuje przepływ informacji poprzez rozesłanie planu produkcyjnego do każdego działu, aby ten produkował i przemieszczał



materiał. W systemie ssącym proces ten jest odwrócony. Każdy ruch materiału generuje sygnał informacyjny (kanban). Kiedy części są „zasysane”, sygnał by produkować więcej jest wysyłany do etapu dostarczającego.

System ssący to metoda sterowania przepływem produkcji, który jest regulowany rzeczywistym zużyciem materiałów oraz zsynchronizowanymi sygnałami o potrzebie uzupełnienia materiału.



Rys. 8. System „ssący”
Źródło: materiały PP

Reasumując system „push” polega na wytwarzaniu produktów a następnie ich sprzedaży, czyli „wypychaniu” do następnego procesu. W systemach „push” gromadzi się duże zapasy, niekiedy zbędne. Produkcja opiera się na prognozowaniu, a nie na realnym zapotrzebowaniu. W systemie „pull” ma miejsce uzupełnianie materiału czy produktu tylko w sytuacji, gdy zasoby zostaną wykorzystane. Zestawienie różnic pomiędzy systemami sterowania produkcją przedstawiono w tabeli 1.

Tabela 4. Różnice w systemach sterowania produkcją

Push	Pull
– wysoki poziom zapasów, – duże zapasy produkcji w toku, – zarządzanie w stylu „gaszenia pożarów” – duże partie produkcyjne, – większe problemy jakościowe, – estymowanie wielkości produkcji.	– niski poziom zapasów, – niskie zapasy produkcji w toku, – zarządzanie przez monitoring, – małe partie produkcyjne, – zredukowane problemy jakościowe, – sprecyzowana wielkość produkcji.

Źródło: opracowanie własne

Wskazać należy również na przykłady stosowania ww. systemów w wybranych przedsiębiorstwach produkcyjnych (tab. 2).



Tabela 5. System push i pull w wybranych przedsiębiorstwach

Push	Pull
Procter & Gamble (P&G) P&G, gigant branży dóbr konsumpcyjnych, używa systemu Push do produkcji swoich popularnych produktów, takich jak środki czystości czy pieluchy. Firma przewiduje popyt na te produkty na podstawie analiz rynkowych, danych o sprzedaży i trendów. Produkcja jest planowana z wyprzedzeniem, a wyprodukowane towary trafiają do magazynów i sklepów, gdzie są dostępne dla klientów. Dzięki temu P&G może korzystać z efektu skali i zapewniać ciągłość dostępności produktów w sklepach.	Dell Computers Dell jest jednym z najgłośniejszych przykładów firmy wykorzystującej system Pull w produkcji komputerów. W tym przypadku produkcja komputerów rozpoczyna się dopiero po złożeniu zamówienia przez klienta. Klient może wybierać specyfikacje swojego komputera, takie jak procesor, ilość pamięci RAM czy typ dysku twardego. Dopiero wtedy, gdy zamówienie zostanie przyjęte, Dell uruchamia produkcję, co pozwala na dostosowanie produktów do indywidualnych potrzeb klienta. Dzięki temu Dell eliminuje koszty związane z przechowywaniem gotowych komputerów, produkując tylko te modele, które są faktycznie zamówione.
Coca-Cola Coca-Cola wykorzystuje system Push w swoim procesie produkcji napojów. Firma przewiduje zapotrzebowanie na różne warianty napojów na podstawie danych sprzedaży i sezonowości, a następnie produkuje je z wyprzedzeniem. Produkcja odbywa się w dużych ilościach i towary są dystrybuowane do sklepów, gdzie czekają na konsumentów.	Toyota – System "Just-in-Time" (JIT) Toyota jest pionierem w stosowaniu systemu Pull, zwłaszcza poprzez swój system Just-in-Time (JIT). W tym systemie produkcja samochodów zaczyna się dopiero, gdy zapotrzebowanie na dany model jest potwierdzone. Na przykład, części do samochodów są dostarczane na linię montażową dokładnie w momencie, gdy są potrzebne, eliminując konieczność przechowywania zapasów. Produkcja w Toyocie jest ściśle skoordynowana z rynkiem, a każdy element w procesie produkcji jest "ciągnięty" przez zapotrzebowanie, co pozwala zminimalizować zapasy i czas oczekiwania.
Nike Nike używa systemu Push w swojej produkcji obuwia i odzieży. Firma przewiduje zapotrzebowanie na swoje produkty na podstawie analiz rynkowych, prognoz sezonowych i danych o sprzedaży, a następnie produkuje je z wyprzedzeniem. Produkty są następnie dystrybuowane do magazynów i sklepów, gdzie są gotowe do sprzedaży. Dzięki systemowi Push Nike jest w stanie utrzymać stałą dostępność swoich produktów w sklepach na całym świecie.	BMW (produkcja samochodów) BMW stosuje system Pull w produkcji swoich samochodów. W przypadku pojazdów z wyższej półki, klienci mają możliwość wyboru wielu opcji konfiguracji, takich jak kolor nadwozia, rodzaj tapicerki czy dodatkowe akcesoria. Produkcja samochodów rozpoczyna się dopiero po złożeniu zamówienia przez klienta. Dzięki temu firma unika nadprodukcji i może dostarczać auta dostosowane do indywidualnych potrzeb klienta.

Źródło: opracowanie własne

Warto również wskazać na inne sposoby sterowania przepływem w przedsiębiorstwie. Jednym z takich sposobów jest koncepcja **Just in Time**, której głównym celem jest dostarczenie materiałów i innych dóbr w ściśle określonych ilościach oraz dokładnie w takim czasie, w jakim są one potrzebne w przedsiębiorstwie. Pozwala to na minimalizację kosztów



zapasów i marnotrawstwa w systemie logistycznym. Koncepcja ta opiera się na czterech założeniach:

- a) zero zapasów,
- b) krótkie serie produkcyjne,
- c) minimalizacja kolejek,
- d) krótkie cykle realizacji zamówienia,
- e) wysoka jakość.

Koncepcja Just in Time ma zatem na celu eliminację zbędnych zapasów oraz minimalizację czasu oczekiwania kolejnych linii produkcyjnych.

Nadrzędnym celem koncepcji Just in Time (JiT) jest dostarczanie właściwych materiałów, części, półproduktów czy produktów gotowych zgodnie z zasadą 5W, tj. we właściwym czasie, w miejsce przeznaczenia, w odpowiedniej jakości i ilości oraz do odpowiedniego odbiorcy. Zazwyczaj dostawa materiałów niezbędnych do realizacji procesu wytwórczego odbywa się częściej w niewielkich partiach, co wymaga właściwego planowania i sterowania materiałami. Nadrzędnym elementem tej koncepcji jest redukcja zapasów, co przedkłada się na redukcję kosztów działalności przy zachowaniu wysokiej jakości. Just in Time, szczególnie w kontekście logistyki produkcji, niezależnie od wielkości przedsiębiorstwa, charakteryzują następujące determinanty:

- produkt powinien być zaprojektowany pod kątem modularności, z dużą łatwością jego wytwarzania i eliminowania zbędnej złożoności,
- koniecznym jest zastosowanie form potokowych w organizacji produkcji, co sprawia, iż system produkcyjny funkcjonuje płynnie, zapewniając synchronizację procesów produkcyjnych,
- należy dążyć do eliminowania źródeł marnotrawstw, które powstają w procesie produkcyjnym,
- koniecznym jest wyposażenie systemu produkcyjnego w możliwie uniwersalne wyposażenie,
- należy zapewnić wysoką jakość i terminowość dostaw, co jest możliwe dzięki wdrożeniu kompleksowego sterowania jakością, w ramach którego pracownicy produkcyjni odpowiadają za jakość wyrobu,



- wszystkie elementy organizacji powinny być dostosowane do ciągłych usprawnień, a system informatyczny na tyle silnie rozwinięty, aby można było w czasie rzeczywistym podejmować właściwe decyzje (Kosieradzka, Lis, 1996).

Podstawowym zadaniem systemu Just in Time jest produkcja w **czasie taktu**⁷. Czas taktu to czas, który upływa pomiędzy zakończeniem produkcji jednej jednostki a rozpoczęciem produkcji kolejnej jednostki, a jego celem jest utrzymanie tempa produkcji zgodnego z wymaganym popytem.

Czas taktu oblicza się, dzieląc dostępny czas pracy przez zapotrzebowanie na produkty. Wzór na obliczenie czasu taktu to:

$$\text{Czas taktu} = \frac{\text{Dostępny czas produkcji}}{\text{Zapotrzebowanie na produkcję (np. ilość jednostek)}}$$

Przykład 5

Firma XYZ zajmuje się produkcją mebli na zamówienie. W ciągu dnia roboczego (8 godzin) firma jest w stanie wyprodukować różne modele mebli, jednak ze względu na zróżnicowane potrzeby klientów, zapotrzebowanie na poszczególne modele jest inne. Firma ma także możliwość wydłużenia swojej pracy do 10 godzin dziennie, jeśli zapotrzebowanie na produkty rośnie w szczytowym okresie.

Obecnie firma planuje produkcję dwóch modeli mebli:

- Model A: Zapotrzebowanie wynosi 120 sztuk dziennie.
- Model B: Zapotrzebowanie wynosi 80 sztuk dziennie.

Zakład dysponuje 8 godzinami roboczymi dziennie, ale istnieje możliwość, że w przyszłości czas pracy zostanie wydłużony do 10 godzin dziennie, aby sprostać wzrastającemu zapotrzebowaniu.

Zadanie

1. Oblicz czas taktu dla obu modeli mebli, zakładając, że firma pracuje 8 godzin dziennie.
2. Oblicz czas taktu dla obu modeli mebli, zakładając, że firma wydłuży czas pracy do 10 godzin dziennie.

⁷ Czas taktu (ang. takt time) to pojęcie stosowane w zarządzaniu produkcją, które oznacza czas, jaki upływa między rozpoczęciem produkcji kolejnej jednostki produktu, aby sprostać wymaganiom popytu na produkty. Jest to podstawowy wskaźnik, który pozwala określić tempo pracy produkcji w celu dostosowania jej do potrzeb rynku i zapotrzebowania klientów.



3. Porównaj obie sytuacje i wyjaśnij, jak wydłużenie czasu pracy wpływa na tempo produkcji oraz efektywność.
4. Załóż, że zakład może produkować jeden mebel w ciągu 15 minut (średni czas produkcji). Określ, ile sztuk obu modeli firma może wyprodukować w ciągu dnia w sytuacji, gdy czas pracy jest wydłużony do 10 godzin.

Rozwiązanie

1. Czas taktu przy 8 godzinach pracy dziennie:

- Dostępny czas pracy = 8 godzin = 480 minut.
- Zapotrzebowanie na Model A = 120 sztuk dziennie.
- Zapotrzebowanie na Model B = 80 sztuk dziennie.

Należy obliczyć czas taktu dla każdego modelu:

- Model A:

Czas taktu dla Modelu A = $480 \text{ minut} / 120 \text{ sztuk} = 4 \text{ minuty}$ na jednostkę

- Model B:

Czas taktu dla Modelu B = $480 \text{ minut} / 80 \text{ sztuk} = 6 \text{ minut}$ na jednostkę

2. Czas taktu przy 10 godzinach pracy dziennie:

- Dostępny czas pracy = 10 godzin = 600 minut.
- Zapotrzebowanie na Model A = 120 sztuk dziennie.
- Zapotrzebowanie na Model B = 80 sztuk dziennie.

Należy obliczyć czas taktu dla każdego modelu:

- Model A:

Czas taktu dla Modelu A = $600 \text{ minut} / 120 \text{ sztuk} = 5 \text{ minut}$ na jednostkę

- Model B:

Czas taktu dla Modelu A = $600 \text{ minut} / 80 \text{ sztuk} = 7,5 \text{ minuty}$ na jednostkę

3. Porównanie obu sytuacji:

- Przy 8 godzinach pracy:
 - Model A: Produkcja jednej sztuki zajmuje 4 minuty.
 - Model B: Produkcja jednej sztuki zajmuje 6 minut.
- Przy 10 godzinach pracy:
 - Model A: Produkcja jednej sztuki zajmuje 5 minut.
 - Model B: Produkcja jednej sztuki zajmuje 7,5 minuty.



Wnioski: Wydłużenie czasu pracy o 2 h powoduje wydłużenie czasu taktu dla obu modeli mebli. Produkcja staje się mniej intensywna, ponieważ firma ma więcej czasu do realizacji zamówień, co pozwala na lepsze dostosowanie się do rosnącego zapotrzebowania. Wydłużenie godzin pracy może jednak pozwolić na większą elastyczność w produkcji i lepsze zarządzanie obciążeniem pracowników.

4. Ilość produkcji przy 10 godzinach pracy:

- Średni czas produkcji na jednostkę = 15 minut.

Należy obliczyć, ile sztuk można wyprodukować w ciągu dnia przy 10 godzinach pracy (600 min):

- Model A:

Liczba sztuk dla Modelu A = $600 \text{ minut} / 15 \text{ minut na sztukę} = 40 \text{ sztuk}$

- Model B:

Liczba sztuk dla Modelu B = $600 \text{ minut} / 15 \text{ minut na sztukę} = 40 \text{ sztuk}$

Firma może wyprodukować 40 sztuk każdego modelu w ciągu dnia roboczego, jeśli wydłuży czas pracy do 10 godzin.

Koniec przykładu 5

Czas taktu w produkcji ma kluczowe zadania, do których zaliczyć należy:

- równoważenie obciążenia produkcyjnego,
- optymalizacja procesów,
- planowanie zasobów,
- zarządzanie efektywnością.

Czas taktu to kluczowy wskaźnik, który pomaga w synchronizacji produkcji z popytem rynkowym. Umożliwia on ustalenie optymalnego tempa produkcji, które zapewnia, że produkcja odbywa się zgodnie z zapotrzebowaniem, minimalizując nadprodukcję oraz braki w dostępnych produktach. Analizując czas taktu warto wspomnieć o czasie cyklu. Czas cyklu to czas, który upływa od momentu rozpoczęcia produkcji jednostki do momentu jej zakończenia. Czas taktu odnosi się natomiast do tempa produkcji w kontekście zapotrzebowania rynkowego, a nie samego czasu wytworzenia jednostki. Czas taktu i czas cyklu powinny być ze sobą zrównoważone. Jeśli czas cyklu jest dłuższy niż czas taktu, oznacza to, że produkcja nie nadąża za zapotrzebowaniem i mogą wystąpić opóźnienia. Jeśli czas cyklu jest krótszy, może dochodzić do nadprodukcji.



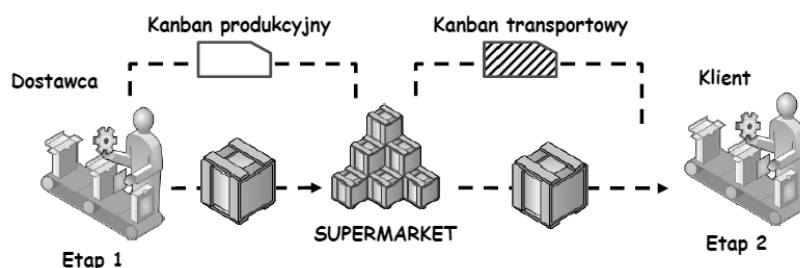
Kolejnym narzędziem umożliwiającym sterowanie przepływem materiałów jest system **kanban**. Kanban w szczupłym wytwarzaniu jest specjalnym narzędziem sterowania informacją i regulowania przemieszczeń materiałów pomiędzy procesami produkcyjnymi (termin ten w języku japońskim oznacza „znak” lub „szyld”). Kanban sprzężony z czasem taktu, procesami o przepływie ciągłym, produkcją w systemie ssania i poziomowanym harmonogramowaniem umożliwia realizację w strumieniu wartości zasady Just in Time. Zwykle kanban jest wykorzystywany do sygnalizowania, kiedy produkt zostaje zużyty (skonsumowany) przez proces w dole strumienia. W najprostszym przypadku, zdarzenie to generuje sygnał o potrzebie uzupełnienia produktu przez proces w górze strumienia (www.lean.org.pl; 18.03.2025).

Kanban to skuteczny system zarządzania produkcją, który pomaga utrzymywać procesy produkcyjne zgodne z rzeczywistym zapotrzebowaniem, minimalizując marnotrawstwo, nadprodukcję i zapasy. Jego główne zasady to wykorzystanie kart kanban, zasada "pull" oraz kontrolowanie liczby jednostek roboczych w procesie produkcji. Wymaga to zaawansowanego planowania i dobrej organizacji, ale w zamian oferuje wiele korzyści w postaci redukcji kosztów, zwiększenia elastyczności i poprawy jakości.

Niezbędnym elementem odpowiednie przepływu materiałów w procesach produkcyjnych w organizacji są karty kanban. Wyróżnia się następujące rodzaje kart kanban:

- **kanban (karta) produkcyjny**, który wykorzystywany jest do zlecenia wykonania pewnej liczby przedmiotów,
- **kanban (karta) transportowy**, który stanowi dokument pobrania z poprzedniego stanowiska produktu i umożliwia przemieszczanie się pojemnika z jednej linii produkcyjnej na drugą (Skowronek, Sarjusz-Wolski, 2008).

Kanban produkcyjny reguluje zatem produkcję części, transportowy z kolei nakazuje dokonać dostawy części (rys. 8).



Rys. 9. Kanban transportowy i produkcyjny
Źródło: materiały PP



Warto dodać, iż każda karta kanban (rys. 9) zawiera określone informacje. Są one niezbędne do zarządzania zapasami i procesami produkcyjnymi. Najczęściej na karcie kanban znajduje się numer referencyjny (kod lub numer materiału czy produktu), ilość, opis produktu, miejsce dostawy, data, poziom zapasu, kod dostawcy oraz podpis osoby odpowiedzialnej (Liker, 2004).

Kanban Produkcyjny				
Poprzedni proces Obróbka skrawaniem		Bieżący proces Powlekanie		
Nazwa części Oprawki		Nr części VDI DIN 68880		
Wielkość partii 20		Nr karty 1/10		
Parametry				
d	20mm	D	50mm	l
				70mm

Rys. 10. Przykładowa karta kanban
Źródło: leanactionplan.pl (19.03.2025)

Karta przyczepiana jest do pojemnika, z którego pracownik produkcyjny pobiera materiał do produkcji. W momencie, gdy pojemnik jest pusty, karta odkładana jest do zbiorczej skrzynki. Pracownik magazynu zbiera wszystkie karty w ściśle określonych odstępach czasu i przy kolejnej dostawie uzupełnia części na linii produkcyjnej. System ten jest wspierany infrastrukturą IT.

Głównym celem i zadaniem organizacji produkcji według systemu kanban jest minimalizacja kosztów poprzez eliminację wszelkich strat. Często eliminację strat określa się mianem „siedem razy zero”, a więc realizacja hasła: zero braków, zero opóźnień, zero zapasów, zero kolejek, zero bezczynności, zero zbędnych operacji technologicznych i kontrolnych oraz zero zbędnych przemieszczeń (Borkowski, Ulewicz, 2008).

System kanban działa właściwie, jeżeli spełnione są następujące warunki:

- do każdego pojemnika przypisana jest wyłącznie jedna karta,
- klient wewnętrzny inicjuje przepływ materiału ze stanowiska poprzedzającego,
- ilość i rodzaj dostarczonego materiału są identyczne z wytycznymi na karcie,
- żadna karta nie może opuścić obiegu,
- w przepływie kart kanban musi funkcjonować zasada FIFO (Szymonik, 2012).

Ta koncepcja ma wiele korzyści, do których zaliczyć trzeba uporządkowany przepływ materiałów w całym systemie produkcyjnym oraz ścisłą kontrolę zapasów. Kanban znajduje



zastosowanie, gdy produkcja jest powtarzalna, a plany produkcyjne są stabilne. Bardzo ciężko jest wdrożyć system, gdy zamówienia są nieprzewidywalne co do wielkości i różnorodności partii.

Aby system kanban funkcjonował efektywnie, konieczne jest spełnienie szeregu warunków organizacyjnych i technologicznych. Przede wszystkim należy stworzyć odpowiednią strukturę procesów produkcyjnych, zapewnić efektywną komunikację między działami, zdefiniować odpowiednie poziomy zapasów oraz zaangażować pracowników w działanie systemu. Tylko wówczas kanban będzie w stanie zrealizować swoje cele: optymalizację zapasów, redukcję marnotrawstwa i poprawę efektywności procesów produkcyjnych. Aby system kanban był efektywny, procesy produkcyjne i przepływ materiałów muszą być szczegółowo zaplanowane i standaryzowane. Standaryzacja procesów produkcyjnych pozwala na precyzyjne określenie, kiedy i jakie materiały powinny być dostarczane do produkcji, a także kiedy należy je zamawiać od dostawców. W takim systemie konieczne jest zrozumienie i określenie wszystkich etapów produkcji, co umożliwia sprawne funkcjonowanie karty kanban jako sygnału do uruchamiania kolejnych procesów produkcyjnych (Amano, 2015). System kanban funkcjonuje na zasadzie tzw. „ciągnięcia” (pull system), gdzie produkcja i zamówienia materiałów są uruchamiane dopiero wtedy, gdy zapotrzebowanie na materiał wystąpi w kolejnym etapie produkcji. Kluczowym elementem jest określenie minimalnego poziomu zapasów (tzw. poziom kanban), poniżej którego powinno zostać uruchomione zamówienie nowych materiałów. Ustalenie odpowiedniego poziomu zapasów zapewnia płynność procesów produkcyjnych i zapobiega zarówno nadprodukcji, jak i niedoborom materiałów (Monden, 2012). Optymalizacja liczby kart kanban jest niezbędna do zapewnienia efektywności systemu. Zbyt mała liczba kart kanban może powodować opóźnienia w produkcji z powodu braku materiałów, podczas gdy nadmiar kart może prowadzić do nadmiernych zapasów i niepotrzebnych kosztów magazynowania. Właściwa liczba kart powinna być określona na podstawie analizy zapotrzebowania i czasu realizacji zamówień w danym systemie produkcyjnym (Liker, 2004). System kanban powinien być zintegrowany z systemem zarządzania jakością, aby zapewnić, że materiały i produkty spełniają wymagania jakościowe na każdym etapie produkcji. Kontrola jakości powinna obejmować nie tylko kontrolowanie produktów gotowych, ale także monitorowanie jakości procesów produkcyjnych, co pozwala na wczesne wykrywanie problemów i minimalizowanie ryzyka błędów w produkcji. W efekcie system kanban działa nie tylko na rzecz optymalizacji zapasów,



ale także na rzecz poprawy jakości procesów produkcyjnych (Szczepański, 2015). Współczesne systemy kanban mogą być wspierane przez systemy informatyczne, takie jak ERP (Enterprise Resource Planning) i WMS (Warehouse Management Systems), które automatyzują procesy związane z zarządzaniem zapasami, monitorowaniem kart kanban oraz analizą danych. Technologie te pozwalają na szybszą i bardziej precyzyjną wymianę informacji, co zwiększa efektywność i elastyczność systemu (Szczepański, 2015).

Zakończenie

Planowanie i sterowanie produkcją są kluczowymi elementami zarządzania przedsiębiorstwem, mającymi na celu zapewnienie płynności procesów produkcyjnych, optymalizację wykorzystania zasobów oraz dostosowanie produkcji do zmieniającego się popytu. Skuteczne planowanie produkcji wymaga precyzyjnego prognozowania popytu, alokacji zasobów, oraz zapewnienia ciągłości w dostawach surowców i komponentów. Z kolei sterowanie produkcją polega na monitorowaniu postępu produkcji i podejmowaniu odpowiednich działań korygujących, aby zachować zgodność z planami i zoptymalizować efektywność. Harmonogramowanie produkcji to kluczowy proces w zarządzaniu produkcją, który pozwala na efektywne planowanie i realizację zadań produkcyjnych. Dzięki odpowiedniemu harmonogramowi możliwe jest zoptymalizowanie czasu, zasobów i kosztów, co ma bezpośredni wpływ na jakość produktów oraz rentowność przedsiębiorstwa.

Bibliografia

- Charytonow, J., Nowakowski, M. (2014). *Zarządzanie produkcją i jakością*. Warszawa: Wydawnictwo Naukowe PWN.
- Filipiak, J. (2003). *Zarządzanie produkcją*. Warszawa: Wydawnictwo Naukowe PWN.
- Vollmann, T.E., Berry, W.L., Whybark, D.C. (2012). *Manufacturing Planning and Control for Supply Chain Management*.
- Heizer, J., Render, B. (2014). *Operations Management* (wyd. 11). Pearson Education.
- Jacobs, F. R., Chase, R. B. (2011). *Operations and Supply Chain Management* (13th ed.). McGraw-Hill Education.
- Karpus, G. (2017). *Zbiór zadań z logistyki*. Warszawa: WSiP.
- Kosieradzka, A., Lis, S. (1996). *Produktywność. Metody analizy oceny i tworzenia programów poprawy*. Warszawa: Oficyna Wydawnicza PW.



- Liker, J. K. (2004). *The Toyota Way: 14 Management Principles from the World's Greatest Manufacturer*. New York, NY: McGraw-Hill.
- Monden, Y. (2012). *The Toyota Production System: An Integrated Approach to Just-In-Time* (3rd ed.). Boca Raton: CRC Press.
- S. Borkowski, S., R. Ulewicz, R. (2008). *Zarządzanie produkcją. Systemy produkcyjne*. Sosnowiec: Oficyna Wydawnicza Humanitas.
- Sierżant, W. (2000). *Planowanie i sterowanie produkcją*. Katowice: Wydawnictwo Politechniki Śląskiej.
- Skowronek, Cz., Sarjusz-Wolski, Z. (2008). *Logistyka w przedsiębiorstwie*. Warszawa: PWE.
- Szczepański, M. (2015). *Zarządzanie produkcją i logistyką w przedsiębiorstwie*. Warszawa: Wydawnictwo Naukowe PWN.
- Szołtysek, J., Jaroszyński, J. (2009). *Decyzje logistyczne w przedsiębiorstwie. Przykłady i zadania*. Wałbrzych: PWSZ.
- Szymonik, A. (2012). *Logistyka produkcji. Procesy. Systemy. Organizacja*, Warszawa: Difin.