

Analiza modulacji analogowych, modulacji cyfrowych, transimpedancji, układów BalUn i UnUn

Korepetycje Radiotechnika

Modulacje sygnałów radiowych

Modulacja sygnału polega na zmianie jednej z jego cech charakterystycznych (amplitudy, częstotliwości lub fazy) zgodnie z informacją zawartą w sygnale modulującym. W niniejszej sekcji omówimy modulacje amplitudowe.

Modulacje amplitudowe

W modulacjach amplitudowych nośna zmienia swoją amplitudę zgodnie z sygnałem modulującym.

Zwykła modulacja AM

W klasycznej modulacji amplitudowej (AM) sygnał modulowany ma postać: $s(t) = \big(A + m(t) \big) \cos(2\pi f_c t)$, gdzie:

- A – amplituda nośnej,
- $m(t)$ – sygnał modulujący,
- f_c – częstotliwość nośna.

Spektrum sygnału AM składa się z fali nośnej oraz dwóch wstęg bocznych oddalonych o częstotliwość sygnału modulującego.

Widmo sygnału AM

Wykres sygnału AM

Modulacja DSB-AM

Modulacja dwuwstęgowa (DSB-AM) jest jedną z form modulacji amplitudy, w której obie wstęgi boczne są przesyłane bez fali nośnej. Wyrażenie na sygnał modulowany w tej formie jest opisane wzorem:

$$s(t) = m(t) \cos(2\pi f_c t)$$

gdzie:

- $s(t)$ to sygnał wyjściowy, który jest wynikiem modulacji,
- $m(t)$ to sygnał informacyjny (modulujący),
- f_c to częstotliwość nośną,
- $\cos(2\pi f_c t)$ to funkcja nośna o częstotliwości f_c .

W tej modulacji:

- Sygnał nośny $\cos(2\pi f_c t)$ jest pomnożony przez sygnał $m(t)$, który reprezentuje informacje, które mają być przesyłane.
- Ponieważ w tej wersji DSB-AM nie przesyła się fali nośnej (zatem nie ma składnika stałego nośnej w sygnale), sygnał jest wyłącznie kombinacją wstęg bocznych powstałych w wyniku modulacji.

Widmo sygnału DSB-AM:

Widmo sygnału DSB-AM

Modulacja SSB-AM

Modulacja jednowstęgowa (SSB-AM) pozwala na transmisję tylko jednej wstęgi bocznej, oszczędzając pasmo. Wyrażenie na sygnał modulowany w tej formie jest opisane wzorem:

$$s(t) = \frac{1}{2} [m(t) + j \hat{m}(t)] e^{j2\pi f_c t} + c.c.$$

gdzie:

- $s(t)$ to sygnał wyjściowy,
- $m(t)$ to sygnał informacyjny (modulujący),
- $\hat{m}(t)$ to transformata Hilberta sygnału $m(t)$,
- f_c to częstotliwość nośna,
- $e^{j2\pi f_c t}$ to składnik nośny w postaci zespolonej,
- $c.c.$ oznacza składnik sprzężony zespolony (complex conjugate).

W tej modulacji, zamiast dwóch wstęg bocznych, przesyłana jest tylko jedna wstęga boczna (przesyłana w postaci części rzeczywistej i zespolonej), co pozwala zaoszczędzić pasmo.

Widmo sygnału SSB-AM

Transformata Hilberta to operacja matematyczna, która dla sygnału $m(t)$ generuje sygnał $\hat{m}(t)$, zwany sygnałem hilbertowskim. Jest to transformacja, która przekształca sygnał rzeczywisty na sygnał o tej samej amplitudzie, ale przesunięty w fazie o 90° .

Matematycznie transformata Hilberta $\hat{m}(t)$ jest definiowana jako całka splotowa:

$$\hat{m}(t) = \frac{1}{\pi} \text{P.V.} \int_{-\infty}^{\infty} \frac{m(\tau)}{t - \tau} d\tau$$

gdzie P.V. oznacza wartość główną Cauchy'ego, a $m(t)$ to sygnał wejściowy. Transformata Hilberta jest szeroko stosowana w analizie sygnałów, w tym w modulacji SSB-AM, gdzie pozwala na uzyskanie sygnału, który w połączeniu z oryginalnym sygnałem $m(t)$ tworzy jedną wstęgę boczną w procesie modulacji.

Całka splotowa jest operacją matematyczną, która łączy dwa sygnały w jeden nowy sygnał. Jest to podstawowy koncept w analizie sygnałów i systemów, szczególnie w teorii filtrów i przetwarzania sygnałów.

Dla dwóch funkcji $f(t)$ i $g(t)$, całka splotowa $(f * g)(t)$ jest definiowana jako:

$$(f * g)(t) = \int_{-\infty}^{\infty} f(\tau) g(t - \tau) d\tau$$

W kontekście transformacji Hilberta, całka splotowa jest używana do obliczenia sygnału $\hat{m}(t)$,

który jest transformowaną Hilberta sygnału $m(t)$. Matematycznie jest to zapisywane jako:

$$\hat{m}(t) = \frac{1}{\pi} \int_{-\infty}^{\infty} \frac{m(\tau)}{t - \tau} d\tau$$

W tym przypadku, $m(\tau)$ to sygnał wejściowy, a $\frac{1}{\pi(t - \tau)}$ to funkcja, którą używamy do obliczenia transformacji Hilberta. Warto zauważyć, że to wyrażenie jest formą splotu z funkcją $\frac{1}{\pi t}$, znaną również jako funkcja Hilberta.

Całka splotowa w tym przypadku „przesuwa” sygnał $m(t)$ w czasie, tworząc nowy sygnał $\hat{m}(t)$, który jest fazowo przesunięty o 90° względem $m(t)$, co jest kluczową cechą transformacji Hilberta. Dzięki temu powstaje sygnał, który jest wykorzystywany w modulacji SSB-AM do przesyłania jednej wstęgi bocznej.

Składnik nośny w postaci zespolonej to wyrażenie $e^{j2\pi f_c t}$, które reprezentuje nośną w postaci zespolonej. W tej formie nośna jest zapisana jako liczba zespolona, gdzie j to jednostka urojona, a f_c to częstotliwość nośna. Składnik $e^{j2\pi f_c t}$ opisuje falę nośną, która ma częstotliwość f_c i jest wyrażona w postaci funkcji wykładniczej. Dzięki postaci zespolonej łatwiej jest manipulować fazą i amplitudą sygnału, co jest szczególnie przydatne w analizie i modulacji sygnałów.

Natomiast *c.c.* to skrót od angielskiego terminu „complex conjugate” (sprzężenie zespolone). Oznacza to, że oprócz składnika $e^{j2\pi f_c t}$ w wyrażeniu na sygnał $s(t)$, dodaje się jego sprzężenie zespolone, czyli $e^{-j2\pi f_c t}$. Sprzężenie zespolone polega na zmianie znaku przy jednostce urojonej j . W kontekście modulacji SSB-AM, składnik sprzężony zespolony zapewnia, że sygnał będzie miał rzeczywistą wartość, ponieważ suma składników zespolonych i ich sprzężonych daje rzeczywisty wynik.

Modulacja VSB-AM

Modulacja VSB-AM (Vestigial Side Band) stosowana jest w transmisji telewizyjnej, gdzie część jednej wstęgi bocznej jest tłumiona, aby zaoszczędzić pasmo, zachowując jednocześnie możliwość odbioru pełnej informacji. Jest to rodzaj modulacji amplitudy, w której tylko część jednej z wstęg bocznych (zwykle wstęgi dolnej) jest przesyłana, podczas gdy reszta wstęgi jest tłumiona przez odpowiedni filtr. Dzięki temu pasmo sygnału jest mniejsze niż w tradycyjnej modulacji AM, co jest korzystne w transmisji telewizyjnej, gdzie efektywność pasma jest kluczowa.

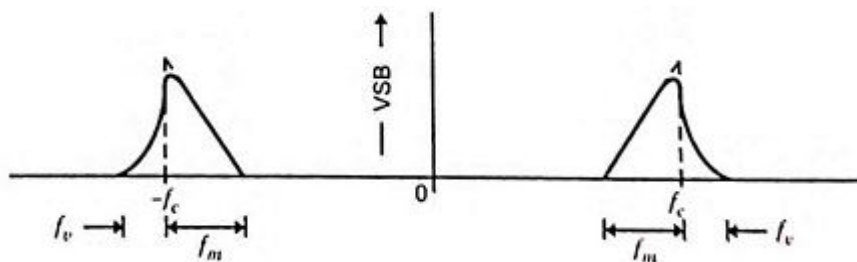
Wyrażenie na widmo sygnału $S(f)$ w modulacji VSB-AM jest zapisane jako:

$$S(f) = M(f) H(f),$$

gdzie:

- $M(f)$ to widmo sygnału informacyjnego, - $H(f)$ to funkcja przenoszenia filtru, który tłumí część jednej wstęgi bocznej, pozostawiając jedynie „resztkową” część tej wstęgi.

Filtr $H(f)$ jest odpowiedzialny za usuwanie nadmiarowych części wstęgi bocznej, co pozwala na zaoszczędzenie pasma, ale w sposób, który nie prowadzi do utraty istotnych informacji.



Wykres widma sygnału VSB-AM

Modulacja częstotliwości

Modulacja częstotliwości (FM) jest techniką, w której częstotliwość nośnej jest modulowana przez sygnał informacyjny $m(t)$. W odróżnieniu od modulacji amplitudy, gdzie zmienia się amplituda nośnej, w modulacji częstotliwości zmienia się jej częstotliwość w zależności od wartości sygnału modulującego. Wzór na sygnał FM można zapisać jako:

$$s(t) = A \cos(2\pi f_c t + \Delta f \sin(2\pi f_m t))$$

gdzie:

- A to amplituda nośnej, - f_c to częstotliwość nośna, - Δf to maksymalne odchylenie częstotliwości (częstotliwość dewiacji), - f_m to częstotliwość sygnału modulującego $m(t)$.

W widmie sygnału FM, w zależności od wartości Δf , pojawia się wiele wstęp bocznych rozmieszczonych wokół częstotliwości nośnej f_c . Wzór na widmo sygnału FM jest bardziej złożony, jednak jego ogólną postać można zapisać jako:

$$S(f) = \sum_{n=-\infty}^{\infty} J_n(\beta) \cdot \delta(f - f_c - n f_m)$$

gdzie:

- $J_n(\beta)$ to funkcja Bessela pierwszego rodzaju z indeksem n , a $\beta = \frac{\Delta f}{f_m}$ jest indeksem modulatora (współczynnikiem dewiacji), - $\delta(f)$ to delta Diraca, reprezentująca wstęgi boczne w widmie.

Dla małych wartości β (modulacja wąskopasmowa), widmo składa się głównie z pierwszej wstęgi bocznej, podczas gdy dla dużych β (modulacja szerokopasmowa) pojawia się wiele wstęp bocznych.

Wykres widma sygnału FM

Widmo sygnału FM z wstęgami bocznymi

W przypadku modulacji częstotliwości widmo sygnału FM rozciąga się na dużą szerokość pasma, szczególnie w przypadku dużych wartości Δf . Oznacza to, że sygnał FM jest bardziej odporny na zakłócenia w porównaniu do sygnałów AM i SSB-AM, jednak wymaga większego pasma transmisyjnego.

Modulacja Fazowa

Modulacja fazowa (PM) jest techniką, w której to faza nośnej jest modulowana przez sygnał informacyjny $m(t)$. W odróżnieniu od modulacji amplitudy (AM), w której zmienia się amplituda nośnej, w modulacji fazowej zmienia się jej faza w odpowiedzi na sygnał modulujący. Sygnał wyjściowy w modulacji fazowej jest opisany wzorem:

$$s(t) = A \cos(2\pi f_c t + \phi(t))$$

gdzie:

- A to amplituda nośnej, - f_c to częstotliwość nośna, - $\phi(t)$ to funkcja fazy, która zależy od sygnału modulującego $m(t)$.

Zwykle faza jest związana z sygnałem modulującym przez wzór:

$$\phi(t) = k_p \cdot m(t)$$

gdzie k_p to współczynnik wzmocnienia fazowego, który określa, jak bardzo sygnał modulujący wpływa na zmianę fazy nośnej.

Widmo sygnału PM jest podobne do widma modulacji częstotliwości, z tym że amplitudy wstępnych bocznych zależą od pierwszych funkcji Bessela, podobnie jak w przypadku modulacji częstotliwości (FM). Zmiana fazy sygnału powoduje przesunięcia w widmie, które reprezentują różne wstęgi boczne o amplitudach $J_n(\beta)$, gdzie β jest współczynnikiem dewiacji fazowej (określającym jak silnie zmienia się faza nośnej w zależności od sygnału modulującego).

Zatem, widmo sygnału PM wyraża się jako:

$$S(f) = \sum_{n=-\infty}^{\infty} J_n(\beta) \cdot \delta(f - f_c - n f_m)$$

gdzie $J_n(\beta)$ to funkcje Bessela pierwszego rodzaju z indeksem n , a $\beta = k_p \cdot m_{\text{max}}$, gdzie m_{max} to maksymalna wartość amplitudy sygnału modulującego $m(t)$.

Podobnie jak w przypadku modulacji częstotliwości, w modulacji fazowej dla małych wartości β pojawiają się tylko pierwsze wstęgi boczne, a dla większych wartości β widmo staje się szersze, rozprzestrzeniając się na wiele wstępnych bocznych.

Wykres widma sygnału PM

Widmo sygnału FM z wstęgami bocznymi

Modulacje cyfrowe

Modulacje cyfrowe są stosowane w transmisji sygnałów cyfrowych, gdzie informacja jest przekazywana przez zmianę parametrów nośnej, takich jak amplituda, częstotliwość czy faza. Dzięki tym technikom możliwa jest efektywna transmisja danych w systemach telekomunikacyjnych.

ASK

FSK

PSK

Modulacje cyfrowe ASK, FSK, PSK dla zakodowanego ciągu 1010

Modulacja ASK (Amplitude Shift Keying)

Modulacja amplitudy (ASK) jest jedną z podstawowych modulacji cyfrowych, w której amplituda nośnej jest modulowana w zależności od wartości cyfrowych bitów. W modulacji ASK mamy dwie możliwe amplitudy:

$$s(t) = \begin{cases} A_1 \cos(2\pi f_c t) & \text{dla bitu 1} \\ A_2 \cos(2\pi f_c t) & \text{dla bitu 0} \end{cases}$$

gdzie A_1 i A_2 to różne amplitudy odpowiadające stanowi logicznemu 1 i 0, a f_c to częstotliwość nośna.

W przypadku tej modulacji sygnał ma tylko jedną wstęgę boczną, której amplituda jest zmieniana zgodnie z wartością przesyłanego bitu.

Modulacja FSK (Frequency Shift Keying)

Modulacja częstotliwości (FSK) polega na zmianie częstotliwości nośnej w zależności od wartości bitu. W przypadku dwufazowej modulacji FSK mamy dwie częstotliwości f_1 i f_2 , które odpowiadają wartościom bitów 0 i 1. Wzór na sygnał w tej modulacji to:

$$s(t) = \begin{cases} A \cos(2\pi f_1 t) & \text{dla bitu 1} \\ A \cos(2\pi f_2 t) & \text{dla bitu 0} \end{cases}$$

gdzie f_1 i f_2 to różne częstotliwości nośne, a A to amplituda nośnej.

Modulacja FSK jest odporniejsza na zakłócenia i szumy w porównaniu do ASK, ponieważ zmiana częstotliwości jest mniej wrażliwa na zmiany amplitudy sygnału.

Modulacja PSK (Phase Shift Keying)

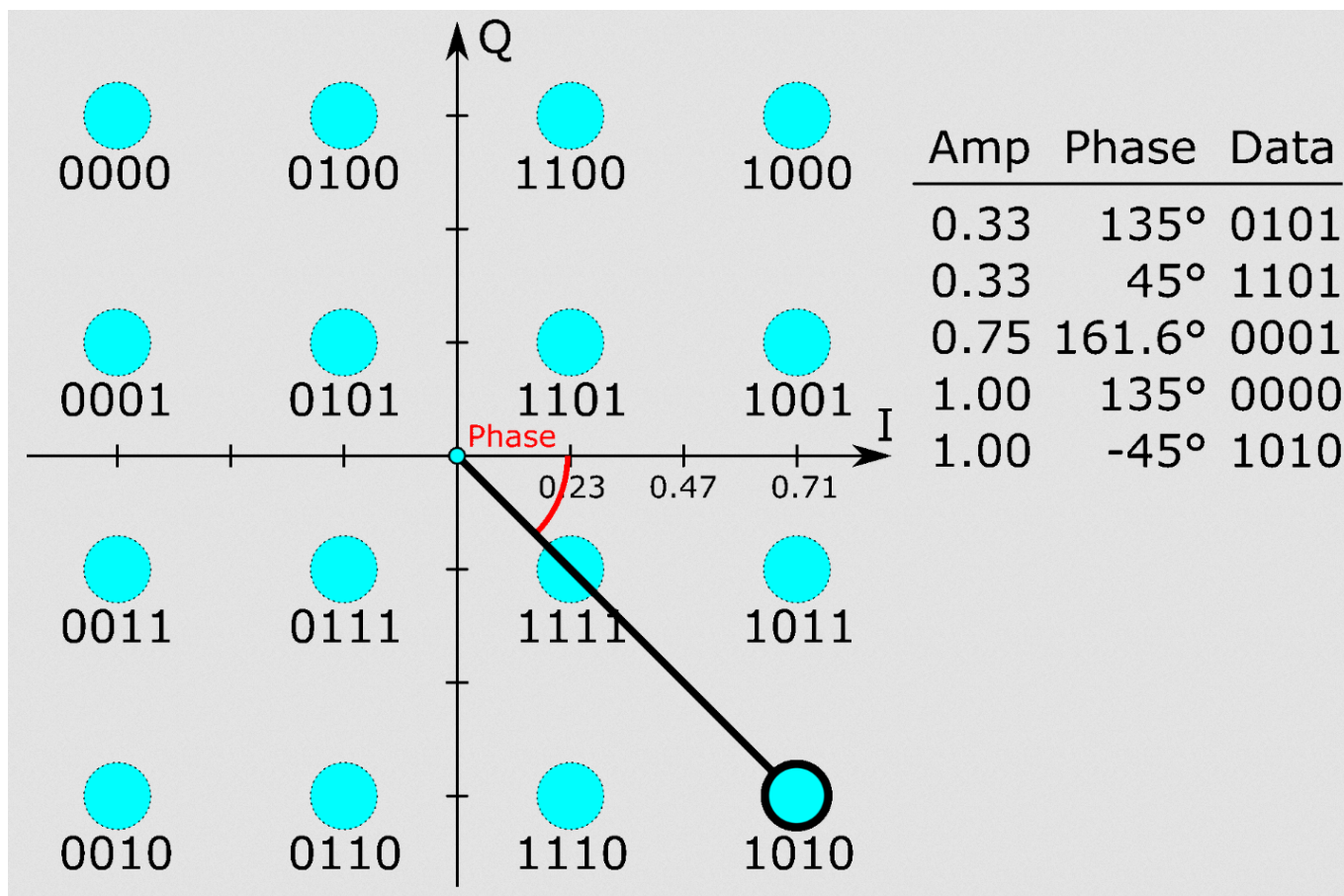
Modulacja fazowa (PSK) polega na zmianie fazy nośnej w zależności od przesyłanych bitów. W najprostszej wersji PSK (BPSK, Binary PSK) mamy dwie możliwe fazy, np. 0 i π , odpowiadające bitom 0 i 1. Wzór na sygnał PSK to:

$$s(t) = A \cos(2\pi f_c t + \phi_k)$$

gdzie $\phi_k \in \{0, \pi\}$ jest fazą nośnej, która zmienia się w zależności od wartości bitu (0 lub 1).

W przypadku PSK możliwa jest również rozszerzona wersja, jak QPSK, gdzie modulujemy cztery różne fazy (np. $0, \frac{\pi}{2}, \pi, \frac{3\pi}{2}$).

Modulacje kwadraturowe



Wykres amplitudowo-fazowy sygnału 16QAM

Modulacje kwadraturowe to technika, w której informacje są transmitowane poprzez modulację zarówno amplitudy, jak i fazy nośnej. Najpopularniejsze przykłady to QAM (Quadrature Amplitude Modulation), gdzie zarówno amplituda, jak i faza są zmieniane w sposób zrównoważony.

W przypadku QAM możemy mieć np. QPSK (Quadrature Phase Shift Keying), gdzie na jednej nośnej transmitowane są dwie oddzielne składowe fazowe. Sygnał wyjściowy QPSK może być opisany równaniem:

$$s(t) = A_1 \cos(2\pi f_c t + \phi_1) + A_2 \cos(2\pi f_c t + \phi_2)$$

gdzie A_1 i A_2 to amplitudy, a ϕ_1 i ϕ_2 to fazy, które mogą przyjmować różne wartości w zależności od kombinacji bitów (np. dla QPSK cztery różne kombinacje fazowe).

Modulacje kwadraturowe, takie jak QAM i QPSK, oferują większą efektywność w wykorzystaniu pasma, ponieważ na każdą nośną można przesłać więcej informacji.

Modulacje Impulsowe

Modulacje impulsowe są technikami, w których informacje są przesyłane za pomocą impulsów o określonych właściwościach, takich jak amplituda, czas trwania czy pozycja w czasie. Wykorzystywane są w różnych aplikacjach komunikacyjnych, w tym w telekomunikacji, systemach radarowych oraz

cyfrowych systemach transmisji danych. Celem tych modulacji jest uzyskanie efektywnej transmisji informacji w ograniczonym paśmie oraz zapewnienie odporności na zakłócenia. Do głównych typów modulacji impulsowych należy: PAM, PWM, PPM i PCM.

Modulacja PAM (Pulse Amplitude Modulation)

Opis

Modulacja amplitudy impulsu (PAM) jest jedną z najprostszych form modulacji, w której amplituda impulsów jest zmieniana w zależności od informacji, którą chcemy przesłać. W tej metodzie amplituda kolejnych impulsów jest proporcjonalna do wartości sygnału informacyjnego.

Wzór matematyczny

Matematycznie, sygnał w modulacji PAM może być zapisany jako: $s(t) = \sum_{n=0}^{N-1} m_n \cdot p(t - nT)$, gdzie: - m_n - amplituda impulsu na n -tej próbce, - T - okres próbkowania, - $p(t)$ - funkcja impulsu (np. funkcja prostokątna).

Zasada działania

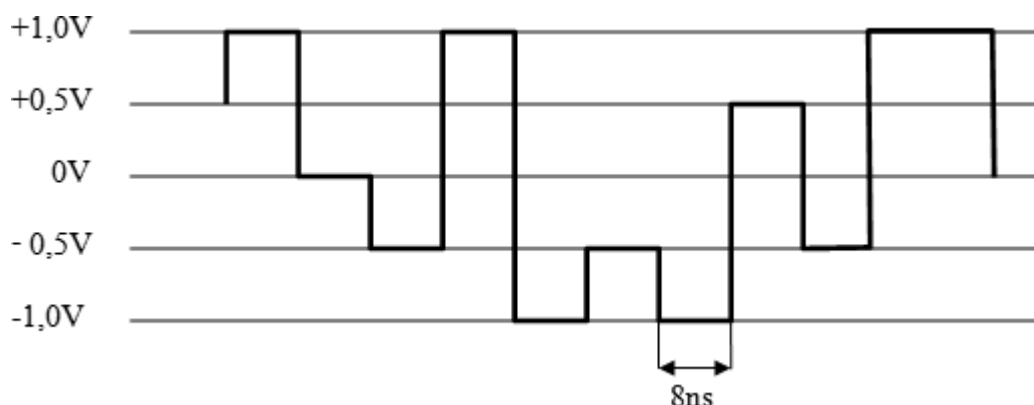
W modulacji PAM amplituda impulsów jest bezpośrednio zależna od wartości sygnału informacyjnego. Zmiana amplitudy impulsów może być realizowana w sposób dyskretny (np. dla sygnałów binarnych) lub ciągły (w przypadku sygnałów analogowych).

Zastosowanie

Modulacja PAM jest stosowana głównie w transmisji sygnałów cyfrowych i analogowych w systemach, gdzie szerokość pasma nie jest ściśle ograniczona, takich jak systemy telekomunikacyjne.

Wykres

Sygnał PAM w domenie czasu to seria impulsów o zmiennej amplitudzie. Poniżej przedstawiono wykres przykładowego sygnału PAM:



Wykres modulacji PAM w czasie

Modulacja PWM (Pulse Width Modulation)

Opis

Modulacja szerokości impulsu (PWM) polega na zmianie czasu trwania impulsu w zależności od wartości sygnału informacyjnego. Częstotliwość impulsów jest stała, a zmienia się tylko ich szerokość.

Wzór matematyczny

W przypadku PWM sygnał opisujemy wzorem: $s(t) = \sum_{n=0}^{N-1} m_n \cdot u(t - nT)$,
gdzie: - m_n - amplituda impulsu, zależna od szerokości impulsu, - $u(t)$ - funkcja prostokątna o zmiennym czasie trwania.

Zasada działania

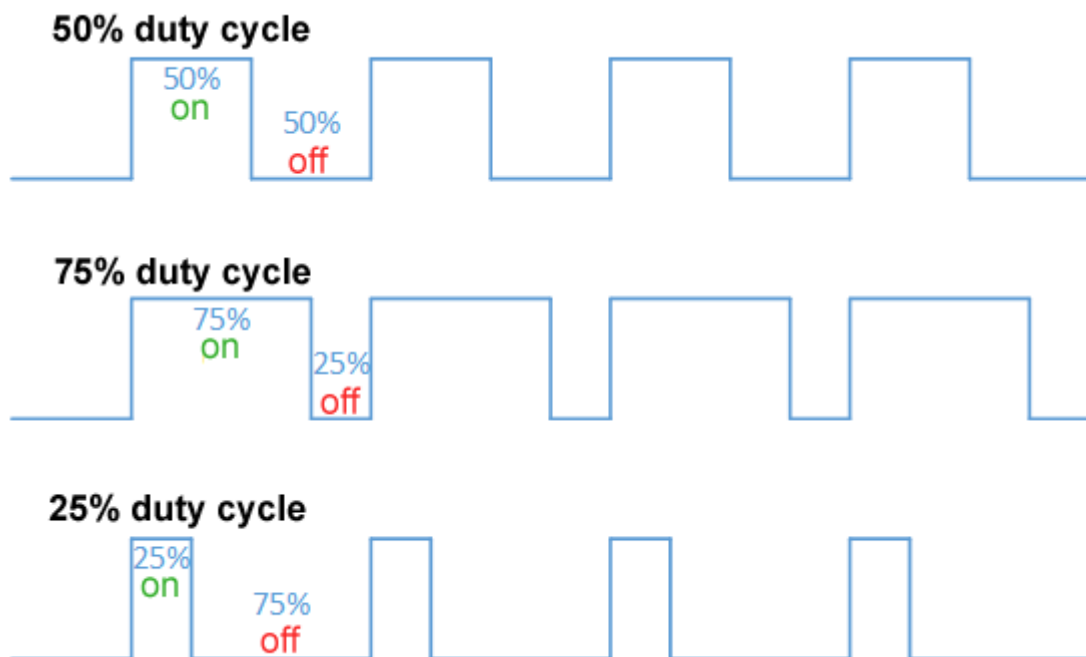
W modulacji PWM szerokość impulsów jest funkcją wartości sygnału informacyjnego. Zmieniając szerokość impulsów, można kodować różne poziomy sygnału, co pozwala na bardziej precyzyjną transmisję.

Zastosowanie

PWM jest powszechnie wykorzystywana w systemach sterowania silnikami, a także w elektronice do regulacji mocy (np. w zasilaczach). Ponadto, modulacja PWM znajduje zastosowanie w transmisji audio i wideo.

Wykres

Sygnał PWM w domenie czasu to seria prostokątnych impulsów o zmieniającej się szerokości:



Wykres modulacji PWM o zmiennych długościach impulsu

Modulacja PPM (Pulse Position Modulation)

Opis

Modulacja pozycji impulsu (PPM) polega na przesuwaniu czasu wystąpienia impulsu w zależności od wartości sygnału informacyjnego. W tej metodzie amplituda impulsu pozostaje stała, ale jego pozycja w czasie jest zmieniana.

Wzór matematyczny

Sygnał w modulacji PPM można zapisać jako: $s(t) = \sum_{n=0}^{N-1} m_n \cdot p(t - nT - \Delta t_n)$, gdzie: - m_n - amplituda impulsu, - Δt_n - opóźnienie impulsu zależne od sygnału informacyjnego.

Zasada działania

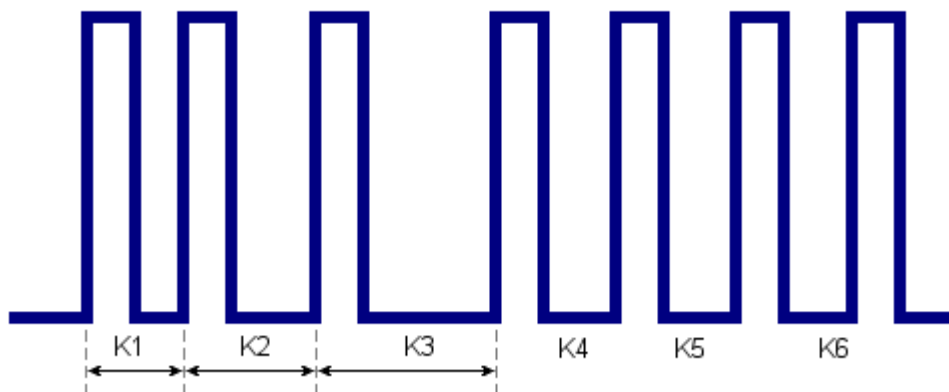
W PPM, zmieniając czas wystąpienia impulsu (pozycję), kodujemy informację. Czas pomiędzy impulsami jest stały, natomiast ich pozycje są funkcją sygnału.

Zastosowanie

Modulacja PPM jest wykorzystywana w systemach, które muszą przesyłać dane w sposób odporny na zakłócenia, takich jak systemy komunikacji optycznej.

Wykres

Sygnał PPM w domenie czasu to seria impulsów, których pozycja w czasie zmienia się w zależności od informacji.



Wykres modulacji PPM w czasie

Modulacja PCM (Pulse Code Modulation)

Opis

Modulacja kodowania impulsów (PCM) to technika cyfrowej modulacji, która polega na próbkowaniu sygnału analogowego i przekształcaniu go w dyskretną sekwencję impulsów. Każdy impuls reprezentuje określoną wartość próbki sygnału analogowego.

Wzór matematyczny

Sygnał PCM zapisujemy jako: $s(t) = \sum_{n=0}^{N-1} m_n \cdot \delta(t - nT)$, gdzie: - m_n - wartość próbki sygnału, - $\delta(t)$ - funkcja Diraca (impuls jednostkowy).

Zasada działania

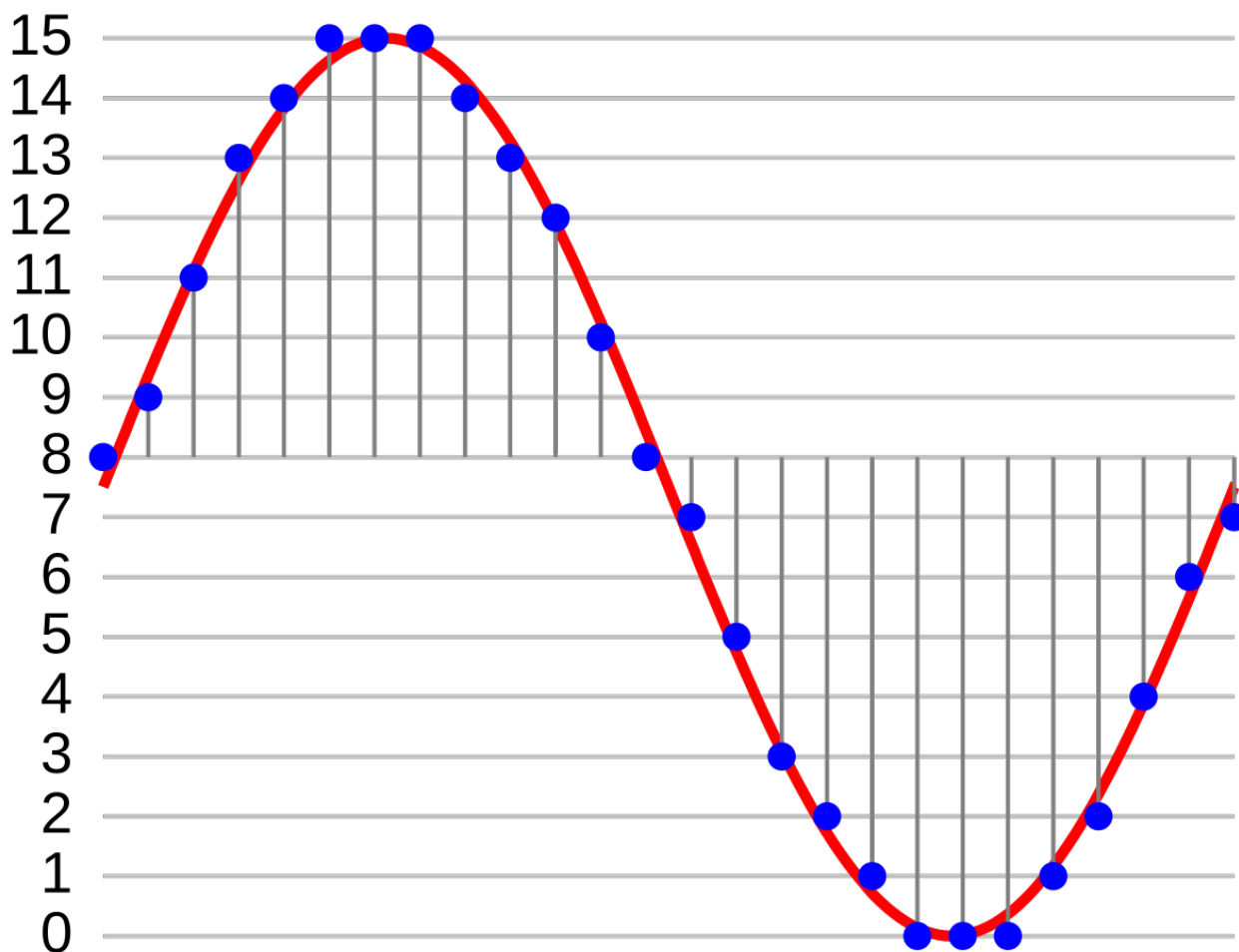
W PCM sygnał analogowy jest próbkowany w określonych odstępach czasu, a następnie każda próbka jest kodowana w postaci cyfrowej (najczęściej w postaci binarnej). Transmitowane są cyfrowe impulsy, które reprezentują poziomy sygnału w danym czasie.

Zastosowanie

PCM jest szeroko stosowana w cyfrowych systemach komunikacji, w tym w telefonii cyfrowej, kompresji audio (np. w formacie MP3) oraz w systemach dźwiękowych.

Wykres

Sygnał PCM to ciąg impulsów, które reprezentują cyfrowe próbki sygnału analogowego:



Wykres modulacji PCM w czasie (sygnał po procesie kwantyzacji)

Transimpedancja: Zastosowanie, Działanie i Zjawiska Fizyczne

Transimpedancja jest istotnym pojęciem w elektronice, szczególnie w kontekście wzmacniaczy transimpedancyjnych (TIA – ang. Transimpedance Amplifier). Jest to parametr opisujący konwersję prądu wejściowego na napięcie wyjściowe. Formalnie, transimpedancja Z_T jest definiowana jako:

$$Z_T = \frac{V_{out}}{I_{in}}$$

co oznacza, że jednostką transimpedancji jest om (Ω).

Zastosowanie

Wzmacniacze transimpedancyjne znajdują szerokie zastosowanie w układach przetwarzających sygnały z czujników prądowych, takich jak:

- Fotodiody w układach optoelektronicznych (np. odbiorniki światłowodowe, detektory optyczne).
- Sondy prądowe w oscyloskopach.
- Systemy biometryczne i sensory chemiczne.

Działanie i Modelowanie Matematyczne

Najczęściej stosowanym układem wzmacniacza transimpedancyjnego jest wzmacniacz operacyjny z rezystorem sprzężenia zwrotnego R_f , który zapewnia odpowiednią konwersję sygnału. Schemat układu przedstawia się następująco:

```
(0.5,0) node[op amp] (opamp) (opamp.+) - (-2,-0.5) node[ground] (opamp.-) - (-1,0.5) to[short, -*]
(-1,1.5) to[R, l=$R_f$] (2,1.5) to[short, -o] (3,1.5) node[right]$V_{out}$ (-1,1.5) to[short] (-1,3) to[I,
l=$I_{in}$] (-1,5) node[above] (opamp.out) - (1.70,1.5);
```

Z analizy węzłowej na wejściu wzmacniacza operacyjnego wynika, że napięcie na wejściu odwracającym wynosi $0V$ (idealny model wzmacniacza operacyjnego). Stosując prawo Ohma do rezystora sprzężenia zwrotnego:

$$V_{out} = -I_{in} R_f$$

co pokazuje, że transimpedancja wynosi:

$$Z_T = -R_f$$

Znaczenie znaku minus przy R_f

Znak minus w wyrażeniu $Z_T = -R_f$ oznacza, że wzmacniacz transimpedancyjny dokonuje odwrócenia fazy sygnału. Oznacza to, że jeśli prąd wejściowy I_{in} jest dodatni, napięcie wyjściowe V_{out} będzie ujemne i odwrotnie. Jest to konsekwencja działania wzmacniacza operacyjnego w konfiguracji odwracającej, w której napięcie na wejściu odwracającym jest równe zeru (masie wirtualnej), a sprzężenie zwrotne przez R_f powoduje odwrótność biegunowości sygnału.

Zjawiska Fizyczne i Ograniczenia

W praktycznych zastosowaniach należy uwzględnić kilka istotnych efektów:

- **Szum Johnsona-Nyquista** – rezystor sprzężenia zwrotnego generuje szum termiczny, który może ograniczać czułość układu.
- **Pojemność wejściowa** – fotodiody i inne sensory prądowe mają swoją pojemność własną C_d , co wpływa na pasmo przenoszenia układu.

- **Pasmo przenoszenia** – w rzeczywistości wzmacniacz operacyjny posiada skończoną szybkość narastania napięcia (slew rate) i ograniczoną szerokość pasma, co wpływa na odpowiedź układu na sygnały zmienne w czasie.

Aby poprawić stabilność i pasmo przenoszenia wzmacniacza transimpedancyjnego, często stosuje się kompensację biegunów poprzez dodanie kondensatora równolegle z rezystorem R_f , co redukuje oscylacje i zwiększa stabilność układu.

Podsumowanie

Wzmacniacze transimpedancyjne są kluczowym elementem wielu układów przetwarzających sygnały z sensorów prądowych. Ich projektowanie wymaga uwzględnienia zarówno parametrów elektronicznych, jak i efektów fizycznych, takich jak szумы i ograniczenia częstotliwościowe. Poprawna analiza matematyczna pozwala na optymalizację układów i zapewnienie ich stabilnej pracy w rzeczywistych zastosowaniach.

Układy BalUn (włas. Symetryzatory)

Układ BalUn (Balanced to Unbalanced) to transformator impedancyjny stosowany do konwersji sygnałów między układami zrównoważonymi (balanced) a niezrównoważonymi (unbalanced). Jego głównym zadaniem jest zapewnienie odpowiedniej transformacji impedancji między tymi dwoma typami układów.

Historia nazwy

Nazwa BalUn pochodzi od angielskich słów *Balanced* (zrównoważony) oraz *Unbalanced* (niezrównoważony). Układy te są szeroko stosowane w komunikacji radiowej oraz transmisji sygnałów, gdzie zrównoważone linie (np. kabel koncentryczny) są wykorzystywane do przesyłania sygnałów w systemach antenowych, a układy niezrównoważone (np. układy z impedancją 50 Ω) są powszechnie wykorzystywane w sprzęcie radiowym.

Znaczenie transimpedancji w BalUn

Transimpedancja w układzie BalUn jest kluczowym parametrem, ponieważ określa zdolność układu do konwersji sygnału z jednej formy na drugą, jednocześnie zachowując odpowiednią impedancję. W praktyce, transimpedancja BalUn zależy od rodzaju zastosowanego transformatora i jego charakterystyki. Współczesne BalUn-y wykorzystywane w urządzeniach RF (Radio Frequency) muszą być zaprojektowane tak, by zminimalizować straty sygnału, zapewniając niskie wartości transimpedancji przy odpowiednich parametrach.

Obliczanie parametrów BalUn

Parametry BalUn są zależne od kilku czynników, takich jak stosunek impedancji (np. 50 Ω do 200 Ω), częstotliwość sygnału oraz kształt transformatora. Aby obliczyć transimpedancję, wykorzystuje się wzory zależne od typu transformatora:

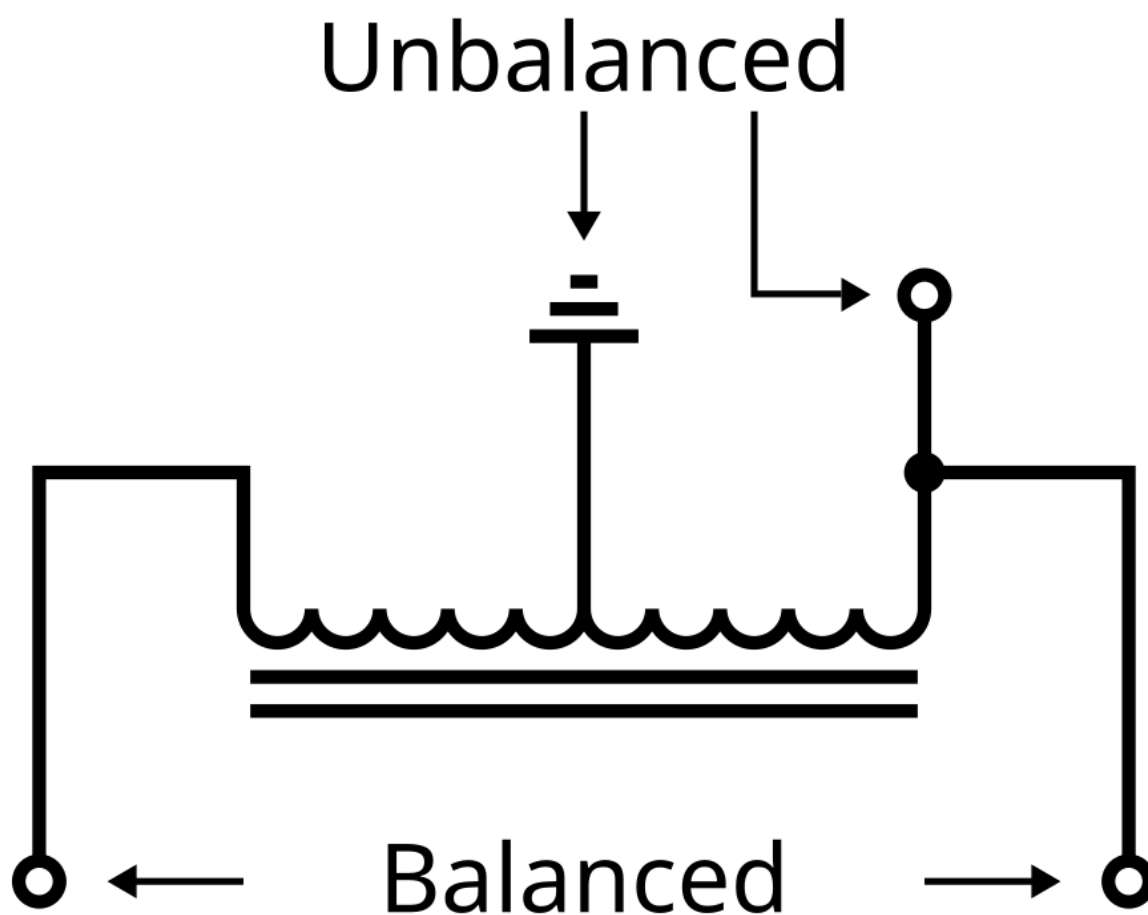
$$Z_{\text{balun}} = \sqrt{Z_{\text{in}} \cdot Z_{\text{out}}}$$

gdzie: - Z_{in} to impedancja wejściowa, - Z_{out} to impedancja wyjściowa.

Dla układu BalUn transformacja impedancji odbywa się w sposób proporcjonalny, a transformator pełni rolę obniżania lub podwyższania impedancji w zależności od konstrukcji układu.

Schemat układu BalUn

Poniżej przedstawiono prosty schemat układu BalUn:



Schemat układu BalUn

Układy UnUn

Układ UnUn (Unbalanced to Unbalanced) jest również układem transformującym impedancję, jednak w tym przypadku oba końce układu są nie zrównoważone. Układy te są stosowane głównie do dopasowywania impedancji w obwodach, gdzie oba elementy systemu są zbudowane na linii nie zrównoważonej.

Historia nazwy

Podobnie jak w przypadku układu BalUn, nazwa UnUn pochodzi od angielskich słów *Unbalanced* oraz *Unbalanced*. W tym układzie, głównym celem jest dopasowanie impedancji między dwoma niezerównoważonymi układami, np. między dwoma urządzeniami, które działają na różnych poziomach impedancji (np. 75 Ω i 50 Ω). Układ ten jest powszechnie stosowany w sieciach telekomunikacyjnych oraz w systemach audio.

Znaczenie transimpedancji w UnUn

Transimpedancja w układzie UnUn pełni podobną rolę jak w układzie BalUn, ponieważ pozwala na skuteczną konwersję sygnału między różnymi impedancjami niezerównoważonymi. Jest to istotne w kontekście dopasowania impedancji dla poprawnej transmisji sygnałów oraz unikania strat mocy.

Obliczanie parametrów UnUn

Podobnie jak w przypadku BalUn, obliczanie parametrów układu UnUn zależy od impedancji wejściowej i wyjściowej. Wzór na transimpedancję w przypadku układu UnUn jest następujący:

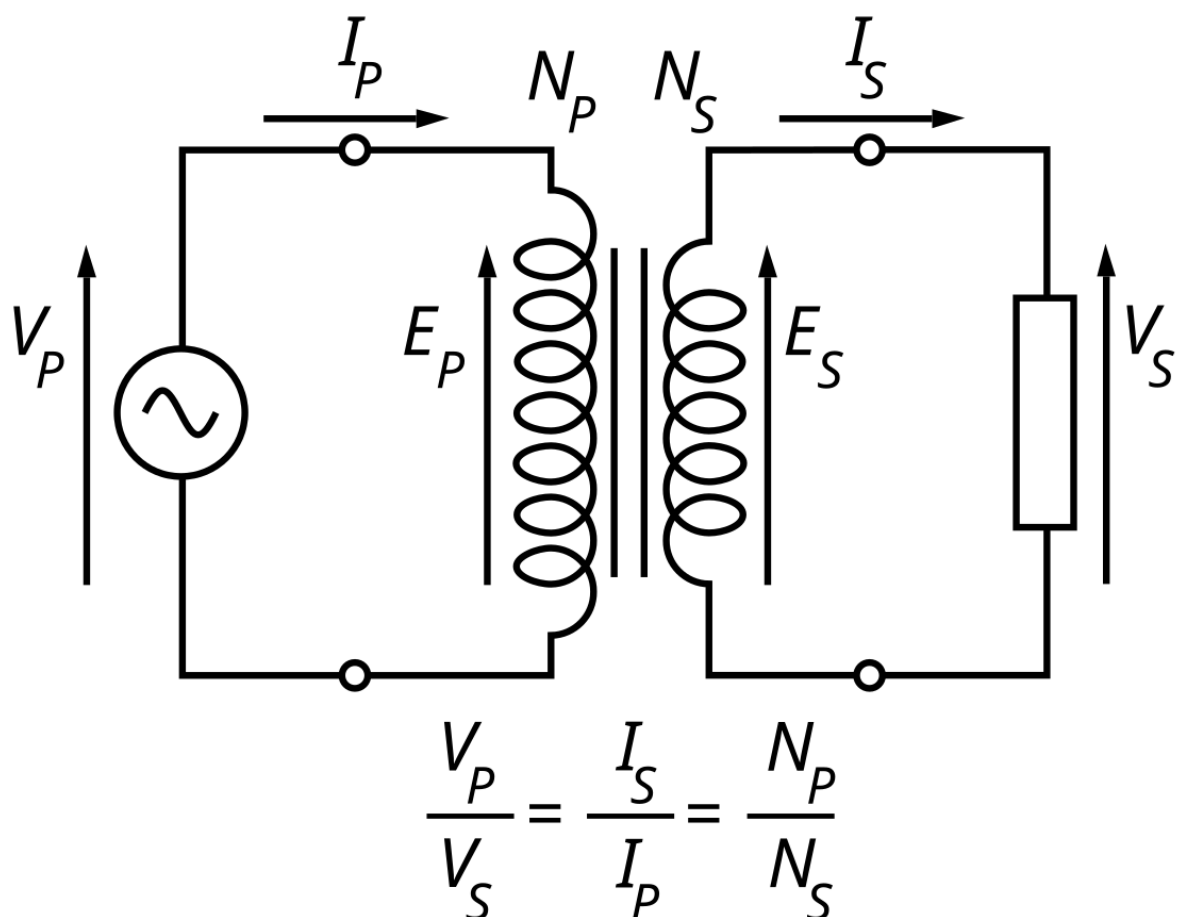
$$Z_{\text{unun}} = \frac{Z_{\text{in}} \cdot Z_{\text{out}}}{Z_{\text{in}} + Z_{\text{out}}}$$

gdzie: - Z_{in} to impedancja wejściowa, - Z_{out} to impedancja wyjściowa.

Układy UnUn mogą być wykorzystywane do łączenia urządzeń o różnych impedancjach, co sprawia, że są one szczególnie użyteczne w systemach audio i telekomunikacyjnych.

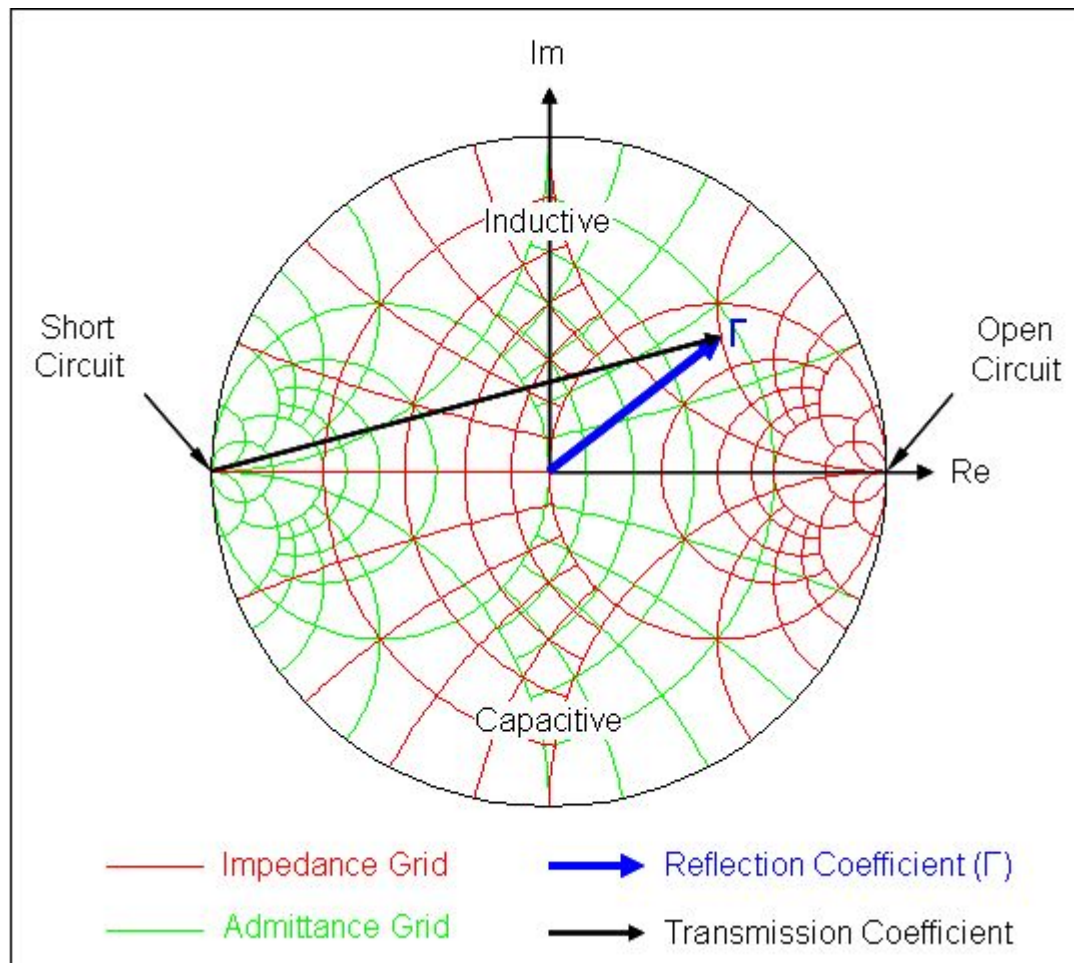
Schemat układu UnUn

Poniżej przedstawiono przykładowy schemat układu UnUn:



Schemat układu UnUn

Wykresy Smitha i ich zastosowanie w dopasowaniu impedancyjnym



Ilustracja ułatwiająca zrozumienie działania wykresu smitha

Wykres Smitha jest narzędziem graficznym wykorzystywanym w telekomunikacji oraz elektronice do analizy impedancji i dopasowywania impedancyjnego w układach radiowych, mikrofalowych i innych obwodach RF. Jest to szczególny przypadek wykresu zespolonego, na którym reprezentowane są zarówno impedancje, jak i admitancje. Dzięki swojej konstrukcji, wykres Smitha pozwala na łatwe przedstawienie właściwości impedancyjnych obwodów, co ułatwia projektowanie układów RF.

Podstawowe pojęcia

Impedancja Z obwodu elektrycznego jest wielkością zespoloną i ma postać: $Z = R + jX$ gdzie:

- R – rezystancja,
- X – reaktancja,
- j – jednostka urojona.

Wykres Smitha przedstawia impedancje w układzie zespolonym, gdzie osie odpowiadają za komponenty rezystancyjne i reaktancyjne. Główne elementy wykresu to okręgi odpowiadające wartościom rezystancji, oraz okręgi odpowiadające wartościom reaktancji.

Zastosowanie wykresu Smitha do dopasowania impedancyjnego

Dopasowanie impedancyjne jest procesem dostosowywania impedancji źródła i obciążenia, aby zminimalizować straty energii i zapewnić maksymalny transfer mocy. Wykres Smitha jest narzędziem umożliwiającym szybkie i intuicyjne dokonanie tego dopasowania. Zastosowanie wykresu Smitha do dopasowania impedancyjnego obejmuje następujące kroki:

1. **Określenie impedancji źródła i obciążenia:** Na wykresie Smitha przedstawiamy impedancje źródła i obciążenia jako punkty na odpowiednich okręgach, które reprezentują wartość rezystancyjną oraz reaktancyjną.
2. **Obliczenie współczynnika dopasowania:** Następnie wyznaczamy współczynnik dopasowania (np. współczynnik odbicia) na podstawie odległości między punktami reprezentującymi impedancje źródła i obciążenia oraz centrum wykresu. Wartość ta wskazuje na jakość dopasowania impedancyjnego.
3. **Dobór elementu dopasowującego:** Wykres Smitha pozwala na wybranie odpowiednich elementów dopasowujących, takich jak dławiki, kondensatory lub inne elementy pasywne, które pozwolą na przesunięcie punktu impedancyjnego w stronę centrum wykresu. Przesunięcie to zmniejsza współczynnik odbicia i poprawia efektywność dopasowania.

Ważną cechą wykresu Smitha jest to, że umożliwia on również wizualizację zmian impedancji w zależności od długości linii transmisyjnej. Na przykład, zmieniająca się impedancja wzdłuż linii transmisyjnej może być przedstawiona jako krzywa na wykresie, co ułatwia ocenę efektów zmiany długości linii na dopasowanie impedancyjne.

Matematyka leżąca u podstaw wykresu Smitha

Wykres Smitha jest narzędziem, które opiera się na matematyce zespolonej i transformacjach impedancji. Jego celem jest przedstawienie impedancji i admitancji w sposób graficzny, co pozwala na intuicyjne zrozumienie ich właściwości oraz łatwe dopasowanie impedancyjne. W tej subsekcji omówimy matematyczne podstawy wykresu Smitha, w tym sposób reprezentacji impedancji oraz zależności między impedancją, admitancją a współczynnikiem odbicia.

Impedancja i jej reprezentacja na płaszczyźnie zespolonej

Impedancja Z jest wielkością zespoloną, którą można zapisać w postaci: $Z = R + jX$ gdzie:

- R – część rzeczywista, odpowiadająca rezystancji,
- X – część urojona, odpowiadająca reaktancji,
- j – jednostka urojona, $j^2 = -1$.

Wykres Smitha opiera się na odwzorowaniu tej impedancji na jednostkowym okręgu w płaszczyźnie zespolonej. Impedancje normalizowane względem charakterystycznej impedancji linii transmisyjnej Z_0 są przedstawiane jako punkt na tym okręgu, gdzie: $z = \frac{Z}{Z_0}$. Reprezentacja impedancji na wykresie Smitha uwzględnia zarówno rezystancję, jak i reaktancję, a okręgi na wykresie odpowiadają różnym wartościom tych parametrów.

Admitancja

Admitancja to odwrotność impedancji w obwodach prądu zmiennego. Jest to wielkość zespolona, która opisuje łatwość przepływu prądu przez elementy obwodu. Admitancję Y definiuje się jako odwrotność impedancji Z : $Y = \frac{1}{Z} = \frac{1}{R + jX}$ gdzie:

- R – rezystancja,
- X – reaktancja,
- j – jednostka urojona.

Admitancja jest również wielkością zespoloną, której część rzeczywista to konduktancja (G - odwrotność rezystancji), a część urojona to susceptancja (B): $Y = G + jB$. W obwodach elektrycznych, konduktancja (G) mierzy zdolność przewodzenia prądu, a susceptancja (B) odpowiada za reakcję elementów obwodu, takich jak kondensatory i dławiki, na zmienne pole elektryczne. Podobnie jak impedancja, admitancja jest wykorzystywana w analizach linii transmisyjnych i w dopasowaniu impedancyjnym.

Transformacja impedancji do admitancji

Na wykresie Smitha przedstawiane są nie tylko impedancje, ale również admitancje, które są odwrotnością impedancji. Admitancję Y oblicza się jako: $Y = \frac{1}{Z} = \frac{1}{R + jX}$. Zaletą wykresu Smitha jest to, że umożliwia on jednoczesne przedstawienie zarówno impedancji, jak i admitancji, co jest szczególnie przydatne w analizach linii transmisyjnych i dopasowania impedancyjnego.

Wartości admitancji można wyrazić w postaci: $Y = G + jB$ gdzie:

- G – konduktancja, część rzeczywista admitancji,
- B – susceptancja, część urojona admitancji.

Na wykresie Smitha dla admitancji, osie odpowiadają podobnym wartościom, ale reprezentują one konduktancję i susceptancję. Okręgi, które reprezentują admitancje, są symetryczne względem osi rzeczywistej.

Współczynnik odbicia

Kolejnym ważnym elementem matematycznym związanym z wykresem Smitha jest współczynnik odbicia Γ , który opisuje stopień, w jakim sygnał zostaje odbity od obciążenia. Współczynnik odbicia jest definiowany jako stosunek amplitudy fali odbitej do amplitudy fali przychodzącej i może być obliczany ze wzoru: $\Gamma = \frac{Z_L - Z_0}{Z_L + Z_0}$ gdzie:

- Z_L – impedancja obciążenia,
- Z_0 – charakterystyczna impedancja linii transmisyjnej.

Na wykresie Smitha współczynnik odbicia jest przedstawiany jako odległość punktu reprezentującego impedancję obciążenia od centrum wykresu. Wartość $\Gamma = 0$ oznacza brak odbicia (impedancja dopasowana), podczas gdy $\Gamma = 1$ oznacza pełne odbicie (impedancja źródła i obciążenia są całkowicie niedopasowane).

Mapowanie impedancji na wykresie Smitha

Podstawową ideą wykresu Smitha jest mapowanie impedancji normalizowanej $z = \frac{Z}{Z_0}$ na okrągłej siatce. Okręgi te mogą być interpretowane w kontekście rozkładu wartości R (rezystancja) i X (reaktancja). W zależności od wartości z , impedancje mogą zostać przedstawione na wykresie jako:

- Punkty na okręgach odpowiadających rezystancji stałej (wartości R),
- Punkty na okręgach odpowiadających reaktancji stałej (wartości X),
- Linie łączące punkty na wykresie reprezentują zmiany w impedancji w wyniku zmian długości linii transmisyjnej.

Wszystkie te transformacje umożliwiają wizualizację i manipulację impedancjami i admitancjami w sposób umożliwiający łatwe przeprowadzanie dopasowania impedancyjnego i rozwiązywanie problemów związanych z maksymalnym transferem mocy.

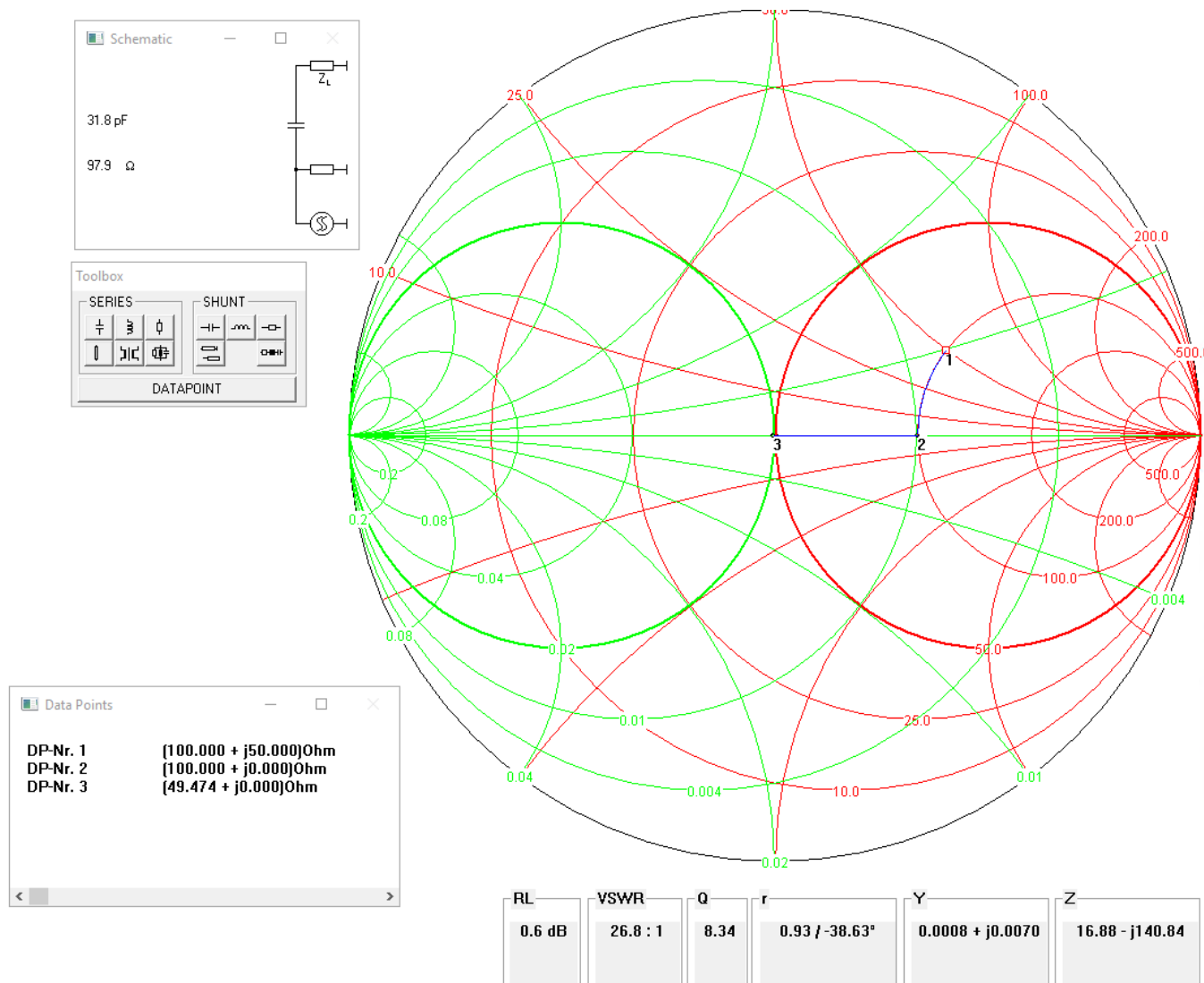
Przykład

Założmy, że mamy obwód z impedancją źródła $Z_s = 100 + j50 \, \Omega$. Aby dopasować impedancję, możemy na wykresie Smitha zaznaczyć punkt i znaleźć odpowiedni element dopasowujący, na przykład kondensator lub dławik, który przesunie impedancję w stronę centrum wykresu.

Program wykorzystany do prezentacji poniżej można pobrać pod linkiem:

<http://filevista.ardugeek.ovh/public/yq/wykres-smitha.exe> link jest zabezpieczony hasłem: *Radio23022025*

Program niestety jest w wersji demo.



Zastosowanie wykresu smitha

Po wykonaniu obliczeń przez program możemy zauważyć że żeby wyrównać impedancję do wartości 50 Ω musimy zastosować rezystor równolegle do źródła sygnału o wartości około 100 Ω oraz kondensator szeregowo do źródła o wartości około 30 pF

99

Vestigial Sideband Modulation System (VSB).

<https://www.eeeguide.com/vestigial-sideband-modulation-system-vsbs/>.

Wikipedia free online encyclopedia. <https://www.wikipedia.org/>.

Microwaves101.com The world's microwave information source since 2001.

<https://www.microwaves101.com/encyclopedias/smith-chart-basics>.

From:

<https://wiki.ostrowski.net.pl/> - Kacper's Wiki

Permanent link:

<https://wiki.ostrowski.net.pl/doku.php?id=notatki:radiotechnikamodulacje&rev=1746610233>

Last update: **2025/05/07 11:30**

