



Fundusze Europejskie
dla Wielkopolski

Dofinansowane przez
Unię Europejską



SAMORZĄD
WOJEWÓDZTWA
WIELKOPOLSKIEGO



Mechanika i budowa maszyn

**Skrypt dla studentów
studiów niestacjonarnych**

Rok III

redakcja

Robert Cieślak

Konin 2025

Tytuł

Mechanika i budowa maszyn
Skrypt dla studentów studiów niestacjonarnych
Rok III

Redakcja naukowa

Robert Cieślak

Autorzy

Kamil Łodygowski
Adam Mazgajczyk
Edward Pająk
Paweł Rutecki

Projekt pn.

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



Fundusze Europejskie
dla Wielkopolski

Dofinansowane przez
Unię Europejską



SAMORZĄD
WOJEWÓDZTWA
WIELKOPOLSKIEGO



Wydawca

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



Spis treści

1. Eksploatacja maszyn i diagnostyka (Kamil Łodygowski)	4
2. Napędy współczesnych pojazdów (Kamil Łodygowski)	21
3. Systemy sterowania CNC (Edward Pająk)	38
4. Technologie łączenia (Adam Mazgajczyk)	86
5. English for Mechanical Engineering Paweł Rutecki	120



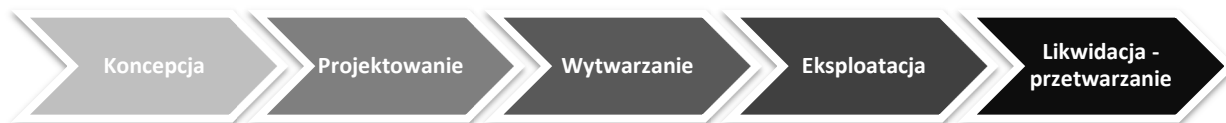
1. Eksploatacja maszyn i diagnostyka (dr inż. Kamil Łodygowski)

1.1. Eksploatacja maszyn

Maszyny i urządzenia tworzą istotną część majątku każdego konsorcjum. Dodatkowo w przypadku zakładów produkcyjnych odpowiedzialne są za generowanie zysku i w momencie zatrzymania odpowiadają bezpośrednio za często nie małe straty. Właściwe wykorzystanie wspomnianych środków trwałych wymaga opracowania i stosowania optymalnych warunków ich eksploatacji.

Eksploatacja jest jedną z pięciu etapów cyklu życia każdego obiektu technicznego:

- generowanie koncepcji technicznego zadania (faza wartościowania). Ocena i końcowy wybór właściwego rozwiązania problemu na tym etapie stanowi istotny wkład w efektywność realizacji zamierzonego celu. Błędy, które zostaną tu popełnione są trudne a często nawet niemożliwe do wyeliminowania w kolejnych fazach.
- Projektowanie. Na tym etapie realizuje się proces syntezy maszyny tak by były spełnione wymagania i funkcje wynikające z przyjętych rozwiązań na etapie wartościowania.
- Wytwarzanie, czyli proces materialnego wykonania maszyny o wymaganych właściwościach, zapewniających realizację celów użytkowych przewidzianych w projekcie.
- Eksploatacja, etap, w której maszyna realizuje założone cele, dla których została wymyślona, zaprojektowana i wytworzona.
- Likwidacja - ostatnia faza w życiu każdego obiektu technicznego. Oddziaływanie na środowisko w czasie procesu likwidacji maszyny poddawane jest analizie już w fazie generowania koncepcji oraz projektowania technicznego działania.



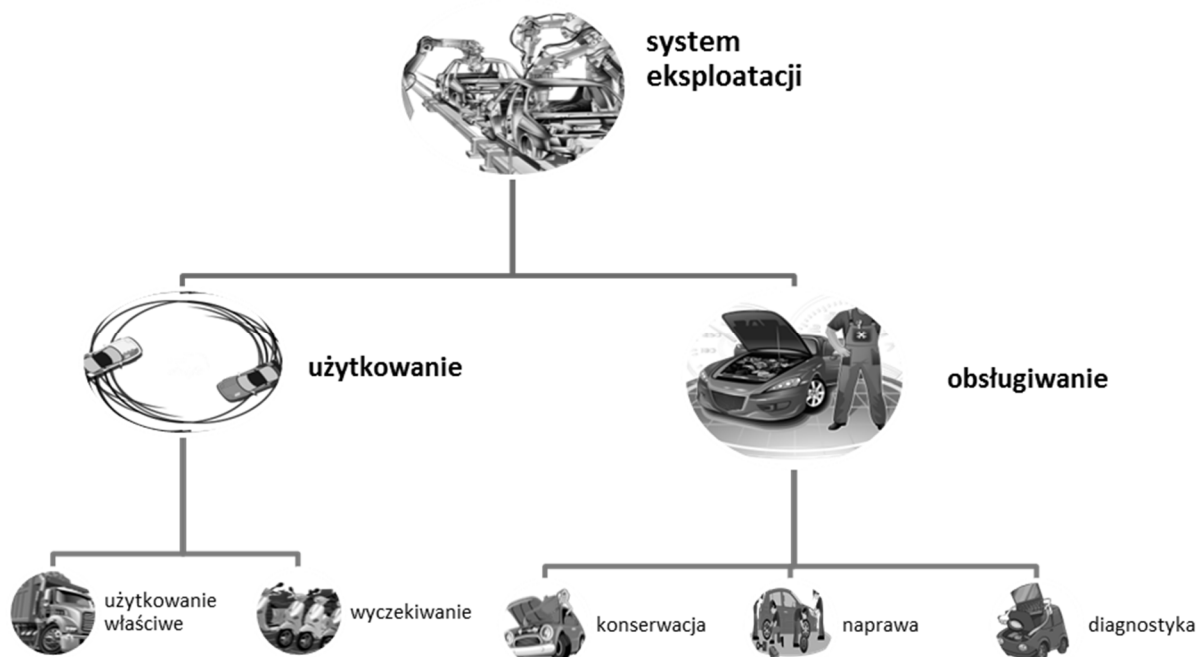
Rys. 1.1 Etapy życia obiektu technicznego

Źródło: opracowanie własne.

Eksploatacja jest zespołem celowych, organizacyjno-technicznych i ekonomicznych działań ludzi z obiektem technicznym oraz są to wzajemne relacje, występujące pomiędzy tym obiektem od chwili uruchomienia po przejęciu do wykorzystania zgodnie z przeznaczeniem, aż do likwidacji¹.

Podczas okresu eksploatacji obiektu technicznego wyróżnić można dwa procesy: użytkowanie oraz obsługiwanie. Procesem **użytkowania** maszyny nazywamy cykl zdarzeń związanych z działaniem technicznie zdatnej maszyny, również wszystkie zdarzenia związane ze sterowaniem i diagnostyką stanu technicznego, prowadzoną podczas użytkowania przez operatora maszyny. W przypadku maszyn z grupy automatów, zdarzenia związane z ich sterowaniem i bieżącą kontrolą stanu technicznego, jak i poprawnością realizowanych przez te maszyny procesów przebiegają bez udziału operatora. Użytkowanie jest głównym i najbardziej pożądanym procesem eksploatacji. Dąży się do tego by ten proces był jak najdłuższy.

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



Rys. 1.2 Schemat systemu eksploatacji

Źródło: opracowanie własne

Procesem **obsługiwania** maszyn nazywa się czynności związane z kontrolą, utrzymaniem oraz przywracaniem stanu technicznej sprawności maszyny. W procesie tym realizowane są różne czynności obsługi urządzenia wyłączonego na czas obsługi z użytkowania takie jak czyszczenie, konserwacja, przeglądy, naprawy okresowe, a także wszelkiego rodzaju remonty. Działania związane z obsługą urządzenia mogą mieć charakter planowy (systematyczny) lub doraźny (poawaryjny). Dąży się do tego by okres obsługiwania był jak najkrótszy. Tym samym analizując wyżej przedstawioną definicję procesu obsługiwania, praca operatora maszyny w czasie użytkowania nie zalicza się do czynności obsługowych, ponieważ stanowi nieodzowny element procesu użytkowania danej maszyny.

Użytkowane maszyny osiągają właściwą wydajność tylko wtedy, kiedy ich wszystkie mechanizmy będą miały zapewnione warunki pracy zgodne z założeniami i właściwościami konstrukcyjnymi. Zmiana tych czynników zmienia charakter pracy całego urządzenia, powodując przyspieszone zużycie mechanizmów, elementów, a w konsekwencji nawet uszkodzenie. Nie wystarcza zatem maszynę lub urządzenie nowocześnie zaprojektować i wytworzyć, z zachowaniem najwyższego standardu jakości i przy spełnieniu wymogów ekologicznych - należy ją także nowocześnie eksploatować.



Rys. 1.3 Zależność procesu zużywania się części od czasu

Źródło: M. Cymerys, Użytkowanie i obsługiwane maszyn i urządzeń 723[02].Z2.03. Poradnik dla ucznia. ITE PiB. Radom 2007

Analizując zużywanie się części w funkcji czasu podczas procesu eksploatacji można wyróżnić trzy okresy. Pierwszy okres, nazywany docieraniem trwa dość krótko, jednak jest bardzo ważny dla prawidłowego działania urządzenia. Następuje wówczas dopasowywanie się współpracujących powierzchni. W początkowej fazie ubytki materiału są dość intensywne. Z kolei pod koniec okresu rzeczywista realna powierzchnia styku obu elementów powiększa się, tym samym maleje intensywność zużycia oraz następuje stabilizacja naprężeń oraz odkształceń w warstwie wierzchniej.

Okres drugi jest to najdłuższy przebieg, podczas którego następuje normalna praca elementów maszyny. Zachodzące zjawiska charakteryzuje się powolną prawie stałą intensywnością zużycia. Trwałość poszczególnych części maszyn określana jest właśnie na podstawie tego okresu.

W momencie, kiedy następuje przekroczenie dopuszczalnego luzu pary trącej rozpoczyna się ostatni, trzeci okres. Następuje nadmierne nagrzewanie się, ulega zmniejszeniu sprawności mechanicznej, z kolei następuje wzrost zużycia materiałów smarnych wynikające bezpośrednio z zakłócenia normalnej współpracy części. Ciągłe



użytkowanie w takich okolicznościach powoduje całkowite zniszczenie lub awarię poszczególnych części².



Rys. 1.4 Procesy niszczenia maszyn i urządzeń

Źródło: opracowanie własne

Techniczne procesy zużycia, czyli fizyczne starzenie elementów obiektów technicznych zachodzące w materiałach konstrukcyjnych w wyniku oddziaływania makro- i mikrootoczenie, powodują bezpośrednio zmiany ich właściwości użytkowych (rys. 3).

Oprócz normalnego procesu starzenia podczas właściwej pracy maszyna jest także narażona na oddziaływanie różnych czynników wymuszających (rys. 4).

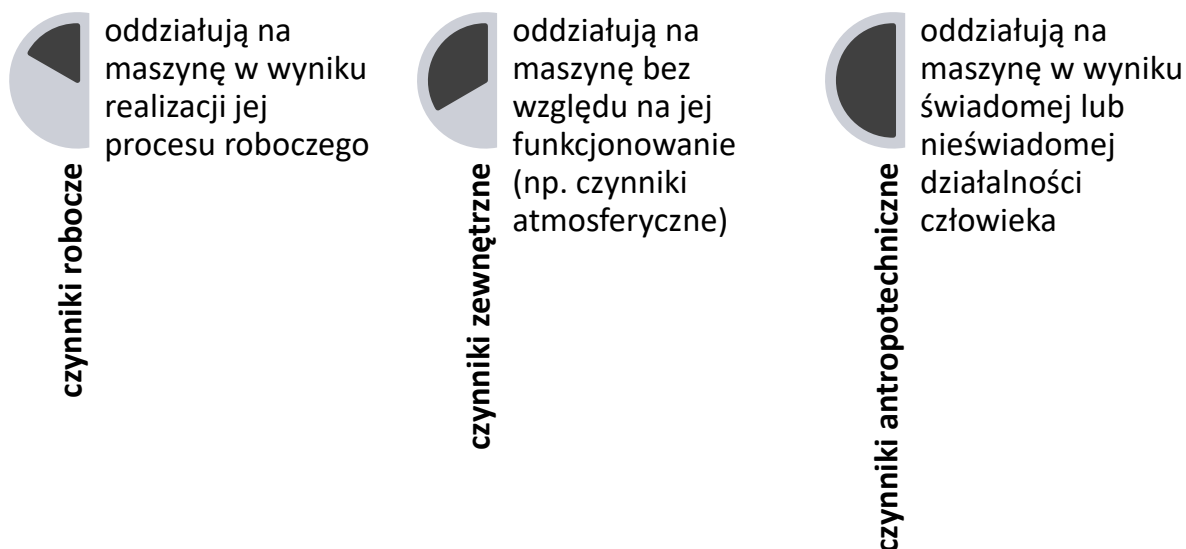
- czynniki robocze, które oddziałują na maszynę w wyniku realizacji jej procesu roboczego,
- czynniki zewnętrzne, które oddziałują na maszynę bez względu na jej funkcjonowanie (np. czynniki atmosferyczne),

² M. Cymerys, *Użytkowanie i obsługa maszyn i urządzeń 723[02].Z2.03. Poradnik dla ucznia*, ITE PiB, Radom 2007.



- czynniki antropotechniczne, które oddziałują na maszynę w wyniku świadomej lub nieświadomej działalności człowieka.

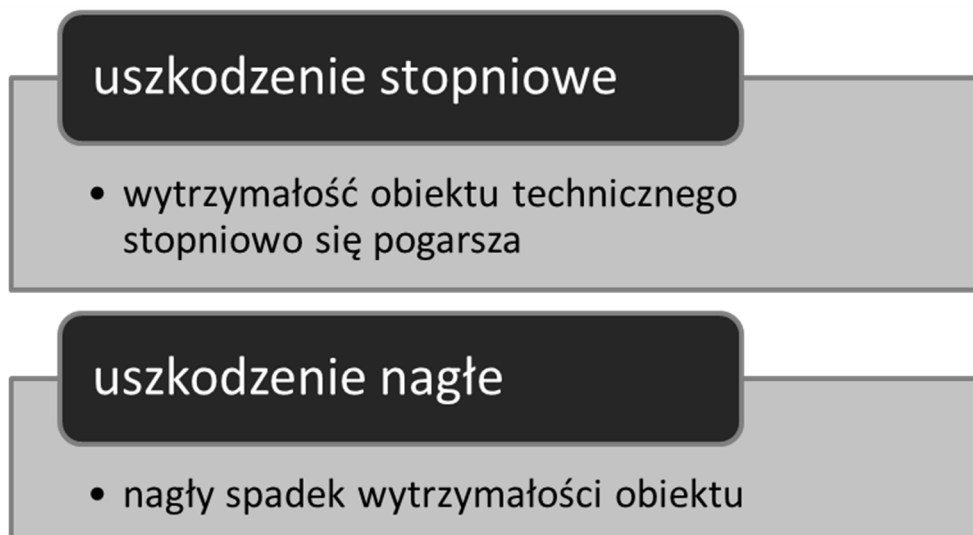
Innym typem utraty normalnych właściwości pracy przez maszynę jest **uszkodzenie** obiektu technicznego, czyli zdarzenie losowe, które powoduje, że maszyna czasowo lub na stałe traci stan zdatności i przechodzi do stanu niezdatności. Uszkodzeniem nazywamy stan, kiedy wartości parametrów określających obciążenie obiektu, zespołu, podzespołu lub elementu przekroczą graniczną wartość wytrzymałości. Uszkodzenie jest, więc zdarzeniem niezamierzonym, nie wliczając uszkodzeń celowych.



Rys. 1.5 Klasyfikacja czynników wymuszających

Źródło: opracowanie własne

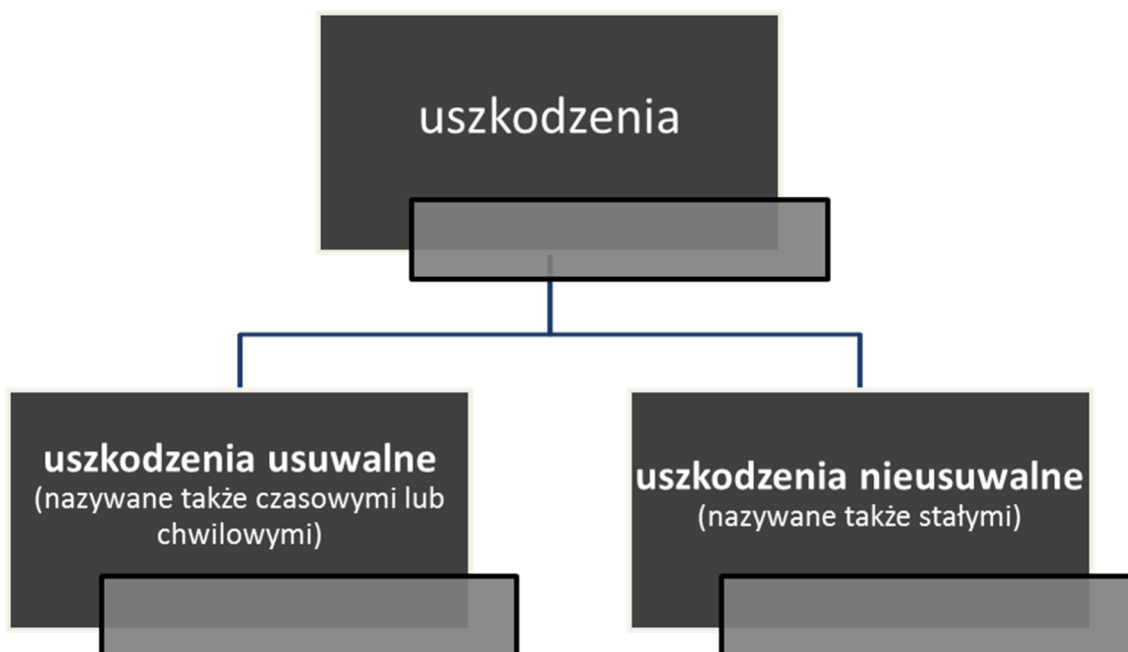
Klasyfikacja uszkodzeń została przedstawiona na rysunku 6.



Rys. 1.7 Klasyfikacja uszkodzeń ze względu na efekt końcowy

Źródło: opracowanie własne

Z kolei uszkodzenia ze względu na efekt końcowy zilustrowano na rysunku 7.



Rys. 1.7 Klasyfikacja uszkodzeń ze względu na efekt końcowy

Źródło: opracowanie własne



Właściwości urządzeń mechanicznych opisane odpowiednimi wielkościami, stanowią istotne elementy uwzględniane w procesie użytkowania maszyny. Rozróżnia się przy tym:

- wielkości krytyczne - których odchylenia od normy uniemożliwiają dalsze użytkowanie maszyny, a ich przekroczenie (zignorowanie) przez użytkownika jest niedopuszczalne, ze względu na zagrożenia dla samej maszyny lub człowieka z nią współpracującego (operatora) - np. złamane narzędzie obrabiarki, przekroczenie dopuszczalnej prędkości obrotowej ściernicy, spadek ciśnienia poniżej dopuszczalnej granicy w hydraulicznych układach mocowania maszyny, itp.
- wielkości ważne - to takie, których odchylenie od normy w sposób istotny zmniejsza sprawność i efektywność maszyny i może spowodować zagrożenie dla maszyny i człowieka jak np. nadmierny luz między czopem wału korbowego a panewką, ponadnormatywne stępienie narzędzia, zanieczyszczony filtr, zbyt niskie ciśnienie w oponach kół samochodowych itp.
- wielkości mało ważne - to takie, których odchylenia od normy nie powodują istotnych zakłóceń w pracy maszyny i mają charakter odwracalny jak np. mechaniczne uszkodzenia elementów osłonowych maszyny, przepalenie żarówki w instalacji oświetleniowej itp.
- wielkości pomijalne - są to takie wielkości, które nie mają żadnego wpływu na sprawność techniczną i efektywność pracy maszyny; zaliczają się do nich np. cechy estetyczne urządzenia, które doznając uszczerbku z jakiegokolwiek powodu, nie zmniejszają w najmniejszym stopniu jego technicznej sprawności (np. ubytki lakieru).

Stan eksploatacyjny obiektu technicznego przedstawia stadium podczas eksploatacji.

Wyróżnia się następujące podstawowe stany eksploatacyjne:

- naprawa główna,
- naprawa średnia,
- naprawa bieżąca,
- obsługa bieżąca,
- likwidacja,
- transport.

Naprawa jest działaniem, którego efektem ma być przywrócenie wartości użytkowej (funkcjonalności, sprawności techniczno-ekonomicznej) obiektu technicznego.



Racjonalna gospodarka remontowa powinna nie tylko być ukierunkowana na szybkie wykonywanie naprawy, ale także nie dopuścić do uszkodzenia awaryjnego obiektu, np. poprzez stosowanie systemu utrzymania planowo-zapobiegawczego.

Proces technologiczny naprawy jako zbiór operacji technologicznych, wykonywanych w określonej kolejności przez odpowiednio przygotowanych pracowników, na odpowiednio wyposażonych stanowiskach roboczych, ma za zadanie przywrócenie stanu zdatności obiektu technicznego. **Podatność naprawcza** jest z kolei prawdopodobieństwem, że wymagający obsługi obiekt techniczny zostanie w określonym czasie naprawiony. **Naprawialność** natomiast to działania niezbędne do przeprowadzenia obsługi. Te dwie cechy charakteryzują zarówno konstrukcje obiektu, jak i jego niezawodność – mają istotny wpływ na sumaryczne koszty eksploatacji.

Gotowość obiektu naprawialnego, 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.
- 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 pełni³.

Mając określone warunki eksploatacji danego obiektu technicznego i chcąc uzyskać informację czy w określonym czasie będzie on pracował prawidłowo mówimy wtedy o jego niezawodności. Niezawodność obejmuje różne wymagania opisane charakterystykami technicznymi, ekonomicznymi i socjologicznymi takich obiektów.

Wyróżnia się:

- niezawodność techniczną, która uwzględnia charakterystyki techniczne,
- niezawodność techniczno-ekonomiczną, która uwzględnia charakterystyki techniczne i ekonomiczne,

³ Ewald Macha, *Niezawodność maszyn*, Politechnika Opolska, Opole 2001.



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

Niezawodność obiektu jest to jego zdolność do spełnienia stawianych mu wymagań. Wielkością charakteryzującą zdolność do spełnienia wymagań jest przede wszystkim prawdopodobieństwo spełniania wymagań. Pojęcia trwałości i niezawodności są często stosowane zamiennie. Należy jednak je rozgraniczyć. Trwałość wskazuje okres, kiedy to obiekt nie wykazuje znaczącej utraty początkowego poziomu jakości. Tym samym określa jak długo spełnione będą stawiane mu wymagania w przydzielonych warunkach użytkowania. Natomiast niezawodność wskazuje na możliwość spełnienia przez obiekt techniczny stawianych mu oczekiwań w określonym czasie i w określonych warunkach użytkowania.

Jeżeli głównym wymaganiem jest to, aby wyrób był zdatny w przedziale czasu $(0, t)$, miarą tu może być ilość wykonanej pracy, czynności jak również czas itp. wtedy: „niezawodność obiektu jest to prawdopodobieństwo, że obiekt jest zdatny 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 inna⁴.

Każdy obiekt techniczny charakteryzuje się określoną niezawodnością (R), trwałością (T) i gotowością (G). Mając takie dane dotyczące danego wyrobu można sprawniej i odpowiedzialniej podejmować decyzje związane z jego produkcją oraz ich długotrwałą eksploatacją. Ze względu na rodzaj charakterystyki, która jest istotna dla danego obiektu technicznego, obiekty dzieli się na 8 klas:

1. Obiekty typu I (ilość), to obiekty, którym nie stawia się wymagań jakościowych związanych z ich niezawodnością, trwałością i gotowością.
2. Obiekty typu R (niezawodność), od których wymaga się dużej niezawodności. Typowymi przykładami takich obiektów są obiekty specjalnego przeznaczenia (np. samoloty, helikoptery).

⁴ E. Macha, *Niezawodność maszyn*.



3. 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.
4. 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.
5. Obiekty typu G (gotowość), od których wymaga się dużej gotowości. Są to głównie pogotowia - medyczne (karetka reanimacyjna), straż pożarna (wóz strażacki) i inne.
6. 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.
7. Obiekty typu TG (trwałość i gotowość), od których wymaga się głównie dużej trwałości i gotowości, np. statki ratownictwa morskiego i inne.
8. Obiekty typu RTG (niezawodność, trwałość, gotowość). Są to różnego rodzaju obiekty pogotowia, charakteryzujące się dużą trwałością i niezawodnością⁵.

1.2. 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 – scharakteryzowanie bieżącego stanu technicznego,

geneza – określenie przyczyn zaistnienia obecnego stanu,

⁵ E. Macha, *Niezawodność maszyn*.



prognoza – scharakteryzowanie granic czasowych przyszłych zmian stanów, wtedy czas pracy maszyny określony jest jej obiektywnym stanem technicznym, a nie statystyką – niezawodnością.

Przeanalizujmy diagnostyczny system eksploatacji (rys. 3). 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 pozyskiwania 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.

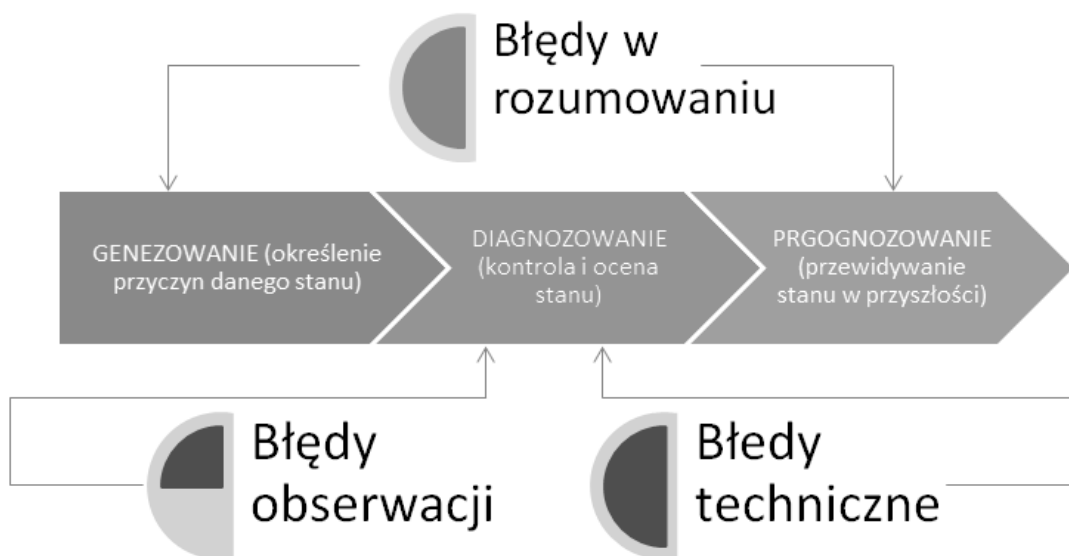
Podczas procesu diagnozowania należy przed wszystkim ocenić stan obiektu w danej chwili oraz określić niezgodności parametrów regulacyjnych. W przypadku stanu niezdatności należy zlokalizować uszkodzenia i wskazać przewidywany czas naprawy. Ostatnim etapem jest określenie terminu następnego diagnozowania.



Rys. 1.8 Algorytm postępowania podczas procesu diagnostycznego

Źródło: opracowanie własne.

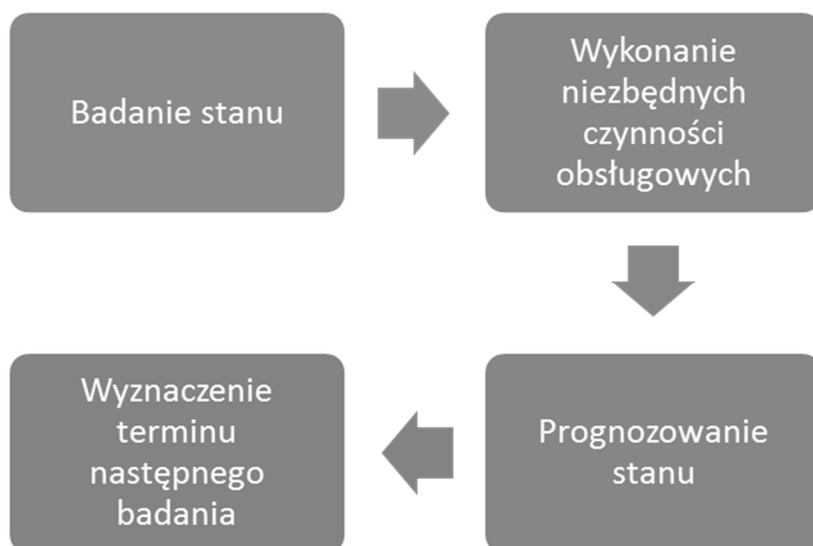
Proces postępowania z obiektem technicznym uzależniony jest od określenia jego stanu. W przypadku, kiedy maszyna jest w stanie zdatnym, postępowanie odbywa się wg algorytmu przedstawionego na rysunku 4. Jeżeli obiekt jest w stanie niezdatności, wtedy procedura wygląda jak na rysunku 5.



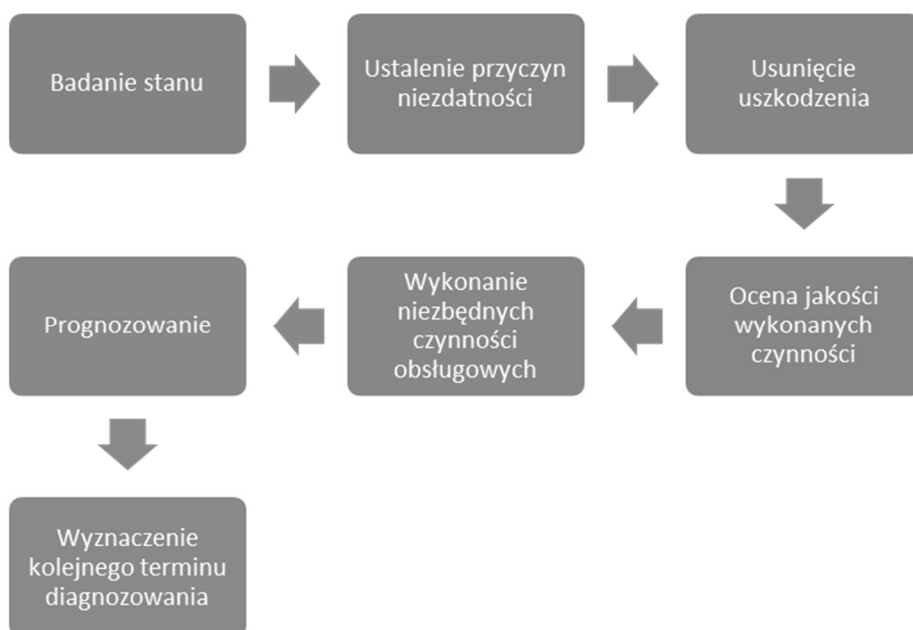
Rys. 1.9 Diagnostyczny system eksploatacji

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 ciągłe 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. 1.10 Proces postępowania diagnostycznego z obiektem technicznym zdatnym



Rys. 1.11 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)⁶.

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 (rys. 1.9) 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.

1.3. Zadania

1. Dokonać analizy wybranego zespołu (opis warunków pracy, realizowane zadania, ogólna budowa, zdjęcie lub schemat przedstawiający badany zespół z zaznaczonymi elementami (częściami zespołu), umiejscowienie, sposób kontroli i demontażu).

⁶ S. Legutko, *Eksploatacja maszyn*.



2. Dokonać analizy wybranego elementu (części) zespołu (opis, zadania, budowa, zdjęcie lub schemat przedstawiający badany element, umiejscowienie wybranego elementu w zespole).
3. Przedstawić diagnostykę zespołu (procedura diagnozowania zespołu, potencjalne objawy poszczególnych uszkodzeń) np. tabelarycznie.

Literatura

Cymerys M., *Użytkowanie i obsługiwane maszyn i urządzeń 723[02].Z2.03. Poradnik dla ucznia*, ITE PiB, Radom 2007

Kokociński J., *Wibroakustyczna diagnostyka maszyn*, „Chemia Przemysłowa” 2012, nr 3.

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

Macha E., *Niezawodność maszyn*, Politechnika Opolska, Opole 2001.



2. Napędy współczesnych pojazdów (dr inż. Kamil Łodygowski)

Wstęp

W XX wieku dominacja silników spalinowych została ugruntowana dzięki ich relatywnie wysokiej gęstości energii oraz rozwiniętej infrastrukturze paliwowej. Jednak już w drugiej połowie stulecia zaczęto dostrzegać ograniczenia tego rozwiązania, szczególnie w kontekście emisji zanieczyszczeń i efektywności energetycznej. Kryzysy naftowe lat 70. oraz rosnąca świadomość ekologiczna przyczyniły się do intensyfikacji badań nad alternatywnymi źródłami napędu.

Na przełomie XX i XXI wieku nastąpił przełom technologiczny związany z rozwojem elektroniki oraz materiałów, co umożliwiło praktyczne zastosowanie napędów hybrydowych i elektrycznych. Współczesne układy napędowe wykorzystują zaawansowane algorytmy sterowania energią oraz systemy odzyskiwania energii, takie jak rekuperacja. Jednocześnie rozwój technologii magazynowania energii, szczególnie akumulatorów litowo-jonowych, otworzył drogę do popularyzacji pojazdów elektrycznych.

Obecnie badania koncentrują się na dalszym zwiększaniu sprawności energetycznej oraz redukcji emisji gazów cieplarnianych, zgodnie z założeniami polityki klimatycznej i zasadami zrównoważonego rozwoju. W tym kontekście szczególne znaczenie zyskują także napędy wodorowe oparte na ogniwach paliwowych oraz koncepcje paliw syntetycznych. Analiza historyczna pokazuje, że ewolucja napędów samochodowych nie jest procesem liniowym, lecz wynikiem złożonej interakcji czynników technologicznych, ekonomicznych i społecznych, które wspólnie kształtują kierunki rozwoju współczesnej motoryzacji.

2.1 Napędy spalinowe

Do drugiej dekady XXI wieku pojazdy samochodowe były wyposażone głównie w silniki spalinowe tłokowe. W skład takiego silnika wchodziły zespoły mechanizmów (układy) oraz szkielet silnika zwany kadłubem. Układy silnika odpowiedzialne były za doprowadzenie paliwa w odpowiednim czasie i w odpowiedniej ilości, a także za smarowanie owego urządzenia, chłodzenie, usuwanie spalin i wywołanie zapłonu.

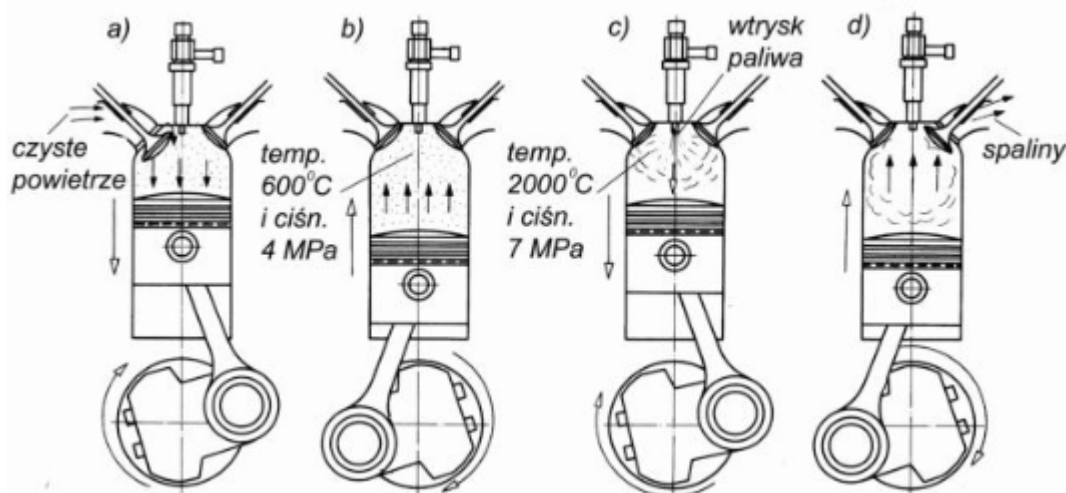


- Układ korbowy zamienia ruch tłoka posuwisto-zwrotny w cylindrze na ruch obrotowy wału korbowego.
- Układ zasilania doprowadza mieszkankę powietrza i paliwa do cylindra, która poddaje się spalaniu w cylindrze silnika.
- Układ rozrządu ma za zadanie zamykanie i otwieranie kanałów dolotowych w określonym czasie, które dostarczają do cylindra substancje konieczne do przeprowadzenia spalania. Z kolei kanały wylotowe usuwają spaliny na zewnątrz.
- Układ chłodzenia ma na celu stworzenie optymalnych warunków cieplnych, tak aby umożliwić silnikowi ekonomiczną pracę.
- Układ smarowania rozprowadza olej pomiędzy współdziałające z sobą elementy silnika, aby zmniejszyć opory tarcia oraz oczyszczać powierzchnie trące.

W zależności od sposobu zapłonu można rozróżnić silniki z zapłonem iskrowym oraz z zapłonem samoczynnym. W silniku z zapłonem iskrowym podstawowym paliwem są lżejsze węglowodory (benzyna) bądź ich mieszanki (np. z etanolem lub eterami). Przygotowana sprężona mieszanka paliwa i powietrza zapala się poprzez iskrę elektryczną. Z kolei silnik z zapłonem samoczynnym (wysokoprężny) pracuje na oleju napędowym, który również podlega mieszanii z biokomponentami - estrami. Czyste powietrze zostaje zassane do cylindra, które pod wpływem dużego sprężania nagrzewa się. Wtrysnięte paliwo do komory spalania zapala się samoczynnie.

Silniki możemy podzielić na dwusuwowe oraz czterosuwowe. Mogą pracować zarówno z zapłonem iskrowym, jak i samoczynnym. Istnieje również podział na liczbę obrotów silnika. Są to silniki wolnoobrotowe, średnioobrotowe i szybkoobrotowe.

Poniższy rysunek przedstawia schemat działania silnika czterosuwowego z zapłonem samoczynnym.



Rys. 1. Schemat działania czterosuwowego silnika z zapłonem samoczynnym
Źródło: <http://www.pcez-bytow.pl/download/plk/2.-dzialanie-tlokowych-siln-07.02.22.pdf>

Oczyszczone powietrze wpływa do cylindra poprzez otwarty zawór dolotowy w czasie przesunięcia się tłoka od GMP (górne martwe położenie) w kierunku wału korbowego. Jest to tak zwany suw ssania. Poprzez zmianę kierunku ruchu tłoka w DMP (dolne martwe położenie) zamyka się zawór dolotowy i następuje suw sprężania, który trwa przez kolejne pół obrotu wału korbowego. Następuje przejście tłoka z pozycji DMP do GMP. Zachodzi wtedy sprężanie powietrza w cylindrze do 3-5 MPa. Wskutek sprężania temperatura powietrza zwiększa się do około 600 stopni Celsjusza. Pod koniec suwu sprężania do cylindra zostaje wtrysnięte rozpylone paliwo. Mieszając się z powietrzem odparowuje i zapala się samoczynnie. Następuje spalanie, temperatura gazów sięga 2000-2500 °C, a ciśnienie rośnie do 5-12 MPa. Pod wpływem wysokiego ciśnienia gazów tłok zmienia położenie z górnego martwego położenia do dolnego martwego położenia. Wykonuje w ten sposób suw pracy.

Objętość przestrzeni roboczej, czyli chwilowa objętość cylindra nad tłokiem stopniowo powiększa się, a znajdujące się w niej gazy spalinowe ulegają rozprężeniu. Pod koniec suwu pracy otwiera się zawór wylotowy i gazy spalinowe uchodzą na zewnątrz silnika¹.

Przesuwający się tłok w cylindrze z DMP do GMP usuwa resztę spalin i wykonuje suw wydechu. Po wypchnięciu spalin zamyka się zawór wylotowy. Cylinder zostaje opróżniony i gotowy do przyjęcia nowej porcji powietrza, natomiast tłok jest w górnym martwym położeniu i może wykonać kolejny suw ssania.

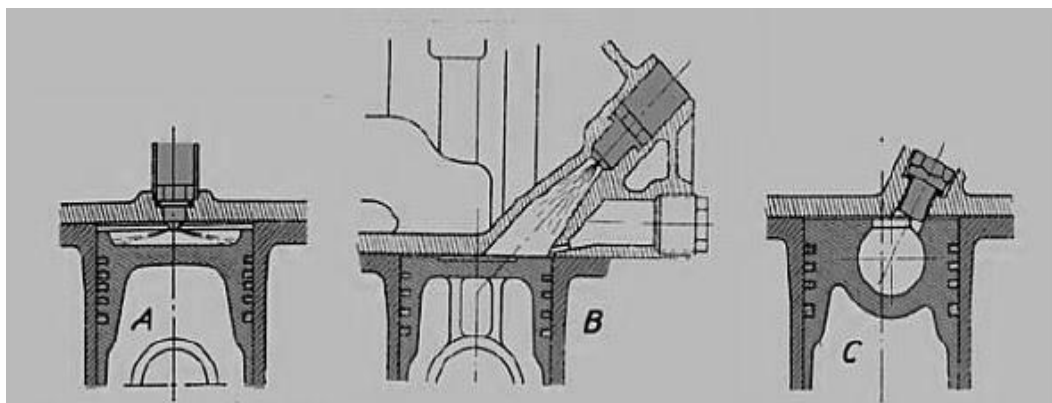
¹ Dąbrowski S., Kozłowska D., *Maszyny i ciągniki rolnicze*, Państwowe Wydawnictwo Rolnicze i Leśne, 1976, s. 189.

W silniku czterosuwowym cykl (obieg) pracy zostaje zamknięty w czterech suwach tłoka, wał korbowy obraca się wtedy dwa razy. Najistotniejszym suwem jest suw pracy, który jest źródłem energii mechanicznej silnika. Pozostałe suwy wykorzystują energię powstałą podczas suwu pracy.

Celem układu zasilania jest stworzenie mieszanki powietrzno-paliwowej zapewniającej jak najbardziej ekonomiczny i precyzyjny przebieg spalania. Do takiego układu zalicza się urządzenia doprowadzające paliwo i powietrze do cylindrów silnika, a także szereg urządzeń oczyszczających ten układ.

Na prawidłową pracę silnika duży wpływ ma kształt komory spalania. W silnikach z zapłonem samoczynnym stosuje się komory spalania z bezpośrednim wtryskiem paliwa z zawirowaniami powietrza, a także bez zawirowań. Spotyka się również komory wirowe w głowicy silnika.

Poniższy rysunek przedstawia rodzaje komór spalania.



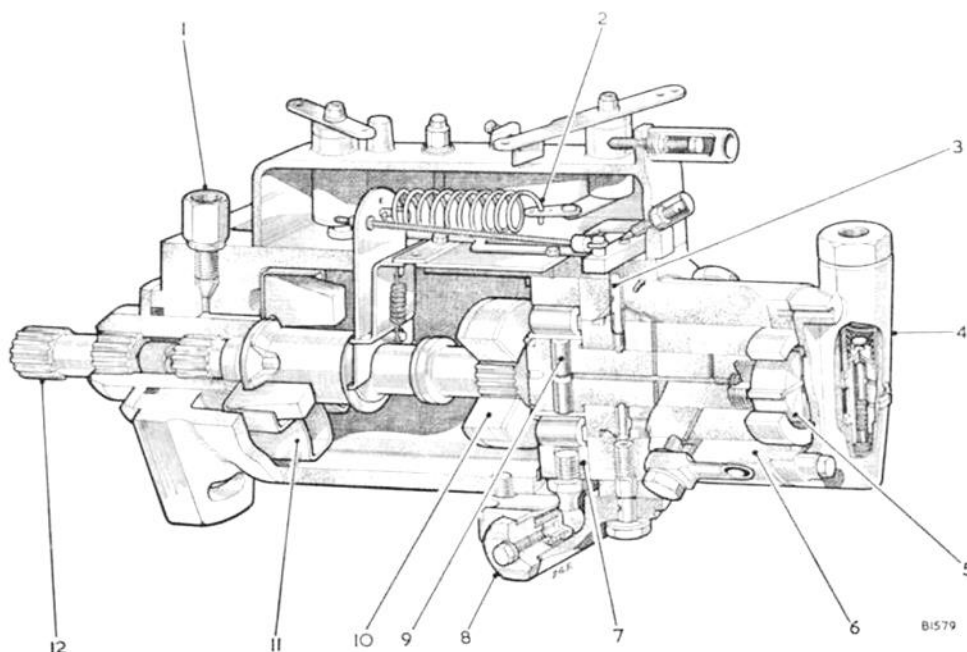
Rys. 2. Komory spalania w silnikach z zapłonem samoczynnym: a - bezpośredni wtrysk bez zawirowań powietrza, b - komora wirowa w głowicy silnika, c - bezpośredni wtrysk z zawirowaniem powietrza

Źródło: wszystkodlawarsztatu.pl

2.2 Przykładowa konstrukcja pompy wtryskowej

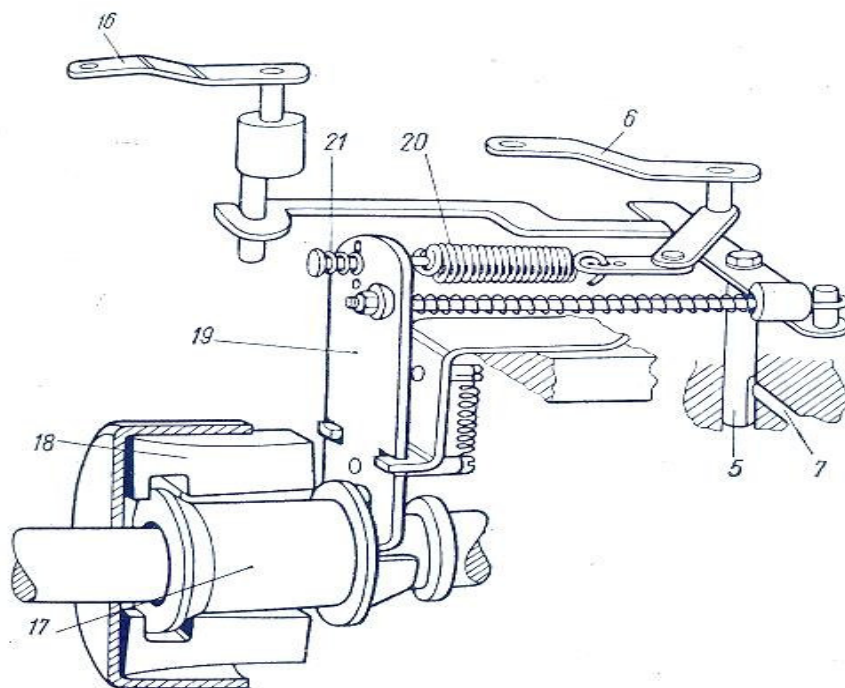
Przykładem pompy wtryskowej stosowanej w silnika o zapłonie samoczynnym jest pompa DPA co oznacza Distributor type fuel injection pump size A (pompa wtryskowa wielkości A typu rozdzielczego) z mechanicznym regulatorem obrotów. Jest to pompa wtryskowa rozdzielaczowa. W porównaniu z pompą wtryskową rzędową, pompa rozdzielaczowa posiada tylko jedną sekcję tłoczącą. Dodatkowym elementem jest

zastosowanie rozdzielacza, który doprowadza odpowiednie dawki paliwa do króćców pompy, a następnie do wtryskiwaczy. Układ wtryskowy uzyskuje wysokie ciśnienie przez zamianę ruchu obrotowego krzywek na ruch posuwisto-zwrotny tłoka lub układu tłoków. Pompa smarowana jest olejem napędowym przepływającym przez pompę i z powrotem do filtrów paliwowych. Konstrukcja ta pozbawiona jest sprężyn popychacza, kół zębatach oraz łożysk tocznych, co wydłużyło jej okres eksploatacji.



Rys. 3 Pompa DPA z regulatorem mechanicznym: 1- króciec wlotu paliwa, 2- sprężyna główna regulatora, 3- zawór dawkujący, 4- króciec przewodu zasilającego, 5- wirnik pompy przetłaczającej, 6- pompa skrzydełkowa, 7- występ popychacza (ślizgacz), 8- regulator dawki, 9- rolki popychacza, 10- ciężarek regulatora, 11- zespół ciężarków regulatora, 12- wałek wielowypustowy

Źródło: Falkowski H., Janiszewski T., *Aparatura paliwowa silników wysokoprzężnych – budowa i działanie*, Wydawnictwo Komunikacji i Łączności, Warszawa 1977.

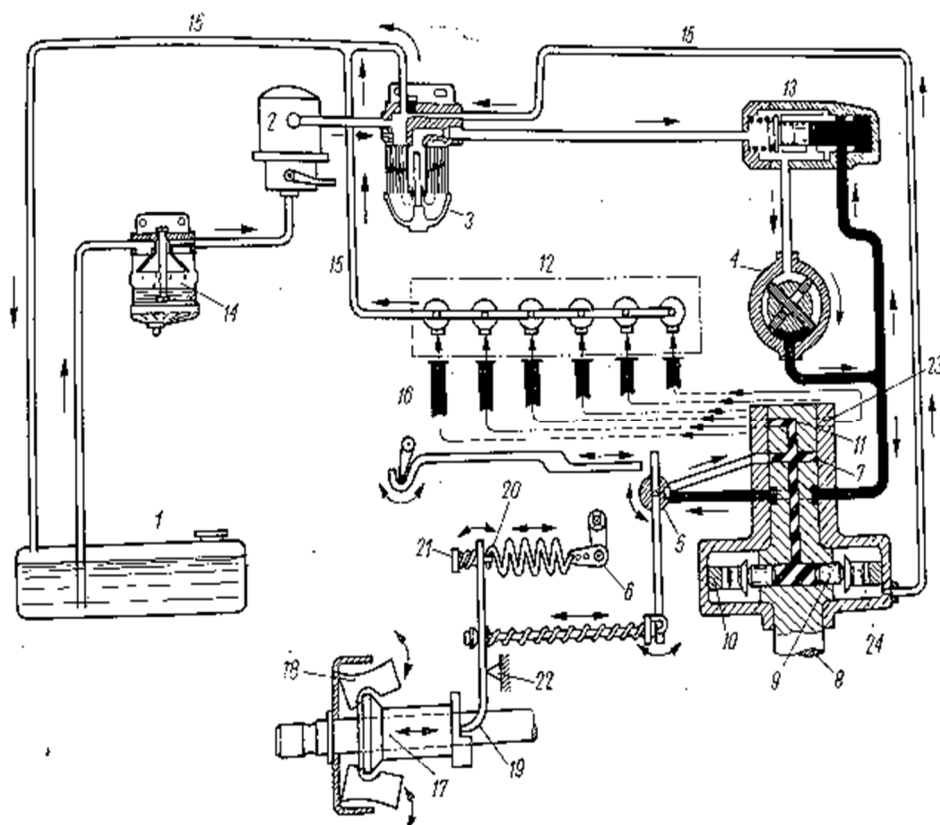


Rys. 4 Regulator mechaniczny pompy DPA: 5- zawór dawujący, 6- dźwignia regulacyjna, 7- otwory zasilające, 16- dźwignia „stop”, 17- tuleja oporowa, 18- ciężarek regulatora, 19- główna dźwignia regulatora, 20- sprężyna główna regulatora, 21- sprężyna wolnych obrotów

Źródło: Falkowski H., Janiszewski T., Aparatura paliwowa silników.

Zmiana tłoczonej dawki paliwa przez pompę następuje poprzez dławienie dopływu paliwa do rozdzielacza. Zawór dawujący (5) posiada obrotową przepustnicę, która przymyka kanał dozujący (7) poprzez obrót, tym samym powstrzymując przepływ paliwa. Przepustnica może zostać obrócona za pomocą dwuramiennej dźwigni, na jednym ramieniu zamocowane jest ciężko dźwigni „stop” (16); następuje wtedy zamknięcie zaworu dawującego, na drugiej natomiast ciężko ciężarkowego regulatora mechanicznego umieszczonego na wałku napędowym pompy. Regulator mechaniczny posiada 6 ciężarków (18) przytwierdzonych bez sworzni w blaszanej osłonie, które w czasie ruchu obrotowego wpływają na przesuwany pierścień (17), przenosząc ruch poosiowy na dźwignię regulatora (19), następnie do przepustnicy zaworu dawującego. Główna sprężyna regulatora (20) wytwarza siłę przeciwstawiającą się działaniu ciężarków, którą kierowca może regulować za pomocą dźwigni regulacyjnej (6), zmieniając w ten sposób zakres obrotów silnika. Wpływ na dźwignię

regulatora ma również sprężyna biegu jałowego (21) pracująca na ściskanie. Podczas wykonywania pracy na maksymalnych obrotach sprężyna jest całkowicie ściśnięta².



Rys. 5 Schemat układu paliwowego z pompą rozdzielczą DPA i regulatorem mechanicznym;
1 - zbiornik paliwa, 2 - pompa zasilająca, 3 - filtr dokładnego oczyszczania, 4 - pompa przetłaczająca,
5 - zawór dawkujący, 6 - dzwignia regulacyjna, 7 - otwory zasilające, 8 - wirnik z rozdzielaczem,
9 - tłoki, 10 - pierścień krzykowy, 11 - otwór rozdzielacza, 12 - wtryskiwacze, 13 - zawór regulacyjny,
14 - separator wody, 15 - przewody odprowadzające przeciek paliwa, 16 - dzwignia „stop”, 17 - tuleja
oporowa, 18 - ciężarek regulatora, 19 - głowna dzwignia regulatora, 20 - sprężyna głowna regulatora,
21 - sprężyna głownych obrotów, 22 - wspornik, 23 - otwór wtryskowy, 24 - dwie rolki popychacza

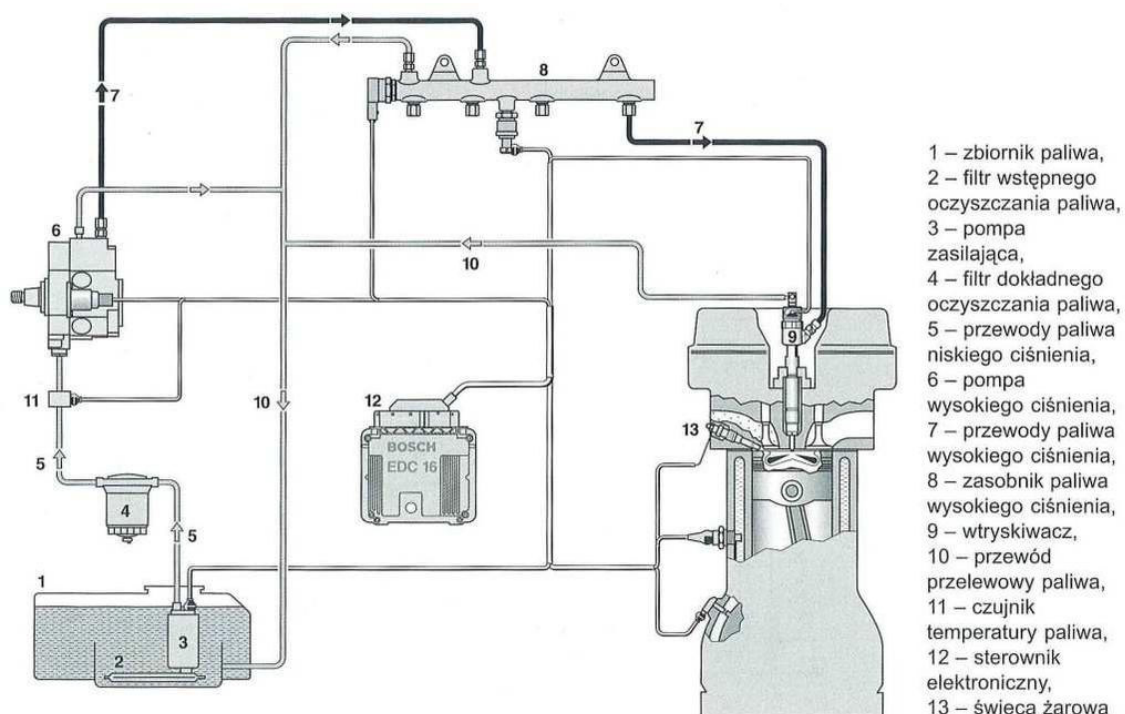
Źródło: Falkowski H., Janiszewski T., *Aparatura paliwowa silników wysokoprężnych*.

2.3 Nowoczesne konstrukcje silników spalinowych

Nowoczesne silniki o zapłonie samoczynnym stosowane w pojazdach samochodowych odbiegają w znacznym stopniu od tych, jakie były stosowane w połowie XX w. Podstawowa cecha, czyli samoczynny zapłon bez udziału urządzeń zapłonowych się nie zmieniła. Zaszły

² Falkowski H., Janiszewski T., *Aparatura paliwowa silników*, s. 36-44.

natomiast istotne zmiany w układzie zasilania paliwem. Stale zaostrzane normy emisji spalin przyczyniły się do poszukiwania coraz to nowszych rozwiązań poprawy spalania przy ograniczonych emisjach związków toksycznych i niskiego zużycia paliwa. Jedną z współczesnych możliwości poprawy przebiegu spalania jest zastosowanie aparatury wtryskowej, która umożliwia wtrysk paliwa bezpośredni przy wysokich ciśnieniach dochodzących do 250 MPa, z indywidualnym podziałem dawki na każdy z cylindrów silnika. Aktualnie takim wymaganiom może sprostać układ wtryskowy Common Rail. Układ ten, z początku stosowany w samochodach osobowych, coraz częściej stosowany jest w ciągnikach rolniczych.



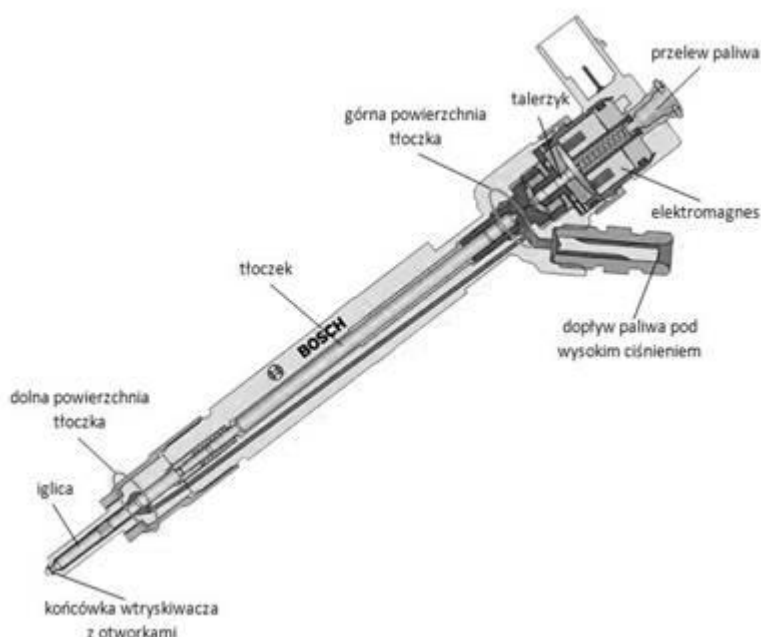
Rys. 6 Schemat układu Common Rail

Źródło: <https://docplayer.pl/49051902-Napedy-i-uklady-napedowe.html>

Zasada działania układu wtryskowego Common Rail jest dość prosta. W zbiorniku paliwa umieszczona jest pompa zasilająca. Paliwo transportowane jest ze zbiornika za pomocą pompy zasilającej poprzez filtry wstępnego i dokładnego oczyszczania aż do pompy wysokiego ciśnienia. Napędzana ruchem obrotowym silnika pompa spręża olej napędowy, transportując go do zasobnika paliwa wysokiego ciśnienia. Następnie paliwo pod wysokim

ciśnieniem trafia do wtryskiwaczy. Wtryskiwacze sterowane są poprzez komputer, w odpowiednim momencie otwierają elektromagnetyczne zawory i rozpylają paliwo w komorze spalania. Dodatkowym elementem jest świeca żarowa wpływająca na poprawę spalania i rozruchu silnika.

Pierwsze układy Common Rail posiadały wtryskiwacze elektromagnetyczne pracujące na ciśnieniu 1000-1300 bar. Wraz z postępem i zwiększeniem ciśnienia do 2000 bar zastąpiono je wtryskiwaczami piezoelektrycznymi.



Rys. 7 Wtryskiwacz elektromagnetyczny firmy BOSCH

Źródło: http://sjsutst.polsl.pl/archives/2015/vol86/065_ZN86_2015_KoniecznyAdamczykAdamczyk

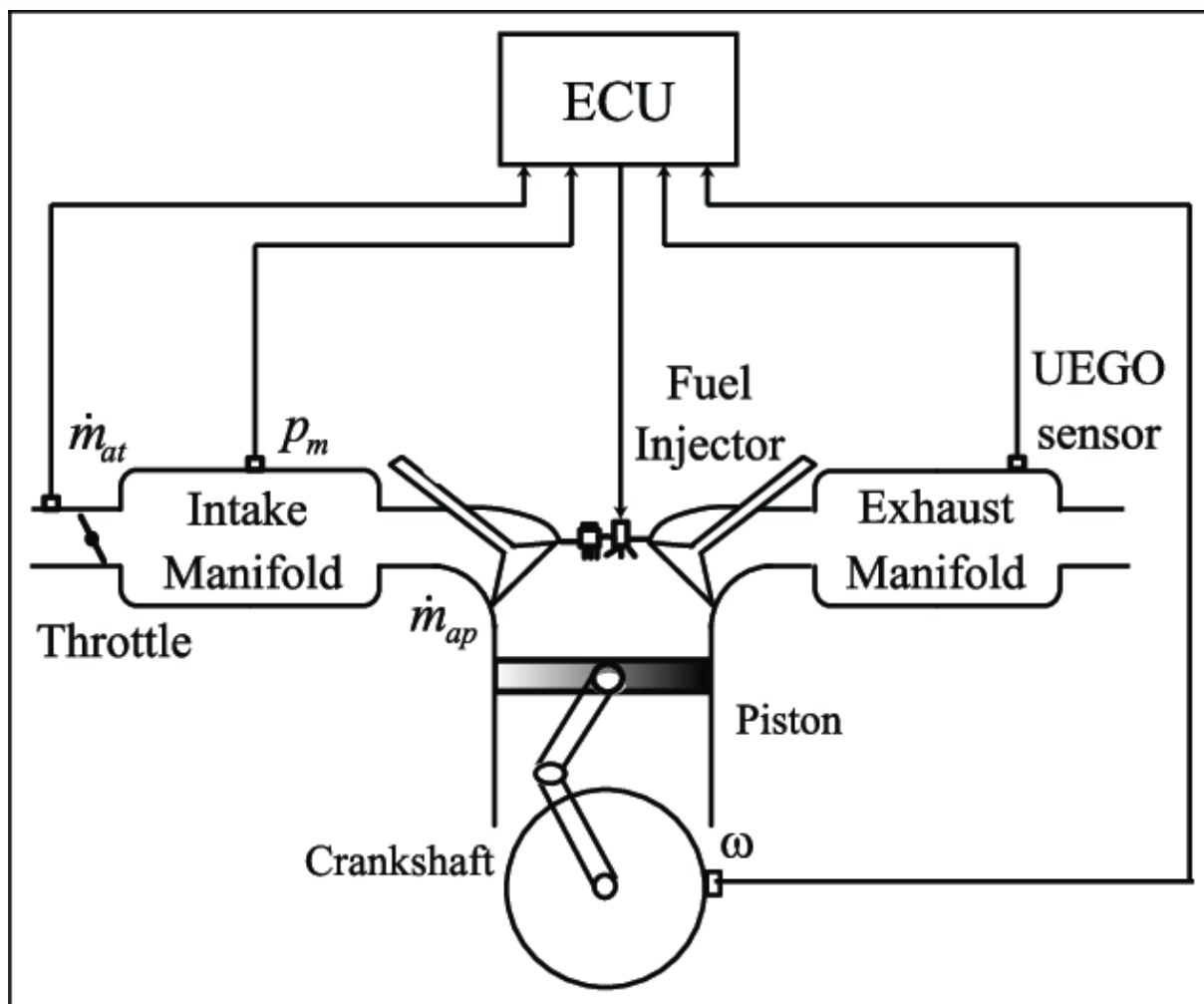
Wtryskiwacz elektromagnetyczny podzielony jest na część górną i dolną. Zarówno w części górnej jak i dolnej znajduje się komora spalania. W obu przypadkach ciśnienie jest takie samo jak w zasobniku paliwa wysokiego ciśnienia, z tym, że w dolnej komorze średnica tłoczka jest mniejsza. Skutkuje to mniejszą siłą nacisku na tłoczek mimo takiego samego ciśnienia. Napierające ciśnienie na dolną powierzchnię tłoczka doprowadza do zamknięcia iglicy i zaprzestaniem podawania paliwa do komory spalania. Gdy sterownik elektryczny nakazuje podanie dawki paliwa, elektromagnes umieszczony w górnej części wtryskiwacza unosi talerzyk zamykający przestrzeń wysokiego ciśnienia. Poprzez utratę ciśnienia w górnej komorze spalania następuje wzrost ciśnienia paliwa w dolnej komorze. Wysokie ciśnienie powoduje przesunięcie się iglicy uwalniając paliwo do cylindra.



Zasada działania wtryskiwacza piezoelektrycznego jest podobna. Iglica również znajduje się w dolnej części wtryskiwacza, nie ma natomiast tłoczka i trzpienia. Główną różnicą jest wypełnienie komory wysokiego i niskiego ciśnienia materiałem piezoelektrycznym. Przepływający prąd powoduje wydłużenie się tego materiału zmieniając tym samym ciśnienie w komorach przez poruszającą się iglicę. Dodatkowo do wtryskiwacza został zamontowany siłownik hydrauliczny w celu zwiększenia skoku materiału piezoelektrycznego. Układ Common Rail jest bardzo wrażliwy na paliwo złej jakości. Należy regularnie wymieniać filtry paliwowe i stosować paliwo dobrej jakości.

W przypadku silników spalinowych o zapłonie iskrowym dzisiejsza konstrukcja to również silnik z bezpośrednim wtryskiem paliwa (GDI – Gasoline Direct Injection) stanowi zaawansowaną odmianę klasycznego silnika benzynowego, w której paliwo jest wtryskiwane bezpośrednio do komory spalania, a nie – jak w starszych układach MPI – do kolektora dolotowego. Rozwiązanie to pozwala na precyzyjniejsze sterowanie procesem tworzenia mieszanki paliwowo-powietrznej, co przekłada się na wyższą sprawność, większą moc jednostkową oraz redukcję emisji spalin.

Budowa silnika GDI obejmuje typowe elementy silnika tłokowego: blok cylindrów, głowicę, tłoki, korbowody, wał korbowy oraz układ rozrządu (najczęściej DOHC). Kluczową różnicą jest układ zasilania paliwem. Składa się on z pompy niskiego ciśnienia (w zbiorniku), pompy wysokiego ciśnienia (napędzanej mechanicznie, często od wałka rozrządu), przewodów paliwowych o wysokiej wytrzymałości oraz wtryskiwaczy umieszczonych bezpośrednio w głowicy cylindrów. Wtryskiwacze GDI pracują przy ciśnieniach rzędu 50–350 barów, co umożliwia bardzo drobne rozpylanie paliwa.



Rys. 8 Schemat budowy silnika o ZI GDI

Źródło: https://www.researchgate.net/figure/Schematic-architecture-of-a-GDI-engine_fig1_309695187

Istotnym elementem jest także układ dolotowy powietrza, często wyposażony w przepustnicę sterowaną elektronicznie oraz system zmiennej geometrii kanałów dolotowych (tzw. tumble lub swirl), które wpływają na ruch powietrza w cylindrze. Współczesne silniki GDI wykorzystują również czujniki (np. ciśnienia, temperatury, składu spalin) oraz zaawansowany sterownik ECU, który zarządza czasem i dawką wtrysku oraz momentem zapłonu.

Zasada działania silnika GDI opiera się na czterosuwowym cyklu pracy: ssania, sprężania, spalania (pracy) i wydechu. W fazie ssania do cylindra zasysane jest powietrze (bez paliwa). Następnie, w zależności od trybu pracy, paliwo jest wtryskiwane bezpośrednio do cylindra w określonym momencie:



- W trybie jednorodnym (homogenicznym) paliwo jest wtryskiwane wcześniej, podczas suwu ssania, co umożliwia dokładne wymieszanie z powietrzem. Powstaje jednorodna mieszanka o współczynniku λ bliskim 1, spalana po zapłonie iskrowym.
- W trybie warstwowym (stratyfikowanym) paliwo wtryskiwane jest późno, pod koniec suwu sprężania. Tworzy się lokalnie bogata mieszanka w pobliżu świecy zapłonowej, podczas gdy reszta objętości cylindra zawiera ubogą mieszankę lub samo powietrze. Pozwala to na pracę przy bardzo ubogich mieszankach ($\lambda > 1$), co obniża zużycie paliwa.

Zapłon następuje dzięki iskrze generowanej przez świecę zapłonową. Po spaleniu mieszanki powstaje wysoka temperatura i ciśnienie gazów, które powodują ruch tłoka w dół (suw pracy), a następnie spaliny są usuwane w suwie wydechu.

Zalety technologii GDI obejmują zwiększoną sprawność cieplną, lepszą dynamikę silnika, możliwość pracy na ubogich mieszankach oraz redukcję emisji CO₂. Jednocześnie występują wyzwania techniczne, takie jak emisja cząstek stałych (PM), konieczność stosowania filtrów GPF oraz większa złożoność układu paliwowego.

2.4 Napędy elektryczne

Elektryczne pojazdy osobowe (EV)

Elektryczne pojazdy osobowe to pojazdy, które są przeznaczone do transportu osób. Zasilane są akumulatorami litowo-jonowymi lub innymi nowoczesnymi akumulatorami, które zapewniają ich autonomię na poziomie od 100 do 600 km na jednym ładowaniu, w zależności od modelu i warunków eksploatacji. Do najpopularniejszych typów pojazdów osobowych należą:

- Pojazdy elektryczne w pełni elektryczne (BEV, Battery Electric Vehicle) – pojazdy, które są całkowicie napędzane energią elektryczną, zasilane wyłącznie z akumulatorów. Przykłady: Tesla Model 3, Nissan Leaf, BMW i3.
- Pojazdy elektryczne z zasięgiem rozszerzonym (EREV, Extended Range Electric Vehicle) – pojazdy elektryczne wyposażone w silnik spalinowy, który działa jako generator, gdy akumulator się wyczerpie. Przykład: Chevrolet Volt.



Elektryczne pojazdy hybrydowe (HEV)

Pojazdy hybrydowe elektryczne (HEV) to pojazdy, które łączą silnik elektryczny i spalinowy. Silnik elektryczny wspomaga silnik spalinowy, co pozwala na zmniejszenie zużycia paliwa i emisji spalin. Pojazdy te mogą działać w trybie elektrycznym (na krótkich dystansach) lub hybrydowym, gdzie oba silniki współpracują w celu optymalizacji efektywności energetycznej. Przykłady: Toyota Prius, Honda Insight.

Pojazdy elektryczne plug-in (PHEV)

Pojazdy elektryczne typu Plug-in Hybrid (PHEV) to pojazdy, które łączą funkcje pojazdu elektrycznego i hybrydowego, jednak różnią się tym, że można je ładować za pomocą zewnętrznego źródła energii (gniazdka elektrycznego lub stacji ładowania). Dzięki temu pojazdy te mogą poruszać się w pełni elektrycznie na krótkich trasach i przełączać na silnik spalinowy na dłuższe odległości. Przykłady: Mitsubishi Outlander PHEV, BMW i8, Audi Q5 TFSI e.

Elektryczne pojazdy użytkowe (EV – Light Commercial Vehicles)

Elektryczne pojazdy użytkowe to pojazdy wykorzystywane do transportu towarów, które również wykorzystują napęd elektryczny. Są to głównie dostawcze pojazdy, które zapewniają niskie koszty eksploatacji oraz redukcję emisji spalin. Przykłady: Mercedes-Benz eSprinter, Nissan e-NV200, Renault Kangoo Z.E.

2.5 Budowa pojazdu elektrycznego

Pojazdy elektryczne wykorzystują napęd elektryczny, w którym energia elektryczna jest magazynowana w akumulatorach (lub innych systemach magazynowania energii) i wykorzystywana do napędu silnika elektrycznego. W zależności od typu pojazdu, silnik elektryczny może być wykorzystywany w połączeniu z różnymi systemami magazynowania energii, takimi jak akumulatory litowo-jonowe, akumulatory kwasowo-ołowiowe czy superkondensatory. Pojazdy te charakteryzują się zerową emisją spalin w trakcie eksploatacji, co przyczynia się do poprawy jakości powietrza w miastach.

Samochody elektryczne (EV) różnią się od tradycyjnych pojazdów spalinowych pod względem konstrukcji, ponieważ zamiast silnika spalinowego wykorzystują silnik elektryczny



do napędu. Ich budowa opiera się na kilku kluczowych elementach, które współpracują w celu zapewnienia efektywności energetycznej, komfortu jazdy i bezpieczeństwa.

1. Akumulator (Bateria)

Akumulator to najważniejszy element każdego samochodu elektrycznego, który gromadzi energię elektryczną potrzebną do napędu silnika. Współczesne pojazdy elektryczne korzystają najczęściej z akumulatorów litowo-jonowych (Li-ion), które charakteryzują się wysoką gęstością energetyczną, długą żywotnością i szybkim czasem ładowania.

Funkcja: Przechowywanie energii elektrycznej, która jest wykorzystywana przez silnik elektryczny do napędu pojazdu.

Wielkość: Zależna od modelu pojazdu; akumulatory mogą mieć pojemność od 20 kWh do ponad 100 kWh, co determinuje zasięg pojazdu na jednym ładowaniu.

Lokalizacja: Akumulator zazwyczaj umieszczony jest w dolnej części pojazdu (np. w podłodze), co poprawia stabilność samochodu.

2. Silnik elektryczny

Silnik elektryczny to element napędowy samochodu elektrycznego, który zamienia energię elektryczną z akumulatora na energię mechaniczną, przekazując ją na koła.

Funkcja: Napędzanie pojazdu, czyli przekształcanie energii elektrycznej w energię ruchu.

Rodzaje: W samochodach elektrycznych stosuje się głównie silniki prądu stałego (DC) lub prądu przemiennego (AC). Silniki indukcyjne (AC) są powszechnie używane, ponieważ oferują prostą konstrukcję, dużą niezawodność i wysoką wydajność.

Wydajność: Silniki elektryczne charakteryzują się wysoką sprawnością energetyczną, często przekraczającą 90%, co przekłada się na mniejsze straty energii.

3. Inwerter

Inwerter to urządzenie, które konwertuje prąd stały (DC) z akumulatora na prąd przemienny (AC), który zasila silnik elektryczny. Ponadto, inwerter kontroluje obroty silnika i umożliwia precyzyjne sterowanie momentem obrotowym.

Funkcja: Przekształcanie prądu stałego w prąd przemienny oraz kontrolowanie pracy silnika elektrycznego (regulacja obrotów).



Sterowanie: Dzięki inwerterowi możliwe jest płynne przyspieszanie, hamowanie odzyskiwanie energii i optymalizacja pracy silnika.

4. Układ przeniesienia napędu

Układ przeniesienia napędu w samochodzie elektrycznym jest prostszy niż w pojazdach spalinowych, ponieważ nie wymaga skomplikowanej skrzyni biegów. Zamiast tego, pojazdy elektryczne często wyposażone są w jednobiegowy układ przekładni.

Funkcja: Przekazywanie mocy z silnika elektrycznego na koła pojazdu.

Konstrukcja: Ze względu na charakter pracy silnika elektrycznego, samochody elektryczne nie potrzebują wielostopniowej skrzyni biegów. Silnik elektryczny zapewnia odpowiednią moc w szerokim zakresie obrotów.

5. System zarządzania energią (BMS)

BMS (Battery Management System) to system zarządzający akumulatorem w samochodzie elektrycznym. Monitoruje on stan naładowania, temperaturę, napięcie oraz inne parametry akumulatora, dbając o jego optymalną pracę i bezpieczeństwo.

Funkcja: Zapewnienie bezpieczeństwa akumulatora, monitorowanie jego stanu oraz zapewnianie efektywnego zarządzania energią (ładowanie, rozładowywanie).

Bezpieczeństwo: BMS zapobiega przegrzewaniu się, przeładowaniu i nadmiernemu rozładowaniu akumulatora, co mogłoby prowadzić do uszkodzenia ogniwa.

6. System hamulcowy i regeneracyjny układ hamulcowy

System hamulcowy w samochodach elektrycznych działa podobnie jak w pojazdach spalinowych, ale dodatkowo większość nowoczesnych samochodów elektrycznych wyposażona jest w hamulce regeneracyjne. Tego typu hamulce wykorzystują silnik elektryczny do spowalniania pojazdu, jednocześnie odzyskując część energii kinetycznej i przesyłając ją z powrotem do akumulatora.

Funkcja: Hamowanie i odzyskiwanie energii. Dzięki hamulcom regeneracyjnym zmniejsza się zużycie tradycyjnych hamulców, a pojazd może naładować akumulator podczas zwalniania.



Zalety: Hamulce regeneracyjne pozwalają na zwiększenie efektywności energetycznej i wydłużenie zasięgu pojazdu.

2.6. Zadanie

Porównać wymiary układu wydechowego dla seryjnego układu rozrządu ze zmodyfikowanym

1. Obliczyć średnicę wewnętrzną rur:

$$D = 2 \sqrt{\left(i \cdot \sqrt{(H - (R - r))^2 + (R - r)^2 \cdot (R + r)} \right) [\text{mm}]}$$

gdzie

D = średnica wewnętrzna rury w mm

R = promień zewnętrznej przylgni zaworowej w gnieździe ssącym w mm (31.5 mm)

r = promień wewnętrznej przylgni zaworowej zaworu ssącego w mm (25,8 mm)

H = maksymalny wznios zaworu w mm

i = ilość zaworów ssących w cylindrze

2. Obliczyć długość rur: $L = \frac{K \cdot 4250}{n \cdot x}$

gdzie:

L = długość rury mierzona od gniazda wydechowego

K = kąt między osiami krzywek (kąt zawarty pomiędzy początkiem otwarcia wydechu, a początkiem otwarcia ssania) mierzony na wale korbowym (otwarcie zaworu dolotowego: 7st przed GMP, zamknięcie 35st po DMP; otwarcie zaworu wylotowego: 37st przed DMP, zamknięcie 5st po GMP)

n = obroty początku działania falowego

x = ilość odbicie fali

3. Obliczyć pojemność całego układu wydechowego. Pojemność układu wydechowego ma duże znaczenie, ponieważ potęguje działanie rur. Podobnie jak w instrumentach muzycznych pudła rezonansowe potęgują działanie strun i podobnie jak w instrumentach muzycznych duże pojemności wzmacniają niskie częstotliwości, zaś



małe – wysokie. Dla jednocylindrowego silnika, pojemność układu wydechowego powinna być równa pojemności rury o wyliczonej średnicy i długości dla pierwszego odbicia fali. Oczywiście w silnikach wielocylindrowych pojemność całego układu wydechowego powinna być sumą pojemności rur wyliczonych j/w.

$$V = K \cdot 13351,77 \cdot \left(\frac{D}{2}\right)^2 \cdot \left(\frac{C}{n}\right) [\text{ccm}]$$

gdzie:

V = całkowita pojemność układu w ccm

C = ilość cylindrów

K = kąt między osiami krzywek

D = obliczona średnica rury w cm

n = obroty początku działania falowego wydechu

Literatura

Dąbrowski S., Kozłowska D., *Maszyzny i ciągniki rolnicze*, Państwowe Wydawnictwo Rolnicze i Leśne, 1976, s. 189.

Falkowski H., Janiszewski T., *Aparatura paliwowa silników wysokoprężnych - budowa i działanie*, Wydawnictwo Komunikacji i Łączności, Warszawa 1977.

http://sjsutst.polsl.pl/archives/2015/vol86/065_ZN86_2015_KoniecznyAdamczykAdamczyk
[dostęp: 10.04.2026]

<http://www.pcez-bytow.pl/download/plk/2.-dzialanie-tlokowych-siln-07.02.22.pdf> [dostęp: 10.04.2026]

<https://docplayer.pl/49051902-Napedy-i-uklady-napedowe.html> [dostęp: 10.04.2026]

https://www.researchgate.net/figure/Schematic-architecture-of-a-GDI-engine_fig1_309695187 [dostęp: 10.04.2026]

www.wszystkodlawarsztatu.pl [dostęp: 10.04.2026]



3. Systemy sterowania CNC (dr hab. inż. Edward Pająk, prof. nadzw.)

3.1. Wprowadzenie - zakres tematyczny rozdziału

Jak opisuje teoria początku Wszechświata, powstał on około 14 mld lat temu. Był to moment, w którym czas, przestrzeń, materia i energia zaczęła istnieć i od tamtej chwili Wszechświat się rozszerza. Około 4,5 mld lat temu obłok pyłu kosmicznego i gazu został wprowadzony w obrót przez wybuch sąsiedniej gwiazdy – wówczas formuje się Układ Słoneczny i nasza planeta Ziemia. W ciągu kolejnych milionów lat proste cząsteczki – molekuły (atomy powiązane wiązaniami chemicznymi), łączyły się ze sobą tworząc coraz to bardziej złożone związki. Powstają aminokwasy, które po połączeniu tworzyły białka. Są one kluczową częścią złożonego systemu chemicznego istot żywych. W wyniku ewolucji ok. 300 000 lat temu powstał w Afryce człowiek rozumny (*Homo sapiens*). Około 45 000 lat temu homo sapiens dotarł do Europy. Rozwiązywanie problemów egzystencjalnych wymagało wspólnego działania całego plemienia – pracy zespołowej. A ta wymagała opracowania systemu porozumiewania się. I tak powstał pierwszy język. Pojawiły się kolejne potrzeby, których zaspokojenie było inspiracją do podjęcia prób rozwiązania problemu. Nastąpił proces zmian rozwojowych techniki określony mianem **postępu technicznego** i w jego ramach konkretnych rozwiązań technicznych – **wynalazków**.

Skala zapotrzebowania ludzkości na wyroby spowodowała, że tradycyjne manufaktury w których wyroby wytwarzane były ręcznie przez rzemieślników okazały się niewystarczające. Pod koniec XVIII wieku, James Watt udoskonalił maszynę parową, do tego stopnia, że umożliwiła ona zastąpienie pracy mięśni ludzi lub zwierząt używaną do napędu urządzeń technologicznych, poprzez energię mechaniczną. W dziejach ludzkości rozpoczął się okres **mechanizacji**.

Postępujący w następnych latach rozwój techniki doprowadził w połowie XX wieku do rozwoju **automatyzacji**. Zastępowała ona nie tylko ludzką pracę fizyczną (mechanizacja), ale również umożliwiała wykonywanie przez urządzenia technologiczne określonych czynności bez udziału człowieka. Częścią automatyzacji jest **sterowanie**, które odpowiada za realizację



określonego działania, natomiast automatyzacja odpowiada za samodzielność całego procesu¹. Sterowanie jest więc czynnością, w wyniku której osiągamy pożądany rezultat². Z pojęciem sterowania związany jest szczególny przypadek sterowania - **regulacja**. W wyniku jej działania utrzymywane jest określona wartość regulowanego parametru³. Ogólnie można powiedzieć, że **każda regulacja jest sterowaniem, ale nie każde sterowanie jest regulacją**.

W niniejszym rozdziale przedstawiona zostanie problematyka sterowania w szerszym ujęciu, a więc z uwzględnieniem regulacji. Przedstawiona zostanie inżynierska klasyfikacja sposobów sterowania, przy czym szczególną uwagę zwrócono na sterowanie numeryczne CNC (*Computer Numerical Control*). W końcowej części rozdziału przedstawiono szereg przykładów praktycznego zastosowania układów sterowania w różnych gałęziach gospodarki.

3.2. Sposoby sterowania

Wraz z rozwojem automatyzacji powstawały różne koncepcje sterowania, które związane były zarówno z wymaganiami obiektu sterowania, a także z możliwościami technicznymi w ówczesnym okresie. Stąd sposoby sterowania można uporządkować stosując różne kryteria podziału.

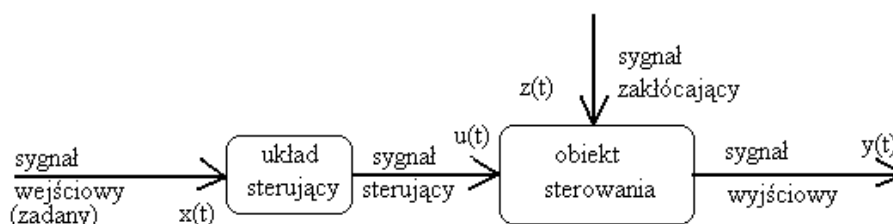
Pierwszy i zarazem najistotniejszy, zarówno z punktu widzenia eksploatacyjnego, jak i technicznego, jest sygnalizowany już w poprzednim rozdziale podział na **sterowanie w układzie otwartym** (bez sprzężenia zwrotnego) oraz **sterowanie w układzie zamkniętym** (ze sprzężeniem zwrotnym **zwane również regulacją** (rys. 1).

¹ Można stwierdzić, że automatyzacja działa na poziomie *makro* i dotyczy całego systemu. Natomiast sterowanie działa na poziomie *mikro*, a więc poszczególnych elementów tego systemu.

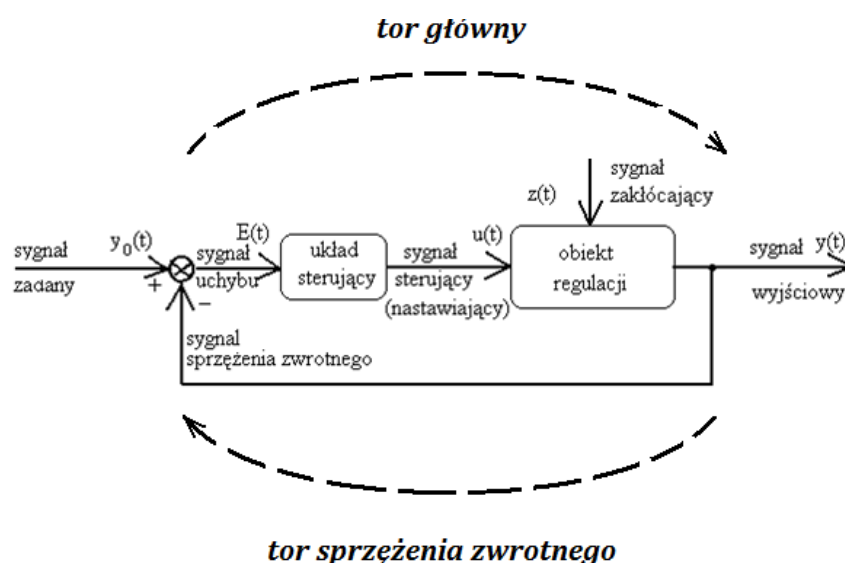
² Przykładowo sterowaniem jest zmiana temperatury w pomieszczeniu. Przykręcając zawór kaloryfera (zdalnie lub manualnie), możemy zmieniać temperaturę w pomieszczeniu (liczymy na to, że w pomieszczeniu będzie utrzymana określona temperatura).

³ Nawiązując do poprzedniego przykładu, w przypadku regulacji temperatury w pomieszczeniu musimy określić jaką konkretnie temperaturę zamierzamy utrzymać. A więc przykładowo ustawiając przy pomocy termostatu (czujnika, sensora) temperaturę 21°C układ ogrzewania będzie zmierzał cały czas do utrzymania zadanej temperatury.

Schemat sterowania w układzie otwartym (bez sprzężenia zwrotnego)



Schemat sterowania w układzie zamkniętym (ze sprzężeniem zwrotnym) - regulacja



Rys. 1. Schemat sterowania w układzie otwartym i zamkniętym (oprac. własne)

Istotną różnicą w przypadku obu sposobów sterowania jest **sprzężenie zwrotne**. Jest to informacja o wyjściu układu, która wraca na jego wejście i wpływa na dalsze działanie układu⁴.

W każdym przypadku określany jest uchyb (błąd), a więc różnica między wartością zadaną, a rzeczywistą. Jeśli w trakcie pracy układu sterowania układ się stabilizuje, a więc zmniejszana jest różnica między wartością zadaną a rzeczywistą mówimy o **sprzężeniu zwrotnym ujemnym** uznawanym za najważniejsze w automatyce. W przypadku sprzężenia

⁴ W przypisie 3 przedstawiono przykład regulacji, a więc sterowania w układzie zamkniętym. W pomieszczeniu (przypis 3), chcemy utrzymać stałą wartość temperatury 21°C. Przez cały czas dokonujemy pomiaru temperatury, a informacje o jej wartości przesyłamy do elektrozaworu zamontowanego na kaloryferze. Jeśli temperatura nie osiąga zadanej wartości (w przykładzie 21°C) zawór jest otwarty. Występuje tzw. uchyb między wartością zadaną a faktyczną. Jeśli temperatura wynosi 21°C elektrozawór jest zamknięty.



zwrotnego dodatniego różnica między wartością zadaną a rzeczywistą wzrasta. Odchylenia są wzmacniane, co może prowadzić do niestabilności układu.

Kolejny podział sposobów sterowania związany jest ze sposobem podejmowania decyzji. W przypadku, jeśli decyzję podejmuje człowiek (np. kierowca, operator maszyny) mówimy o **sterowaniu ręcznym**. Jeśli decyzje podejmuje układ sterowania, który najczęściej wyposażony jest w opracowany algorytm sterowania, mówimy o **sterowaniu automatycznym**. Ze względu na swoją specyfikę, sterowanie CNC zaliczane jest do grupy sterowania automatycznego.

Innym sposobem podziału systemów (sposobów) sterowania jest podział uwzględniający rodzaj algorytmu sterującego. W takim przypadku można mówić o:

- sterowaniu dwustanowym (on/off),
- sterowaniu PID,
- sterowaniu adaptacyjnym,
- sterowaniu predykcyjnym,
- sterowaniu rozmytym (fuzzy).

Sterowanie dwustanowe stanowi najprostszy sposób sterowania, w którym sygnał sterujący może przyjmować tylko dwa stany 0/1 (włącz/wyłącz). Klasycznym przykładem jest omawiany przykład ogrzewania pomieszczenia (przypis 3 i 4). W przypadku, jeśli temperatura jest niższa niż przykładowe 21°C stan „włącz”, jeśli wyższa – stan „wyłącz”. W takim przypadku w pobliżu wartości zadanej następuje co chwila zmiana stanu z 0 na 1 i odwrotnie – a więc włączanie i wyłączanie sterowania. Aby zapobiec tej „skakance” i tym samym ograniczyć możliwość awarii układów sterowania, stosuje się „martwą strefę” w której stan się nie zmienia. Nazywana jest ona **histerezą**. Przykładowo, gdy temperatura w pomieszczeniu $T < 21^{\circ}\text{C}$ stan „1” – włączony. Gdy temperatura $T > 23^{\circ}\text{C}$ stan „0” – wyłączony. W zakresie temperatur $21^{\circ}\text{C} < T < 23^{\circ}\text{C}$ stan układu sterowania się nie zmienia – jest to wspomniana wyżej martwa strefa.

Nazwa **sterowanie PID** pochodzi od trzech członów regulatora: P – proporcjonalny, I – całkujący oraz D – różniczkujący. **Człon proporcjonalny P** reaguje natychmiastowo na stwierdzony uchyb, przy czym im większy błąd tym większa reakcja. **Człon całkujący I** sumuje błędy w czasie, dzięki czemu jego reakcja na uchyb jest stosunkowo „łagodna”. **Człon D – różniczkujący** reaguje na szybkość zmian błędu. W przypadku dużych zmian działa jak



hamulec. Sterowanie PID w porównaniu do sterowania dwustanowego jest płynne i dokładne, co powoduje, że stosowane jest przede wszystkim w układach CNC.

W przypadku sterowania dwustanowego bądź też sterowania PID zakładamy, że zachowanie obiektu jest ściśle określone i stałe. Odwołując się do przykładowego ogrzewania w pomieszczeniu, zakładamy, że temperatura w tym pomieszczeniu ma być stała i wynosić tak jak w przykładzie 21°C . Jednakże niekiedy wymagania są inne. Przykładowo zakładamy, że temperatura 21°C dotyczy pomieszczenia w godzinach od 18 do 24. W godzinach nocnych od 0 do 7 rano ma wynosić 16°C , a od 7 do 18 przykładowo 19°C . Układ ma więc określony model odniesienia, czyli informacje jak ma się zachowywać w czasie. W przykładzie parametry pracy regulatora nie są więc stałe, lecz zmienne w funkcji czasu. Sterowanie takim obiektem odbywa się według zasady **sterowania adaptacyjnego**, a parametry są samoczynnie dostosowane przez układ w trakcie jego pracy na podstawie modelu (tak jak w przykładzie) lub bezpośredniej identyfikacji obiektu. W przypadku bezpośredniej identyfikacji układ sterowania ocenia czy przykładowo w pomieszczeniu przebywają ludzie (wówczas temperatura w pomieszczeniu jest wyższa), lub czy pomieszczenie jest puste (wówczas ogrzewanie jest wyłączane). Sterowanie adaptacyjne można więc potraktować jako „inteligentny” sposób sterowania.

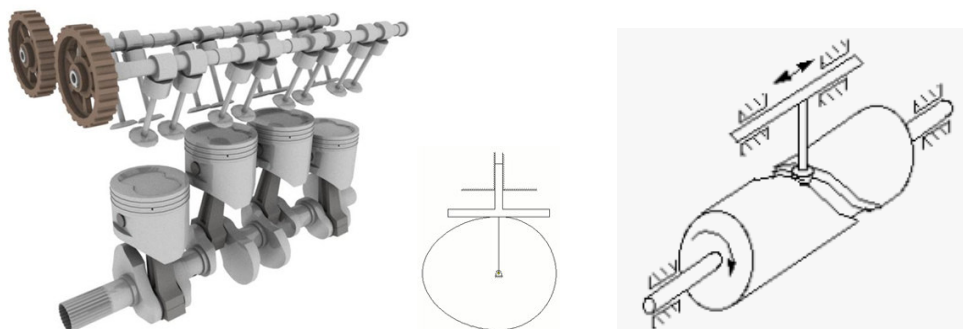
Kolejnym sposobem sterowania jest **sterowanie predykcyjne MPC⁵** (*Model Predictive Control*). System sterowania sprawdza co się stanie za chwilę i już teraz wybierana jest najlepsza decyzja. Jednakże w takim przypadku konieczne jest opracowanie modelu, który przewiduje przyszłe zachowanie obiektu w czasie. Uzupełnieniem sterowania predykcyjnego może być **sterowanie rozmyte (fuzzy)**. Nie tyle przewiduje ono przyszłe zachowanie obiektu sterowania, ile bazuje na regułach podobnych do ludzkiego myślenia: „jeżeli dzieje się to – należy zareagować w taki sposób”. Sterowanie to stosowane jest w wielu urządzeniach gospodarstwa domowego AGD, motoryzacji (komputery pokładowe) i wielu innych obiektach.

Innym kryterium podziału sposobów sterowania może być **rozwiązanie konstrukcyjne układu sterowania** obiektem. W tym aspekcie można mówić o układach sterowania umożliwiających automatyzację sztywną (*fixed automation*) lub automatyzację elastyczną.

⁵ Predykcja to po prostu przewidywanie przyszłych stanów lub zdarzeń na podstawie danych z przeszłości, aktualnych pomiarów lub modelu zjawiska. Najkrócej jest to próba odpowiedzi na pytanie „co będzie dalej”, lecz predykcja nie jest zgadywaniem. Przykładowo jedziesz samochodem – widzisz zakręt i zwalniasz wcześniej nie czekając aż auto wypadnie z drogi. To jest właśnie predykcja.

W przypadku **automatyzacji sztywnej** układ wykonuje ściśle określone, niezmiennie czynności według raz zaprojektowanego schematu. Przykładem może być linia do produkcji butelek. Zawsze produkowana jest ta sama butelka, według tej samej kolejności operacji, a to zapewnia olbrzymią wydajność. Jakakolwiek zmiana np. kształtu butelki jest niemożliwa lub wymaga kosztownej przebudowy. W przypadku **automatyzacji elastycznej** system sterowania można względnie łatwo dostosować do zmian produktu, procesu lub warunków pracy – bez przebudowy sprzętu, a wszystkie zmiany realizuje się programowo. Przykładem może być sterowanie numeryczne centrum obróbkowego CNC. Dzisiaj produkuje ono część „A”, jutro część „B”, a zmiany związane z tą wielką zmiennością produkcji wymagają opracowania a potem wprowadzenia do układu sterowania innego programu. Nie ma natomiast mowy o jakiegokolwiek przebudowie maszyny.

W przypadku automatyzacji sztywnej najczęściej stosowane jest **sterowanie krzywkowe**. Sterowanie jest całkowicie mechaniczne (bez elektroniki), a ruch elementów odbywa się za pomocą **krzywki**, czyli specjalnie ukształtowanego elementu mechanicznego, który podczas obrotu wymusza określony ruch innego elementu np. popychacza (rys. 2).

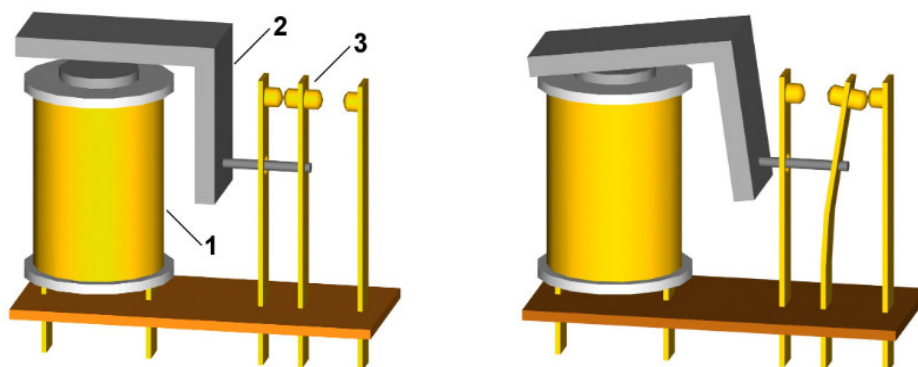


Rys. 2. Przykłady sterowania krzywkowego (oprac. na podst. stron internetowych)

Zaletą sterowania krzywkowego jest bardzo duża dokładność powtarzalnego ruchu. Natomiast wadą brak elastyczności, a zmiana ruchu wymaga zmiany krzywki. Jest to typowe sterowanie programowe mechaniczne, w którym programem jest zarys krzywki. Jest to również sterowanie otwarte – bez sprzężenia zwrotnego.

Etapem pośrednim między sterowaniem mechanicznym (krzywkowym), a nowoczesnym sterowaniem cyfrowym jest **sterowanie elektromechaniczne**. W przypadku tego sterowania sygnał elektryczny z przycisku bądź czujnika pobudza element wykonawczy (np. przekaźnik),

który przekazuje sygnał do elementu mechanicznego wykonującego określony ruch, np. przesunięcie mechanizmu, zamknięcie zaworu, załączanie silnika (rys. 3).



Rys. 3. Przekładnik sterowania elektromechanicznego (oprac. na podst. stron internetowych)

Zaletą sterowania elektromechanicznego jest prosta budowa i wysoka niezawodność układu sterowania, natomiast zasadniczą wadą duże gabaryty szaf sterowniczych. Sterowanie elektromechaniczne urządzeń stworzyło podstawę automatyzacji elastycznej, chociaż elastyczność tych układów była znacznie ograniczona.

Wynikiem rozwoju technologii półprzewodnikowej w latach 50-tych XX wieku, a więc w okresie III rewolucji przemysłowej, było między innymi **sterowanie numeryczne NC** (*Numerical Control*). Jest to metoda sterowania urządzeniem przy pomocy ciągu liczb zapisanych w **programie**, które określają ruchy i operacje wykonywane przez urządzenie. Dalszy rozwój technologii półprzewodnikowej, a szczególnie komputerów skutkował wprowadzeniem **sterowania numerycznego CNC** (*Computer Numerical Control*). Podstawowa różnica między sterowaniem NC i CNC polegała na przygotowaniu programu sterującego. W sterowaniu NC program przygotowany był na taśmie perforowanej i za pomocą odpowiedniego czytnika wprowadzany do sterowanego urządzenia. W przypadku sterowania CNC program wprowadzany jest bezpośrednio do pamięci komputera „pokładowego” urządzenia. Dzięki temu łatwiejsze i szybsze jest przeprogramowanie sterowanego urządzenia, przykładowo obrabiarki CNC. To właśnie sterowanie CNC stanowi fundament współczesnego przemysłu i koncepcji Przemysł 4.0.



3.3. Technika cyfrowa

Technika cyfrowa to dziedzina techniki zajmująca się działaniem i architekturą (budową) najczęściej półprzewodnikowych układów cyfrowych. W niniejszym rozdziale poruszono dwa powiązane ze sobą zagadnienia dotyczące:

- materiałów i związków półprzewodnikowych oraz syntetycznie technologii półprzewodników,
- podstawowych elementów techniki cyfrowej⁶

3.3.1. Materiały i związki półprzewodnikowe

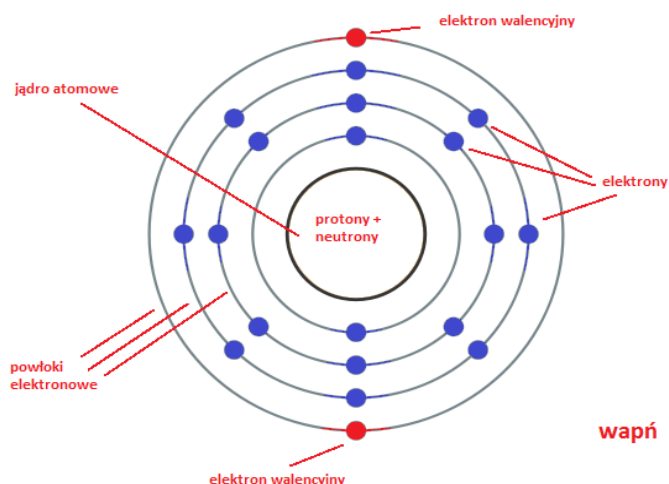
Prąd elektryczny to uporządkowany ruch ładunków elektrycznych. Ilość ładunku elektrycznego przepływającego przez przewodnik w określonym czasie określone jest jako **natężenie prądu** oznaczane literą **I**, a jednostką natężenia prądu jest Amper (A). Im większe natężenie prądu, tym większa ilość przenoszonych ładunków. Lecz aby popłynął prąd elektryczny potrzebne jest źródło napięcia, które wytwarza różnicę potencjałów umożliwiającą ruch ładunku elektrycznego. Tą różnicę potencjałów nazywa się **napięciem U**, a jednostką napięcia jest Volt (V). Jednakże „kluczową kwestią” jest, jak wspomniano na początku ruch ładunków elektrycznych, najczęściej ładunków ujemnych nazwanych **elektronami**.

Podstawowymi elementami, z których zbudowana jest materia⁷ są **atomy**, które składają się z jądra oraz poruszających się wokół niego elektronów mających ładunek ujemny i zajmujące określone poziomy energetyczne (powłoki) – rys. 4⁸.

⁶ Z niektórymi zagadnieniami poruszonymi w niniejszym rozdziale zapoznaliście się już na wcześniejszych zajęciach dydaktycznych. Stąd zaprezentowany materiał będzie po prostu przypomnieniem ułatwiającym zrozumienie zagadnień systemów CNC. Inne zagadnienia stanowią zasygnalizowanie problemów, które rozwinięte zostaną w trakcie dalszych Waszych studiów. Są one jednak również niezbędne do zrozumienia prezentowanego w rozdziale materiału.

⁷ Materiał jest wszystko co ma masę i zajmuje określoną objętość (miejsce w przestrzeni).

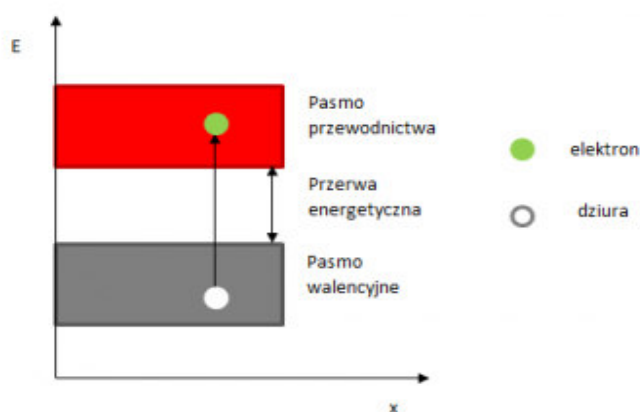
⁸ Współcześnie dzięki osiągnięciom fizyki kwantowej dla niezwykle małych cząsteczek poruszających się z ogromnymi prędkościami, stosuje się kwantowy model atomu. Powstał on dzięki pracy Erwina Schrodingera i Wernera Heisenberga. Założeniem tego modelu jest, że elektrony nie mają określonych orbit (tak jak to zakładał Niels Bohr – rys. 4), lecz znajdują się w chmurze elektronowej i możliwe jest określenie jedynie prawdopodobieństwa ich położenia.



Rys. 4. Model atomu wg Nielsa Bohra (źródło: strony internetowe)

Szczególna rola przypisana jest elektronom znajdującym się na najbardziej zewnętrznej powłoce elektronowej atomu tzw. powłoce walencyjnej. Jest to najdalej położona od jądra powłoka, a elektrony na niej się znajdujące nazywa się **elektronami walencyjnymi**. To one decydują w jaki sposób atom łączy się z innymi atomami (czyli tworzy wiązania chemiczne), jakie ma właściwości chemiczne oraz wpływają na przewodnictwo elektryczne.

Energię elektronów walencyjnych (podobnie elektronów na innych powłokach), opisuje model pasmowy, kluczowy dla zrozumienia elektroniki i półprzewodników (rys. 5).

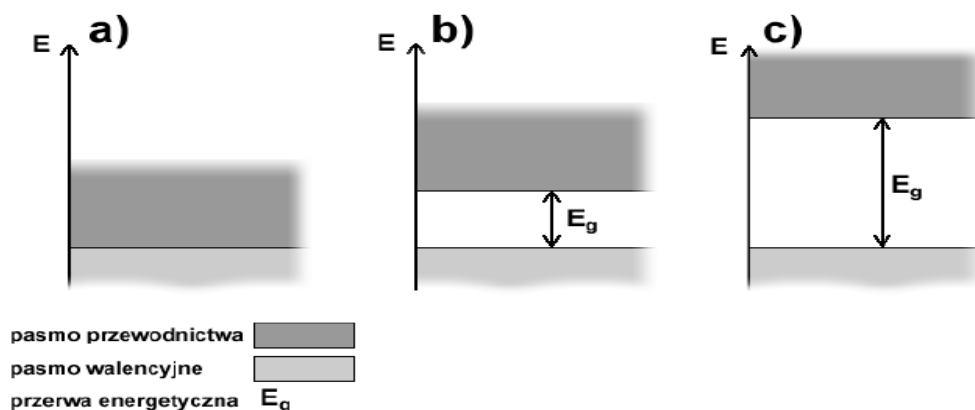


Rys. 5. Energetyczny model pasmowy (źródło: strony internetowe)

Gdy pojedyncze atomy łączą się w ciało stałe ich poziomy energetyczne nakładają się i rozszczepiają, a zamiast pojedynczych poziomów powstają ciągłe zakresy energii zwane **poziomami energetycznymi**. Występują trzy pasma: **pasmo walencyjne** które zawiera elektrony związane z atomami, **pasmo przewodnictwa**, w którym elektrony mogą się

swobodnie poruszać. Odpowiada ono za przewodzenie prądu oraz **pasmo zabronione** zwane również przerwą energetyczną. W tym paśmie elektrony nie mogą się znajdować.

W zależności od nakładania się lub rozszczepiania poziomów energetycznych wyróżnić można: metale, półprzewodniki i izolatory (rys. 6).



Rys. 6. Położenie pasm energetycznych w przypadku metali (a), półprzewodników (b) oraz izolatorów (c)
(źródło: strony internetowe)

W przypadku metali (rys. 6a) nie ma pasma zabronionego, wskutek czego elektrony mogą bez przeszkód przechodzić z pasma walencyjnego do pasma przewodnictwa. Dzięki temu materiały te dobrze przewodzą prąd elektryczny. W przypadku izolatorów (rys. 6c) istnieje pasmo zabronione, a przerwa energetyczna E_g jest duża. Dostarczenie tak dużej energii umożliwiającej „przeskok” elektronów między pasmami walencyjnym i przewodnictwa w praktyce oznacza zniszczenie izolatora. W przypadku półprzewodników (rys. 6b) przerwa energetyczna E_g istnieje, ale jest na tyle mała, że doprowadzenie niewielkiej energii cieplnej, promieniowania świetlnego (fotonów), czy też pola elektrycznego (napięcia) umożliwia „przeniesienie” elektronów z pasma walencyjnego do pasma przewodnictwa, dzięki czemu materiał półprzewodnika zaczyna przewodzić prąd.

Materiały półprzewodnikowe są podstawą do wytwarzania **układów scalonych IC** (*Integrated Circuit*)⁹. Jest to ogólna nazwa, a szczegółowo dotyczy ona procesorów – CPU

⁹ Układ scalony IC to miniaturowy element elektroniczny w którym wiele komponentów (np. tranzystory, rezystory, diody) umieszczonych jest na podłożu, który stanowi materiał półprzewodnikowy. Wszystkie te komponenty tworzą kompletny obwód elektryczny zdolny do samodzielnego działania.



(*Central Processing Unit*), układów pamięci RAM, układów graficznych, wzmacniaczy operacyjnych i innych zwany popularnie **chipami**.

Klasycznym i jednocześnie najtańszym materiałem półprzewodnikowym jest aktualnie **krzem (Si)**. Wykorzystywany jest do produkcji różnych układów scalonych, a jego atutem jest dostępność (to jeden z najpowszechniejszych pierwiastków w przyrodzie), korzystna cena oraz duże doświadczenie technologiczne związane z jego stosowaniem. Istotną wadą jest mniejsza mobilność elektronów, co wpływa na szybkość działania układów scalonych. Alternatywą krzemu są **związki półprzewodnikowe** najczęściej materiałów należących do III i V grupy pierwiastków układu okresowego (rys. 7).

II	III	IV	V	VI
Be	B	C	N	O
Mg	Al	Si	P	S
Zn	Ga	Ge	As	Se
Cd	In	Sn	Sb	Te

Grupa IV: diament, Si, Ge
 Grupy III-V: GaAs, AlAs, InSb, InAs...
 Grupy II-VI: ZnSe, CdTe, ZnO, SdS...

Rys. 7. Fragment tablicy Mendelejewa dotyczący półprzewodników

Z tej grupy materiałów najczęściej stosowany jest **arsenek galu (GaAs)**. Wyższa mobilność elektronów w porównaniu z tradycyjnym krzemem sprawia, że urządzenia oparte na tym związku półprzewodnikowym są szybsze i bardziej efektywne w zastosowaniach technologicznych. Jednak układ scalony wykonany z tego materiału jest znacznie droższy w stosunku do krzemu.

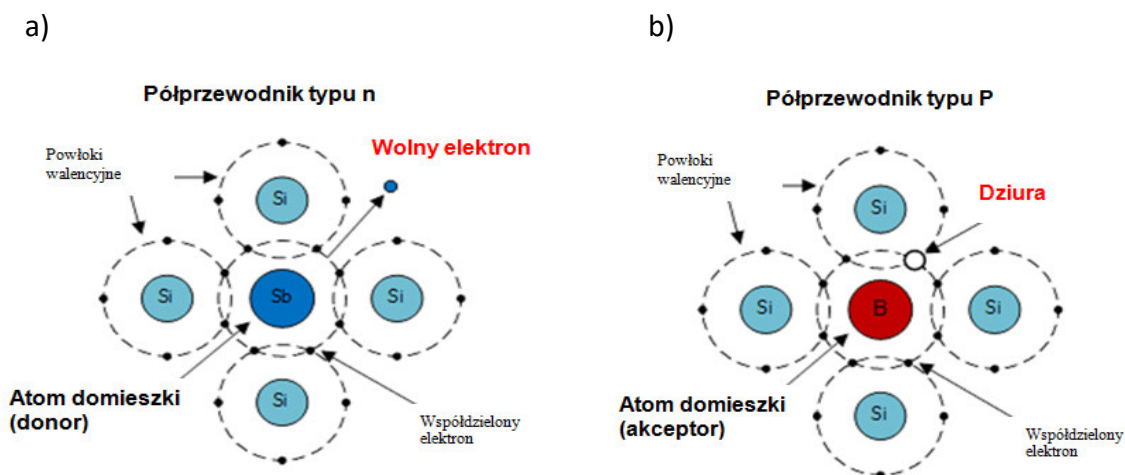
Czysty krzem zwany półprzewodnikiem samoistnym, ma bardzo mało nośników ładunków elektrycznych – słabo przewodzi prąd. Dodając do niego niewielką ilość innych pierwiastków tzw. **domieszek**, tworzymy dodatkowe nośniki ładunku, dzięki czemu przewodnictwo materiału rośnie nawet milion razy. A więcej nośników to łatwiejszy przepływ prądu, a w związku z tym układy działają szybciej przy mniejszych stratach. Proces domieszkowania jest więc procesem kontrolowanego „ulepszania” półprzewodnika.

Podstawowe typy domieszkowania to:

- **typ n** (ujemny), dodajemy atomy z nadmiarem elektronów które są nośnikiem ujemnych ładunków elektrycznych (elektronów),
- **typ p** (dodatni), dodajemy atomy z niedoborem elektronów; powstają tzw. dziury (braki elektronów), które są nośnikami ładunków.

Efektom domieszkowania jest łatwiejsze przejście elektronów, czyli większe przewodnictwo nawet bez dużej energii zewnętrznej. Z tego względu domieszkowanie uważa się za swojego rodzaju energię „wbudowaną” półprzewodnika¹⁰.

Krzem Si jako pierwiastek ma na powłoce walencyjnej 4 elektrony. Jeżeli wprowadzimy do niego domieszkę antymonu (Sb), który na powłoce walencyjnej ma 5 elektronów, to 4 z tych elektronów wykorzystane zostaną w wiązaniach krzemu z antymonem, zaś pozostały elektron stanie się wolnym elektronem stwarzającym warunki do przepływu prądu. W takim przypadku występuje półprzewodnik **typu n**, a atom domieszki nazywany jest **donorem** (rys. 8a).



Rys. 8. Półprzewodniki typu n (a) oraz typu p (b)

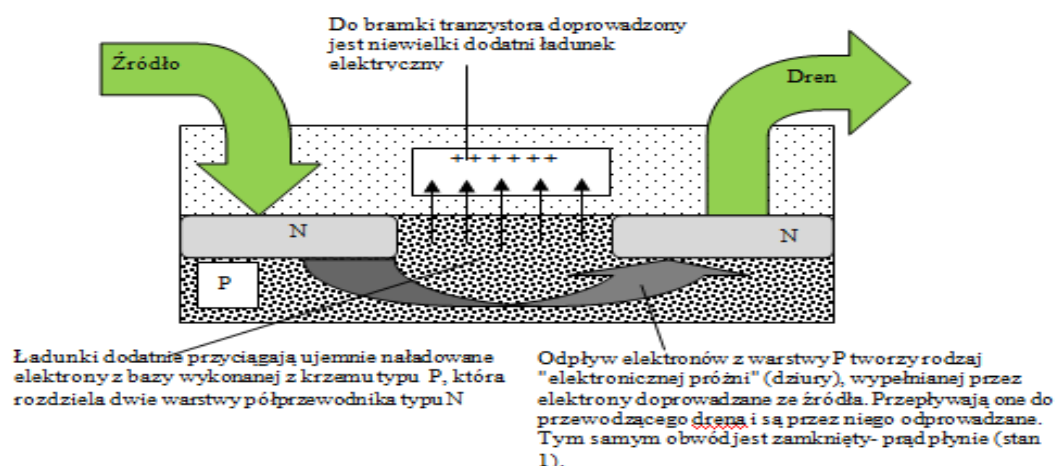
Natomiast jeżeli do krzemu wprowadzimy domieszkę pierwiastka zawierającą na powłoce walencyjnej 3 elektrony walencyjne (np. bor), to w wiązaniu zabraknie 1 elektronu, a w jego miejscu powstanie dziura (rys. 8b). Taki półprzewodnik nazywany jest półprzewodnikiem **typu p**, a atomy domieszki nazywa się **akceptorem**. W półprzewodniku fizycznie przemieszczają się elektrony, jednakże, gdy elektron „przeskakuje” do dziury, to zostawia po sobie kolejną dziurę – czyli sprawia to wrażenie przemieszczania się dziury, która

¹⁰ Wyobraźcie sobie, czysty krzem jako drogę z bardzo małą liczbą samochodów (ładunków). Przemieszczanie między punktami jest więc ograniczone. Domieszkowanie to „dodanie” nowych pojazdów – czyli elektrony są jak nowe auta, a dziury jako wolne miejsca dla poruszających się aut.

zachowuje się jak cząstka o ładunku dodatnim. O ile w półprzewodniku typu n nośnikami ładunku są elektrony, to w półprzewodniku typu p – dziury.

W elektronice szczególne znaczenia mają złącza między dwoma obszarami półprzewodnika (złącza p-n lub n-p). To podstawowy element działania przede wszystkim tranzystorów i innych elementów tworzących układ scalony. Złącze działa jak zawór dla prądu; przepuszcza prąd w jedną stronę, blokuje w drugą¹¹. Właśnie takim zaworem, elementem przełączającym jest w układach elektronicznych **tranzystor**. Istota działania tranzystora przedstawiona jest na rys. 9, a jego głównym zadaniem jest kontrolowanie przepływu prądu elektrycznego – czyli umożliwienie sterowania dużym sygnałem za pomocą małego. Ogólnie mówiąc służy do sterowania, wzmacniania i przełączania sygnałów elektrycznych.

Tranzystor jest podstawowym elementem z którego zbudowane są wszystkie układy scalone. Tworzy on wyłącznie informację binarną: 1 - jeśli prąd płynie, 0 - jeśli prąd nie płynie. Z tych zer i jedynek, zwanych bitami, komputer może utworzyć dowolną liczbę, zakładając że jest wystarczająco wiele zgrupowanych razem tranzystorów. Schemat tranzystora przedstawia rysunek.



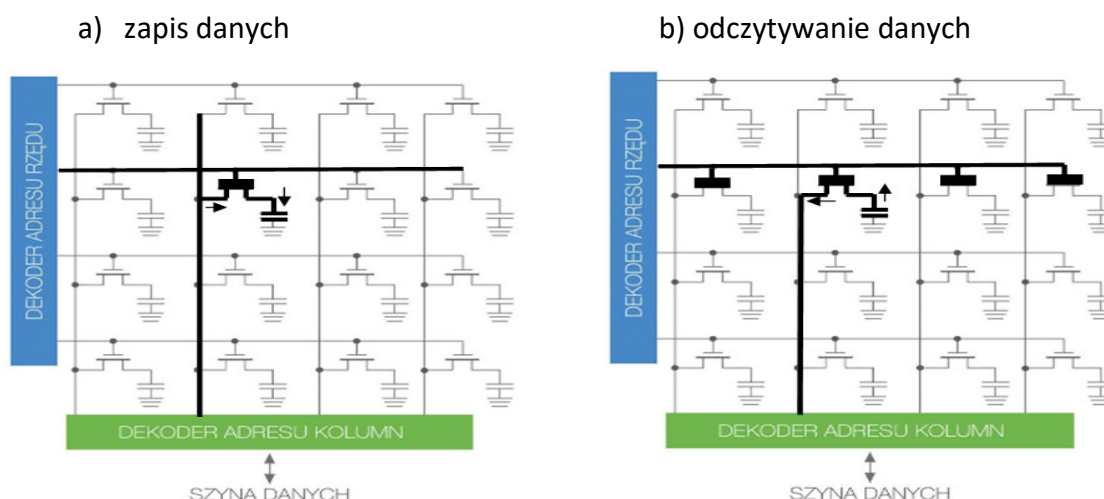
Jeżeli do bramki tranzystora doprowadzony zostanie ładunek ujemny nie powstaje "elektroniczna próżnia" i elektrony źródła są odpychane, a tranzystor jest wyłączony (stan 0).

Rys. 9. Istota działania tranzystora

Tranzystory są podstawowym elementem pamięci operacyjnej komputera RAM (*Random Access Memory*), w której przechowywane są w postaci zero – jedynkowej, tymczasowe dane i programy aktualnie używane przez procesor¹². Schemat zapisywania i odczytywania danych z pamięci operacyjnej przedstawia schemat na rys. 10.

¹¹ Można to przyrównać do zapory na rzece – w jedną stronę przepływ jest możliwy w drugą nie.

¹² Poglądowo – pamięć RAM działa tak jak „biurko robocze” dla procesora. Im większa powierzchnia biurka, tym więcej czynności można na nim wykonać. A więc większa pamięć RAM komputera, tym większe jego możliwości



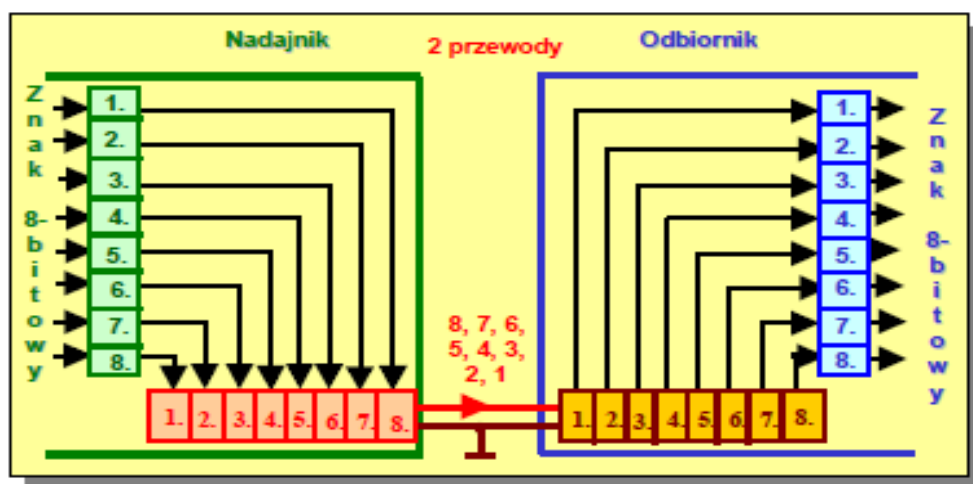
Rys. 10. Istota działania pamięci operacyjnej

Pamięć operacyjną przyrównać można do powszechnie znanego arkusza kalkulacyjnego którego komórkami są miliony, czy dzisiaj nawet miliardy tranzystorów wykonanych z materiałów półprzewodnikowych. Każda komórka przechowuje 0 lub 1 (bit informacji). Dotarcie do określonej komórki arkusza kalkulacyjnego jest możliwe przez podanie jej adresu. Podobnie jest z dostępem do pamięci.

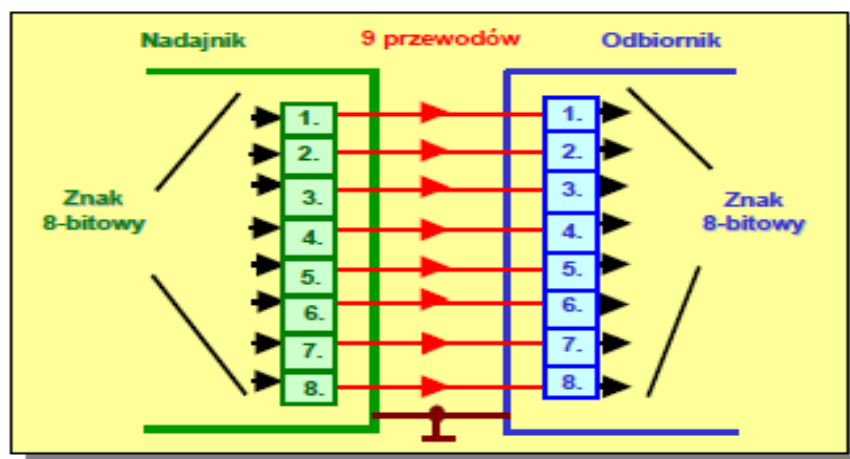
W przypadku zapisywania danych aktywizowana jest przez procesor magistrala adresowa, która wskazuje w którym miejscu ma być zapisana informacja. Równocześnie aktywizowana jest również magistrala danych, którą przekazywane są właściwe dane - zera lub jedynki. Innymi słowy **magistrala adresowa** wskazuje, gdzie będzie coś zapisane, **magistrala danych** natomiast – co zapisujemy. Magistrale te nazywane są również szynami. W przypadku zapisu (rys. 10a – pogrubione linie), **baza** tranzystora spolaryzowana jest dodatnio. Tranzystor przewodzi prąd do drena (**kolektora**), którego zakończenie stanowi kondensator. Magazynuje on ładunek elektryczny „1”. W przypadku pozostałych komórek (tranzystorów), kondensatory nie są ładowane, gdyż albo do bazy, albo do źródła zwanego również **emiterem** nie podłączony jest prąd. Tranzystor nie przekazuje ładunku do kondensatora. W przypadku odczytu danych (rys. 10 b – pogrubione linie), spolaryzowana dodatnio baza powoduje przepływ ładunku z kondensatora do magistrali danych.

Przesyłanie sygnałów magistralami nazywane jest **transmisją sygnałów**. W układach cyfrowych stosowane są dwa podstawowe sposoby przesyłania danych: **transmisja szeregową** (*serial*) oraz **równoległą** (*paralel*).

a)



b)



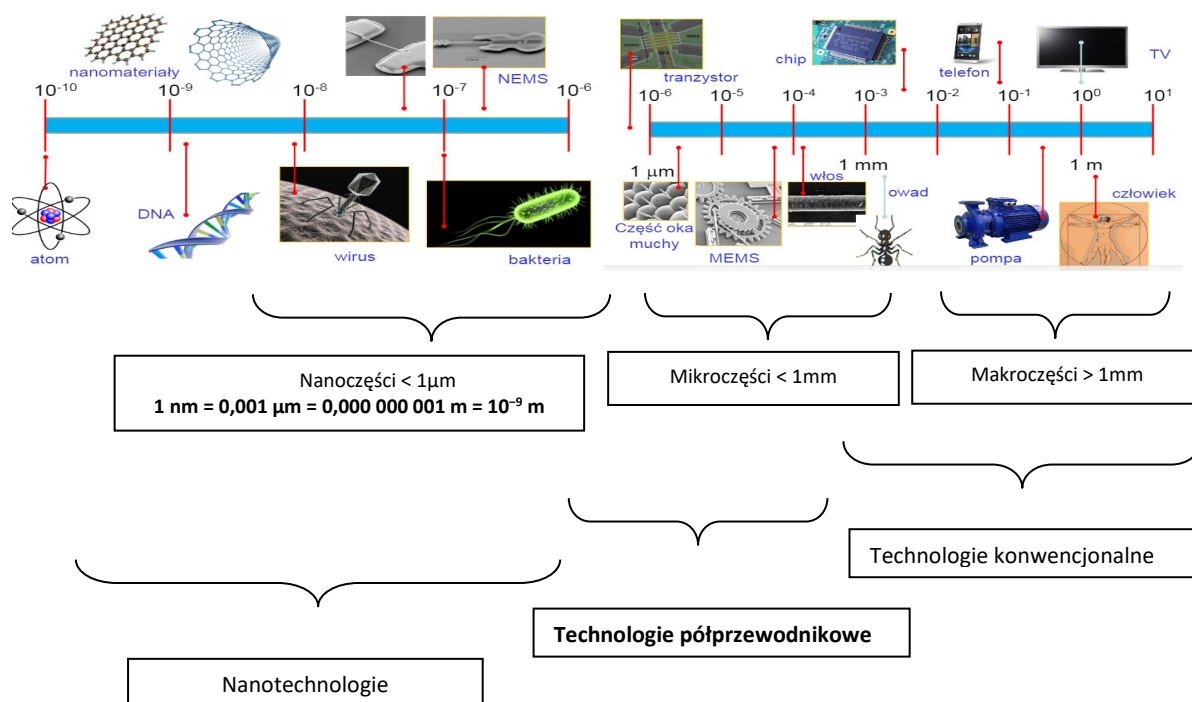
Rys. 11. Transmisja sygnałów szeregowa (a) i równoległa (b)

W przypadku transmisji szeregowej (rys. 11a), dane przesyłane są bit po bicie, jeden za drugim, jedną linią transmisyjną. Odbywa się to małą liczbą przewodów (najczęściej jeden lub dwa przewody), więc teoretycznie przesyłanie informacji jest wolniejsze, ale odległości między nadajnikiem i odbiornikiem mogą być większe. W przypadku transmisji równoległej (rys. 11b), wiele bitów przesyłanych jest jednocześnie, każdy osobną linią (linia zawiera wiele przewodów np. 8, 16, 32). Umożliwia to zwiększenie szybkości przesyłania danych szczególnie na krótkich odległościach.

W nowoczesnych systemach cyfrowych **transmisja szeregowa wyparła równoległą**, bo mimo przesyłania bitów po kolei, osiąga wyższe prędkości dzięki wysokim częstotliwościom i mniejszym zakłóceniom.

3.3.2. Technologia półprzewodnikowa (krzemowa)

Technologię należy rozumieć w uproszczeniu jako „sposób wykonania”. Przykładowo może to być sposób wykonania zespołu, podzespołu czy też części produktu. Zasadniczym problemem w technologii półprzewodnikowej jest wymiar gabarytowy wykonywanego elementu. O ile w klasycznych technologiach mamy do czynienia z elementami dużymi, a więc takimi które możemy zobaczyć, to w przypadku technologii półprzewodnikowej mamy do czynienia z elementami bardzo małymi (np. tranzystory), niewidocznymi gołym okiem. Z tego względu praktycznie niemożliwa była adaptacja technologii klasycznych takich jak przykładowo obróbka skrawaniem, do wykonania elementów półprzewodnikowych. Taką technologię należało „stworzyć od zera”¹³ (rys. 12).



Rys. 12. Wymiary elementów wytwarzanych w ramach poszczególnych technologii

Właśnie to wymiary produkowanych elementów i precyzja ich wykonania spowodowały, że zastosowanie tradycyjnych technologii stało się niemożliwe. Wykonywana część była mała, a do jej obróbki należało stworzyć odpowiednio małą obrabiarkę, a także jeszcze mniejsze narzędzia. Ten kierunek rozwoju technologii okazał się niemożliwy. Należało

¹³ Jeszcze w latach 70-tych XX wieku „szczytem precyzji” był mechanizm zegarka sprężynowego.



opracować zupełnie nowe technologie. Natomiast już nanotechnologia, zmierzająca do wytwarzania jeszcze mniejszych wyrobów, wykorzystuje i rozwija metody obróbki stosowane w technologiach półprzewodnikowych.

Elementy układów scalonych wykonywane są zazwyczaj na podłożu, którym jest materiał półprzewodnikowy. Stąd pierwszą operacją procesu produkcji układów scalonych jest **wykonanie półprzewodnikowych płytek podłożowych** zwanych waflami. Następną operacją jest **wprowadzenie domieszek** zmieniających właściwości płytki podłożowej w określonych jej fragmentach. Następuje kształtowanie elementów układu scalonego według zaprojektowanej uprzednio architektury. Końcowym etapem produkcji jest **mikromontaż**. Polega on na połączeniu struktury półprzewodnikowej z obudową oraz wyprowadzeniami elektrycznymi.

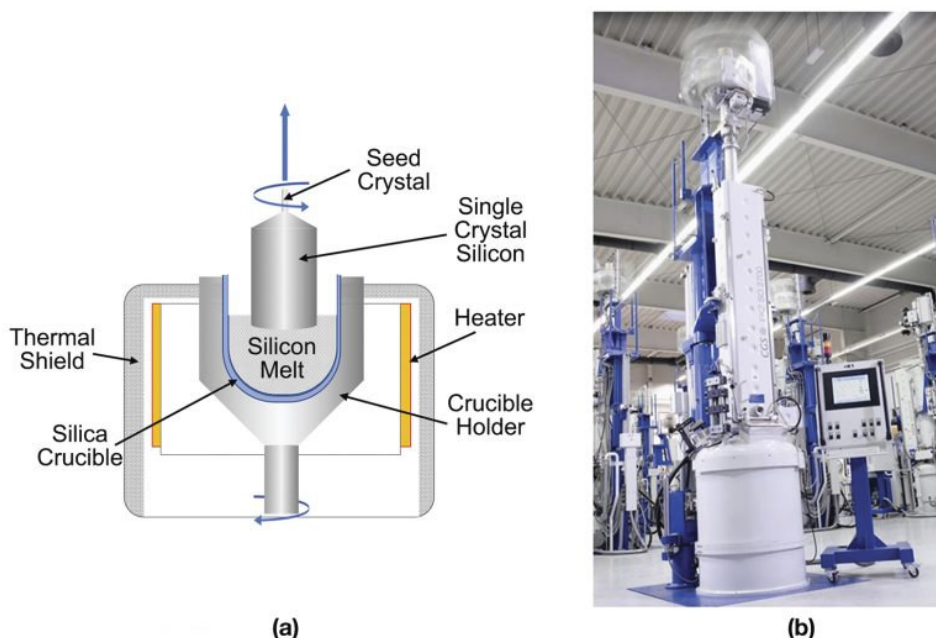
Jak już wspomniano w poprzednim rozdziale w przypadku, kiedy stosowanym do produkcji układu scalonego półprzewodnikiem jest krzem, należy go odpowiednio przygotować¹⁴, gdyż jakość wykonanego układu scalonego zależy od ilości zanieczyszczeń zawartych w półprzewodniku. Opóźniają one ruch elektronów, a to skutkuje gorszymi parametrami pracy układu scalonego. Stąd zabiegi technologiczne związane z wykonaniem płytek podłożowych (wafla) na których „budowane” są układy scalone, rozpoczynają się od przygotowania materiału wyjściowego, którym jest **monokryształ krzemu**.

Monokryształ wytwarzany jest dwuetapowo. Pierwszym etapem jest ekstrakcja (wyodrębnianie) krzemu (Si) z krzemionki (SiO_2), która jest podstawowym składnikiem piasku. Jest to szereg reakcji chemicznych, których końcowym efektem jest wytworzenie krzemu polikrystalicznego zwanego **polikrzemem**. Jest to zlepek małych kryształków krzemu (monokryształów) o wymiarach od 1 do 100 μm . Jego czystość zależy głównie od jakości surowca i przyjętej metody oczyszczana. W przypadku polikryształów krzemu zazwyczaj możliwe jest uzyskanie krzemu o czystości 99,999%. Polikrzem wykorzystywany jest bezpośrednio do produkcji ogniw fotowoltaicznych, a jego czystość jest wystarczająca do osiągnięcia wysokiej efektywności konwersji energii słonecznej, ale nie wystarczająca w przypadku produkcji układów scalonych. W takim przypadku konieczne jest uzyskanie krzemu o wyższej czystości 99,9999%. Czystość monokryształów oraz brak granic między ziarnami ułatwia przepływ elektronów, zapewniając efektywniejszą pracę układu scalonego.

¹⁴ Technologie półprzewodnikowe, które w procesie produkcji wykorzystują półprzewodnik krzemowy, nazywane są również technologiami krzemowymi.

Jednym z najczęściej stosowanych procesów otrzymywania monokryształów krzemu jest **metoda Czochralskiego**¹⁵. Polega na zanurzeniu zarodka kryształu krzemu (stanowi on rodzaj wzorca) w ciekłym krzemie. Zarodek obraca się i jest powoli wyciągany z kąpeli, co pozwala na stopniowe krystalizowanie się materiału wzdłuż osi wyciągania (rys. 13).

Wytworzony metodą Czochralskiego monokryształ zwany również ingotem ma długość ok. 2 metrów i jest cięty na szereg (w granicach 3000 do 5000 szt.) płytek stanowiący podłoże wytwarzanego układu scalonego. Grubość wafli półprzewodnikowych zależy w głównej mierze od ich przeznaczenia, przy czym najczęściej wynosi od 200 do 800 μm (0,2 do 0,8 mm).



Rys. 13. Schemat procesu Czochralskiego (a) oraz urządzenie technologiczne (b)

[źródło: <https://pl.dsolar.com/>]

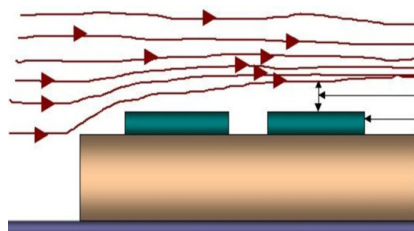
Po dalszych zabiegach związanych z zapewnieniem płytce podłożowej odpowiednich parametrów jakościowych, rozpoczyna się kolejna operacja dotycząca wytworzenia niezwykle małych elementów układu scalonego. Ich wymiary gabarytowe aktualnie oscylują

¹⁵ Jan Czochralski polski chemik i metaloznawca, urodził się w 1895 roku w Kcyni. Do II wojny światowej pracował na uczelniach w Berlinie, potem na Politechnice Warszawskiej zajmując się wytwarzaniem monokryształów. Po wojnie władze PRL negatywnie oceniły działania naukowca. Wycofał się z życia naukowego – zmarł w 1953 roku w Poznaniu. W Dolinie Krzemowej jego nazwisko jest słabo rozpoznawalne, natomiast metoda otrzymywania monokryształów zaadoptowana do wytwarzania monokryształów krzemu, stanowi podstawowy proces wielu współczesnych gigantów technologii krzemowej.

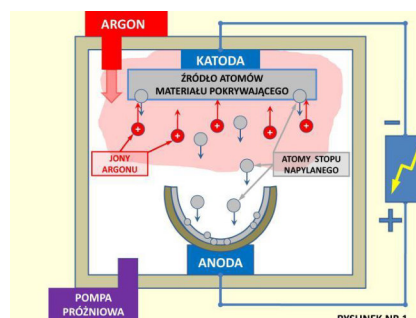
w granicach od kilku do kilkunastu nm. Odbyna się to przez domieszkowanie płytki półprzewodnika. Proces ten polega na wprowadzeniu do czystego krzemu (lub innego półprzewodnika), śladowych ilości innych pierwiastków w celu zmiany jego właściwości elektrycznych. Powstałe, wskutek tego procesu warstwy zawierają głównie domieszki donorowe (typu n) oraz domieszki akceptorowe (typu p). Są to tzw. **cienkie warstwy**, których nałożenie nie powoduje naruszenie struktury podłoża. Technika cienkowarstwowa umożliwia modyfikację powierzchni, a tym samym sterowanie własnościami mikroukładów elektronicznych.

Najczęściej stosowanymi technikami tworzenia cienkich warstw są metody: CVD (*Chemical Vapour Deposition*) - chemiczne metody nakładania warstw oraz PVD (*Physical Vapor Deposition*) - fizyczne metody nakładania warstw. **Metoda CVD** (rys. 14a), polega na osadzeniu warstwy z fazy gazowej przepływającego gazu w temperaturze nieco powyżej 450°C (dla nanowarstw).

a)



b)



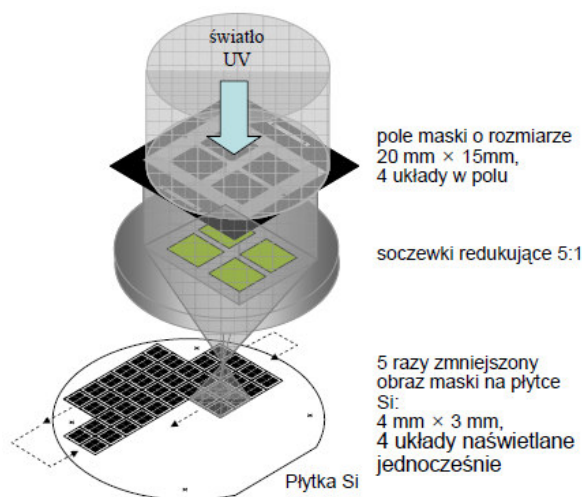
Rys. 14. Schemat metody CVD (a) PVD (b)

Warstwa tworzy się na skutek reakcji chemicznych między płytką podłożową na której warstwa jest osadzana, a przepływającym gazem. W przypadku metody **fizycznego osadzania warstwy PVD** (*Physical Vapor Deposition*), warstwa powstaje z fazy gazowej, a atomy materiału warstwy za pomocą pola elektrycznego lub magnetycznego przemieszczane są w próżni między źródłem atomów (tzw. targetem), a miejscem naniesienia warstwy.

Podstawowym problemem jest precyzyjne określenie miejsca, w którym warstwa powinna być nałożona. Odbyna się to w **procesie litografii**, który uważany jest za najważniejszy proces w cyklu produkcji chipów. Polega on na naniesieniu na powierzchnię płytki podłożowej topografii (wzoru) układu scalonego określonego maską wykonaną na

podstawie opracowanego projektu. Jednakże wzór ten jest odpowiednio pomniejszony, a to powoduje, że na jednym waflu znajdują się tysiące lub nawet setki tysięcy układów scalonych. W układzie scalonym znajdują się natomiast niekiedy miliardy maleńkich tranzystorów, bramek elektrycznych które włączają się i wyłączają w celu wykonania obliczeń. To właśnie w nanometrach określa się wymiary pojedynczego tranzystora. Przykładowo w roku 2025 firma TSMC z Taiwanu przygotowuje się do wdrożenia technologii 2nm¹⁶, jednakże sukces tego przedsięwzięcia zależy od Holenderskiej firmy ASML – wytwórcy urządzeń do tak precyzyjnej litografii.

Schemat procesu litografii przedstawia rys. 15.



Rys. 15. Schemat procesu litografii

Proces ten wymaga wykonania tzw. maski, a więc zaprojektowanego wzorca (topografii), który ma być odwzorowany na odpowiednio przygotowanej krzemowej płytce podłożowej. Maską jest to najczęściej płytka szklana z naniesionym na jednej stronie wzorem. Wzór jest to warstwa metalu obrazująca ścieżki i okna mikrosystemu. Wymiary maski są tak dobrane, aby można było na niej szczegółowo odwzorować wymagania dotyczące określonych miejsc, które następnie będą przykładowo domieszkowane (w przykładzie dydaktycznym na rys. 15 maska ma wymiary 20 x 15 mm). Tak przygotowana maska jest naświetlana np. promieniami ultrafioletowymi. Układ optyczny urządzenia do litografii zmniejsza obraz maski. Przykład na rys. 15 przedstawia 5-krotne zmniejszenie wymiarów maski, a więc odwzorowanie nastąpi na

¹⁶ Dla przykładu wymiar atomu to ok. 0,1 nm, cząsteczka białka 5 do 10 nm, wirus 30 do 100 nm, a komórka 10000 nm, czyli 10 μ m. Generalnie rząd 10 nm to poziom, na którym zaczyna się **nanotechnologia**.

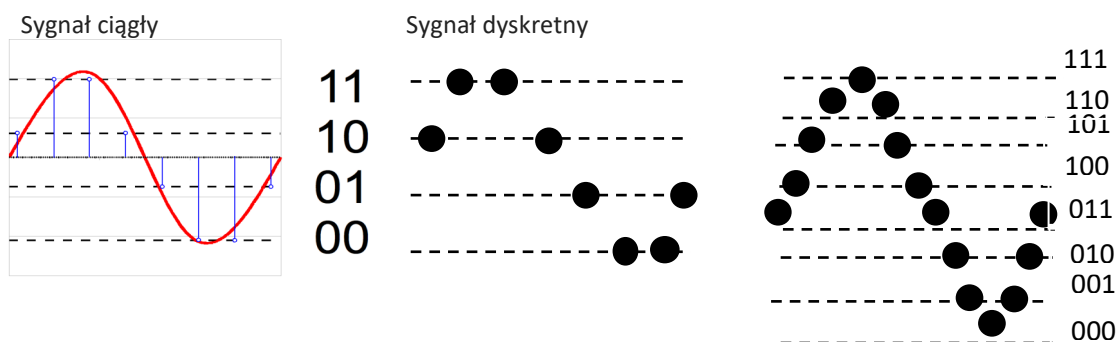


powierzchni 4 x 3 mm. Ponieważ jednocześnie odwzorowane były 4 elementy, stąd wymiary pojedynczego odwzorowanego elementu to 1,0 x 0,75 mm. Przesuwając płytkę podłożową i powtarzając naświetlanie, można umieszczonym na masce wzorem pokryć całą płytkę wafla. A więc na jednej płytce podłożowej wykonywanych są setki tysięcy a nawet miliony układów scalonych. Opisany proces "przenoszenia" wzorca maski na płytkę podłożową jest niepełny. Energia promieniowania np. ultrafioletowego jest zbyt mała, aby na twardej płytce krzemowej utrwalić wzorzec maski. Z tego względu przed naświetlaniem płytka podłożowa musi być pokryta warstwą materiału zwanego **fotorezystemem**, który na skutek tego naświetlania zmienia właściwości (np. łatwiej go rozpuścić). Naświetlony fotorezyst jest usuwany w procesie trawienia, a nienaświetlony pozostaje na płytce stanowiąc ochronę przed dostępem atomów domieszki doprowadzanych w procesie domieszkowania.

Kończącym etapem procesu wytwarzania układów scalonych jest **proces dicingu**, zwany potocznie kostkowaniem. Jak wspomniano na płytce podłożowej znajduje się wiele wytworzonych układów scalonych. Z tego względu płytka ta musi zostać pocięta na indywidualne chipy, które następnie oprawiane są w metalową obudowę posiadającą doprowadzenia i wyprowadzenia prądowe umożliwiające wykorzystanie układu scalonego zgodnie z jego przeznaczeniem.

3.3.3. Podstawowe elementy techniki cyfrowej

Technika cyfrowa to dziedzina zajmująca się przetwarzaniem, przesyłaniem i przechowywaniem informacji w postaci **sygnałów dyskretnych**, czyli takich, które przyjmują określone wartości w układzie binarnych 0 i 1. Tym samym odróżnia się zasadniczo od techniki analogowej, która zajmuje się przetwarzaniem **sygnałów ciągłych**, a więc takich które mogą przyjmować dowolne wartości w pewnym zakresie (rys. 16).

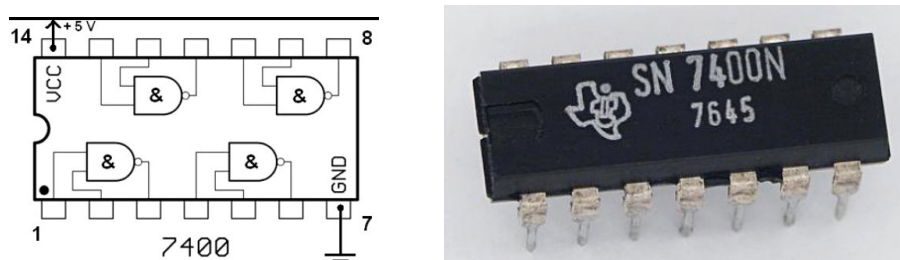


Rys. 16. Porównanie sygnału ciągłego i dyskretnego

Na rys. 16 przedstawiono **proces kwantyzacji**, a więc przekształcania sygnału analogowego na cyfrowy. Polega ona na przypisaniu wartości do określonych poziomów. Zawsze to powoduje pewien błąd, ale im więcej poziomów (bitów), tym dokładniejsze odwzorowanie sygnału. Przykładowo dla 2 bitów mamy $2^2=4$ poziomy, dla 3 bitów mamy $2^3=8$ poziomów (rys. 16), dla 8 bitów wyznaczane są wartości na 256 poziomach, a dla 10 bitów na 1024 poziomach itd. Tak więc sygnał ciągły jest coraz dokładniej odwzorowany sygnałem cyfrowym. Jednocześnie sygnał cyfrowy jest bardziej odporny na zakłócenia.

Przykłady zastosowania techniki cyfrowej przedstawiono już w rozdz. 3.3.1, głównie w odniesieniu do pamięci cyfrowych, których każda komórka pamięci stanowi **przerzutnik**, a więc element cyfrowy mogący osiągnąć dwa stabilne stany 0 lub 1. Można więc go traktować jako „miniaturową pamięć”.

Podstawowym elementem elektroniki cyfrowej są natomiast **bramki logiczne**, które wykonują operacje logiczne na sygnałach binarnych 0/1. Są to zazwyczaj układy scalone (rys. 17), realizujące pewną prostą funkcję logiczną, której argumenty (zmienne logiczne) oraz sama funkcja mogą przyjmować jedną z wartości 0 lub 1.



Rys. 17. Układ cyfrowy 7400 zawierający 4 bramki logiczne NAND

Mówiąc prościej bramka przyjmuje jeden lub więcej sygnałów binarnych (0 – brak napięcia, 1 – obecność napięcia) i na ich podstawie „podejmuje decyzję” jaki ma być sygnał wyjściowy. Najważniejsze bramki logiczne przedstawione są na rys. 18.

Bramki logiczne

(rodzaj, funkcja logiczna, symbol, tablica prawdy)

NOT (negacja):

$$f = \bar{a}$$

a	f
0	1
1	0

EXOR (Exclusive OR, wyłączna suma logiczna):

$$f = a \oplus b$$

a	b	f
0	0	0
0	1	1
1	0	1
1	1	0

AND (iloczyn):

$$f = a \cdot b$$

a	b	f
0	0	0
0	1	0
1	0	0
1	1	1

NAND (NOT-AND, negacja iloczynu):

$$f = \overline{a \cdot b}$$

a	b	f
0	0	1
0	1	1
1	0	1
1	1	0

OR (suma):

$$f = a + b$$

a	b	f
0	0	0
0	1	1
1	0	1
1	1	1

NOR (NOT-OR, negacja sumy):

$$f = \overline{a + b}$$

a	b	f
0	0	1
0	1	0
1	0	0
1	1	0

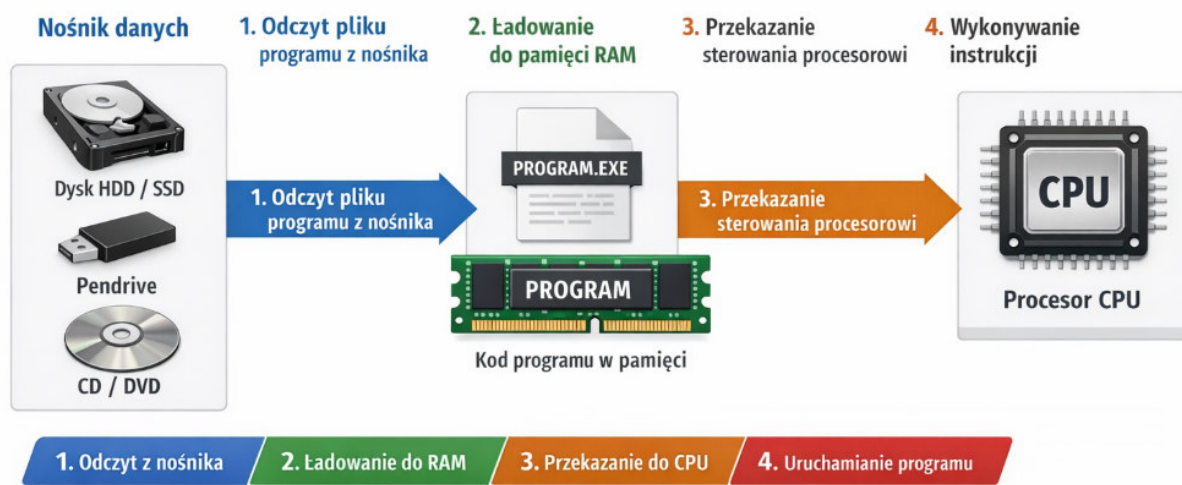
Rys. 18. Tablica podstawowych bramek logicznych

Przykładowo, jeśli na wejściu „a” jest prąd „1” oraz na wejściu „b” załączony jest włącznik „1” – to bramka AND określi, że na wyjściu popłynie prąd „1”. W każdym innym przypadku na wyjściu prąd nie popłynie „0”. Natomiast w przypadku bramki OR wystarczy, że na jednym z wejść „a” lub „b” pojawia się sygnał to wystarczy on do tego, aby na wyjściu popłynął prąd „1”. Układ scalony jest to w zasadzie wiele bramek połączonych ze sobą w funkcjonalną całość (rys. 17). Bramki logiczne są fundamentem działania nie tylko komputerów, ale także wszelkich układów sterowania i automatyki.

3.4. Układy sterowania urządzeń CNC

Współczesne układy sterowania większości urządzeń działają w oparciu o sterowanie w układzie zamkniętym. Układ sterowania musi być wyposażony w pamięć operacyjną RAM, do której wprowadzony jest program sterowania. Aktualnie najczęściej odbywa się to

z poziomu nośników zewnętrznych (dysk, pendrive, CD/DVD) lub poprzez ładowanie z sieci (*network boot*) – rys. 19.



Rys. 19. Wprowadzanie programu do systemu sterowania CNC

W wielu układach sterowania w „podręcznej” pamięci operacyjnej pamięci ROM (*Read Only Memory*) zapisane jest oprogramowanie uruchamiane bezpośrednio po włączeniu zasilania – BIOS (*Basic Input/Output System*). BIOS sprawdza poprawność działania procesora (CPU), pamięci (RAM) i urządzeń wejścia/wyjścia.

3.4.1. Nośniki informacji stosowane w systemach sterowania CNC

Pierwszym nośnikiem informacji znanym już w XIX wieku i stosowanym w różnych urządzeniach była **taśma dziurkowana (perforowana)** odczytywana w specjalnej konstrukcji czytnikach (rys. 20a). Szybkość wprowadzania programu z taśmy perforowanej była bardzo ograniczona, a ponadto ten nośnik informacji w warunkach przemysłowych często bywał uszkodzony. Późniejszym sposobem zapisywania, a następnie wprowadzania programu do układu sterowania, był **zapis na taśmie magnetycznej** (rys. 20b). Taśma pokryta warstwą materiału ferromagnetycznego przesuwana się wzdłuż głowicy zapisującej, która wytwarza pole magnetyczne. Na taśmie powstają namagnesowane obszary (bity 0 i 1). W przypadku odczytu danych, głowica wykrywa zmiany pola magnetycznego. Zaletą tego zapisu jest bardzo duża pojemność i niski koszt przechowywania danych. W warunkach przemysłowych (pył, kurza) tego typu nośniki danych obecnie stosowane są sporadycznie. Natomiast

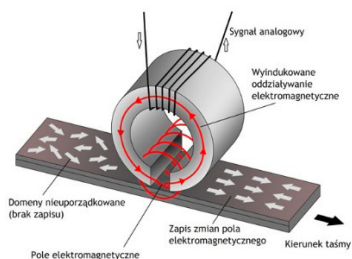
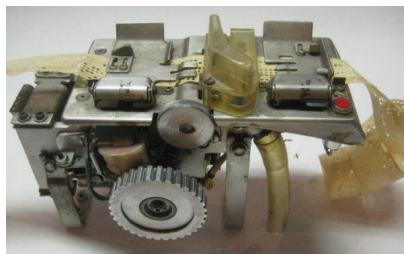
nowoczesnym standardem jest **LTO** (*Linear Tape-Open*), sposób używany do profesjonalnych systemów przechowywania danych, czy też systemów backupu.

O ile LTO nadaje się przykładowo do archiwizacji danych (sporadyczny dostęp do danych), to w przypadku codziennej pracy stosowane są **dyski HDD** (*Hard Disk Drive*), szeroko stosowany nośnik danych zapewniający szybki dostęp do opracowanego programu. Dysk HDD składa się z kilku kluczowych elementów: talerzy (tzw. plater) – metalowe lub szklane dyski pokryte warstwą magnetyczną, głowice odczytujące i zapisujące dane, które zamocowane są na ramieniu poruszającym się nad powierzchnią talerzy oraz silnika wprawiającego talerze w ruch z prędkością 5400 do 15000 obr/min – rys.20c.

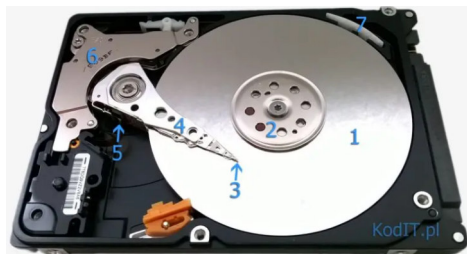
a.



b.



c.



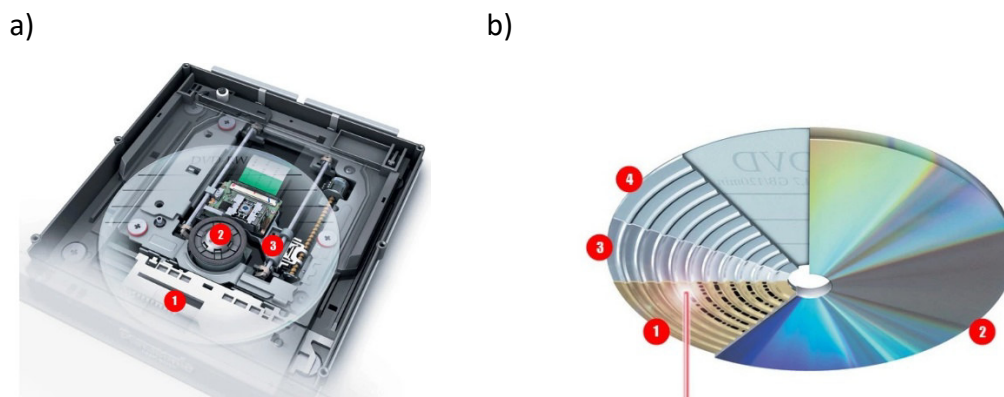
Rys. 20. Zapis programu na taśmie perforowanej (a), nośniku magnetycznym (b), twardym dysku (c)

Innym nośnikiem konkurującym z tzw. twardym dyskiem HDD jest **dysk SSD** (*Solid State Drive*) zwany często dyskiem półprzewodnikowym. Nie posiada on ruchomych części, a dane zapisywane są w postaci ładunków¹⁷.

Najczęściej stosowanym obecnie nośnikiem informacji są **pendrive** (inaczej pamięć USB lub *flash drive*) – małe przenośne urządzenia wykorzystujące pamięć typu flash (czyli nieulotną). Oznacza to, że zapisywane dane nie znikają po odłączeniu zasilania i można je wielokrotnie zapisywać i usuwać. Podłącza się go przez port USB (*Universal Serial Bus*), dzięki czemu działa od razu bez konieczności instalowania dodatkowych programów.

¹⁷ Taka pamięć omawiana była w poprzednich rozdziałach pracy.

Równie często stosowanym nośnikiem danych zapisywanych i odczytywanych za pomocą lasera są **nośniki optyczne** (płytki) **DVD** (*Digital Versatile Disc*). Dane zapisywane są w postaci mikroskopijnych zagłębień na powierzchni płyty (rys. 21).



Rys. 21. Napęd DVD (a) oraz zapis na nośniku DVD (b)

Kiedy tacka wsunie się do urządzenia (rys. 21a), zespół zębatek (1) przesunie napęd (2) pod płytkę. Jednocześnie zespół lasera (3) zamocowany na sankach, przesuwając laser do miejsca rozpoczęcia (nagrywania/ odtwarzania). Nagrywanie rozpoczyna się od ścieżek wewnętrznych po spirali na zewnątrz. Silnik rozpędza płytkę do prędkości 1475 obr/min. W miarę przesuwania się lasera na zewnątrz prędkość obrotowa zmniejsza się do 575 obr/min. Zapewnia to stałą szybkość przepływu danych. Laser zapisuje zera i jedynki cyfrowych danych (bity) na wytłoczonej wcześniej spiralnej ścieżce. W uproszczeniu wygląda to tak, że dla każdej jedynki laser jest włączany, a dla każdego zera wyłączany. Włączony laser zaczernia farbą (1) pod przezroczystą warstwą ochronną (2). Poniżej znajduje się metaliczny reflektor (3) umieszczony na warstwie nośnej z nadrukiem (4). Podczas odczytu napęd rozpoznaje jedynki po tym, że w zaczernionych punktach nie następuje odbicie światła.

Do przesyłania programów różnych urządzeń technologicznych coraz częściej wykorzystywany jest internet bądź przez network boot lub streaming. W przypadku **network boot** (bootowania z sieci) karta sieciowa pobiera adres IP, odnajduje serwer startowy z którego pobierany jest program. Przykładem może być program sterujący obrabiarką CNC znajdujący się na serwerach wspomagających projektowanie konstrukcji i technologii (CAD/CAM). W przypadku **streamingu** przesyłanie danych odbywa się na bieżąco przez internet, bez konieczności wcześniejszego pobierania całego pliku (dane wysyłane są w małych fragmentach bezpośrednio odtwarzanych przez układ sterowania). Aktualnie ten



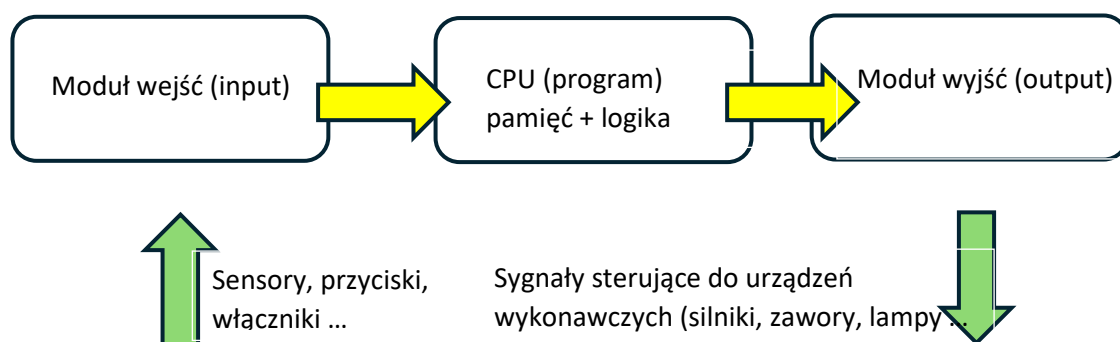
sposób nie jest stosowany do przesyłania programów sterujących do urządzeń CNC, gdyż wymaga on bardzo szybkiego internetu. Niespełnienie tego warunku może spowodować „zacinanie się” programu i niewłaściwą pracę sterowanego urządzenia.

3.4.2. Sterownik programowalny, sensory i aktuatory

Niezależnie od nośnika, na którym zapisany jest program, zostaje on załadowany do pamięci RAM i następnie wykonywany według instrukcji przez CPU (*Central Processing Unit*), który pracuje w cyklu: pobranie instrukcji z pamięci RAM -> dekodowanie (zrozumienie polecenia) -> wykonanie polecenia.

Współczesne układy sterowania CNC wyposażone są w wielu przypadkach w różnego rodzaju sensory (czujniki), analizujące otoczenie, w którym działa urządzenie technologiczne. Podobnie urządzenia sterowane numerycznie wyposażone są w układy wykonawcze (aktuatory), które fizycznie umożliwiają wykonanie poleceń zawartych w programie sterującym. Można w takim przypadku stwierdzić, że urządzenia te posiadają znamiona sztucznej inteligencji.

W układach sterowania CNC najczęściej używany jest specjalny komputer przemysłowy, zwany **sterownikiem PLC** (*Programmable Logic Controller*). Odbiera on sygnały z różnych urządzeń, przetwarza je w mikroprocesorze powszechnie nazywanym CPU (*Central Processing Unit*), według programu i wysyła sygnały sterujące do innych elementów (rys. 22).



Rys. 22. Schemat sterownika PLC

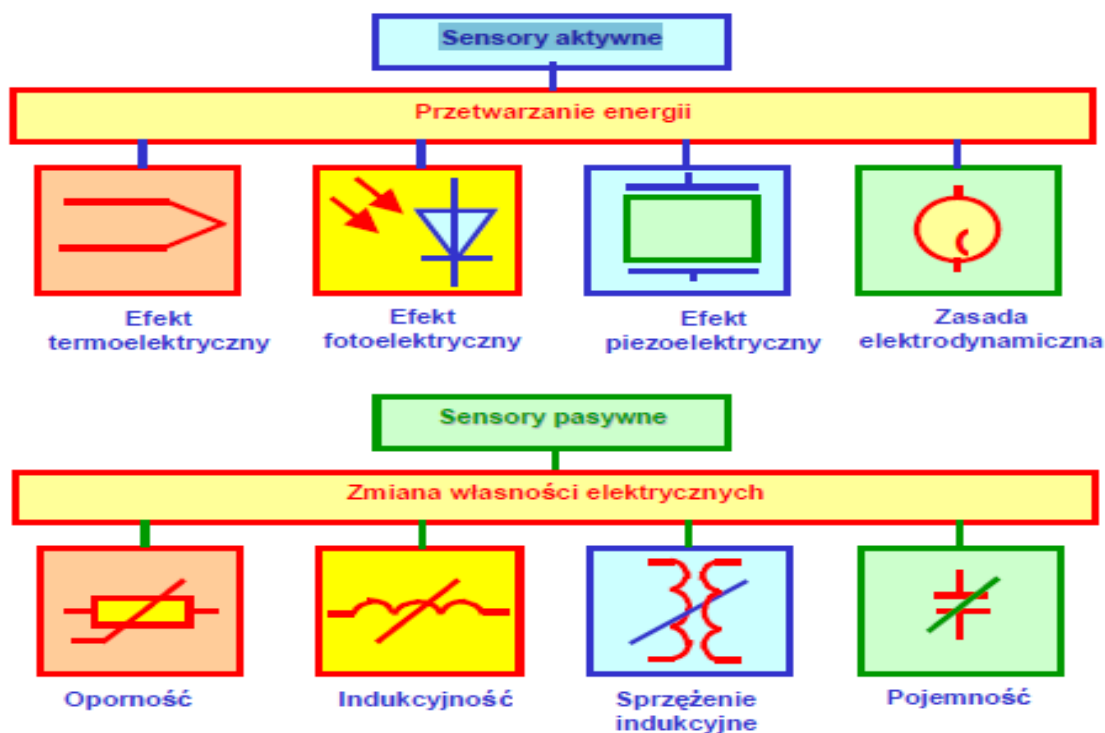
W istocie sterownik PLC od zwykłego komputera różni się przede wszystkim urządzeniami wejścia/wyjścia. Sterowniki PLC koncentrowane są na odbieraniu

i przekazywaniu sygnałów sterujących, natomiast klasyczne komputery wyposażone są w monitory, drukarki, skanery i inne urządzenia ułatwiające użytkownikowi (człowiekowi) porozumiewanie się z maszyną.

W układach sterowania CNC informacje dotyczące sygnałów wejściowych pochodzą z sensorów rozmieszczonych w obszarach istotnych dla działania sterowanego urządzenia. Są one dobierane w zależności od doprowadzanego do sterownika parametru. Na ogół stosowane są dwa typy sensorów: aktywne i pasywne.

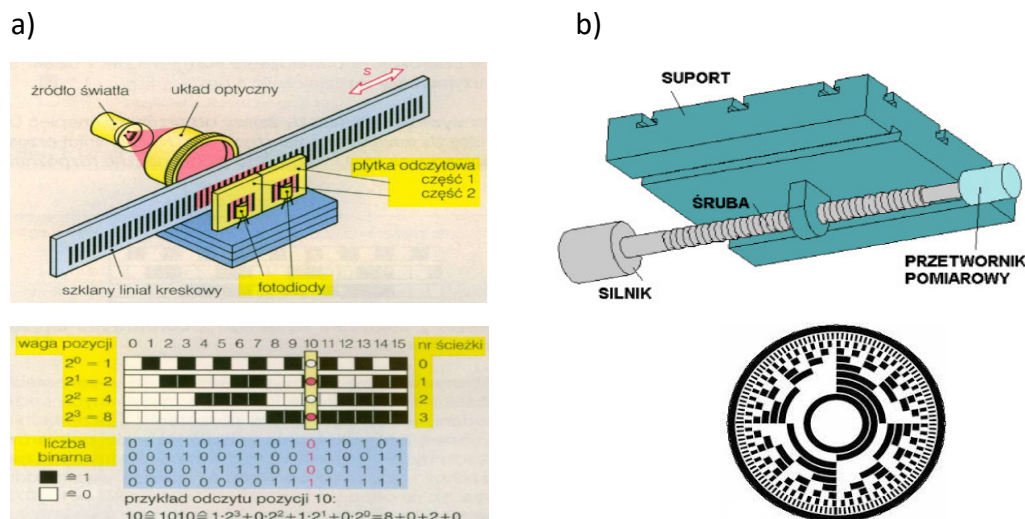
Sensory aktywne same wytwarzają sygnał elektryczny w odpowiedzi na zmianę parametrów mierzonego zjawiska fizycznego. Nie potrzebują zewnętrznego zasilania do wygenerowania sygnału pomiarowego, a energia pochodzi bezpośrednio z badanego zjawiska. Przykładem sensorów aktywnych jest termopara, czujnik piezoelektryczny, mikrofon dynamiczny i inne. Sensory aktywne generują niewielkie sygnały, które zazwyczaj wymagają wzmocnienia.

Sensory pasywne nie wytwarzają samodzielnie sygnału elektrycznego, lecz zmieniają swoje własności elektryczne pod wpływem mierzonego zjawiska. Aby działały, muszą być zasilane z zewnętrznego źródła energii. Przykładem mogą być fotorezystory zmieniające rezystancję pod wpływem światła, sensory pojemnościowe i inne. Przykłady sensorów przedstawiono w tablicy na rys. 23.



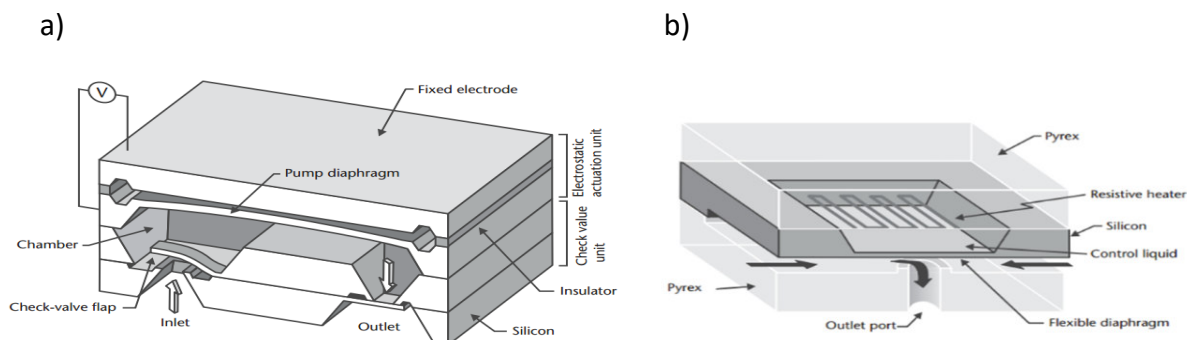
Rys. 23. Przykłady sensorów najczęściej stosowanych w układach sterowania CNC

W obrabiarkach sterowanych numerycznie, bardzo istotną informacją doprowadzaną do układu sterowania jest informacja o aktualnym położeniu suportów lub innych elementów wykonawczych obrabiarki. Rolę tą spełniają sensory położenia: linały kodowe i tarcze kodowe (rys. 24).

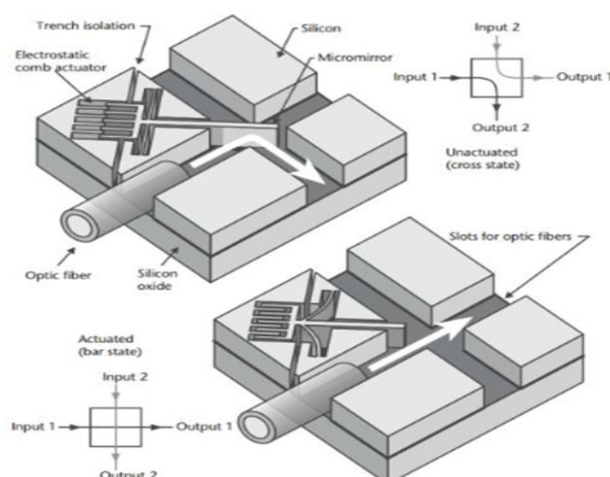


Rys. 24. Sensory położenia: linał kodowy (a) i tarcza kodowa (b)

Doprowadzone do układu sterowania sygnały wejściowe, przetwarzane są zgodnie z wprowadzonym programem, a wynikiem tego jest sygnał sterujący przesyłany do **aktuatora**. Jest to element wykonawczy, który zamienia sygnał sterujący na działania fizyczne np. ruch, ciepło, siłę lub światło. Przykładem mogą być silniki elektryczne, siłowniki pneumatyczne czy też hydrauliczne, zawory i inne (rys. 25).

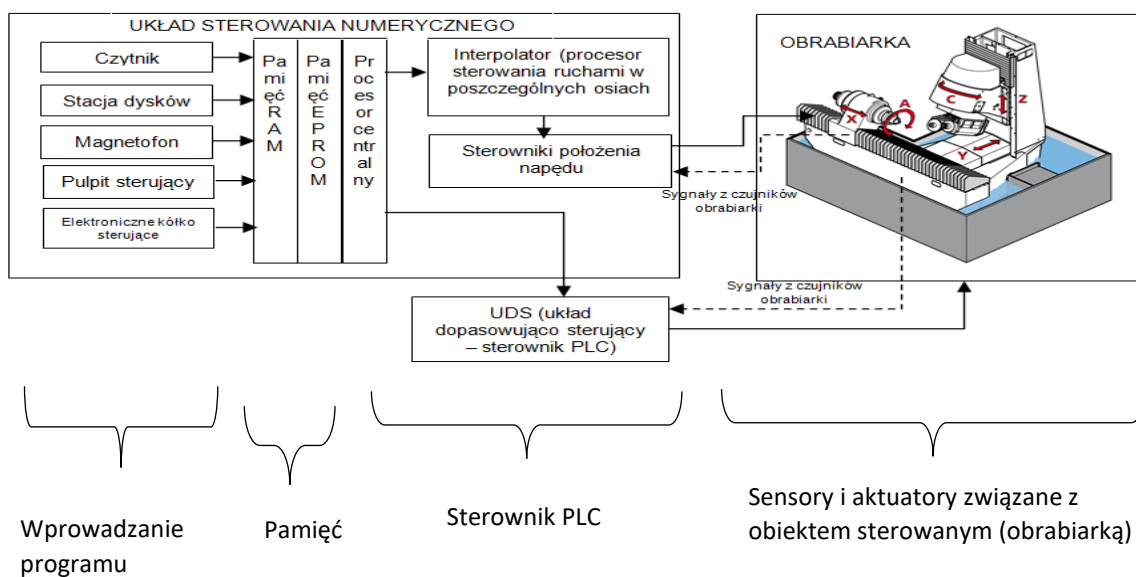


c)



Rys. 25. Aktuatory o wymiarach poniżej 0,01 mm: mikropompa (a), mikrozawór (b),
przełącznik optyczny (c)

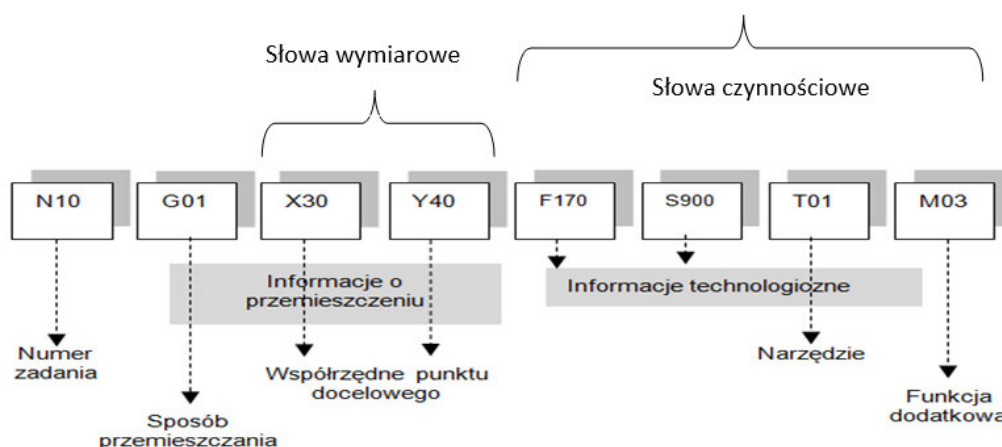
Poglądowy schemat układu sterowania CNC obrabiarki przedstawiono na rys. 26.



Rys. 26. Schemat układu sterującego obrabiarki

Program, czyli grupa poleceń wprowadzona do układu CNC obrabiarki, określa prowadzenie narzędzia wzdłuż wymaganego toru (linii, łuków za które odpowiada interpolator) oraz czynności jakie muszą wykonywać zespoły obrabiarki. Są one zapisane w postaci symbolicznej – w języku

układu sterowania obrabiarki¹⁸. Poszczególne linie programu, zwane blokami, budowane są w oparciu o składnię języka programowania. Składają się z szeregu instrukcji (słów) – rys. 27.



Rys. 27. Przykładowy blok programu wprowadzany do układu sterującego obrabiarki CNC

Słowa wymiarowe związane są z przemieszczeniami narzędzi i suportów obrabiarki (np. przemieszczenie liniowe narzędzia z punktu A do B), natomiast **słowa czynnościowe** zawierają informację o sposobie działania układów obrabiarki (prędkość posuwu F, prędkość obrotowa wrzeciona S, numer narzędzia T01 czy uruchomienie funkcji dodatkowej M03).

Program wprowadzony jest najpierw do pamięci RAM układu sterowania obrabiarki i blok po bloku pobierany do rejestrów mikroprocesora. Zgodnie ze słowami programu do aktuatorów obrabiarki przekazywane są polecenia programu, których wykonanie potwierdzone jest przez sensory obrabiarki (najczęściej sensory położenia). To potwierdzenie jest sygnałem do rozpoczęcia realizacji kolejnego bloku programu.

W podobny sposób działają inne urządzenia technologiczne CNC, przykładowo roboty przemysłowe, czy też środki transportu AVG (*Automated Guided Vehicle*), chociaż coraz wyraźniej widoczne jest wprowadzanie do tych urządzeń układów sterowania, których częścią są programy **sztucznej inteligencji AI** (*Artificial Intelligence*).

¹⁸ Budowa programu ujęta jest w normie DIN 66025.



3.5. Inteligentne systemy sterowania CNC

Inteligentne systemy sterowania są kolejnym etapem rozwoju układów sterowania, wykraczającym daleko poza automatyzację. Poprzedzająca automatyzację **mechanizacja**, umożliwiła zastąpienie pracy człowieka przez maszynę. **Automatyzacja** natomiast pozwoliła na samodzielną pracę maszyny lub tylko z minimalnym udziałem człowieka. Zawsze jednak musiał istnieć opracowany przez człowieka program sterujący maszyną. Rozwój współczesnej automatyzacji zmierza w kierunku **inteligentnych opartych o AI systemów**, które ograniczają udział lub niekiedy nawet wykluczają człowieka w sterowaniu procesem¹⁹. Uzyskanie tego efektu wymaga jednak spełnienia kilku warunków nazywanych niekiedy „wielką piątką” cyfryzacji:

- posiadania dużej ilości danych umożliwiających uczenie się systemu i analizowanie informacji (**big data**),
- przetwarzanie danych w **chmurach obliczeniowych CC** (*cloud computing*) oraz przechowywanie ich w **chmurach danych CD** (*cloud data*),
- środków technicznych umożliwiających łączenie między sobą różnych urządzeń i przedmiotów; aktualnie jest to łączenie głównie z internetem – tzw. **internet rzeczy IoT** (*internet of things*), chociaż w przypadku mniejszych odległości między obiektami stosowane jest również do łączenia z internetem Wi-Fi (aktualnie zasięg kilkudziesięciu metrów w budynku) oraz Bluetooth (zasięg 10 -20 m),
- posiadania **sieci telekomunikacyjnych** umożliwiających przesyłanie danych z dużą prędkością oraz odporną na spadki wydajności związane z dużą liczbą uczestników korzystających z sieci w tym samym czasie (**sieć 5G**); aby internet, czyli globalna sieć komputerów mogła być ze sobą połączona w celu przesyłania danych – musi istnieć sieć telekomunikacyjna (kable światłowodowe, stacje bazowe itp.),
- opracowania **algorytmów** (zestawu reguł i metod obliczeniowych), które pozwolą **sztucznej inteligencji** analizować dane, uczyć się zależności i podejmować decyzje.

¹⁹ Jest to obszar IV rewolucji przemysłowej, której umowny początek to rok 2011, znanej pod hasłem Przemysł 4.0.



Idea sztucznej inteligencji pojawiła się wraz z rozwojem w latach 50. XX wieku, medycyny a konkretnie neurobiologii – nauki o działaniu mózgu²⁰. Stwierdzono wówczas, że mózg istoty żywej jest siecią komórek neuronowych, które wysyłają impulsy elektryczne do różnych części mózgu, odpowiedzialnych za „funkcjonowanie” istoty żywej. Prawie od razu znalazła się grupa entuzjastów inżynierów – ludzi, którzy postanowili przy pomocy elementów mechanicznych odwzorować działanie mózgu istoty żywej. Pracę zapoczątkował angielski informatyk Alan Turing. Już w połowie XX w. zasugerował on, że maszyny podobnie jak ludzie, mogą wyciągać logiczne wnioski w celu rozwiązywania problemów lub podejmowania decyzji. Lecz prawdziwym przewrotem okazała się zorganizowana w 1956 roku konferencja w Dartmouth College w New Hampshire (USA), która zgromadziła „śmietankę” młodych entuzjastów z całego świata dyskutujących o czymś, czego nie nazwano jeszcze wówczas sztuczną inteligencją – tych, którzy postanowili kontynuować i rozwijać koncepcję zmarłego w 1954 roku Alana Turinga. Warsztaty w Dartmouth zapoczątkowały badania nad sztuczną inteligencją, a ona sama stała się poważną dyscypliną naukową. W latach 60. i 70. XX wieku powstały pierwsze systemy, które wykorzystywały wiedzę ekspertów do podejmowania decyzji w konkretnych dziedzinach. Były to **systemy ekspertowe**. W latach 80. i 90. skupiono się na tworzeniu systemów uczących się, takich jak **sieci neuronowe**. Stało się to możliwe w związku z rozwojem technologii komputerowych, które umożliwiły wzrost mocy obliczeniowej komputerów (*computer power*), czyli „siły” komputera do przetwarzania danych.

3.5.1. Systemy ekspertowe (eksperckie)

Systemy ekspertowe zwane również systemami eksperckimi stanowiły jedną z pierwszych rozpoczętych w latach 70-tych i 80-tych XX wieku, prób naśladowania przez maszynę zdolności ludzi (ekspertów) w rozwiązywaniu konkretnych problemów. Stało się to możliwe dzięki postępowi w obszarze technologii półprzewodnikowych, który umożliwił zwiększenie poziomu **skali integracji** produkowanych układów scalonych²¹. Efektem tego

²⁰ Badania zapoczątkował hiszpański neurobiolog Santiago Ramon Cajal.

²¹ Skala integracji odnosi się do liczby elementów elektronicznych takich jak tranzystory, kondensatory i inne, które są zintegrowane w pojedynczym układzie scalonym. Im więcej tych elementów, tym większa skala integracji i tym bardziej zaawansowane technologie stosowane muszą być podczas produkcji takich układów scalonych. Pierwsze układy scalone budowane były na poziomie SSI (*Small-Scale Integration*) i zawierały



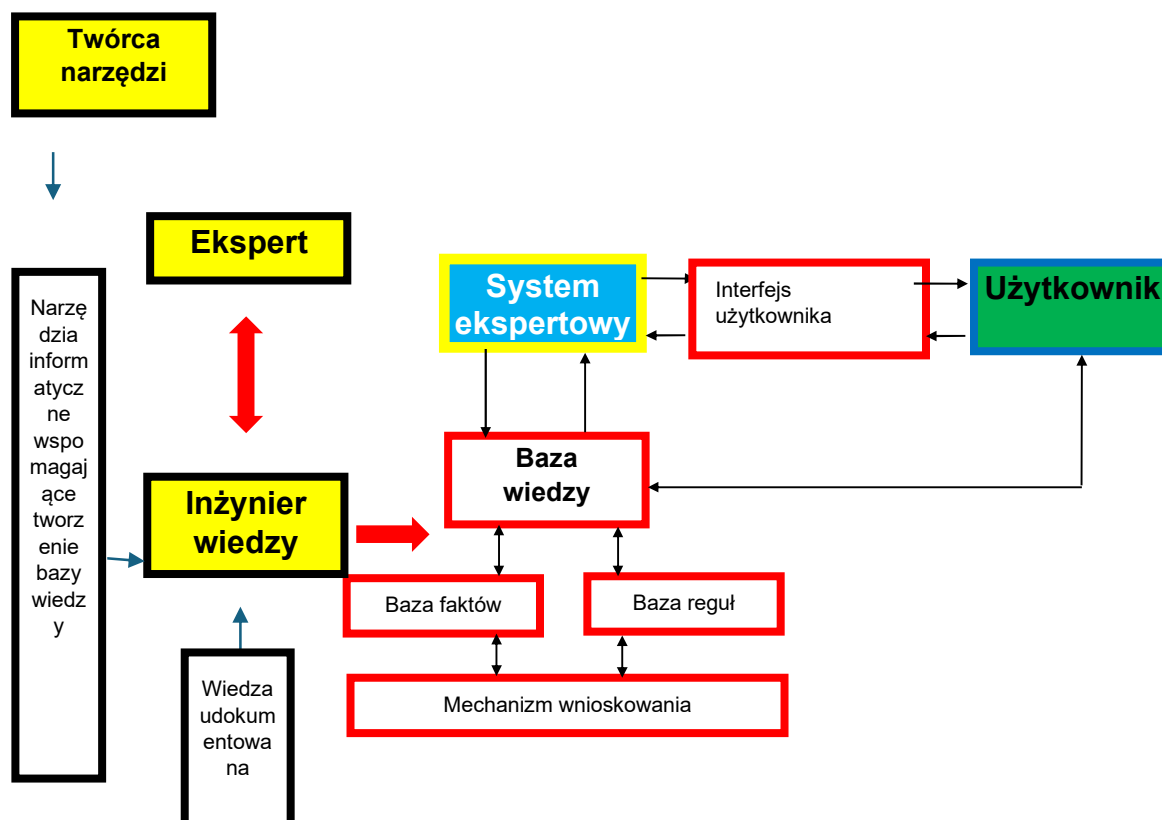
była możliwość zwiększenia ilości informacji przechowywanej w układach scalonych (pamięciach RAM, ROM) oraz szybkości jej wyszukiwania – przy akceptowalnych kosztach tych układów.

Podstawą tych systemów była wiedza ekspertów, a więc osób posiadających wiedzę i umiejętności, zdobywane przez naukę i doświadczenie, które czyniły je autorytetami w danym obszarze. Wiedza ekspertów zapisywana była w odpowiedniej formie w pamięci komputera i w miarę potrzeby wykorzystywana przez osoby nawet „luźno” związane z danym obszarem nauki. Dane zapisane w pamięci są podstawą do wyciągania przez system ekspertowy wniosków w sposób symulujący proces rozumowania człowieka – eksperta w danej dziedzinie.

Zastosowanie systemów ekspertowych spowodowało, że wiedza ekspertów stała się tania i łatwo dostępna. Jednak była ona „wiedzą martwą” nie inspirującą, nie zawierającą elementów zdroworozsądkowych bądź intuicyjnych. Obecnie znaczenie tych systemów jest ograniczone przez inne systemy z grupy inteligentnych, chociaż w niektórych dziedzinach, gdzie wiedza ekspercka jest kluczowa, systemy ekspertowe nadal znajdują zastosowanie.

Klasyczny system ekspertowy obejmuje kilka kluczowych komponentów, które współpracują ze sobą, aby system mógł skutecznie analizować informacje, wnioskować na ich podstawie i udzielać ekspertyzy w danej dziedzinie. Najważniejszym elementem tego systemu jest **baza wiedzy** (*Knowledge Base*). Jest to centralna część systemu ekspertowego, w której przechowywana jest wiedza ekspertów w postaci reguł, faktów, procedur czy innych form. Baza wiedzy dostarcza informacji niezbędnych do podejmowania decyzji (rys. 28).

kilkadziesiąt elementów elektronicznych. W latach 80. XX wieku budowano układy scalone w skali VLSI (*Very Large Scale Integration*). Zawierały one już miliony elementów elektronicznych. Aktualnie budowane są układy GSI (*Gigantic Scale Integration*) zawierające miliardy elementów elektronicznych.



Rys. 28. System ekspertowy (oprac. własne na podstawie: E. Pająk, Zarządzanie produkcją, PWN, Warszawa 2006)

Skuteczność systemu ekspertowego zależy głównie od „ilości” wiedzy zgromadzonej w bazie. Im większa złożoność bazy wiedzy określana liczbą reguł w nich zawartej, tym trafniejsza sugestia decyzji, lecz tym dłuższy czas przeszukiwania bazy wiedzy celem znalezienia „właściwej” reguły. Z tego względu kolejnym ważnym elementem systemu ekspertowego jest przyjęty sposób przeszukiwania bazy wiedzy zwany w wielu pracach **silnikiem wnioskowania** (*Inference Engine*). Generalnie odpowiada on za przetwarzanie informacji, generuje nowe informacje lub rozwiązania wykorzystując reguły i algorytmy przetwarzania informacji. Praca tego „mechanizmu przeszukiwania” bazy wiedzy (silnika), przekazywana jest przez **interfejs** do **użytkownika**. Może on przyjmować różne formy, np. tekstowy, graficzny lub inny. Stanowi drogę komunikacji użytkownika z systemem.

Jednakże pierwszym i zasadniczym krokiem zmierzającym do zbudowania systemu ekspertowego jest **pozyskanie wiedzy i utworzenie z niej bazy wiedzy**. Rola ta najczęściej powierzana jest osobie nazwanej powszechnie „inżynierem wiedzy”. Jego zadaniem jest współpraca z ekspertem lub ekspertami oraz ujęcie ich wiedzy najczęściej w postaci reguł

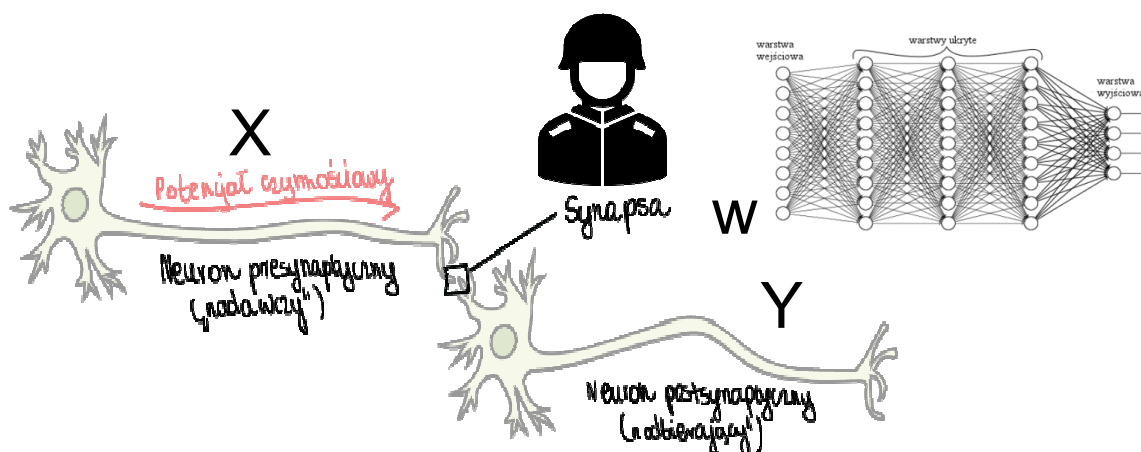
zapisywanych w bazie wiedzy. **Reguły** są warunkowymi stwierdzeniami o istnieniu pewnych zależności między obiektami. Opisują one reakcje obiektu na przedstawione lub napływające fakty.

IF fakt THEN wniosek AND/OR działanie

Przykład wykorzystania systemu ekspertowego przedstawiony został w następnym rozdziale i dotyczy wykorzystania systemu ekspertowego do „nadzorowania” pracy silnika spalinowego różnych pojazdów (działanie tzw. komputera pokładowego).

3.5.2. Sztuczne sieci neuronowe

Mózg człowieka (*cerebrum*) składa się z miliardów neuronów²², które komunikują się ze sobą za pomocą impulsów elektrycznych, umożliwiając mózgowi przetwarzanie informacji i podejmowanie decyzji. Jest to do tej pory niedościgniony wzór dla systemów sztucznej inteligencji. Neurony mózgu komunikują się ze sobą dzięki połączeniom nazywanym **synapsami** (rys. 29).

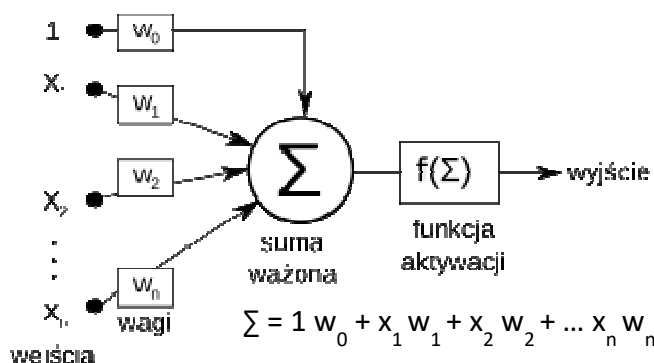


Rys. 29. Biologiczna i sztuczna sieć neuronowa

Synapsy spełniają nieustannie ważną rolę, gdyż albo wzmacniają przesyłany z neuronu do neuronu sygnał, względnie go osłabiają lub w ogóle blokują.

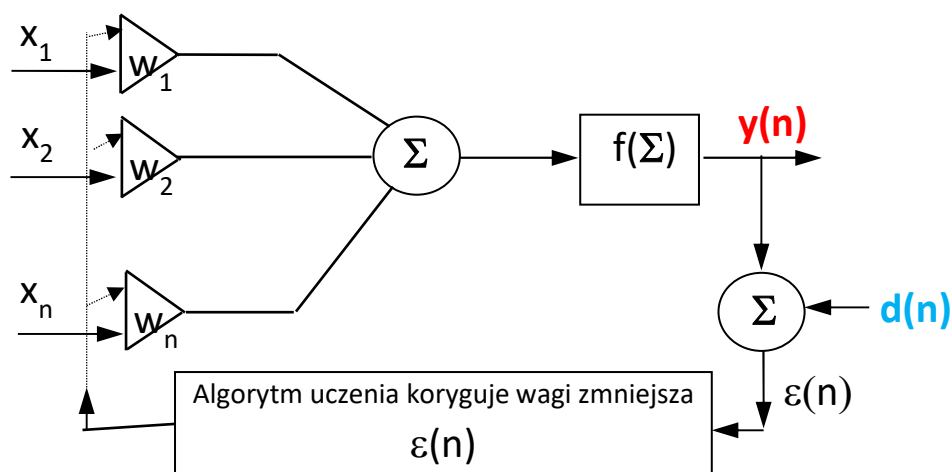
²² Szacuje się, że mózg człowieka zawiera od 30 do 100 miliardów neuronów.

Opracowanie przez inżynierów modelu neuronu (rys. 30), stanowiło istotny krok rozwoju sztucznych sieci neuronowych SSI.



Rys. 30. Model matematyczny neuronu wg McCullocha-Pittsa (źródło <https://pl.wikipedia.org>)

Z matematycznego punktu widzenia problem był prosty. Wartość „x” stanowiąca wyjście z innych neuronów sieci, lub też sygnał z sensorów, mnożymy przez wartość wagi „w” i sumujemy w jądrze neuronu. Funkcja aktywacji przekształca wyliczoną wartość i to ona stanowi wyjście „naszego” neuronu, które staje się wejściem do kolejnego neuronu sieci. Niby nic trudnego, ale ... nigdzie nie określono jakie wartości przyjmują wagi, a przecież to dla „działania” modelu neuronu jest szczególnie istotne! W punkcie wyjścia wartości tych wag można przyjąć zupełnie losowo, ale już kolejne kroki nie mogą być dowolne. Jakie one będą zależy od przyjętego sposobu trenowania sieci neuronowej (nauczania sieci). Dydaktyczny przykład trenowania pojedynczego neuronu sieci przedstawiono na rys. 31. Jest to przykład „**trenowania z nauczycielem**”.



Rys. 31. Dydaktyczny przykład trenowania pojedynczego neuronu sieci

(źródło <https://pl.wikipedia.org>)

Warunkiem rozpoczęcia trenowania jest konieczność posiadania przez nauczyciela obszernego zbioru danych do trenowania zawierających informacje jaka jest prawidłowa wartość wyjściowa $d(n)$, dla określonego zbioru sygnałów x doprowadzanych do synaps. Przykładowo, jeśli dla danych wejściowych x_1 do x_n przy założonych wstępnie wagach synaptycznych w_1 do w_n uzyskaliśmy na wyjściu wartość $y(n)$, to musimy ją porównać z zestawem danych trenujących posiadanym przez nauczyciela. Jeśli oczekiwana wartość wynosiła $d(n)$ i jest ona różna od otrzymanej $y(n)$ to powstaje błąd $\varepsilon(n) = y(n) - d(n)$. Korzystając ze stosowanego przez nauczyciela algorytmu, następuje w kolejnym kroku zmiana wartości wag synaptycznych. Powtarza się procedurę, ale tym razem dla skorygowanych wag synaptycznych. Takie korekty prowadzi się do czasu, kiedy nauczyciel określi, że wartość błędu $\varepsilon(n)$ jest na tyle mała, że można przyjąć, że neuron został wytrenowany. Problemem jest jednak to, że omawiany przykład trenowania dotyczy nie całej sieci, a **tylko jednego neuronu**. W przypadku, jeśli sieć posiada np. 39 neuronów, liczba powiązań neuronów wynosi 1500. W takim przypadku trenowanie sieci jest procesem niesłychanie czasochłonnym i kosztownym, zniechęcającym do korzystania z sieci – szczególnie wobec faktu, że **sieć musi być wytrenowana** – a jeśli tak nie jest, nie wolno jej stosować (nie daje prawidłowych odpowiedzi). To spowodowało, że po pierwszej fali entuzjazmu w latach 70-90. XX wieku, prace nad sztucznymi sieciami neuronowymi zostały ograniczone.



Problem trenowania sztucznych sieci neuronowej wrócił wraz z rozwojem internetu²³. Właśnie to dzięki zasobom internetu możliwe stało się szybsze trenowanie sieci chociaż ważne, aby mieć świadomość, że niekoniecznie lepsze.

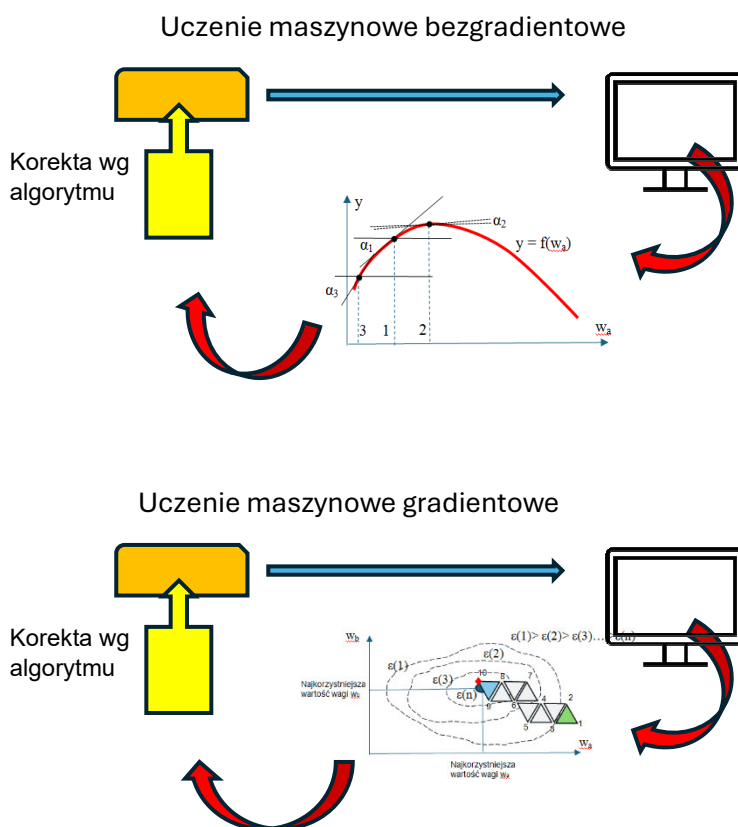
Możliwe stało się wprowadzenie tzw. **uczenia maszynowego**, a więc takiego w której sieć neuronowa uczy się na podstawie danych i doświadczenia, a nie na podstawie sztywnych reguł – tak jak to ma miejsce w przypadku systemów ekspertowych. W uproszczeniu ten sposób trenowania sieci w przypadku rozpoznawania pojazdów wyglądałby następująco: wprowadzamy do systemu wiele zdjęć obrazujących nadwozia pojazdów np. samochodu osobowego, ciężarówki, autobusów itd. Algorytm uczenia sieci analizuje cechy wprowadzonych obrazów i po treningu potrafi sam ocenić nowe zdjęcie i rozpoznać jaki pojazd przedstawiony jest na zdjęciu. W praktyce istnieje wiele sposobów, które można zaliczyć do obszaru uczenia maszynowego. Wszystkie one mają zasadniczy cel – zmniejszenie liczby danych niezbędnych do wytrenowania sieci i tym samym do skrócenia czasu treningu. W tym kierunku zmierza metoda **uczenia maszynowego nienadzorowanego** (*unsupervised learning*), w których system sam szuka podobieństw w danych. Jest to bardzo wygodna i szybka metoda trenowania sieci, lecz może być obarczona błędem. Nie zawsze sieć prawidłowo rozpoznaje dane, a to prowadzi do błędnego działania sieci.

Przykładem innego podejścia trenowania sieci jest **uczenie maszynowe przez wzmacnianie** (*reinforcement learning*). W przypadku **uczenia maszynowego gradientowego** do systemu wprowadzany jest model, który wskazuje czy podjęte zmiany np. wag synaptycznych sieci prowadzą do zmniejszenia błędu, czy też odwrotnie. Na tej podstawie następuje korekta parametrów wag sieci neuronowych. Przedstawiony sposób przypomina zaprezentowany wcześniej sposób uczenia sieci z nauczycielem, przy czym trenowanie sieci przez uczenie maszynowe metodą gradientową odbywa się znacznie szybciej. Podstawą opracowanego modelu są badania doświadczalne, co znacznie zwiększa jakość wytrenowania sieci. Podobnie działa **uczenie maszynowe bezgradientowe**. Jednakże w takim przypadku podstawą korekty wag są opracowane przez człowieka (nauczyciela) modele matematyczne. Dla parametrów wejściowych sieci neuronowej system oblicza z modelu prawidłowe parametry wyjścia sieci neuronowej i następnie porównuje je do parametrów rzeczywistego

²³ Prawdziwy rozwój internetu dla zwykłych użytkowników nastąpił w 1991 roku, kiedy Tim Berners-Lee opracował sieć WWW (*World Wide Web*). Powstały wtedy strony internetowe, przeglądarki i system linków. Od 2000 roku następuje dynamiczny rozwój internetu na cały świat.

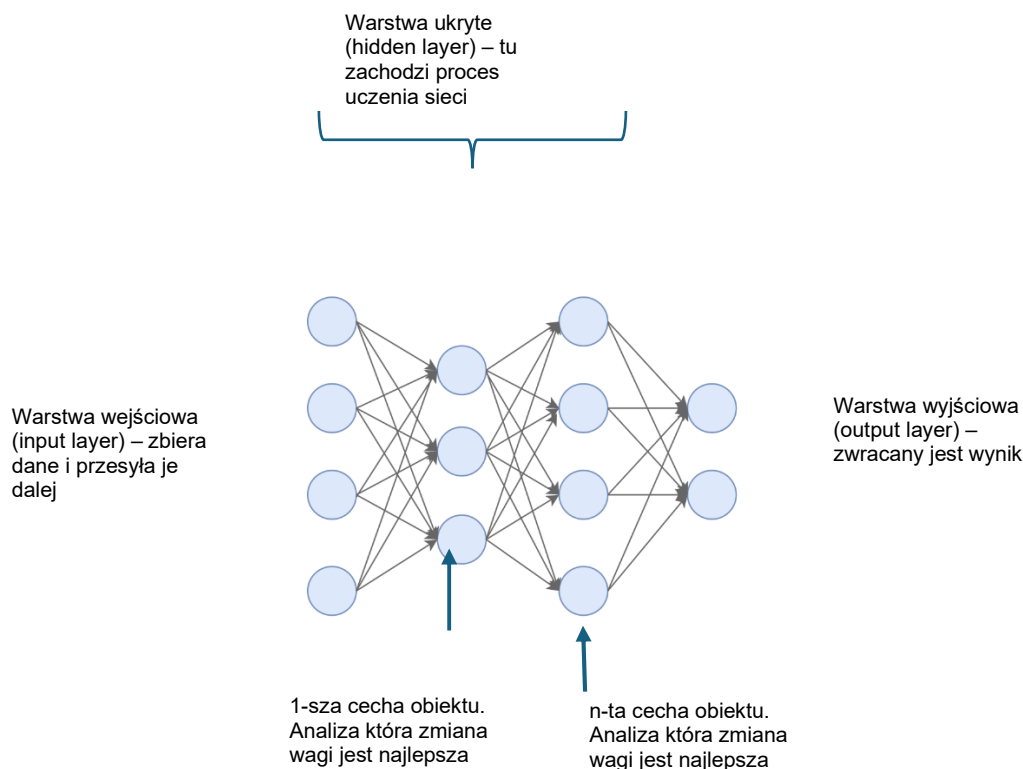
„wyjścia” SSN. W zależności od różnicy między nimi (błędu), dokonywana jest korekta wag synaptycznych sieci zgodnie z ustalonym algorytmem.

Metody gradientowe i bezgradientowe uczenia maszynowego, są w pewnym sensie łącznikiem między trenowaniem z nauczycielem, a uczeniem maszynowym nienadzorowanym. Z jednej strony skraca czas trenowania sieci, z drugiej poprawa jakość wytrenowania (rys. 32).



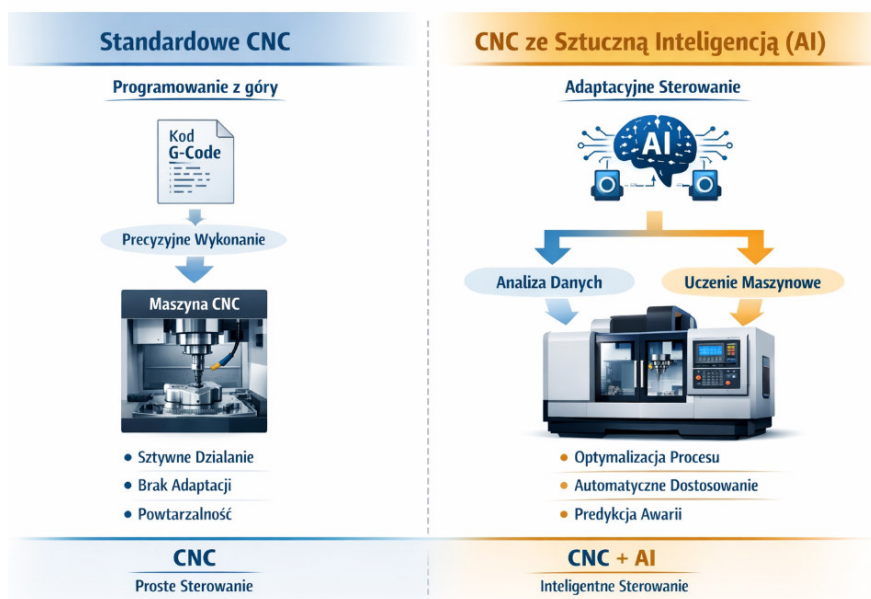
Rys. 32. Schemat uczenia maszynowego sztucznej sieci neuronowej

Odmianą uczenia maszynowego jest tzw. **uczenie głębokie** (rys. 33).



Rys. 33. Schemat uczenia głębokiego

Istota uczenia głębokiego polega na tym, że komputer uczy się rozpoznawać wzorce za pomocą wielowarstwowych sieci neuronowych. Każda warstwa przetwarza dane dotyczące fragmentu obiektu (cechy – *feature engineering*) i przekazuje wynik do następnej warstwy. Przykładowo w trakcie analizy twarzy jedna warstwa będzie analizowała cechy nosa, druga oczy, a trzecia usta. Tak więc przykładowo ustalony wzorzec nosa jest przekazywany do kolejnej warstwy, która rozpoznaje oczy, a sumaryczny wynik pracy obu tych warstw przekazywany jest do kolejnej warstwy „zajmującej się” ustami. Uczenie polega na dostosowywaniu wag połączeń między neuronami tak, aby wynik był jak najbardziej poprawny. Im więcej warstw w sieci, tym bardziej złożone zależności może ona wykrywać. Tak więc sieć uczy się reprezentacji danych w kolejnych warstwach, zamiast otrzymywać od człowieka gotowe reguły.

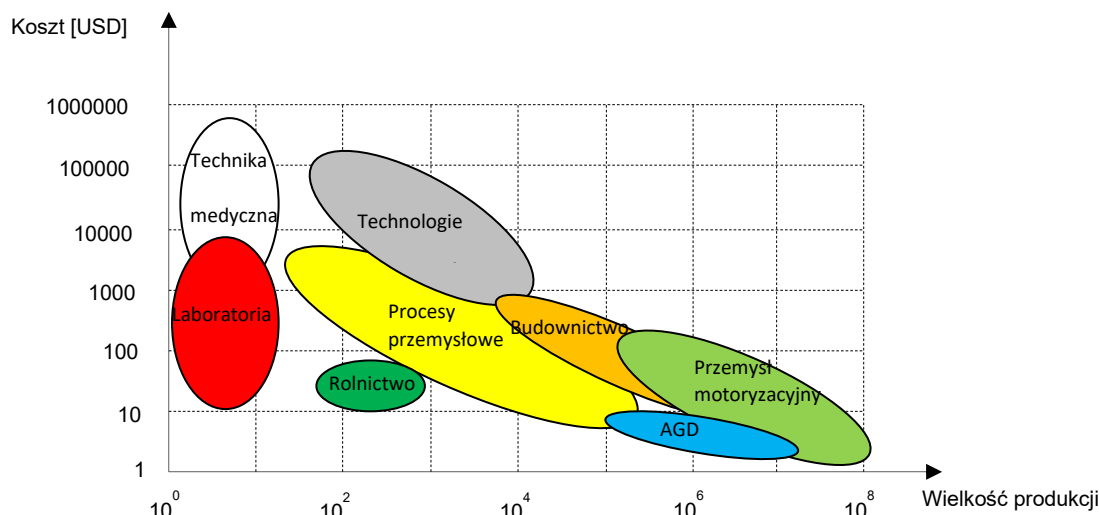


Rys. 34. Sterowanie CNC oraz CNC+AI

Rzeczywistość pokazuje, że coraz częściej sterowanie numeryczne CNC jest integrowane ze sztuczną inteligencją AI w **inteligentnych maszynach i urządzeniach technologicznych** (rys. 34). Popularnie można powiedzieć, że CNC to precyzyjne „ręce” maszyny, a AI to jej „mózg”.

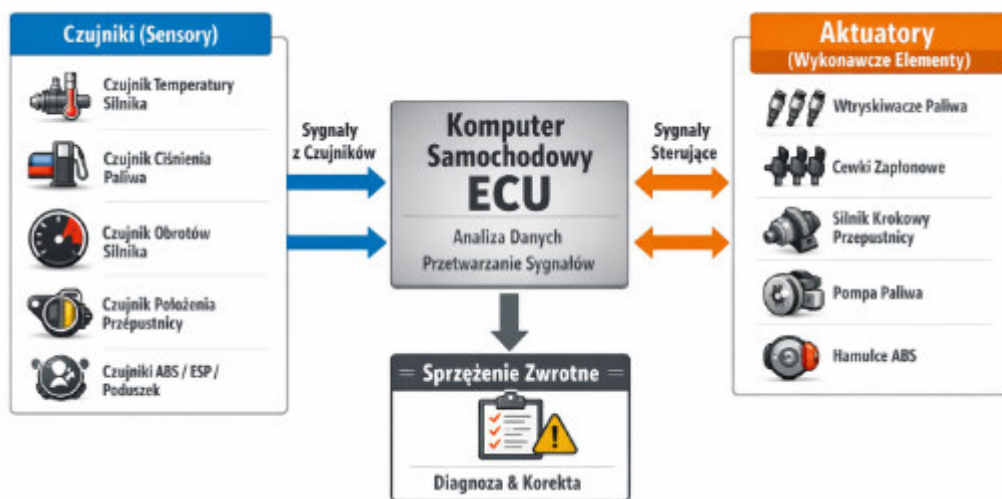
3.6. Przykłady inteligentnego sterowania CNC - AI

Założeniem IV rewolucji przemysłowej określanej w skrócie jako Przemysł 4.0 jest szeroko rozumiana współpraca człowieka z maszyną, przy czym wiele zadań wykonywanych dotychczas tylko przez człowieka, przejmowanych jest przez maszynę pracującą pod nadzorem człowieka. Ten kierunek stał się dominujący w sytuacji szybkiego rozwoju technologii półprzewodnikowych i tym samym rozszerzania możliwości zastosowania układów scalonych (rys. 35).



Rys. 35. Obszar zastosowania układów scalonych (oprac. własne na podstawie danych z internetu)

Zastosowanie układów scalonych w przemyśle motoryzacyjnym przyniosło konkretne efekty. Aktualnie każdy współczesny samochód „naszpikowany” jest elektroniką nadzorującą pracę wszystkich ważniejszych układów pojazdu (rys. 36).

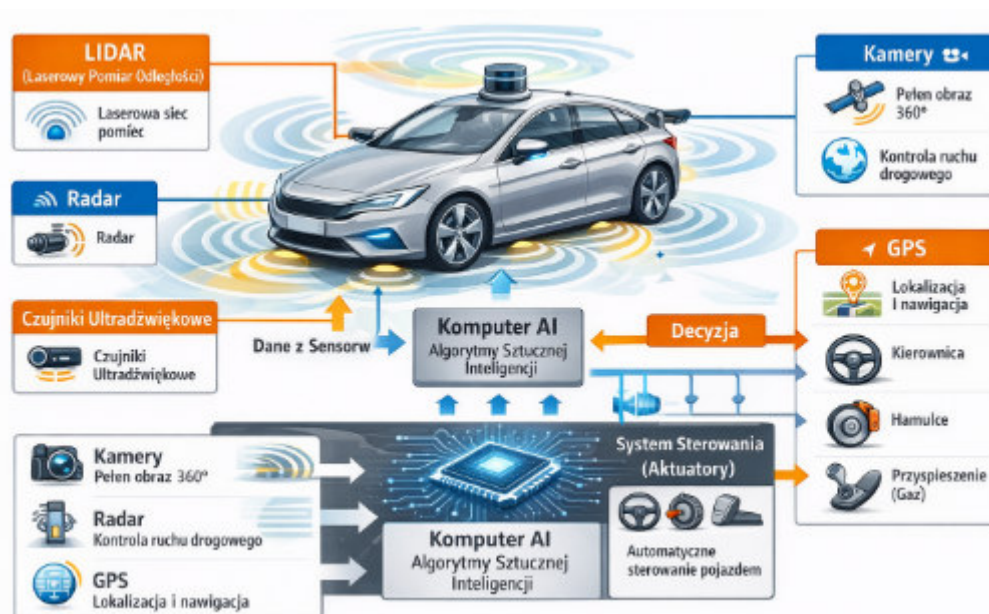


Rys. 36. Schemat działania komputera samochodowego ECU

Sensory pozyskują informację z różnych układów pojazdu i przekazują je do ECU (komputera pokładowego), w którym następuje ich przetwarzanie. Aktualnie najczęściej wykorzystane w tym celu są systemy ekspertowe. Przetworzone sygnały w przypadku konieczności dokonania korekty parametrów pracy układu, przekazywane są do aktuatorów

– czyli elementów wykonawczych korygujących pracę tych układów. Efektem działania tego systemu jest przede wszystkim obniżenie zużycia paliwa przy jednoczesnej poprawie innych parametrów eksploatacyjnych, bezpieczeństwo podróży oraz ochrona środowiska. Wprowadzone do bazy wiedzy reguły, reagują zmieniając parametry pracy układu samochodu, jednocześnie przekazując stosowny komunikat kierowcy (baza wiedzy okresowo jest uzupełniana przez producenta o kolejne nowe reguły, stąd konieczne jest co pewien czas aktualizacja oprogramowania).

Wskazane wyżej rozwiązania stanowią, jak już wspomniano, standard współczesnej motoryzacji. Ale warto w tym miejscu wspomnieć również o znacznych osiągnięciach w zakresie **pojazdów autonomicznych**, a więc takich które potrafią poruszać się samodzielnie bez udziału kierowcy. Obok jednostek ECU posiadać one muszą sensory zbierające informację o otoczeniu pojazdu. Stosując programy sztucznej inteligencji komputer pokładowy (zazwyczaj inny niż ECU), podejmuje decyzje o bezpiecznym poruszaniu się pojazdu (rys. 37).



Rys. 37. Schemat budowy pojazdu autonomicznego

Pojazdy autonomiczne już dawno mają za sobą okres „raczkowania”. Skala poziomów autonomii pojazdów ma 6 stopni (od 0 do 5). Większość pojazdów znajdujących się w sprzedaży to zazwyczaj pojazdy 2 i 3 stopnia w podanej wyżej skali. Nadmienić jednak



należy, że producenci takich pojazdów (Tesla, Toyota, Mercedes-Benz, Waymo), wkrótce zaoferują klientom samochody z wyższym niż dotychczas stopniem autonomii.

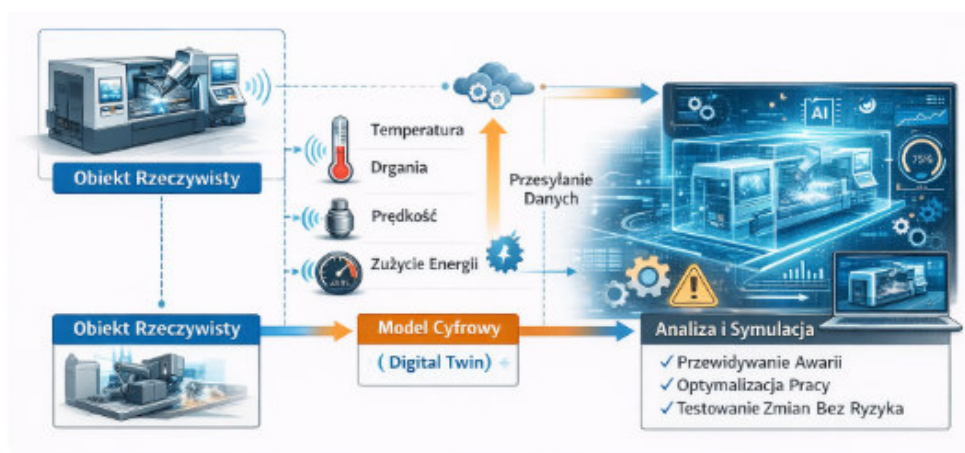
W zakresie pojazdów szynowych warto wspomnieć o całkowicie automatycznej linii kolejowej Yurikamome, którą oddano do użytku w Tokio w 1995 r. Składała się z 11 stacji i była pierwszym w Tokio w pełni automatycznym systemem transportu pasażerskiego (rys. 38).



Rys. 38. Całkowicie autonomiczna linia kolejowa łącząca centrum Tokio z wyspą Odaiba

(źródło: Yurikamome – Wikipedia, wolna encyklopedia)

Innym przykładem coraz szerzej stosowanym w praktyce jest **technologia Digital Twin** – cyfrowych bliźniaków. Cyfrowe bliźniaki to cyfrowa kopia rzeczywistego obiektu, maszyny lub systemu, która istnieje w komputerze i jest połączona z rzeczywistym obiektem poprzez dane z sensorów. Dzięki temu można monitorować, analizować i symulować działanie urządzenia w czasie rzeczywistym (rys. 39).



Rys. 39. Schemat technologii Digital Twins

Przykładem działania tej technologii może być produkcja łopatk turbiny lotniczej: maszyna CNC obrabia część, czujniki połączone internetem rzeczy (IoT) zbierają dane o temperaturze i drganiach występujących w trakcie obróbki. Dane te przesyłane są do komputera, na którym znajduje się „cyfrowy bliźniak” symulujący proces obróbki. Sztuczna inteligencja wykrywa zbliżenie się okresu intensywnego zużycia narzędzia (przekroczenie dopuszczalnego zużycia) i dokonuje korekty parametrów obróbki CNC, jednocześnie przekazując odpowiedni komunikat obsłudze.

W roku 2024 w fabrykach na całym świecie pracowało ok. 4,7 mln robotów przemysłowych, tj. około 9% więcej niż w roku poprzednim. Największy park robotów na świecie, bo ponad 2 mln robotów mają Chiny. Jednakże można prognozować, że coraz częściej stosowane będą w przemyśle **coboty** (*collaborative robot*), czyli roboty współpracujące z ludźmi – zaprojektowane do pracy bezpośrednio obok ludzi, w przeciwieństwie do klasycznych robotów przemysłowych, które ze względu na bezpieczeństwo działają zwykle w odizolowanych strefach. Wyposażone są one w czujniki siły, kolizji i obecności człowieka, dzięki którym zatrzymują ruch przy kontakcie z człowiekiem. Głównymi firmami produkującymi coboty są: Universal Robots (Dania), FANUC Corporation (Japonia), ABB Group (Szwajcaria) – rys. 40.



Rys. 40. Przykłady cobotów firmy Kawasaki Robotics (źródło: strona internetowa Kawasaki)

W perspektywie rozwoju robotyki w II połowie XXI wieku, są roboty pracujące w ramach kierunku nazwanego **autonomiczna robotyką**. Schemat stanowiska autonomicznej robotyki przedstawia rys. 41.

Stanowiska składa się z obszaru „klasycznej” sztucznej inteligencji AI, której zadaniem jest diagnoza problemu. Drugi obszar to sztuczna inteligencja agentyczna (agentowa), której zadaniem jest opracowanie technologii naprawy oraz trzeci obszar - robotyki, w której następuje fizyczna naprawa uszkodzenia.



Rys. 41. Przykład stanowiska autonomicznej robotyki (oprac. własne)

Przedstawiony przykład wskazuje, że w niedalekiej nawet perspektywie udział w procesach zarówno produkcyjnych jak i usługowych sztucznej inteligencji będzie wzrastał.



Sztuczna inteligencja dyskryminacyjna działa na podstawie danych wyuczonych w trakcie treningu. AI nie wykracza więc poza wzorce i reguły wprowadzone do bazy wiedzy. W takim przypadku sztuczna inteligencja jest pod pełną kontrolą człowieka. **Sztuczna inteligencja generatywna** samodzielnie analizuje duże zbiory danych i na tej podstawie tworzy „nowe treści”. Z tej racji kontrola człowieka nad wynikami analizy dokonywanej przez AI jest ograniczona. **Sztuczna inteligencja agentowa** samodzielnie podejmuje decyzje dotyczące sposobu działania jak i wykonuje podjęte decyzje. W tym przypadku nie ma żadnej kontroli człowieka nad podjętymi przez AI działaniami.

Warto w tym miejscu przytoczyć słowa nieżyjącego już polskiego pisarza powieści gatunku science fiction Stanisława Lema: „Sztuczna inteligencja może osiągnąć poziom tak złożony, że stanie się niezrozumiała dla ludzi. Porzuca kontakt z ludźmi, uznając ludzkość za zbyt prymitywną”.



4. Technologie łączenia (dr Adam Mazgajczyk)

4.1. Techniki łączenia metali

Technika łączenia to zbiór metod i sposobów stosowanych do łączenia elementów w celu utworzenia trwałej i stabilnej całości. Jest to dziedzina inżynierii i technologii, która obejmuje różnorodne techniki zapewniające mocne, niezawodne i efektywne połączenia między materiałami, takimi jak metale, tworzywa sztuczne, drewno czy ceramika¹.

1. Celem techniki łączenia jest zapewnienie trwałości, wytrzymałości i funkcjonalności połączenia, a także możliwość demontażu lub trwałego związania elementów.
2. Rodzaje technik łączenia dzielimy na:
 - mechaniczne: np. spawy, śruby, gwoździe, kołki, zatrzaski, złącza klinowe,
 - spawalnictwo: łączenie elementów metalowych przez stopienie ich brzegów, np. spawanie łukowe, MIG/MAG, TIG,
 - zgrzewanie: łączenie elementów poprzez ich rozgrzewanie i wciskanie lub dociskanie, np. zgrzewanie oporowe, ultradźwiękowe,
 - łączenie chemiczne: użycie klejów, żywic, epoksydów, które łączą elementy na poziomie chemicznym,
 - inne metody: np. lutowanie, skręcanie, zaciskanie.
3. Kryteriami doboru techniki łączenia jest:
 - rodzaj i właściwości materiałów,
 - wymagana wytrzymałość i trwałość połączenia,
 - warunki pracy (np. temperatura, obciążenia),
 - możliwość demontażu,
 - koszty i efektywność technologiczna.
4. Znaczenie techniki łączenia:
 - od niej zależy funkcjonalność i bezpieczeństwo wyrobu,
 - wpływa na jakość, trwałość i estetykę produktu,

¹ T. Rydzkowski, I. Michalska-Požoga, „Rozwój technik łączenia materiałów spawanie tworzyw polimerowych”, Welding Technology Review, 87/11/2015, s. 18-21.



- ma kluczowe znaczenie w procesach produkcyjnych i naprawczych.

Technika łączenia to kluczowa dziedzina techniczna służąca do trwałego i efektywnego łączenia elementów w różnych dziedzinach przemysłu i rzemiosła, obejmująca różnorodne metody dostosowane do charakterystyki materiałów i wymagań końcowego produktu.

4.2. Metody i sposoby łączenia metalu

Połączenia metali to sposoby łączenia dwóch lub więcej elementów metalowych w trwałą całość. W zależności od potrzeb konstrukcyjnych, technologicznych i użytkowych, stosuje się różne metody łączenia - zarówno nierozłączne, jak i rozłączne².

1. Połączenia nierozłączne są stosowane tam, gdzie liczy się duża wytrzymałość i sztywność, nie można ich rozłączyć bez uszkodzenia połączenia lub elementów, tworzą trwałą całość³.

Do najczęściej stosowanych należą:

- połączenia spawane metali to technika łączenia elementów metalowych polegająca na miejscowym stopieniu brzegów łączonych elementów i zespolenia ich razem, często z użyciem materiału dodatkowego (spoiwa), który po schłodzeniu tworzy trwałe i wytrzymałe połączenie. Zaletami połączeń spawanych jest duża wytrzymałość, szczelność (trwałe i mocne połączenia), oszczędność materiału, możliwość łączenia różnych materiałów i kształtów, automatyzacja procesu (np. roboty spawalnicze) a wadami jest specjalistyczna wiedza i urządzenia, wpływ na właściwości mechaniczne materiałów (np. możliwość odkształceń, utleniania, pęknięcia, wtrącenia), trudność demontażu, wpływ na strukturę metalu⁴.

Połączenia spawane dzielimy na połączenia:

- pełne (np. złącza czołowe, pachwinowe, czołowe narożnikowe) zapewniające maksymalną wytrzymałość,

² P. Chmiel, „Wykonywanie nierozłącznych połączeń blach 721 [01]. Z1. 06”, Instytut Technologii Eksploatacji – Państwowy Instytut Badawczy, 2007, s. 59.

³ M. Zasada, „Wykonywanie połączeń rozłącznych i nierozłącznych 724 [02]. 01. 05”, Instytut Technologii Eksploatacji – Państwowy Instytut Badawczy, Radom, 2007, s. 59.

⁴ M. Grigorenko, W. Gieorgij, A. Kostin, „Spawalność stali i kryteria jej oceny”, Welding Technology Review, 85/7/2013, s. 11-17.



- cienkościenne stosowane przy elementach o niewielkiej grubości, np. cienkie blachy,
- połączenia punktowe stosowane głównie w spawaniu elektrodą otuloną lub w spawaniu MIG/MAG.

Metody spawania dzielimy na spawanie:

- łukowe (np. MIG, MAG, TIG, elektrodą otuloną) najczęściej stosowane w przemyśle,
- gazowe np. spalanie acetylenu, rzadziej w przemyśle maszynowym,
- plazmowe precyzyjne i głębokie łączenia,
- odpornościowe (np. spawanie punktowe) głównie w przemyśle samochodowym.

Właściwości połączeń spawanych:

- wytrzymałość mechaniczna zależna od rodzaju spawania, materiału i techniki,
- odporność na korozję istotna w warunkach agresywnego środowiska,
- estetyka i jakość powierzchni ważne w przypadku elementów widocznych.

Czynniki wpływające na jakość połączeń spawanych:

- dobór odpowiedniej metody spawania i parametrów,
- przygotowanie elementów (np. czyszczenie, odpowiednie ustawienie),
- kontrola jakości (np. badania ultradźwiękowe, radiograficzne).

Połączenia spawane metali są kluczowym elementem w budowie maszyn, konstrukcji stalowych, przemysłach samochodowym i lotniczym, zapewniając trwałość i wytrzymałość struktur. Ich właściwy dobór i wykonanie są decydujące dla bezpieczeństwa i funkcjonalności końcowego produktu.

- połączenia zgrzewane metali to metoda łączenia elementów metalowych poprzez ich trwałe zespolenie bez użycia spoiw chemicznych, głównie poprzez nagrzewanie stykających się powierzchni do wysokiej temperatury i ich dociśnięcie, wywołanie plastycznej deformacji (np. zgrzewanie punktowe, garbowe, liniowe). Istnieje kilka głównych rodzajów połączeń zgrzewanych, które różnią się techniką wykonania i zastosowaniem (np. cienkie blachy, pręty, elementy karoserii samochodowych)⁵.

⁵ P. Kustroń, J. Leśniewski, B. Białobrzaska, „Nowoczesne metody zgrzewania tarcowego punktowego materiałów konstrukcyjnych”, *Welding Technology Review*, 88/8/2016, s. 43-46.



Połączenia zgrzewane metali dzielimy na:

- zgrzewanie na zimno (zimne zgrzewanie) wykonane przez nacisk elementów metalowych w temperaturze pokojowej, charakteryzuje się brakiem podgrzewania metali, stosowane podczas łączenia cienkich blach, elementów elektronicznych, elementów o niskich wymaganiach wytrzymałościowych,
- zgrzewanie na gorąco wykonane przez podgrzewanie elementów do temperatury zbliżonej do temperatury plastyczności metali, następnie ich łączenie przez nacisk, umożliwia uzyskanie mocnych połączeń, stosowane w przypadku metali o wysokiej odporności na zginanie i rozciąganie,
- zgrzewanie oporowe wykonane przez przepływ prądu elektrycznego przez styeczne powierzchnie elementów, co powoduje ich nagrzewanie się i złączenie pod naciskiem, typowe dla produkcji elementów samochodowych, rur, profili, dzielimy na zgrzewanie oporowe punktowe, liniowe, wycinanie, zgrzewanie z dociskiem,
- zgrzewanie ultradźwiękowe wykonane przez wykorzystanie wysokich częstotliwości drgań ultradźwiękowych do wywołania tarcia i nagrzewania powierzchni złączanych elementów stosowane w elektronice, łączeniu cienkich foli metalowych, komponentów elektronicznych,
- zgrzewanie metodą szczypcową i metodami specjalistycznymi obejmuje zgrzewanie plazmowe, laserowe, ultradźwiękowe, które pozwalają na precyzyjne i szybkie łączenie elementów.

Charakterystyka połączeń zgrzewanych:

- trwałość i wytrzymałość zależą od rodzaju metali, grubości, techniki zgrzewania oraz warunków obróbki,
- połączenia zgrzewane są często bez widocznych spoin, co poprawia estetykę i funkcjonalność,
- metoda ta jest szybka, automatyzowana i ekonomiczna w dużych seriach produkcyjnych,
- wymaga odpowiednich urządzeń i precyzyjnego doboru parametrów procesu.

Połączenia zgrzewane metali są szeroko stosowane w przemyśle motoryzacyjnym, budowlanym, elektronicznym i wielu innych dziedzinach ze względu na ich trwałość, estetykę i efektywność technologiczną.



- połączenia lutowane metali to technika łączenia elementów metalowych za pomocą stopu lutowniczego (spoiwo, lut) o niższej temperaturze topnienia niż metal łączony, stop lutowniczy jest topiony w miejscu styku, bez nadtapiania elementów łączonego metalu. Rozróżniamy **lutowanie miękkie** (poniżej 450°C, np. cyna), **lutowanie twarde** (powyżej 450°C, np. mosiądz, srebro). Zaletami połączeń lutowanych jest: niska temperatura lutowania w porównaniu z innymi metodami spawania, szybkość i prostota wykonania, możliwość naprawy i demontażu (w przeciwieństwie do spawania). Wadami połączeń lutowanych jest: mniejsza wytrzymałość mechaniczna w porównaniu do spawania, narażenie na uszkodzenia termiczne i korozję, jeśli nie zastosuje się odpowiednich materiałów zabezpieczających, czystość powierzchni do uzyskania dobrego połączenia. Stosowane w instalacjach hydraulicznych, pneumatycznych, chłodniczych, przemyśle elektrycznym, elektrotechnicznym i mechanicznym, elektrotechnice, elektronice (lutowanie elementów na płytkach PCB)⁶.
 - proces lutowania polega na rozgrzaniu miejsca połączenia do temperatury topnienia stopu lutowniczego, stop lutowniczy wnikając między powierzchnie, tworzy trwałe połączenie,
 - najczęściej stosowane rodzaje stopów lutownicznych to cyny z ołowiem (np. 60/40, 63/37) lub bez ołowiu (np. cyna z miedzią, srebrne i cynowe), dobór stopu lutowniczego zależy od wymagań dotyczących wytrzymałości, odporności na korozję i temperaturę pracy,
 - połączenia lutowane jest w pełni przewodzące elektrycznie i termicznie, jest elastyczne i może tłumić drgania, choć nie jest tak wytrzymałe mechanicznie jak spawanie czy nitowanie, połączenie lutowane jest odporne na warunki atmosferyczne, o ile zastosowano odpowiedni stop lutowniczy i odpowiednią technikę.

Połączenia lutowane metali są powszechnie stosowane ze względu na ich prostotę, szybkość oraz dobre przewodnictwo elektryczne i cieplne, choć ich wytrzymałość mechaniczna jest ograniczona w porównaniu z innymi metodami łączenia.

⁶ Z. Mirski, T. Wojdat, „Połączenia lutowane aluminium z miedzią stalą niestopową i stopową wykonane spoiwami cynkowymi”, *Welding Technology Review*, 85/4/2013, s. 2-8.



- połączenia klejone metali to technologia łączenia elementów metalowych za pomocą klejów specjalistycznych metalowych, które tworzą trwałą spoinę chemiczną po utwardzeniu, zapewniając trwałość i wytrzymałość układów. Zaletami połączeń klejonych metali jest: brak konieczności wiercenia i gwintowania, możliwość łączenia materiałów różnych typów (np. metali z tworzywami), lżejsza konstrukcja w porównaniu do spawania czy nitowania, możliwość łączenia cienkich i delikatnych elementów. Wadami połączeń klejonych są: wymagania co do powierzchni (konieczne staranne przygotowanie), ograniczona odporność na wysokie temperatury (w zależności od kleju), mniejsza wytrzymałość na wysokie obciążenia w porównaniu do spawania, czas utwardzania i konieczność odpowiednich warunków. Połączenia klejone metali stosuje się do łączenia elementów o małych obciążeniach, w przemyśle lotniczym, elektronicznym⁷.

Rodzaje klejów stosowanych do metali:

- kleje epoksydowe charakteryzują się wysoką wytrzymałością mechaniczną, dobrą odpornością na czynniki chemiczne i termiczne,
- kleje cyjanoakrylowe (superglue) charakteryzują się szybkim wiązaniem, stosowane przy niewielkich obciążeniach i dla elementów o małych powierzchniach styku,
- kleje poliuretanowe są elastyczne, odporne na wstrząsy i warunki atmosferyczne,
- kleje akrylowe charakteryzują się dobrymi właściwościami przy łączeniu różnych materiałów, szybkie wiązanie,
- kleje silikonowe stosuje się głównie do uszczelniania, ale również stosowane jako połączenia elastyczne.

Proces klejenia metali:

- przygotowanie powierzchni to oczyszczenie z tłuszczów, rdzy, starej farby, odtłuszczenie i ewentualne szlifowanie,
- aplikacja kleju to równomierne rozprowadzenie na powierzchni,
- dociskanie elementów to zapewnienie kontaktu i odpowiedniego nacisku podczas utwardzania,

⁷ M. Karny, „Połączenia klejone w strukturach kompozytowych-metodyka badań”, Prace Instytutu Lotnictwa, 3/244/2016, s. 97-108.



- utwardzenie to czas i warunki (temperatura, wilgotność) zależne od rodzaju kleju.

Cechy połączeń klejonych metali:

- rozkład naprężeń, kleje rozkładają naprężenia na dużej powierzchni, co zmniejsza ryzyko pęknięć,
- estetyka, brak widocznych elementów mocujących, co poprawia wygląd,
- elastyczność, niektóre kleje zapewniają elastyczność, co jest ważne przy materiałach pracujących pod obciążeniem dynamicznym,
- odporność na korozję, odpowiedni dobór kleju i przygotowanie powierzchni zapobiegają powstawaniu korozji.

Połączenia klejone metali stanowią nowoczesną i coraz bardziej popularną metodę łączenia elementów metalowych, szczególnie tam, gdzie konieczne jest zachowanie estetyki, minimalizacja masy lub łączenie różnych materiałów. Właściwy dobór kleju i odpowiednie przygotowanie powierzchni są kluczowe dla uzyskania trwałego i wytrzymałego połączenia.

Zalety i wady połączeń nierozłącznych metali

- duża wytrzymałość,
- sztywność konstrukcji,
- mniejsza masa (brak śrub itp.),
- brak możliwości demontażu,
- trudniejsze naprawy,
- często wymagają specjalistycznego sprzętu.

Połączenia rozłączne stosowane tam, gdzie potrzebny jest demontaż, naprawa lub regulacja, można je wielokrotnie rozłączać i łączyć bez niszczenia części, wymagają zwykłe elementów łączących⁸.

Do najczęściej stosowanych należą:

- połączenia gwintowe (śrubowe) metali to popularny i szeroko stosowany sposób łączenia elementów metalowych, oparty na zastosowaniu śrub, nakrętek (nakładek z gwintem) i podkładek. Łatwe do montażu i demontażu, uniwersalne, możliwość

⁸ M. Zasada, „Wykonywanie połączeń rozłącznych i nierozłącznych 724 [02]. 01. 05”, Instytut Technologii Eksploatacji – Państwowy Instytut Badawczy, Radom, 2007, s. 59.



luzowania się połączenia, konieczność stosowania zabezpieczeń. Zaletami połączeń gwintowanych jest: możliwość wielokrotnego demontażu, precyzyjne mocowanie, uniwersalność (dostępne w różnych rozmiarach i typach). Wadami połączeń gwintowanych jest: zużycie gwintu przy wielokrotnym demontażu, ryzyko poluzowania pod wpływem drgań (wymaga stosowania nakładek zabezpieczających np. podkładek sprężynowych), możliwość uszkodzenia gwintu podczas montażu/demontażu. Połączenia gwintowe stosuje się w konstrukcjach stalowych i metalowych, maszynach i urządzeniach, elementach mechanicznych i budowlanych, instalacjach przemysłowych⁹.

Budowa i elementy składowe:

- śruby, elementy z gwintem zewnętrznym, które wkręca się w element z gwintem wewnętrznym,
- nakładki (nakrętki), element z gwintem wewnętrznym, służący do mocowania śrub,
- podkładki, elementy rozkładające nacisk i chroniące powierzchnie łączonych elementów.

Rodzaje gwintów:

- gwint metryczny (np. M6, M10) najczęściej stosowany w Europie,
- gwint calowy (UNC, UNF) stosowany głównie w krajach anglojęzycznych,
- inne specjalne gwinty np. trapezowe, stożkowe.

Charakterystyka połączeń gwintowanych:

- łatwość montażu i demontażu umożliwia szybkie i wielokrotne łączenie elementów,
- możliwość regulacji np. dokręcanie i odkręcanie pozwalają na regulację siły mocowania,
- precyzyjne przenoszenie obciążeń szczególnie w połączeniach śrubowych o dużej wytrzymałości.

Właściwości mechaniczne:

- wytrzymałość na rozciąganie i ścinanie zależy od rodzaju użytego metalu, średnicy i gwintu,

⁹ W. Kaczmarek, „Analiza czysto-tarciowego modelu połączenia gwintowego”, Biuletyn Wojskowej Akademii Technicznej, 52.5/6/2003, s. 123-139.



- odporność na odkształcenia i zużycie odpowiednio dobrany materiał i obróbka zapewniają trwałość.

Połączenia gwintowane metali to niezawodny i elastyczny sposób łączenia elementów, który dzięki różnorodności dostępnych rodzajów i rozmiarów spełnia wymagania wielu dziedzin przemysłu i budownictwa.

- połączenia nitowane metali są popularną metodą łączenia elementów metalowych, za pomocą nitów, które po wbiciu lub zaciśnięciu tworzą trwałe połączenie, charakteryzującą się wytrzymałością i trwałością. Zaletami połączeń nitowanych jest: trwałość i odporność na drgania oraz wibracje, brak konieczności stosowania spawania czy klejenia, możliwość łączenia cienkich lub trudnych do spawania materiałów, łatwość naprawy i demontażu w niektórych przypadkach. Wadami połączeń nitowanych metalu jest: mniejsza elastyczność w modyfikacji połączenia w porównaniu do innych metod, wymóg precyzyjnego wiercenia otworów, specjalistycznych narzędzi do montażu. Połączenia nitowane metali stosowane są w budownictwie (np. konstrukcje stalowe, mosty), przemyśle lotniczym i kosmicznym, produkcji pojazdów i maszyn, elektronice i drobnych urządzeniach mechanicznych, podczas łączenia cienkich blach, w miejscach narażonych na drgania.¹⁰

Budowa i zasada działania:

- połączenie nitowe składa się z nitów, elementów metalowych, które przebijają łączone materiały i są następnie rozginane lub wyciskane, aby zapewnić trwałe zespolenie,
- nity mają zwykle głowicę z jednej strony i część rozprężną (np. trzpień i pierścień rozprężny) po drugiej stronie, co umożliwia ich zaciskanie.

Rodzaje nitów:

- nity pełne, jednolite, z jednego kawałka metalu, zapewniające najwyższą wytrzymałość,
- nity z główką (np. z główką stożkową lub walcową), używane głównie do połączeń, gdzie istotne jest estetyczne wykończenie,

¹⁰ A. Rudawska, M. Błaziak, „Analiza porównawcza siły niszczącej połączenia klejowe, klejowo-nitowe oraz nitowe stopu tytanu”, Technologia i Automatyzacja Montażu, 4/2011, s. 40-44.



- nity rozprężne (np. nity sprężynowe), po wbiciu rozprężają się, tworząc mocne połączenie,
- nity aluminiowe, stalowe, miedziane w zależności od wymagań dotyczących odporności na korozję i wytrzymałości.

proces łączenia:

- nity są wprowadzane przez otwory w łączonych elementach,
- po wbiciu trzpienia lub końcówki, jest on wyciskany lub rozginany, co powoduje powstanie głowicy lub rozprężenie, zapewniając trwałe połączenie.

Połączenia nitowane są niezawodnym i trwałym sposobem łączenia metali, szczególnie tam, gdzie wymagana jest wysoka wytrzymałość i odporność na drgania.

- połączenia wpustowe, klinowe, wielowypustowe służą do łączenia wałów z kołami zębatymi, pasowymi, umożliwiają przenoszenie momentu obrotowego¹¹.
 - połączenia wpustowe są to połączenia, w których elementy łączą się za pomocą wpustów, występów i odpowiadających im wycięć lub rowków. Zaletą połączeń wpustowych jest: zapewnienie dokładnego pozycjonowania elementów, duża wytrzymałość na obciążenia osiowe i momenty skręcające. Wadą połączeń wpustowych jest: konieczność dokładnego wykonania i dopasowania, trudności w montażu i demontażu. Połączenia wpustowe najczęściej stosuje się w przekładniach, wałach, osiach i elementach, gdzie wymagana jest precyzyjna transmisja momentu obrotowego,
 - połączenia klinowe opierają się na użyciu klinów, elementów w kształcie trójkąta lub trapezu, które służą do łączenia elementów poprzez wtopienie się między nie. Zaletą połączenia klinowego jest łatwość montażu i demontażu, możliwość regulacji położenia. Wadą połączenia klinowego jest możliwość zużywania się, wymóg regularnej konserwacji, mniejsza precyzyjność niż połączenia wpustowego. Połączenia klinowe często stosowane do mocowania lub napinania elementów, np. w przekładniach, kołach zębatych, łożyskach,
 - połączenia wielowypustowe to połączenia, w których elementy łączone mają wiele równoległych wypustów, które współpracują z odpowiadającymi im rowkami lub kanałami. Zaletą połączeń wielowypustowych jest: wysoka

¹¹ M. Łoboda, „Oprogramowanie wspomagające obliczenia konstrukcyjne połączeń kształtowych”, Inżynieria Rolnicza, 16/2012, s. 225-233.



wytrzymałość na obciążenia, dokładne pozycjonowanie, minimalne luz. Wadą połączeń wielowpustowych jest: skomplikowalność w wykonaniu i montażu, konieczność precyzyjnego dopasowania elementów. Połączenia wielowpustowe stosowane są głównie w przekładniach, wałach i kołach zębatych, gdzie wymagana jest duża wytrzymałość i precyzja w przenoszeniu momentu.

Wybór odpowiedniego rodzaju połączenia metalowego zależy od wymagań dotyczących wytrzymałości, precyzji, łatwości montażu i demontażu oraz warunków pracy.

Zalety i wady połączeń rozłącznych metali

- łatwy montaż i demontaż,
- możliwość wymiany części,
- uniwersalność,
- mogą się luzować,
- zwykle mniejsza wytrzymałość niż połączenia nierozdzielne,
- większa liczba elementów.

2. Kryteria doboru rodzaju połączenia

Przy wyborze metody łączenia metali bierze się pod uwagę:

- rodzaj i własności łączonych materiałów,
- grubość elementów,
- obciążenia działające na połączenie,
- wymagania szczelności i estetyki,
- możliwość demontażu,
- warunki eksploatacji (temperatura, środowisko pracy).

Tabela porównawcza połączeń rozdzielnych i nierozdzielnych metali:

Cecha	Połączenia rozdzielne	Połączenia nierozdzielne
Możliwość rozłączenia	Tak, bez uszkodzenia elementów	Nie, tylko przez zniszczenie połączenia



Cecha	Połączenia rozdzielne	Połączenia nierozdzielne
Trwałość połączenia	Mniejsza	Bardzo duża
Demontaż	Łatwy, wielokrotny	Brak możliwości demontażu
Wytrzymałość	Zwykle mniejsza	Zwykle większa
Zastosowanie	Elementy wymagające naprawy lub regulacji	Konstrukcje stałe
Przykłady	Śruby, wkręty, sworznie, kołki	Spawanie, nitowanie, lutowanie, klejenie
Liczba elementów	Większa (elementy łączące)	Mniejsza
Koszt naprawy	Niski	Wysoki lub niemożliwy

4.3. Rodzaje rur

Rury to elementy konstrukcyjne i instalacyjne, które służą do transportu cieczy, gazów, a także jako elementy konstrukcyjne w różnych gałęziach przemysłu. Wyróżniamy różne rodzaje rur, które różnią się materiałem, sposobem wykonania, kształtem oraz zastosowaniem¹²:

- rury stalowe
 - rury stalowe czarne (węglowe) wykonane z żelaza węglowego, stosowane głównie w instalacjach wodociągowych, gazowych oraz przemysłowych, charakteryzują się dużą wytrzymałością mechaniczną,
 - rury stalowe ocynkowane, pokryte warstwą cynku, co zapewnia odporność na korozję, używane w instalacjach wodociągowych i gazowych,
 - rury nierdzewne, wykonane ze stali nierdzewnej, odporne na korozję i wysokie temperatury, stosowane w przemysłach spożywczym, chemicznym, medycznym.
- rury aluminiowe lekkie, odporne na korozję, stosowane w przemyśle lotniczym, motoryzacyjnym, a także w instalacjach energetycznych i przemysłowych,

¹² M. Kowacki, „Rury i ich zastosowanie w budownictwie”, Nowoczesne Budownictwo Inżynieryjne, 1/2020, s. 56-64.



- rury miedziane charakteryzują się dobrą przewodnością cieplną i elektryczną, odpornością na korozję, często używane w instalacjach grzewczych, chłodniczych, wodociągowych oraz w elektrotechnice,
- rury plastikowe wykonane z tworzyw sztucznych to popularny rodzaj elementów instalacyjnych używanych w różnych gałęziach przemysłu, budownictwie oraz w instalacjach wodno-kanalizacyjnych, przemysłowych i energetycznych.

Charakteryzują się one szeregiem cech, które wyróżniają je na tle rur wykonanych z innych materiałów, takich jak:

- lekkość, rury z tworzyw sztucznych są znacznie lżejsze od rur metalowych (stalowych czy żeliwnych), co ułatwia ich transport, montaż i obsługę,
- odporność na korozję, tworzywa sztuczne są odporne na działanie korozji, chemikaliów i czynników środowiskowych, co zapewnia długą żywotność instalacji i minimalizuje konieczność konserwacji,
- łatwość montażu, dzięki lekkości i elastyczności, rury te można łatwo ciąć, łączyć i instalować nawet w trudnodostępnych miejscach,
- różnorodność materiałów, dostępne są różne rodzaje tworzyw sztucznych, takie jak polietylen (PE), polipropylen (PP), polichlorek winylu (PCV), poliuretan (PUR) czy polibutadien (PB), z różnymi właściwościami chemicznymi i mechanicznymi,
- dobre właściwości izolacyjne, tworzywa sztuczne mają niską przewodność cieplną, co pomaga w utrzymaniu temperatury przesyłanych mediów,
- odporność na działanie czynników chemicznych, wiele tworzyw jest odporne na działanie kwasów, zasad i innych substancji chemicznych, co czyni je odpowiednimi w różnych branżach przemysłowych.

Podstawowe rodzaje rur z tworzyw sztucznych:

- rury z polietylenu (PE) charakteryzują się dużą elastycznością, odpornością na uderzenia i chemikalia, są szeroko stosowane w instalacjach wodociągowych, gazowych oraz systemach przemysłowych, występują w wersjach wysokiej i niskiej gęstości (HDPE, LDPE),
- rury z polipropylenu (PP) są odporne na wysokie temperatury i chemikalia, stosowane w instalacjach ciepłej i zimnej wody, systemach grzewczych oraz przemysłowych,



- rury z polichlorku winylu (PCV), są twarde, sztywne i odporne na chemikalia, znajdują zastosowanie głównie w instalacjach wodociągowych, kanalizacyjnych, a także w systemach odwadniających oraz systemach przemysłowych, wyróżnia je niska cena i łatwość obróbki,
- rury chlorowane polichlorki winylu (CPVC), odporne na wysokie temperatury, stosowane w przemysłowych systemach wodociągowych,
- rury z poliuretanu (PUR), używane głównie jako element izolacji lub w specjalistycznych instalacjach,
- rury z polibutadienu (PB), są wykorzystywane w instalacjach budowlanych, głównie do ciepłej i zimnej wody oraz ogrzewania, dzięki swojej odporności na korozję, mróz, wysokie ciśnienie i chemię, a także wyjątkowej elastyczności i trwałości,
- rury z polietylenu o dużej gęstości (PEX-Al.-PEX), używane głównie w instalacjach wodnych i ogrzewaniu podłogowym.

Rury z tworzyw sztucznych są niezwykle wszechstronne, lekkie, odporne na korozję i chemikalia, co czyni je idealnym wyborem do szerokiego zakresu zastosowań. Ich właściwości pozwalają na szybki i ekonomiczny montaż, a różnorodność materiałów umożliwia dopasowanie do specyficznych wymagań technicznych i środowiskowych.

- rury ceramiczne, wykonane z ceramiki, używane głównie w przemysłach chemicznym i elektronicznym, do transportu agresywnych substancji lub w warunkach wysokiej temperatury,
- rury betonowe, wykonane z betonu zbrojonego lub niezbrojonego, stosowane głównie w kanalizacji, drenażu oraz w konstrukcjach podziemnych.

Wybór rodzaju rury zależy od jej przeznaczenia, wymagań dotyczących odporności na korozję, temperaturę, ciśnienie oraz warunki środowiskowe. Każdy z tych rodzajów ma swoje zalety i ograniczenia, co pozwala na ich optymalne dopasowanie do konkretnego zastosowania.

4.4. Metody i sposoby łączenia rur

Metody łączenia rur to techniki stosowane w celu trwałego i szczelnego połączenia elementów rurociągów. Wybór odpowiedniej metody zależy od rodzaju materiału,



zastosowania, ciśnienia, temperatury oraz wymagań dotyczących szczelności i wytrzymałości¹³.

- zgrzewanie metoda polega na stopieniu powierzchni łączonych elementów, które następnie łączą się na skutek schłodzenia. W przypadku metali najczęściej stosuje się zgrzewanie oporowe, łukowe, czy metodę zgrzewania termicznego. Zalety: trwałe, szczelne połączenie, brak elementów dodatkowych. Wady: wymaga specjalistycznego sprzętu i umiejętności, nie zawsze możliwe w przypadku różnych materiałów,
- spawanie łączenie rur poprzez stopienie ich krawędzi i ich zespawanie, np. spawanie MIG/MAG, TIG, czy spawanie łukowe. Zalety: mocne, trwałe i szczelne połączenie, szerokie zastosowanie. Wady: wymaga specjalistycznego sprzętu i doświadczenia, proces czasochłonny, konieczność ochrony przed zanieczyszczeniami,
- złączki i nakrętki (połączenia mechaniczne) polega na użyciu elementów takich, jak złączki, kolanka, trójniki, nakrętki, zatrzaski, które mocuje się na rurach. Zalety: szybki montaż, możliwość demontażu, stosunkowo niskie koszty. Wady: mogą wymagać uszczelnień, mniej wytrzymałe na wysokie ciśnienia i temperatury w porównaniu do zgrzewania czy spawania,
- lutowanie i spawanie (techniki termiczne)
lutowanie łączenie rur za pomocą stopienia materiału lutowniczego, który wypełnia szczeliny, najczęściej stosowane w instalacjach wodociągowych i gazowych,
spawanie łączenie rur za pomocą stopienia materiału spawalniczego, który wypełnia szczeliny, stosowane głównie w metalach,
- złącza typu szybkozłącza i kolanka specjalne systemy szybkozłączy, które pozwalają na szybkie i łatwe łączenie rur bez użycia narzędzi termicznych czy mechanicznych. Zalety: szybki montaż i demontaż, dobra szczelność w odpowiednich warunkach. Wady: mogą mieć ograniczenia co do wytrzymałości i ciśnienia.

Wybór metody łączenia rur zależy od materiału, przeznaczenia instalacji, wymagań dotyczących trwałości i szczelności. Metody mechaniczne (złączki, nakrętki) są szybkie i łatwe w montażu, natomiast metody termiczne (zgrzewanie, spawanie) zapewniają trwałe i szczelne połączenia, ale wymagają specjalistycznego sprzętu i umiejętności.

¹³ A. Wróblewska, „Ocena możliwości stosowania w instalacjach gazowych systemów rur wielowarstwowych z tworzyw sztucznych”, Nafta-Gaz, 66/2010, s. 597-601.



1. Metody łączenia rur stalowych można podzielić na kilka podstawowych grup, w zależności od zastosowanych technik i wymagań dotyczących wytrzymałości, szczelności oraz trwałości połączenia¹⁴.

Do najczęściej stosowanych metod należą:

- spawanie polega na stopieniu się brzegów rur i ich zespoleniu w miejscu łączenia. Zalety połączeń spawanych to: trwałość, szczelność, możliwość łączenia rur o różnych grubościach i rozmiarach. Wady połączeń spawanych: wymóg specjalistycznej wiedzy, sprzętu i odpowiednich warunków,
 - metody spawania:
 - spawanie metodą MIG/MAG (Metal Inert Gas / Metal Active Gas) szybkie i wydajne, stosowane w przemysłowych zastosowaniach,
 - spawanie metodą TIG (Tungsten Inert Gas) precyzyjne, idealne do cienkich rur i wymagających połączeń wysokiej jakości,
 - spawanie łukowe elektrodą otuloną popularne w pracach warsztatowych i na placu budowy.
- lutowanie i spawanie miękkie, łączenia wykonywane przez topienie materiału lutowniczego, który wnika między elementy, zapewniając szczelne połączenie. Stosowane do łączenia małych rur, instalacji hydraulicznych i gazowych o niskim ciśnieniu. Zalety połączeń lutowanych to: szybkość wykonania, prostota wykonania, nie wymaga wysokich temperatur. Wady połączeń lutowanych: nie wytrzymują wysokich ciśnień i obciążeń mechanicznych,
- złączki i elementy rozłączalne, stosowanie specjalnych złączek, nakrętek, obejm i innych elementów umożliwiających rozłączanie i ponowne łączenie rur. Zalety: łatwość montażu i demontażu, możliwość modyfikacji układu. Wady: mogą wymagać regularnej konserwacji i skracają czas użytkowania układu,
 - metody:
 - złączki gwintowane, rurki z gwintowanymi końcami, które są skręcane z złączkami,
 - złączki z zaciskami, obejm, zaciski sprężynowe, szybkozłączki,
 - złączki skręcane na wkładki lub kołnierze.

¹⁴ M. Piekarek, „Montaż instalacji z rur stalowych 713 [02] Z1/2/3/4.02”, Instytut Technologii Eksploatacji – Państwowy Instytut Badawczy, Radom, 2006, s. 49



- metody mechaniczne (łączenia na wcisk, zaciski, śruby), łączenia rur za pomocą elementów mechanicznych, które mocują je ze sobą bez konieczności spawania. Zalety metody mechanicznej to: szybki montaż, możliwość rozłączenia. Wady metody mechanicznej to: mniejsza szczelność, ograniczona wytrzymałość mechaniczna,
 - metody:
 - wkładki i tuleje,
 - zaciski sprężynowe,
 - śruby i nakrętki.
- metody łączenia na kołnierze, rurki końcowe wyposażone w kołnierze, które są mocowane do siebie za pomocą śrub i nakrętek. Metoda stosowana w instalacjach przemysłowych, wysokociśnieniowych systemach rurowych. Zalety metody łączenia na kołnierze to: łatwość rozłączania, możliwość wymiany elementów. Wady metody łączenia na kołnierze to: konieczność zastosowania uszczelek, które zapewniają szczelność.

Wybór metody łączenia rur stalowych zależy od wymagań projektowych, warunków pracy, rodzaju instalacji oraz oczekiwanej wytrzymałości i szczelności. Spawanie jest najtrwalszą i najbardziej uniwersalną metodą, natomiast złączki i mechaniczne metody łączą się z łatwością montażu i demontażu, ale mogą mieć ograniczenia w zakresie wytrzymałości.

2. Metody łączenia rur aluminiowych są kluczowe dla zapewnienia trwałości, szczelności i wytrzymałości konstrukcji¹⁵.

Do najczęściej stosowanych metod należą:

- lutowanie (lutowanie miękkie i twarde) polega na łączeniu rur poprzez stopienie cienkiej warstwy cyny lub innego stopu lutowniczego, który wnikać między powierzchnie, tworzy trwałe połączenie. Stosowane głównie w instalacjach wodnych, gazowych oraz w układach chłodniczych. Zalety połączeń lutowanych: szybkie i proste, dobra szczelność, nieduże wymagania sprzętowe. Wady połączeń lutowanych: wymagają czystości powierzchni, mniejsza wytrzymałość na duże obciążenia.

¹⁵ A. Ambroziak, P. Białucki, W. Derlukiewicz, A. Lange, „Właściwości aluminiowych złączy lutowanych lutami na osnowie cynku”, *Welding Technology Review*, 88/9/2016, s. 56-60.



- zgrzewanie metoda łączenia poprzez podgrzewanie i dociskanie rur do momentu ich zespolenia. Stosowane w przemyśle, szczególnie w dużych średnicach rur aluminiowych. Zalety metody zgrzewania: solidne i trwałe połączenia, brak materiałów dodatkowych. Wady metody zgrzewania: konieczność specjalistycznego sprzętu i wykwalifikowanej obsługi.
 - typy:
 - zgrzewanie oporowe,
 - zgrzewanie termiczne,
 - zgrzewanie doczołowe.
- spawanie (np. TIG, MIG/MAG) łączenie rur poprzez stopienie ich krawędzi i zespolenie ich za pomocą łuku elektrycznego. Zalety metody spawanej: zapewnia mocne, trwałe i hermetyczne połączenia, możliwość spawania cienkich i grubych rur. Wady metody spawanej: wymaga specjalistycznej wiedzy, sprzętu i odpowiednich parametrów spawania.
- złączki i elementy mechaniczne stosowanie specjalnych złączek, kolanek, mufek, śrub i nakrętek. Zalety: szybki montaż, łatwa rozbiórka, brak konieczności spawania czy lutowania. Wady: mogą wymagać uszczelnień, mniejsza szczelność przy dużych ciśnieniach.
- metody mechaniczne (np. zaciski, klamry) użycie elementów zaciskowych do tymczasowego lub trwałego łączenia rur. Zalety: szybki montaż, brak konieczności obróbki termicznej. Wady: mogą mieć ograniczone zastosowania w wysokociśnieniowych systemach.

Wybór metody łączenia rur aluminiowych zależy od wymagań technicznych, typu instalacji, jej przeznaczenia oraz dostępnego sprzętu. Najczęściej stosowane to lutowanie, spawanie oraz złączki mechaniczne, przy czym najważniejsze jest zapewnienie odpowiedniej szczelności i wytrzymałości połączenia.

3. Metody łączenia rur miedzianych są kluczowe dla zapewnienia szczelności, trwałości i bezpieczeństwa instalacji wodociągowych, grzewczych czy gazowych¹⁶.

¹⁶ M. Więcek, „Montaż instalacji z rur miedzianych 713 [02]. Z1/2/3/4.03”, Instytut Technologii Eksploatacji – Państwowy Instytut Badawczy, Radom, 2006, ss. 46



Do najczęściej stosowanych metod należą:

- lutowanie miękkie (lutowanie cynowe) jest to najpopularniejsza metoda łączenia rur miedzianych w instalacjach domowych i małych systemach. Polega na podgrzaniu połączenia i wprowadzeniu lutowia cynowego, które topi się i wypełnia szczeliny między rurami. Zalety lutowania miękkiego: prosta, szybka, niskokosztowa, dostępna do samodzielnego wykonania. Wady lutowania miękkiego: wymaga odpowiedniej techniki i czystości powierzchni, nie jest odpowiednia do wysokociśnieniowych instalacji,
 - proces:
 - przygotowanie powierzchni (oczyszczenie, odtłuszczenie i odrdzewienie),
 - podgrzewanie miejsc łączenia za pomocą palnika gazowego,
 - wprowadzenie lutowia cynowego do rozgrzanego połączenia,
 - schłodzenie i sprawdzenie szczelności.
- lutowanie twarde (lutowanie srebrne lub miedziowe) metoda stosowana głównie w przemysłowych instalacjach i w sytuacjach wymagających wysokiej szczelności i odporności na ciśnienie. Lutowanie twarde wykorzystuje materiały lutownicze o wyższej temperaturze topnienia. Zalety lutowania twardego: bardzo trwałe i szczelne połączenia. Wady lutowania twardego: wymaga specjalistycznych narzędzi i doświadczenia.
 - proces:
 - przygotowanie powierzchni,
 - podgrzewanie połączenia do wyższej temperatury,
 - wprowadzenie lutowia twardego (np. srebrnego),
 - schłodzenie i kontrola szczelności.
- złączki i elementy zaciskowe (złącza zaciskowe, kolanka, trójniki), elementy wykonane z miedzi lub innych materiałów, które umożliwiają łączenie rur bez konieczności lutowania. Zalety: szybki montaż, brak konieczności użycia otwartego ognia, możliwość demontażu. Wady: mogą być mniej szczelne, szczególnie przy dużych ciśnieniach, wymagają odpowiedniego doboru elementów.
 - metoda:
 - rury są wprowadzone do złączek lub zacisków,
 - złączki są mocowane za pomocą zacisków, śrub lub specjalnych narzędzi.



- zgrzewanie, rzadziej stosowana metoda, polegająca na łączeniu rur przez podgrzewanie i dociskanie, wymaga specjalistycznego sprzętu i jest bardziej popularny w przemysłowych zastosowaniach. Zalety zgrzewania: trwałe, szczelne połączenia. Wady zgrzewania: wysokie koszty i specjalistyczne narzędzia.
- metoda zaciskowa (złącza szybkozłączki) stosowana głównie w instalacjach przemysłowych i systemach tymczasowych. Zalety metody zaciskowej: szybki montaż, brak konieczności lutowania. Wady metody zaciskowej: ograniczona trwałość przy dużych obciążeniach.
 - proces:
 - rury są wkładane do złączki,
 - złącze jest zaciskane specjalnym narzędziem.

Wybór metody łączenia rur miedzianych zależy od zastosowania, wymagań dotyczących szczelności, ciśnienia, dostępnego sprzętu oraz umiejętności wykonawcy. Lutowanie miękkie jest najczęściej stosowane w instalacjach domowych, natomiast lutowanie twarde i złączki zaciskowe znajdują zastosowanie w instalacjach przemysłowych i specjalistycznych.

4. Metody łączenia rur polietylenowych (PE) to popularna technika łączenia elementów tego typu rur, stosowana głównie w instalacjach wodociągowych, gazowych oraz przemysłowych służąca zapewnieniu szczelności i trwałości instalacji¹⁷.
- zgrzewanie termiczne (zgrzewanie na ciepło):
 - zgrzewanie na ciepło (fusion welding) polega na podgrzaniu końców rur lub kształtek do momentu, do ich miejscowego stopienia (aż materiał stanie się plastyczny), a następnie ich połączeniu pod naciskiem, aby uzyskać trwałe złącze. Po ostygnięciu powstaje jednolita spójna struktura jest najbardziej popularna metoda i uznawana za najtrwalszą dla rur (PE), która zapewnia wysoką trwałość i powtarzalność połączeń. Wykorzystuje zgrzewarki elektrooporowe lub metodę ręczną, lub zgrzewarki automatyczne,

¹⁷ A. Rudzki, M. Kroczyk, „Możliwości wykorzystania rur z PE w gazownictwie”, Gaz, Woda i Technika Sanitarna, 2021, s. 2-5.



- metoda zgrzewania doczołowego (butt fusion) rury są ustawiane na specjalnym stanowisku, końce są podgrzewane do odpowiedniej temperatury, a następnie łączone pod naciskiem,
- metoda zgrzewania bocznego (socket fusion), stosowana głównie do łączenia rur i złączy o mniejszych średnicach, gdzie końcówki są podgrzewane i wstawiane do siebie,
- metoda zgrzewania na gorący warkocz (hot gas welding) używa gorącego gazu (np. propanu), który topi końce rur, a następnie łączy je pod naciskiem, stosowana szczególnie w naprawach lub łączeniach rur w trudnych warunkach.

Wskazówki:

- przed zgrzewaniem konieczne jest odpowiednie przygotowanie końcówek rur (cięcie na prosto, oczyszczenie z zanieczyszczeń),
 - ważne jest dobranie odpowiedniej temperatury i czasu zgrzewania zgodnie z normami i instrukcjami producentów,
 - zawsze należy korzystać z certyfikowanego sprzętu i przestrzegać wytycznych technicznych.
- złącza z elementami złącznymi (mechaniczne), złączki, mufy, kolana i inne elementy złączkowe, które mogą być łączone za pomocą specjalnych złączy mechanicznych, np. z gwintem, zaciskami, obejmami, lub złączami typu szybkozłącza (push-fit), które pozwalają na szybkie i łatwe łączenie rur bez konieczności ich podgrzewania. Zaletą metody jest szybka i prosta instalacja, nie wymaga specjalistycznego sprzętu.
 - złącza typu elektrooporowe (np. zgrzewanie elektrofuzji) (specjalne złączki z wbudowanymi elementami oporowymi) rura i złącze mają wbudowane elementy przewodzące, a złączenie dokonuje się przez podłączenie ich do zgrzewarki i przepuszczenie prądu elektrycznego, co powoduje podgrzewanie złączki i końcówki rury, aż do stopienia się materiału i trwałego połączenia. Ta metoda jest szczególnie popularna w instalacjach wodociągowych i gazowych. Zaletą metody jest powtarzalność, szybki czas zgrzewania, brak konieczności użycia otwartego ognia,
 - metoda zgrzewania ultradźwiękowego, wykorzystywana głównie w przemyśle do łączenia małych odcinków rur lub elementów, gdzie złącza się je za pomocą ultradźwięków generowanych przez specjalne urządzenia,



Najpopularniejszymi i najbardziej powszechnymi metodami łączenia rur (PE) są zgrzewanie na ciepło (zwłaszcza zgrzewanie doczołowe i boczne) oraz złącza elektrooporowe. Dobór metody zależy od średnicy rur, warunków instalacji, wymagań dotyczących szczelności oraz dostępnych narzędzi.

5. Metody łączenia rur polipropylenowych (PP) są kluczowe dla zapewnienia trwałości i szczelności instalacji¹⁸.

Do najczęściej stosowanych metod należą:

- zgrzewanie na gorąco (zgrzewanie termiczne) polega na podgrzaniu końców rur i kształtek do określonej temperatury, a następnie ich połączeniu pod naciskiem, stosowane w instalacjach wodociągowych, centralnego ogrzewania. Zalety metody: silne i trwałe połączenie, odporne na ciśnienie, szczelne i odporne na korozję.
 - proces:
 - wybranie odpowiednich elementów (rura i złączka),
 - podgrzewanie końców za pomocą specjalnego urządzenia (zgrzewarki),
 - po osiągnięciu temperatury końcówki łączone są pod naciskiem, co powoduje stapianie materiału.
- zgrzewanie na zimno (zgrzewanie z użyciem kleju) rzadziej stosowana metoda w przypadku rur PP, głównie w specjalistycznych zastosowaniach. Metoda mniej popularna dla rur PP, bardziej dla systemów z innymi materiałami.
 - proces:
 - użycie specjalistycznych klejów (np. na bazie klejów rozpuszczalnikowych) do łączenia końców rur.
- złącza mechaniczne (złączki, obejmy, kolanka, trójniki) używanie gotowych złązek i elementów łączących, które są mocowane na rurach. Stosowana w mniejszych instalacjach, podczas napraw systemów, gdzie konieczne jest rozłączanie. Zalety metody: szybki montaż, możliwość rozłączenia, nie wymaga specjalistycznego sprzętu.
 - rodzaje:
 - złączki skręcane (np. z nakrętkami i uszczelkami),

¹⁸ A. Rudawska, S. Kurzep, „Wybrane właściwości mechaniczne i szczelność połączeń klejowych rur z tworzyw polimerowych”, Przetwórstwo Tworzyw, 22.4/172/2016, s. 357-364.



- złączki zatrzaskowe,
- kolanka, trójniki, redukcje.

- złączki typu PUSH-fit (wtykowe) rury wciska się w specjalne złączki, które automatycznie zabezpieczają połączenie. Stosowana metoda w systemach wodociągowych i ogrzewania. Zalety: szybki montaż, brak konieczności używania narzędzi czy klejów.
- metoda klejenia (drogą chemiczną) rzadziej stosowana dla PP, głównie dla PE, ale mogą wystąpić specjalistyczne kleje do PP.

Najpopularniejszą i najtrwalszą metodą łączenia rur PP jest zgrzewanie na gorąco, które zapewnia szczelne, trwałe połączenie odporne na ciśnienie i czynniki zewnętrzne. Metody mechaniczne, takie jak złączki i obejmy są wygodne i szybkie, szczególnie w mniejszych instalacjach lub podczas napraw. Wybór metody zależy od rodzaju instalacji, wymagań dotyczących szczelności i dostępności sprzętu.

6. Metody łączenia rur PVC są kluczowe dla zapewnienia szczelności i trwałości instalacji wodociągowych, kanalizacyjnych oraz innych systemów rurowych¹⁹.

Do najczęściej stosowanych metod należą:

- złącza skręcane (złącza gwintowane) rury PVC wyposażone w gwinty, które umożliwiają ich bezpośrednie połączenie za pomocą nakrętek i tulei. Przeznaczone głównie do rur o mniejszych średnicach i systemów, które mogą wymagać demontażu. Zalety: łatwość montażu i demontażu. Wady: wymaga zastosowania uszczelek lub taśm teflonowych dla zapewnienia szczelności.
- złącza wciskowe (złącza typu push-fit) specjalne elementy z tworzywa, które poprzez nacisk lub zatrzaski łączą dwie rury bez konieczności użycia narzędzi specjalistycznych. Zalety: szybki i prosty montaż, zapewniają szczelne połączenie bez konieczności klejenia. Wady: wymagają odpowiedniego doboru średnic i typu rur.
- złącza klejone (złącza na podstawie kleju PVC) połączenia wykonywane przez sklejanie końców rur za pomocą specjalnego kleju do PVC (np. kleju na bazie rozpuszczalników). Zalety: trwałe i szczelne. Wady: wymaga dokładnego oczyszczenia

¹⁹ Z. Nitkiewicz, P. Chmielowiec, H. Stokłosa, M. Mierzwa, „Porównanie właściwości litego i porowatego PVC stosowanego do wyrobu rur”, *Kompozyty*, 4/2004, 421-425.



i odtłuszczenia powierzchni rur przed klejeniem, proces montażu jest czasochłonny i wymaga zachowania odpowiednich warunków (np. czas schnięcia).

- złącza skręcano-zaciskowe (złącza typu clamp lub zaciski) elementy łączące rurę poprzez zaciskanie na jej powierzchni, często stosowane w systemach tymczasowych lub serwisowych. Zalety: łatwe w montażu i demontażu, idealne do tymczasowych połączeń lub w miejscach, gdzie konieczny jest dostęp do instalacji.
- złącza typu mufy i kolana (złączki) elementy w kształcie mufy lub kolana, które służą do zmiany kierunku lub przedłużenia rur. Zalety: montowane za pomocą kleju lub złącz gwintowanych, umożliwiają tworzenie różnych konfiguracji instalacji.

Wybór metody łączenia rur PVC zależy od specyfiki instalacji, wymagań dotyczących trwałości, szczelności, łatwości montażu oraz demontażu. Najczęściej stosowane są złącza klejone ze względu na trwałość i szczelność, jednak w niektórych przypadkach korzysta się ze złączy gwintowanych lub złączy wciskowych.

7. Metody łączenia rur poliuretanowych (PUR) są kluczowe dla zapewnienia szczelności, trwałości i wytrzymałości instalacji²⁰.

Do najczęściej stosowanych metod należą:

- zgrzewanie termiczne (zgrzewanie na gorąco) polega na podgrzaniu końców rur do odpowiedniej temperatury, a następnie ich złączeniu pod naciskiem, co powoduje stopienie powierzchni i trwałe połączenie. Stosowane w instalacjach wymagających wysokiej szczelności i odporności na ciśnienie. Zalety: wysoka trwałość. Wady: konieczność specjalistycznego sprzętu i odpowiedniej temperatury.
- fuzja na zimno (złącza chemiczne) używa się specjalnych złączy lub klejów chemicznych (np. żywic epoksydowych lub klejów na bazie poliuretanu), które wiążą rurę i złącze na zimno. Stosowane do mniejszych średnic rur lub tam, gdzie zgrzewanie termiczne jest trudne. Zalety: łatwa w użyciu, nie wymaga specjalistycznego sprzętu. Wady: może mieć ograniczenia co do wytrzymałości na ciśnienie.
- złącza mechaniczne (ściskowe, gwintowane, szybkozłącza) polega na użyciu specjalnych złączek, nakrętek, obejm, zacisków lub szybkozłączy, które mocują rury

²⁰ A. Roszkowski, „Wytyczne Izby Gospodarczej Wodociągi Polskie odnośnie stosowania rur z tworzyw sztucznych”, Stowarzyszenie PRIK, Toruń, s. 1-14.



razem bez konieczności ich topienia czy klejenia. Stosowana w miejscach, gdzie wymagana jest szybka instalacja lub demontaż. Zalety: łatwa i szybka instalacja, możliwość demontażu. Wady: może być mniej odporne na dużą szczelność w porównaniu do zgrzewania.

- metoda łączenia za pomocą złączek typu push-fit lub typu szybkozłącza, rury wsuwa się do złączek, które zabezpieczają połączenie za pomocą uszczelek lub zacisków. Stosowana w systemach tymczasowych lub instalacjach, gdzie istotna jest szybka wymiana elementów. Zalety: szybka i wygodna. Wady: może wymagać okresowej kontroli szczelności.

Wybór metody łączenia rur poliuretanowych zależy od wymagań technicznych, warunków środowiskowych, potrzeb demontażu oraz dostępnego sprzętu. Zgrzewanie termiczne jest najtrwalsze i najbardziej odporne na ciśnienie. Natomiast metody mechaniczne i chemiczne są bardziej wygodne i szybkie w zastosowaniu, ale mogą mieć ograniczenia co do wytrzymałości i szczelności.

8. Metody łączenia rur polibutadienowych (PB) są kluczowe dla zapewnienia szczelności, trwałości i bezpieczeństwa instalacji²¹.

Do najczęściej stosowanych metod należą:

- złączki gwintowane (złącza gwintowane), rury są wyposażone w gwinty, które umożliwiają ich bezpośrednie połączenie za pomocą nakrętek lub innych elementów gwintowanych. Zalety: łatwość montażu i demontażu, możliwość rozłączania. Wady: ryzyko nieszczelności przy nieprawidłowym dokręceniu, konieczność stosowania uszczelek.
- złącza skręcane (złącza typu skręcanego), rury są łączone za pomocą specjalnych złączek, które można skręcać, często z elementami zaciskowymi lub śrubowymi. Zalety: szybki montaż, brak konieczności stosowania klejów czy lutowania. Wady: ograniczona odporność na wysokie ciśnienia i temperatury, konieczność stosowania odpowiednich elementów.
- złącza klejowe (z użyciem specjalnych klejów), rury i złączki łączy się za pomocą klejów chemicznych przeznaczonych do rur polibutadienowych. Zalety: dobre

²¹ M. Piekarek, „Montaż instalacji z rur z tworzyw sztucznych 713[02].Z1/2.04”, Instytut Technologii Eksploatacji – Państwowy Instytut Badawczy, Radom, 2006, s. 54.



szczelności, estetyczny wygląd, brak konieczności stosowania narzędzi mechanicznych. Wady: długi czas schnięcia, konieczność starannego przygotowania powierzchni, ograniczenia przy wysokich ciśnieniach i temperaturach.

- złączki zaciskowe (z zaciskami lub obejmami), używa się złączek z obejmami, które przez zaciskanie na rurze zapewniają połączenie. Zalety: szybki montaż, możliwość demontażu. Wady: konieczność stosowania narzędzi do zaciskania, ograniczenia w zakresie ciśnień.
- złącza typu push-fit lub szybkozłącza rury wsuwa się w złącze, które automatycznie blokuje rurę, zapewniając szczelne połączenie. Zalety: bardzo szybki montaż, nie wymaga użycia narzędzi ani klejów. Wady: ograniczona odporność na wysokie ciśnienia i temperatury, możliwość rozłączenia przy dużych obciążeniach.

Wybór metody łączenia rur polibutadienowych zależy od wymagań instalacji, takich jak ciśnienie, temperatura, trwałość, łatwość montażu i demontażu oraz koszty. W praktyce często stosuje się kombinację różnych metod, aby optymalnie dostosować system do potrzeb konkretnej instalacji.

9. Metody łączenia rur PEX-AL-PEX (polietylenowe z warstwą aluminiową) są kluczowe dla zapewnienia szczelności i trwałości instalacji²².

Do najczęściej stosowanych metod należą:

- złączki z gwintem (gwintowane połączenia) rura jest łączona z innym elementem za pomocą złączki gwintowanej. Zalety: łatwość montażu, możliwość rozmontowania. Wady: konieczność stosowania uszczelek i odpowiednich narzędzi, ryzyko nieszczelności przy niewłaściwym dokręceniu.
- złączki zaciskowe (np. zaciski typu „Clic” lub „Oetiker”) końcówki rur są wkładane do specjalnej złączki, a następnie zaciskane za pomocą narzędzia. Zalety: szybki montaż, trwałe połączenie. Wady: wymaga specjalnego narzędzia do zaciskania, konieczność stosowania odpowiednich złączek.
- złączki skręcane (np. złączki typu „compression”) końcówka rury jest wkładana do złączki, a następnie dokręca się nakrętkę, która zapewnia szczelne połączenie. Zalety:

²² M. Piekarek, „Montaż instalacji z rur z tworzyw sztucznych 713[02].Z1/2.04”, Instytut Technologii Eksploatacji – Państwowy Instytut Badawczy, Radom, 2006, s. 54.



prosty montaż, możliwość rozmontowania. Wady: konieczność stosowania uszczelek, ryzyko nieszczelności przy niewłaściwym dokręceniu.

- złączki PUSH-fit (złącza zatrzaskowe) rura jest wkładana do złączki, która za pomocą zatrzasku mocuje rurę na miejscu. Zalety: szybki i prosty montaż, nie wymaga narzędzi specjalistycznych. Wady: mogą być droższe, konieczność stosowania kompatybilnych złączy.
- spawanie lub zgrzewanie (rzadziej stosowane w przypadku rur PEX-AL-PEX) wysoce specjalistyczne metody, rzadziej stosowane w instalacjach domowych. Zalety: trwałe i szczelne połączenie. Wady: wymaga specjalistycznego sprzętu i wiedzy.

Najczęściej stosowanymi metodami łączenia rur PEX-AL-PEX są zaciski zaciskowe, złączki PUSH-fit oraz złączki skręcane. Wybór metody zależy od rodzaju instalacji, dostępności narzędzi oraz preferencji instalatora. Ważne jest, aby stosować wysokiej jakości złączki i przestrzegać instrukcji producenta, co zapewni trwałe i szczelne połączenie.

10. Metody łączenia rur karbowanych są kluczowe dla zapewnienia szczelności, trwałości i bezpieczeństwa instalacji. Rury karbowane, ze względu na swoją elastyczność i wytrzymałość na ciśnienie, często stosuje się w instalacjach wodno-kanalizacyjnych, grzewczych czy przemysłowych²³.

Do najczęściej stosowanych metod należą:

- złącza skręcane (połączenia gwintowane i skręcane) rury karbowane mogą być łączone za pomocą złączy gwintowanych lub skręcanych, które mocuje się na końcach rur. Metoda ta jest łatwa i szybka, ale wymaga odpowiedniej precyzji w wykonaniu gwintów. Jest stosowana głównie w instalacjach niskociśnieniowych.
- złącza zaciskowe (złączki zaciskowe) rury wprowadza się do specjalnych złączy, które następnie zaciska się za pomocą narzędzi ręcznych lub pneumatycznych. Zapewniają szczelność i możliwość demontażu, stosowane w instalacjach przemysłowych i domowych.

²³ M. Zasada, „Wykonywanie połączeń rozłącznych i nierozłącznych 724 [02]. 01. 05”, Instytut Technologii Eksploatacji – Państwowy Instytut Badawczy, Radom, 2007, s. 59.



- złącza z opaskami zaciskowymi (kolanowe, prostokątne) rury karbowane są łączone za pomocą opasek zaciskowych, które mocuje się wokół końców rur. Prosta i szybka metoda, stosowana głównie w instalacjach niskociśnieniowych.
- złącza wciskowe – rura jest wciskana do złącza specjalnego, które tworzy szczelne połączenie. Metoda szybka i niezawodna, często stosowana w systemach wodociągowych i ogrzewania podłogowego.
- złącza lutowane (lutowanie miękkie lub twarde) rury i złącza są łączone za pomocą lutowania, które tworzy trwałe i szczelne połączenie. Wymaga specjalistycznych narzędzi i umiejętności, zapewnia wysoką trwałość.
- złącza klinowe i zaciskowe (np. z systemem push-fit) rury wkłada się do złącza, które następnie zaciska się mechanicznie lub klinem. Szybkie w montażu, często stosowane w nowoczesnych instalacjach.

Wybór metody łączenia rur karbowanych zależy od rodzaju instalacji, wymagań dotyczących szczelności, ciśnienia oraz łatwości montażu i demontażu. Ważne jest, aby stosować odpowiednie złącza i narzędzia, co zapewni długotrwałe i bezawaryjne funkcjonowanie całego systemu.

4.5. Metody zgrzewania rur z tworzyw sztucznych

1. Metody zgrzewania rur z tworzyw sztucznych są szeroko stosowane w przemyśle instalacyjnym, wodociągowym, gazowym oraz w branży przemysłowej. Pozwalają one na trwałe i szczelne połączenia elementów wykonanych z różnych materiałów termoplastycznych²⁴.

Do najczęściej stosowanych metod należą:

- zgrzewanie na gorąco (zgrzewanie termiczne) polega na podgrzewaniu powierzchni łączonych elementów do momentu ich stopienia, a następnie łączeniu ich ze sobą pod naciskiem.
 - zgrzewanie na ciepło (butt fusion) końce rur są podgrzewane do temperatury topnienia, a następnie łączone pod naciskiem, stosowana najczęściej do rur z PE, PP, PVDF,

²⁴ M. Piekarek, „Montaż instalacji z rur z tworzyw sztucznych 713[02].Z1/2.04”, Instytut Technologii Eksploatacji – Państwowy Instytut Badawczy, Radom, 2006, s. 54.



- zgrzewanie metodą elektrooporową wykorzystanie fal radiowych lub mikrofal do podgrzewania powierzchni łączonych elementów stosowana głównie w produkcji rur z PE, PP.
 - zgrzewanie na zimno (zgrzewanie na styku) łączenie rur bez konieczności podgrzewania, za pomocą specjalnych złączek lub metod mechanicznych.
 - złącza kielichowe (złączki) włożenie końców rur do złączek, które zapewniają szczelne połączenie, stosowana w instalacjach wodociągowych i gazowych,
 - zgrzewanie na wcisk elementy są dociskane do siebie, często z użyciem specjalnych narzędzi.
 - zgrzewanie ultradźwiękowe wykorzystanie fal ultradźwiękowych do wytwarzania ciepła i stopienia powierzchni rur. Stosowane głównie w produkcji rur z PP, PVDF, w szczególnych zastosowaniach przemysłowych.
 - zgrzewanie metodą dociskową (zgrzewanie na zimno) łączenie elementów przez dociskanie bez podgrzewania, często w przypadku elementów o specjalnych właściwościach. Stosowane w niektórych specjalistycznych instalacjach i materiałach.
 - zgrzewanie przy użyciu specjalistycznych urządzeń np. zautomatyzowanych urządzeń, które precyzyjnie kontrolują parametry procesu (temperaturę, nacisk, czas). Stosowana podczas przemysłowej produkcji rur oraz w instalacjach przemysłowych. Wybór metody zgrzewania zależy od rodzaju tworzywa, wymagań szczelności, warunków pracy oraz dostępnych narzędzi i urządzeń. Zgrzewanie rur z tworzyw sztucznych zapewnia trwałe i szczelne połączenia, które są kluczowe dla bezpieczeństwa i funkcjonalności instalacji.
2. Metoda polifuzyjna zgrzewania rur z tworzyw sztucznych odnosi się do techniki łączenia elementów z tworzyw sztucznych, w której stosuje się zarówno zgrzewanie termiczne, jak i metodę mechanicznego docisku, w celu uzyskania trwałego i szczelnego połączenia.

Charakterystyka metody polifuzyjnej obejmuje:

- Połączenie metod: wykorzystuje zarówno podgrzewanie powierzchni rur i kształtek (np. przez zgrzewanie termiczne), jak i ich docisk mechaniczny, co zwiększa wytrzymałość i szczelność połączenia.

Opis procesu zgrzewania:



- najpierw elementy są odpowiednio ustawione i dopasowane,
- potem następuje podgrzewanie stref złącza, zwykle przez styczne podgrzewanie lub z wykorzystaniem zgrzewania oporowego,
- po osiągnięciu odpowiedniej temperatury, elementy są dociskane do siebie pod określonym naciskiem, co powoduje stopienie się materiału i powstanie trwałego połączenia.

Zalety metody polifuzyjnej:

- możliwość uzyskania bardzo mocnych i szczelnych połączeń,
- krótkie czasy realizacji procesu,
- możliwość łączenia rur o różnych wymiarach i materiałach, pod warunkiem odpowiedniego doboru parametrów.

Zastosowania:

- instalacje wodociągowe, kanalizacyjne,
- przemysł chemiczny i petrochemiczny,
- systemy grzewcze i chłodnicze.

Wymagania techniczne:

- precyzyjne ustawienie parametrów procesu (temperatura, czas, nacisk),
- użycie odpowiednich urządzeń zgrzewających dostosowanych do rodzaju tworzywa i średnicy rur,
- zachowanie odpowiednich warunków czystości w miejscu zgrzewania.

Metoda polifuzyjna zgrzewania rur z tworzyw sztucznych jest skuteczną techniką łączenia, która łączy zalety zgrzewania termicznego i mechanicznego, zapewniając trwałe, szczelne i odporne na obciążenia połączenia w różnych zastosowaniach przemysłowych i inżynierskich.

3. Metoda elektrooporowa zgrzewania rur z tworzyw sztucznych jest jednym z najpopularniejszych i najsukuteczniejszych sposobów łączenia elementów z termoplastów, zwłaszcza rur stosowanych w instalacjach wodociągowych, kanalizacyjnych oraz przemysłowych²⁵.

Charakterystyka metody elektrooporowej:

²⁵ P. Buczak, „Zgrzewanie elektrooporowe–zasady dobrych praktyk wykonawczych”, Wodociągi-Kanalizacja, 1/2015, s. 63-78.



- Zgrzewanie elektrooporowe opiera się na wykorzystaniu zjawiska przewodzenia prądu elektrycznego przez elementy z tworzyw sztucznych pod wpływem podgrzewania oporowego. Rury są łączone przy użyciu specjalnych złączy lub złączy wbudowanych w rurę, które mają elementy przewodzące.
- Podgrzewanie: Podczas procesu do złączy lub specjalnych elektrod przepuszcza się prąd elektryczny, co powoduje podgrzewanie się strefy styku na skutek oporu elektrycznego.
- Topnienie i złączenie: Po osiągnięciu odpowiedniej temperatury, pod wpływem ciepła, następuje stopienie powierzchni łączonych elementów. Po wyłączeniu prądu i ochłodzeniu się, tworzy się trwałe, szczelne połączenie.
- Automatyzacja: Proces jest często zautomatyzowany za pomocą specjalistycznych urządzeń, co zapewnia powtarzalność i wysoką jakość złącza.

Opis procesu zgrzewania rur elektrooporowej:

- Przygotowanie elementów: Rury i złączki są dokładnie oczyszczane z zabrudzeń, kurzu, tłuszczów i innych zanieczyszczeń, które mogą osłabić połączenie.
- Ustawienie elementów w urządzeniu: Rura i złączka są umieszczane w specjalnym urządzeniu elektrooporowym, które ustawia je w odpowiedniej pozycji i utrzymuje podczas procesu.
- Podłączenie do źródła prądu: Urządzenie podłącza elementy do źródła prądu o odpowiednich parametrach (napięcie, natężenie, czas).
- Proces zgrzewania: Przez elektrody lub elementy przewodzące przepływa prąd, powodując podgrzewanie powierzchni styku. Po osiągnięciu zadanej temperatury prąd jest odcinany, a elementy pozostają pod naciskiem, co zapewnia równomierne połączenie.
- Chłodzenie: Po zakończeniu podgrzewania i złączeniu elementów, następuje czas chłodzenia, podczas którego nie wolno przesuwować rur, aby połączenie było trwałe.
- Inspekcja: Po zakończeniu procesu często wykonuje się kontrolę wizualną lub pomiary, by upewnić się o jakości złącza.



Zalety metody elektrooporowej:

- wysoka jakość i trwałość połączeń,
- powtarzalność i wysokie parametry techniczne,
- możliwość automatyzacji procesu,
- krótki czas wykonania złącza,
- brak konieczności stosowania klejów czy innych środków chemicznych.

Metoda elektrooporowa jest skuteczną, szybką i niezawodną techniką łączenia rur z tworzyw sztucznych, szczególnie polietylenu (PE) i polipropylenu (PP). Dzięki automatyzacji i precyzji procesu, znajduje szerokie zastosowanie w instalacjach przemysłowych i komunalnych, gwarantując szczelność i długotrwałość połączeń.

4. Metoda doczołowa zgrzewania rur z tworzyw sztucznych znana również jako zgrzewanie doczołowe (ang. butt welding), polega na połączeniu dwóch końców rur poprzez ich dokładne dopasowanie i stopienie powierzchni styku. Jest to jedna z najpopularniejszych i najbardziej skutecznych metod zgrzewania rur z tworzyw sztucznych, szczególnie wykorzystywana w instalacjach wodociągowych, gazowych czy przemysłowych²⁶.

Charakterystyka i opis metody:

- Przygotowanie elementów: Rury są przycinane do odpowiednich długości i końce są dokładnie oczyszczane z zabrudzeń, zadziorów, tłuszczy i innych zanieczyszczeń. Końcówki rur są formowane tak, aby miały równą i gładką powierzchnię, zwykle poprzez cięcie pod kątem prostym lub specjalnym urządzeniem do zgrzewania.
- Ustawienie rur: Rury są precyzyjnie ustawiane w urządzeniu zgrzewalniczym, tak, aby ich końce stykały się dokładnie w osi, zapewniając równomierne połączenie.
- Podgrzewanie: Końce rur są podgrzewane do odpowiedniej temperatury, zwykle za pomocą specjalnych płyt grzewczych lub elementów grzewczych, które zapewniają równomierne rozprowadzenie ciepła. Temperatura i czas podgrzewania są ściśle określone dla danego rodzaju tworzywa i średnicy rur.

²⁶ J. Słania, G. Raczyński, „Zgrzewanie gazociągu średniego ciśnienia z rur polietylenowych”, Welding Technology Review, 85.5/2013, s. 18-25.



- Złączenie: Po osiągnięciu odpowiedniej temperatury końce rur są natychmiast łączone ze sobą, przyciskając je mocno i równomiernie. Tworzy się trwałe połączenie, które w wyniku chłodzenia tworzy integralną całość.
- Chłodzenie: Połączenie jest utrzymywane w stabilnym położeniu przez określony czas, aż do pełnego ostygnięcia i utwardzenia złącza.

Zalety metody:

- trwałe i szczelne połączenie,
- brak konieczności stosowania dodatkowych materiałów (np. złączek czy klejów),
- możliwość zgrzewania rur o dużych średnicach,
- powtarzalność i wysokie parametry jakościowe połączenia.

Wady metody:

- wymaga specjalistycznego sprzętu i wykwalifikowanej obsługi,
- niewłaściwe ustawienie lub parametry mogą prowadzić do uszkodzeń złącza.

Metoda doczołowa zgrzewania rur z tworzyw sztucznych jest skuteczną i szeroko stosowaną techniką łączenia elementów z tworzyw, zapewniającą trwałe i szczelne połączenia w różnych zastosowaniach przemysłowych i instalacyjnych.

Literatura

- Ambroziak A., Białucki P., Derlukiewicz W., Lange A., „Właściwości aluminiowych złączy lutowanych lutami na osnowie cynku”, *Welding Technology Review*, 88/9/2016.
- Buczak P., „Zgrzewanie elektrooporowe–zasady dobrych praktyk wykonawczych”, *Wodociągi-Kanalizacja*, 1/2015.
- Chmiel P., „Wykonywanie nierozłącznych połączeń blach 721 [01]. Z1. 06”, Instytut Technologii Eksploatacji – Państwowy Instytut Badawczy, Radom, 2007.
- Grigorenko M., Gieorgij W., Kostin A., „Spawalność stali i kryteria jej oceny”, *Welding Technology Review*, 85/7/2013.
- Kaczmarek W., „Analiza czysto-tarciowego modelu połączenia gwintowego”, *Biuletyn Wojskowej Akademii Technicznej*, 52.5/6/2003.
- Karny M., „Połączenia klejone w strukturach kompozytowych-metodyka badań”, *Prace Instytutu Lotnictwa*, 3/244/2016.



- Kowacki M., „Rury i ich zastosowanie w budownictwie”, Nowoczesne Budownictwo Inżynieryjne, 1/2020.
- Kustroń P., Leśniewski J., Białobrzeska B., „Nowoczesne metody zgrzewania tarcowego punktowego materiałów konstrukcyjnych”, Welding Technology Review, 88/8/2016.
- Łoboda M., „Oprogramowanie wspomagające obliczenia konstrukcyjne połączeń kształtowych”, Inżynieria Rolnicza, 16/2012.
- Mirski Z., Wojdat T., „Połączenia lutowane aluminium z miedzią stałą niestopową i stopową wykonane spoiwami cynkowymi”, Welding Technology Review, 85/4/2013.
- Nitkiewicz Z., Chmielowiec P., Stokłosa H., Mierzwa M., „Porównanie właściwości litego i porowatego PVC stosowanego do wyrobu rur”, Kompozyty, 4/2004.
- Piekarek M., „Montaż instalacji z rur stalowych 713 [02] Z1/2/3/4.02”, Instytut Technologii Eksploatacji – Państwowy Instytut Badawczy, Radom, 2006.
- Piekarek M., „Montaż instalacji z rur z tworzyw sztucznych 713[02].Z1/2.04”, Instytut Technologii Eksploatacji – Państwowy Instytut Badawczy, Radom 2006.
- Roszkowski A., „Wytyczne Izby Gospodarczej Wodociągi Polskie odnośnie stosowania rur z tworzyw sztucznych”, Stowarzyszenie PRiK, Toruń 2008.
- Rudawska A., Błaziak M., „Analiza porównawcza siły niszczącej połączenia klejowe, klejowo-nitowe oraz nitowe stopu tytanu”, Technologia i Automatyzacja Montażu, 4/2011.
- Rudawska A., Kurzep S., „Wybrane właściwości mechaniczne i szczelność połączeń klejowych rur z tworzyw polimerowych”, Przetwórstwo Tworzyw, 22.4/172/2016.
- Rudzki A., Krocze M., „Możliwości wykorzystania rur z PE w gazownictwie”, Gaz, Woda i Technika Sanitarna, Warszawa, 2021.
- Rydzkowski T., Michalska-Požoga I., „Rozwój technik łączenia materiałów spawanie tworzyw polimerowych”, Welding Technology Review, 87/11/2015.
- Słania J., Raczyński G., „Zgrzewanie gazociągu średniego ciśnienia z rur polietylenowych”, Welding Technology Review, 85.5/2013.
- Więcek M., „Montaż instalacji z rur miedzianych 713 [02]. Z1/2/3/4.03”, Instytut Technologii Eksploatacji – Państwowy Instytut Badawczy, Radom 2006.
- Wróblewska A., „Ocena możliwości stosowania w instalacjach gazowych systemów rur wielowarstwowych z tworzyw sztucznych”, Nafta-Gaz, 66/2010.
- Zasada M., „Wykonywanie połączeń rozłącznych i nierozłącznych 724 [02]. 01. 05”, Instytut Technologii Eksploatacji – Państwowy Instytut Badawczy, Radom 2007.



5. English for Mechanical Engineering

(dedicated for students of metallurgical specialization)

(Paweł Rutecki; PhD Eng.)

1. Types of metal alloys and the effect of alloying elements on final product properties

Metals are a broad class of materials widely used in numerous industrial fields. It is reasonable to assume that every manufacturing process involves metals in some form, for example as machine components, tools, or moulds and dies. In addition, metals are often the final product themselves. They are widely used due to their mechanical properties, including strength, hardness, ductility, and elasticity, as well as their chemical properties such as corrosion resistance, resistance to acids and alkalis, and overall chemical stability. Metals are also used in industry for their ability to conduct heat and electricity.

However, to be precise, metals are far more commonly used in the form of alloys rather than as pure elements. Metal alloys are materials formed by combining a base metal with one or more additional elements. Compared to pure metals, alloys often exhibit improved performance, for example higher strength, greater hardness, or enhanced corrosion resistance. In some cases, alloying is also applied to reduce the overall material cost, such as in gold and copper alloys.

Copper

Copper is one of the oldest metals used by humans, with its use dating back to ancient civilizations. It was widely used for making tools, weapons, and decorative items due to its malleability and durability. In addition, ancient cultures used copper to create early coins and various everyday objects.

Copper is a non-ferrous material valued primarily for its excellent electrical and thermal conductivity, which makes it essential in electrical engineering and electronics. That is why in the 19th and 20th centuries it became a key industrial material. It exhibits good ductility and malleability, which allow it to be easily formed into wires, sheets, and tubes. In addition, copper has good corrosion resistance. When exposed to atmospheric conditions, it develops a stable patina layer that protects the underlying metal from further degradation.



Copper is rarely used in its pure form for mechanical applications due to its relatively low strength. Instead, it is widely used in the form of alloys such as brass and bronze. Brass, a copper–zinc alloy, is commonly applied in plumbing fittings as well as decorative elements and hardware. Bronze, a copper–tin alloy, is used in applications such as coins, as well as bearings and bushings, where good wear resistance is required.

Steel

Steel is one of the most widely used engineering materials in the world. Its large-scale industrial production began in the 19th century with the development of modern steelmaking processes such as the Bessemer process, which made steel an economically viable material for structural and engineering applications. Steel is an alloy primarily composed of iron and carbon. It can be divided into non-alloy steel and alloy steel. Non-alloy steel consists mainly of iron and carbon, whereas alloy steel, in addition to these two elements, contains other alloying elements that influence its final properties.

Typically, steel contains approximately 0.5-1.5% carbon. An increase in carbon content, and therefore a higher proportion of hard and brittle cementite, results in greater steel hardness. In alloy steels, the effect of carbon on hardness is also associated with the tendency of certain alloying elements, especially chromium, to form compounds with carbon, primarily very hard carbides. The addition of chromium also promotes grain refinement of the steel's crystalline structure and increases hardenability. Nickel facilitates the hardening process by slowing down its rate. Manganese increases resistance to impact, tensile loading, and wear. Tungsten, like manganese, increases strength and wear resistance, but also helps create the fine-grained structure of steel. Vanadium and cobalt increase hardness and promote a fine-grained structure. Another typical additive is molybdenum. It increases hardenability, strength, creep resistance, and reduces brittleness.

Alloy steels are widely used in the production of machine parts and structural components that require high strength and durability. They are commonly applied in the automotive and aerospace industries, as well as in the manufacture of tools, gears, shafts, and bearings. Steel is also widely used in construction for load-bearing structures, bridges, and industrial equipment.



Steel is not inherently resistant to corrosion. When it is exposed to water and oxygen, a chemical reaction occurs and the material gradually oxidizes and forms iron oxides, commonly known as rust. In most modern applications, this issue is addressed through protective coatings. Steel can be covered with a variety of protective materials, such as paint used in automotive applications or enamel coatings applied to household appliances such as refrigerators and pots. Alternatively, corrosion resistance can be improved by alloying steel with elements such as nickel and chromium, resulting in stainless steel, which can reduce the tendency to rust.

Aluminium

Aluminium is the second most plentiful metallic element on Earth. At the end of the 19th century, it became an economically competitive material for engineering applications. Today, aluminium is still widely used in many industrial sectors, mainly in the form of alloys. Different aluminium alloys are applied depending on the required properties, which in turn depend on the alloying elements.

Magnesium decreases the electrical conductivity of aluminium alloys while enhancing their mechanical properties and improving corrosion resistance. Silicon improves mechanical properties and does not affect corrosion resistance when it remains in solid solution. Iron enhances mechanical properties and reduces the corrosion resistance of aluminium alloys. Copper also improves mechanical properties but simultaneously decreases corrosion resistance and does not significantly affect formability. Manganese contributes to improved mechanical properties and increases corrosion resistance. A small addition of zinc improves formability. Titanium refines the grain structure, enhances mechanical properties, and does not reduce corrosion resistance. Chromium improves mechanical properties and reduces the susceptibility of AlMg alloys to stress corrosion cracking. Nickel decreases corrosion resistance, while improving mechanical properties, particularly at elevated temperatures.

Based on the type of alloying element, eight groups of aluminium alloys can be distinguished, as presented in the Table 1 below.

Table 1 Aluminium alloys

Alloy code	Alloying element	Example of application
1XXX	None (pure aluminium)	Primarily used in the electrical and chemical industries.



2XXX	Copper	Widely used in aircraft.
3XXX	Manganese	General-purpose alloy for architectural applications and various products.
4XXX	Silicon	Welding rods and brazing sheet.
5XXX	Magnesium	Boat hulls, gangplanks and other products exposed to marine environments and also food packaging.
6XXX	Silicon + Magnesium	Architectural extrusions.
7XXX	Zinc	Aircraft structural components and other high-strength applications.
8XXX	Others	Bottle caps, packaging foil.

Aluminium alloys are characterized by good corrosion resistance, as a thin, dense layer of aluminium oxide forms spontaneously on their surface. This oxide layer acts as a protective barrier, preventing further oxidation and limiting the penetration of corrosive agents.

Aluminium and steel have partially overlapping applications, but aluminium offers an advantage over steel due to its high strength-to-weight ratio and good formability, enabling the production of strong yet lightweight structures, which is particularly important in the automotive and aerospace industries. A comparison of the weights of aluminium and steel is presented in the table below.

Table 2 Comparison of aluminium and steel

	Steel	Aluminium
Weight of 1 m ² sheet thickness 1.5 mm	~ 11.8 kg	~ 4 kg
Weight comparison (1 m ²) when keeping the same stiffness	Thickness: 0.8 mm Weight: 6.24 kg	Thickness: 1.2 mm Weight: 3.24 kg

As shown in the examples above, metals and their alloys exhibit a wide range of properties, and the appropriate material must be selected depending on the intended application. Therefore, careful material selection is essential to meet specific performance requirements.



EXERCISES

Task 1

Match the words from the text with their Polish equivalents.

- | | |
|---------------------------|--------------------------------|
| 1. alloy | A. konstrukcja nośna |
| 2. corrosion resistance | B. wytrzymałość na rozciąganie |
| 3. grain structure | C. pierwiastek |
| 4. hardenability | D. odporność na korozję |
| 5. tensile strength | E. odporność na ścieranie |
| 6. load-bearing structure | F. hartowność |
| 7. chemical element | G. stop metalu |
| 8. wear resistance | H. struktura ziarna |

Task 2

Make adjectives from the following nouns.

Example: elasticity → elastic

1. corrosion → _____
2. strength → _____
3. stability → _____
4. hardness → _____
5. durability → _____
6. resistance → _____

Task 3

Fill in the missing words in the sentences.

(steel, alloying, corrosion, conductivity, ductility, patina, strength, aluminium)

1. Copper is widely used due to its excellent electrical and thermal _____.
2. Metals are often used in the form of _____ rather than pure elements.
3. Steel contains iron and carbon, and may include additional _____ elements.
4. Copper develops a protective layer called a _____ on its surface.
5. Increasing carbon content increases the _____ of steel.
6. Aluminium is valued for its low weight and good _____.



7. Exposure to oxygen and water causes _____ of steel.
8. Copper is not often used in pure form because of its low _____.

Task 4

Match the terms with their definitions.

- | | |
|---------------------|---|
| 1. Aluminium oxide | A. A process of reducing grain size in a metal structure to improve properties. |
| 2. Oxidation | B. A steel resistant to rust due to alloying elements like chromium. |
| 3. Carbide | C. A compound formed between carbon and a metal. |
| 4. Hardenability | D. A thin protective layer formed on aluminium surface. |
| 5. Stainless steel | E. A chemical reaction with oxygen leading to rust formation. |
| 6. Grain refinement | F. The ability of steel to be hardened through heat treatment. |

2. Casting processes in manufacturing

Casting is a manufacturing process in which molten metal is poured into a mould containing a negative impression of the desired shape. The metal enters the mould through a hollow channel known as a sprue. After the mould is filled, both the metal and the mould are cooled until the metal solidifies. Once solid, the final part (also called a casting) is removed from the mould. Casting process is widely used to produce solid ingots for further processing, but it is also an important method for manufacturing finished and semi-finished components. This process makes it possible to create complex shapes without the need to join separate parts, which could introduce weak points and additional operations. In addition, when producing detailed metal objects, casting is more material-efficient than machining or cutting a component from a larger metal block.

Gravity die casting is one of the oldest techniques used for casting metals and metal alloys. In this process, molten metal is poured into a mould cavity through an opening located on the upper surface. The metal fills the cavity solely due to the force of gravity, without the



use of any external pressure. After the mould is filled, the metal begins to cool and solidify. Once the casting has solidified, it can be removed from the mould.

The main advantage of gravity die casting is that it does not require any external force to push the molten metal into the mould, which helps reduce overall production costs. This method is generally cheaper than other casting techniques. However, manual gravity die casting is slower than many alternative casting processes, which makes it more suitable for small- or medium-scale production.

Another technique is **die casting**. In contrast to gravity casting, in die casting metal is forced into the mould under high pressure. This method is typically used to produce small precision parts such as housings and brackets for the automotive industry. Aluminium and copper alloys are commonly used in die casting, whereas ferrous alloys are more difficult to process in this way due to the fact that the moulds are often made from hardened steel.

This technique offers several advantages, including very low unit cost, high definition and surface finish, as well as good dimensional accuracy, which reduces the number of operations required for further machining. However, die casting is also unsuitable for large castings. Additionally, the process requires significant capital investment, and the alloys used must have a low melting point, often at the expense of other properties such as strength and stiffness.

Another widely used method in the metal casting industry is **direct-chill casting** (DC casting). It is applied to produce billets/slabs, which are cylindrical/rectangular ingots, intended for further processing such as extrusion and rolling. The term direct-chill casting refers to the fact that the solidifying metal is cooled directly with water. The process involves two main stages of cooling. Primary cooling takes place inside the mould, where the metal initially solidifies against the mould wall. Secondary cooling occurs when the billet leaves the mould and its surface is cooled directly by water, typically by intensive water jets or immersion. In some systems, an additional cooling zone is applied at the bottom block to improve temperature control. During casting, molten metal flows into a water-cooled, bottomless mould, where a solid shell forms on the outer surface. This shell provides sufficient strength to support the liquid core while maintaining the shape of the billet.

DC casting is mainly applied to aluminium and its alloys, since ferrous alloys are generally less suitable for this process due to their much higher melting temperatures. A major advantage of this method is controlled solidification, which improves product quality and



reduces defects. However, DC casting requires precise control of cooling conditions, since improper cooling may lead to cracking, segregation, or other internal defects.

An interesting technique that differs significantly from other casting methods is **continuous casting**. It is used to produce a large volume of metal in simple shapes by continuously casting and rolling the material in a single process. In this method, molten metal is continuously fed into the gap between two rotating rolls, where it solidifies and is simultaneously formed into a strip. The resulting slabs are subjected to further processing, in order to obtain more practical and usable products.

This method provides important advantages, such as a high yield for a given volume of liquid metal, mainly due to the absence of a contraction pipe, as well as good surface finish. Furthermore, the method offers a very low unit cost, which results from the very high volume of metal that can be produced. However, continuous casting requires a large capital investment to set up the process, and it is limited to producing only simple shapes with a constant cross-section.

Regardless of the chosen casting method, the process is preceded by the melting of the metal alloy. This stage is of critical importance, as many casting defects can be prevented already at this point, including non-metallic inclusions, gas porosity, and structural inhomogeneity across the casting cross-section. There are a few key methods used to help ensure appropriate quality of the cast, and one of them is refining. It involves the removal of unwanted impurities and dissolved gases from the liquid metal, thereby reducing the risk of defects such as porosity and contamination-related inclusions. Another important aspect of this stage is the removal of slag/dross formed on the surface of the melt, which helps to improve the overall cleanliness of the alloy. Additionally, filtration systems are commonly used to purify the molten metal before casting, where ceramic or other filters retain solid impurities and prevent their transfer into the mould, thus reducing the occurrence of inclusions in the final casting. Each of these stages, including refining, slag removal, and filtration, is crucial for ensuring the high quality of the resulting ingots.



EXERCISES

Task 1

Find in the text words for:

- | | |
|-----------------------------|-----------------------------|
| 1) odlewanie grawitacyjne | 11) koszt jednostkowy |
| 2) odlewanie ciśnieniowe | 12) inwestycja kapitałowa |
| 3) odlewanie ciągłe | 13) rafinacja |
| 4) forma odlewnicza | 14) porowatość |
| 5) uzysk | 15) wtrącenia niemetaliczne |
| 6) ciekły metal | 16) żużel/zgary |
| 7) krzepnięcie / zestalenie | 17) filtracja |
| 8) wlew | 18) dokładność wymiarowa |
| 9) wlew (kanał wlewowy) | 19) obróbka skrawaniem |
| 10) temperatura topnienia | 20) wydajność procesu |

Task 2

Answer the following questions based on the text:

1. What is the main difference between gravity die casting and die casting?

.....
.....

2. What are the main advantages of continuous casting mentioned in the text?

.....
.....

3. Why is filtration used before casting molten metal?

.....
.....

Task 3

Match the two parts of the sentences:

- 1) Casting is a manufacturing process in which molten metal
- 2) The metal enters the mould through a hollow channel
- 3) Gravity die casting is a process in which molten metal
- 4) In gravity die casting, the metal fills the cavity



- 5) Once the casting has fully solidified, it
 - 6) In die casting, metal
 - 7) Continuous casting is used to produce a large volume of metal
 - 8) Filtration systems are commonly used to purify molten metal
-
- a) is forced into the mould under high pressure.
 - b) can be removed from the mould.
 - c) known as a sprue.
 - d) is poured into a mould cavity through an opening.
 - e) solely due to the force of gravity.
 - f) is poured into a mould containing a negative impression of the desired shape.
 - g) in simple shapes with a constant cross-section.
 - h) before casting to retain solid impurities.



3. Metal forming and heat treatment

Metal processing is a complex and multi-stage industrial route. Regardless of the casting method used, an ingot or slab obtained after solidification is not a final product and must undergo further processing in order to achieve the required mechanical properties, dimensions and surface quality. Depending on the final application, the transformation from the slab to semi-products such as sheets, strips, or coils may involve several or even more than a dozen successive operations. Moreover, the complete processing chain, from the first treatment of the casting to the final product, usually does not take place within a single company. The following sections describe typical operations commonly applied in metal processing facilities.

Scalping is a surface machining operation performed mainly on ingots (slabs) produced by gravity casting. In this process, a milling machine removes a thin layer from the surface of the ingot. The primary purpose of scalping is to eliminate the casting skin, which may contain cracks, oxide films, gas porosity, non-metallic inclusions, and a zone of frozen crystals. Since the surface of the cast ingot is relatively rough and structurally non-uniform, this outer layer tends to crack during hot rolling, which may result in defects on the rolled product surface. For this reason, ingots are typically scalped to a depth sufficient to fully remove the defective layer. In the case of aluminium and soft aluminium alloys, the removed layer is commonly in the range of 3 to 6 mm per side, while for harder alloys it may reach up to 10 mm per side. In contrast, in continuous casting the semi-product surface is generally more uniform and does not usually require such extensive surface removal.

Heat treatment, in its broadest sense, refers to a group of heating and cooling operations performed in order to modify the microstructure, mechanical properties, and residual stress state of a metallic product. Heat treatment processes are essential in metal processing, as they enable the adjustment of strength, ductility, and internal structural homogeneity prior to forming operations such as rolling.

Homogenization is a heat treatment process in which ingots are heated to a temperature close to their solidification point. This thermal input increases atomic mobility, allowing alloying elements to diffuse more effectively and promoting the formation of a more uniform solid solution. Once the target temperature is reached, the ingot is maintained at this level for a specific period known as the **soaking** time, during which homogenization progresses toward completion. The objective of this process is to ensure that the final rolled product exhibits



uniform chemical composition and consistent physical properties throughout its volume. After the homogenization stage is completed, ingots are typically transferred to hot rolling. In certain cases, if the homogenization temperature is significantly higher than the rolling temperature, controlled cooling is applied before rolling begins. Depending on the production route of a particular plant, homogenization may be performed as a separate process step, in which case additional pre-heating may be required before hot rolling.

Although both pre-heating and homogenization involve raising the temperature of the ingot, these processes serve different purposes and are performed under different conditions. Homogenization is generally carried out at higher temperatures and includes multiple stages: heating to an elevated temperature, soaking for a defined period, and then cooling down to the rolling temperature. Pre-heating, in contrast, is typically a simpler one-step operation intended only to bring the material to the appropriate temperature required for hot rolling, without the extended soaking stage aimed at microstructural equalization.

Another important heat treatment process is **annealing**, which is commonly used to decrease strength and increase ductility between forming stages. Annealing relieves internal stresses and restores the material's ability to undergo further deformation without cracking. The process usually involves heating the metal to a temperature above its recrystallization point, holding it for a specified time, and then cooling it slowly. Annealing is widely applied in manufacturing because it reverses the effects of work hardening, which is especially important during operations such as cold rolling, where excessive strengthening and reduced ductility may lead to fracture or surface defects.

Rolling is one of the most important metal forming processes and involves plastic deformation of the material between rotating rolls. It is generally classified into hot rolling and cold rolling depending on the process temperature.

Hot rolling is performed at temperatures above the recrystallization temperature of the metal. During deformation, the grains elongate and deform. However, due to the elevated temperature, recrystallization occurs simultaneously or immediately after deformation. This prevents significant work hardening and allows large thickness reductions in a relatively efficient manner. The starting material is usually large pieces of metal, like semi-finished casting products, such as scalped ingots.

Cold rolling is performed at temperatures below the recrystallization temperature, typically at or near room temperature. Its main purpose is to further reduce thickness while



improving mechanical strength, surface finish, and dimensional precision. In the initial stages of cold rolling, the reduction per pass is often relatively high; these early passes are sometimes referred to as breakdown passes. Cold rolling causes elongation of grains in the rolling direction and promotes the fragmentation and redistribution of microstructural heterogeneities, resulting in a more uniform structure. However, cold deformation introduces internal stresses and significantly increases work hardening. These effects may be partially or fully removed through subsequent heat treatment, such as intermediate annealing.

Tension levelling is a finishing process used to improve the flatness of rolled strip or sheet. The process is based on applying controlled tensile stress while the material passes through a series of levelling rolls. The combination of stretching and bending produces slight plastic deformation, which reduces residual stresses and eliminates shape defects such as waviness or coil set. As a result, the product exhibits improved flatness, which is crucial for further forming and surface treatment operations.

Surface cleanliness requirements depend strongly on subsequent finishing operations. Processes such as plating, chromate treatment, or the application of conversion coatings require the metal surface to be free from oils, greases, and other contaminants. One commonly used industrial cleaning method is alkaline cleaning, which is particularly widespread in the processing of aluminium alloys. This method is relatively easy to implement in production environments and does not require highly complex equipment. Aluminium is readily attacked by alkaline solutions. Therefore, cleaning baths are usually maintained at a pH between 9 and 11 and are often inhibited to minimize excessive dissolution of the metal.

A practical method used to verify the effectiveness of degreasing is the ink test. In this procedure, a line is drawn on the metal surface using a dedicated test ink. If the surface has been properly degreased, the ink line remains continuous and stable for at least two seconds without pulling back or breaking. If the surface is still contaminated, the ink tends to shrink, break apart, or form droplets, indicating insufficient cleanliness and the need for further surface preparation.

Slitting is a cutting process in which wide metal strip is divided longitudinally into narrower strips. A slitting line can be designed to perform single-strand or multi-strand slitting in one continuous operation. In industrial practice, the material is not always slit into multiple strips. However, trimming of the edges is generally required to ensure correct final dimensions and acceptable edge quality. In some cases, instead of slitting into strips, the product is cut



into sheets. In such situations, the material undergoes both edge trimming and transverse cutting to obtain sheets of the desired length.

The operations described above represent typical steps in the processing route of metals, particularly in the production of sheets and strips intended for further manufacturing. These processes play a key role in achieving the required mechanical properties, microstructural uniformity, surface quality, and dimensional precision. Although these technological routes are fundamental to modern industry, the complexity and duration of the manufacturing process behind everyday objects (including disposable products) are often underestimated.

EXERCISES

Task 1

Match the terms with their definitions:

- | | |
|----------------------|--|
| 1. Scalping | A. A cutting process in which wide metal strip is divided longitudinally into narrower strips. |
| 2. Homogenization | B. A heat treatment process that reduces strength and increases ductility by heating and slow cooling. |
| 3. Annealing | C. A rolling process carried out above the recrystallization temperature. |
| 4. Hot rolling | D. A surface machining operation that removes a thin outer layer from an ingot. |
| 5. Tension levelling | E. A process used to improve flatness by applying controlled tensile stress. |
| 6. Slitting | F. A heat treatment process in which ingots are heated close to the solidification temperature to improve chemical uniformity. |

Task 2

Answer the following questions based on the text.

1. What are the differences between hot rolling and cold rolling? Give at least two differences.

.....
.....



2. Which process would you apply to improve the flatness of a metal product? Briefly justify your answer.

.....
.....

4. Machines Used in Metal Processing

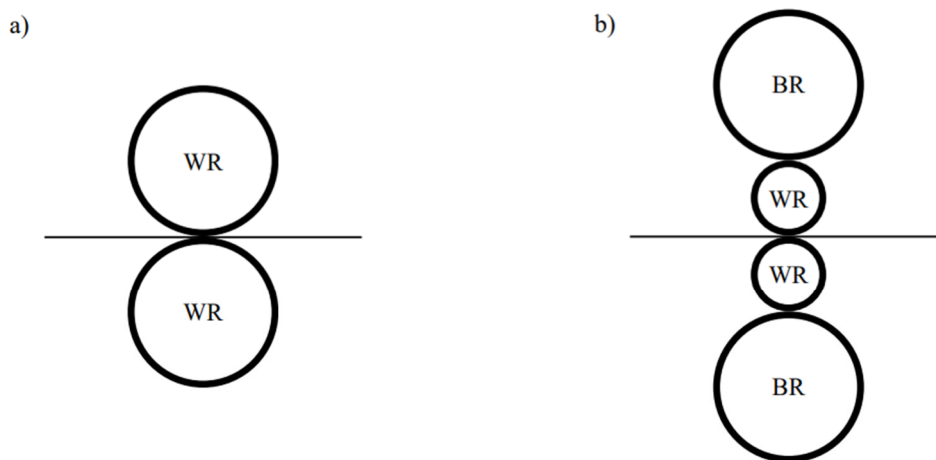
In the previous chapters, only the main operations typically applied in metal processing were described. In practice, however, each of these processes represents a broad technological field that could easily be discussed in a separate, extensive publication. When analysing a manufacturing process, it is also important to consider the equipment that enables its execution. For the purpose of this thesis, the construction and operating principles of rolling mills used in aluminium processing are presented, as rolling is one of the most significant forming methods in the production of aluminium sheets and strips.

At first glance, the design of a rolling mill may appear relatively simple. In reality, however, it is a highly engineered machine, whose performance depends on the precise cooperation of mechanical, hydraulic, and control systems. The rolling process is based on introducing the material between two rolls rotating in opposite directions. Due to the compressive forces applied by the rolls, the thickness of the material is reduced and its length increases. The passage of the metal between a pair of working rolls is referred to as a pass. In industrial practice, it is usually not possible to obtain the final required dimensions in a single pass, therefore rolling is performed through a sequence of multiple passes, each gradually reducing the thickness of the strip.

A key parameter of the rolling mill is the roll gap, defined as the distance between the barrels of the cooperating rolls. The roll gap height is controlled by adjusting the vertical position of the rolls. During rolling, the rolls are subjected to high compressive forces, which leads to elastic deformation of the roll system. As a result, the actual roll gap increases compared to the set value. This phenomenon is known as roll spring, and its magnitude is typically in the range of approximately 0.1% to 0.7% of the roll diameter. The change of the material dimensions during rolling occurs according to the principle of constant volume, meaning that although the thickness decreases and the length increases, the volume of the rolled metal remains unchanged.

Rolling mills can be classified according to the number and arrangement of rolls used in the mill stand. The simplest design is the two-high rolling mill, consisting of two working rolls positioned one above the other. The rolls rotate in opposite directions and the material passes between them. Although this type of mill has a relatively simple construction and is easy to operate, it has limited capability for producing thin products due to roll deflection under high rolling forces. For this reason, two-high mills are mainly used for roughing operations (breakdown mills), where large reductions are applied at higher thickness ranges.

A more commonly used configuration in modern rolling plants is the four-high rolling mill. In this design, two smaller working rolls are supported by two larger backup rolls. The working rolls are in direct contact with the rolled material, while the backup rolls provide mechanical support and prevent excessive bending. This arrangement significantly increases stiffness of the rolling stand and enables the production of thinner sheets and strips with improved thickness uniformity. Four-high mills are widely applied in both hot and cold rolling of aluminium alloys due to their efficiency and improved control of strip geometry.



Legend: WR - Working Roll; BR - Backup Roll

Figure 1 Two-high (a) and four-high (b) rolling mill configurations.

In aluminium hot rolling, proper thermal control of the rolling mill is of particular importance. The surface temperature of the rolls must generally not exceed approximately 250 to 300°C, because at higher temperatures aluminium tends to adhere to the roll surface, which may lead to surface defects and unstable process conditions. For this reason, rolling mills are



equipped with cooling and lubrication systems whose purpose is to prevent overheating of the rolls and reduce thermal stresses. Two basic lubrication and cooling systems are commonly used. The first method involves lubrication with mineral oils or kerosene combined with periodic water spraying. The second method, which is more widely applied in industrial practice, is the use of an emulsion delivered continuously in a large flow rate. Such emulsions typically consist of mineral oil, an emulsifying agent (for example alkaline metal soap), and water. The correct composition and delivery parameters of the rolling emulsion significantly influence friction conditions, roll wear, and the final surface quality of the rolled product.

The reliable operation of a rolling mill requires regular maintenance activities. Roll grinding is one of the most important procedures, as the roll surface condition directly affects the thickness uniformity, flatness, and surface finish of the strip. Over time, rolls develop wear marks and surface damage, which may result in defects transferred to the product. In addition, periodic technical inspections are required to ensure proper operation of the mill stand, bearings, drives, and hydraulic adjustment systems. Another essential aspect is the monitoring of rolling fluid quality. The emulsion properties, such as concentration, contamination level, and temperature, must be controlled, since deterioration of the rolling fluid may lead to insufficient lubrication, increased roll wear, and reduced process stability.

In conclusion, machines used in metal processing, particularly rolling mills, represent advanced and carefully designed systems. Their performance is determined not only by mechanical construction, but also by precise process control, lubrication and cooling systems, and regular maintenance procedures. Understanding the operation of these machines is essential for analysing aluminium processing routes, as equipment condition and operating parameters directly affect the quality of the final rolled products.

EXERCISES

Task 1

Translate the following terms into Polish:

- | | |
|--------------------|-----------------------|
| 1) rolling mill | 7) roll grinding |
| 2) working roll | 8) lubrication system |
| 3) backup roll | 9) rolling emulsion |
| 4) roll deflection | 10) thermal stresses |



5) hot rolling

11) thickness uniformity

6) cold rolling

12) contamination level

Task 2

Write short definitions of the following terms (1–2 sentences each):

- 1) roll diameter
- 2) rolling pass
- 3) rolling emulsion

Task 3

Complete the sentences with a, an, the, or no article (—).

- 1) Rolling is ___ forming process in which metal is reduced in thickness between rotating rolls.
- 2) ___ roll gap must be adjusted precisely to achieve required strip thickness.
- 3) Aluminium may stick to ___ roll surface if the temperature is too high.
- 4) Roll grinding is ___ important maintenance operation in aluminium processing plants.
- 5) Rolling emulsion is usually delivered in ___ continuous flow to ensure proper lubrication and cooling.

Task 4

Answer the following questions based on the text:

- 1) Why are multiple passes usually required in rolling processes?
.....
.....
- 2) What is the function of backup rolls in a four-high rolling mill?
.....
.....
- 3) Give two reasons why cooling and lubrication are necessary in aluminium hot rolling.
.....
.....
- 4) What maintenance activities are important for ensuring stable rolling mill operation?
.....
.....



5. Environmental and sustainability aspects of metal processing

Metals are widely regarded as materials that support sustainable development, mainly due to their high recyclability and the possibility of repeated reprocessing without significant loss of functional properties. In contrast to many other engineering materials, metals can be recovered both from production scrap and from end-of-life products, and then reused in new industrial applications.

Aluminium is a particularly important example, as it is endlessly recyclable, meaning it can be repeatedly remelted and processed while maintaining its usability. Additionally, aluminium offers a highly beneficial strength-to-weight ratio, which makes it an attractive choice for industries such as automotive, aerospace, and rail transport. Reducing the mass of structural components leads directly to lower fuel consumption and reduced greenhouse gas emissions during the operational phase of vehicles. As a result, the use of recyclable aluminium contributes not only to improved sustainability at the production stage but also to environmental benefits throughout the entire life cycle of manufactured products.

Recycling plays a crucial role in environmental protection because it reduces the demand for primary raw materials and limits the amount of waste sent to landfills. The extraction and processing of metal ores is associated with significant energy consumption and environmental burdens, including emissions and the generation of industrial waste. For this reason, increasing the share of secondary raw materials in metal production is a key principle of the circular economy. Furthermore, the use of recycled metal improves resource efficiency and contributes to greater stability in supply chains by decreasing dependence on mining and imported raw materials, which has both environmental and economic advantages.

One of the most important challenges for modern metal processing industries is decarbonisation, which is the reduction of greenhouse gas emissions across industrial operations and supply chains. A central concept in this context is the carbon footprint, defined as the total amount of greenhouse gas emissions expressed as carbon dioxide equivalent. The carbon footprint can be calculated for a product, a technological process, an industrial plant, or an entire company, and it provides valuable information on which stages generate the highest emissions. Decarbonisation involves not only improving energy efficiency but also replacing fossil fuel-based energy sources with renewable alternatives, implementing low-emission technologies, and increasing the proportion of recycled materials used in production.



These strategies are becoming essential due to stricter regulations and the growing importance of sustainable manufacturing.

To assess emissions accurately, greenhouse gas emissions are commonly divided into three categories referred to as Scope 1, Scope 2, and Scope 3. Scope 1 includes direct emissions from sources owned or controlled by the company, such as fuel combustion in furnaces, boilers, or industrial machinery used in metal processing. Scope 2 refers to indirect emissions associated with the generation of purchased electricity, heat, or steam consumed by the company, meaning that emissions occur outside the facility but are driven by its energy demand. Scope 3 includes all other indirect emissions across the value chain, such as those related to raw material extraction, production of auxiliary materials, transportation of inputs and finished products, product used by customers, and end-of-life treatment including disposal or recycling.

The calculation of a carbon footprint is based on collecting data on energy use, fuel consumption, material flows, and transport activities, and then converting these values into CO₂ equivalent emissions using appropriate emission factors. This approach enables companies to monitor their environmental impact, identify key emission sources, and implement targeted measures to reduce greenhouse gas emissions in line with sustainability goals.

Overall, the environmental performance of metal processing depends strongly on the ability to minimise emissions and maximise material recovery throughout the entire life cycle of products. Increasing recycling rates and implementing effective decarbonisation strategies are therefore essential steps toward reducing the carbon footprint of the metal industry. In the long term, combining circular economy principles with low-emission energy sources will play a key role in achieving sustainable industrial development.

EXERCISES

Task 1

Translate the following terms into Polish:

- | | |
|-----------------------------|---------------------|
| 1) recyclability | 8) carbon footprint |
| 2) strength-to-weight ratio | 9) decarbonisation |
| 3) greenhouse gas | 10) fossil fuel |



- | | |
|---------------------|-----------------------|
| 4) life cycle | 11) renewable |
| 5) landfills | 12) low-emissions |
| 6) emissions | 13) end-of-life scrap |
| 7) industrial waste | 14) circular economy |

Task 2

Fill in the missing words in the sentences.

(carbon footprint, decarbonisation, strength-to-weight ratio, industrial waste, recyclability, greenhouse gases, circular economy)

- 1) Metals are considered environmentally friendly materials due to their high _____ and ability to be reused without significant loss of properties.
- 2) Aluminium is widely used in transport industries because of its excellent _____, which helps reduce vehicle weight.
- 3) The extraction and processing of metal ores generates emissions and large amounts of _____.
- 4) Increasing the use of secondary raw materials is a key principle of the _____.
- 5) Modern metal processing industries focus on _____ in order to reduce emissions across production and supply chains.
- 6) The _____ is defined as the total amount of greenhouse gas emissions expressed as carbon dioxide equivalent.
- 7) Reducing _____ emissions is one of the main environmental goals of sustainable manufacturing.

Task 3

Match the term with its correct description.

- | | |
|------------|---|
| 1) Scope-1 | a) Indirect emissions from the supply chain, transport, and product life cycle. |
| 2) Scope-2 | b) Direct emissions from company-owned sources such as furnaces. |
| 3) Scope-3 | c) Emissions related to purchased electricity and heat. |



Bibliography

Singh, R. V. (2000). *Aluminum Rolling: Process, Principles & Applications*. Warrendale: TMS.

Davis, J. R. (1993). *Aluminum and aluminum alloys*. Materials Park, Ohio: ASM International.