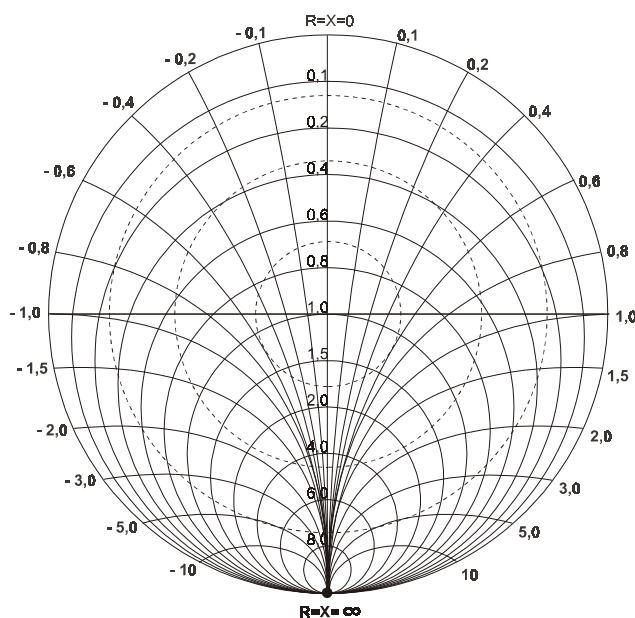


PODSTAWY TECHNIKI MIKROFAL

Mateusz Pasternak



Warszawa 2002

Spis treści

Elementy teorii pola	5
Pojęcie pola	5
Ruch ładunku w polu.....	10
Równania Maxwella w próżni.....	14
Równania Maxwella w postaci całkowej.....	15
Równania falowe.....	16
Twierdzenie Poyntinga.....	17
Płaskie fale elektromagnetyczne	19
Równania Maxwella w postaci zespolonej	22
Równania falowe w postaci zespolonej	23
Rozwiązanie równań Helmholtza poprzez rozdzielenie zmiennych.....	24
Dyskusja	25
Niejednorodne równanie falowe	27
Relatywistyczna postać równań Maxwella	30
Odbicie fal elektromagnetycznych.....	32
Fala pada prostopadle na idealny przewodnik	32
Fala pada prostopadle na idealny dielektryk.....	35
Głębokość wnikania fal elektromagnetycznych.....	38
Wybrane własności fal elektromagnetycznych.....	40
Teoria linii transmisyjnych.....	41
Falowody prostokątne	41
Falowody cylindryczne	45
Miary dopasowania	47
Transformacyjne własności linii przesyłowych	48
Równania telegrafistów.....	50
Wykres Smitha.	55
Wykres Smitha.	55
Odwzorowanie homograficzne	55
Konstrukcja wykresu Smitha	56
Metody dopasowania impedancji.....	60
Transformator ćwierćfalowy	61
Łańcuch transformatorów.....	62
Stroik	63
Macierz rozproszenia	66
Przykład wyznaczania macierzy rozproszenia.....	70
Rezonatory mikrofalowe.....	72
Sposoby sprzęgania rezonatorów z liniami przesyłowymi	74
Dobroć rezonatorów mikrofalowych	74
Trzy rodzaje sprzężeń rezonatorów mikrofalowych	75
Filtry mikrofalowe.....	78

Wyznaczanie elementów filtrów	80
Transformacja częstotliwości	81
Skalowanie wartości elementów	81
Inwertery immitancji	83
Przykłady praktycznych realizacji filtrów.....	86
Sprzęgacze kierunkowe.....	87
Sprzęgacze 3 dB	89
Elementy ferrytowe	91
Wybrane zjawiska fizyczne zachodzące w ferrytach.....	92
Izolatory ferrytowe	94
Cyrkulatory	97
Nieciągłości w liniach transmisyjnych.....	98
Zaginanie linii przesyłowych z kompensacją reaktancji.....	100
Wybrane elementy pasywne toru mikrofalowego.....	101
Przykłady przejść między liniami różnego typu	102
Linie opóźniające	104
Przełączniki nadawczo odbiorcze	104
Elementy akustoelektroniki mikrofalowej	107
Fala Rayleigha	108
Zasada działania przetworników IDT	108
Powierzchniowe własności piezoelektryków	111
Synteza podzespołów z AFP	112
Model funkcji δ	112
Przetworniki jednokierunkowe	114
Rezonatory z AFP	116
Linie dyspersyjne	118
Konwoluty	118
Czujniki	119
Literatura źródłowa	121

Wstęp

Niniejszy skrypt, stanowiąc kompilację informacji pochodzących z różnych źródeł, ma przede wszystkim na celu pomoc w utrwaleniu wiedzy zdobytej przez studentów WEL WAT na wykładach z podstaw techniki mikrofal (ewentualnie teorii pola elektro-magnetycznego). Z tego też powodu informacje w nim zawarte podane są w sposób skrótowy (dla kogoś, kto zetknął się z określonym zagadnieniem w zupełności do jego przypomnienia wystarczy rysunek lub prosty opis) i ich właściwe zrozumienie wymaga czynnego uczestnictwa w zajęciach audytoryjnych.

Rozdziały wykraczające poza ramy programu nauczania należy uznać jako swego rodzaju uzupełnienie i wskazanie możliwych kierunków pogłębiania wiedzy w przedmiotowej dziedzinie. Te zaś fragmenty, które ściśle przylegają do programu nauczania konkretnej grupy studentów, można uznać, co wydaje się dobrym pomysłem, za solidny i łatwy do obudowania w trakcie wykładów szkielet notatek.

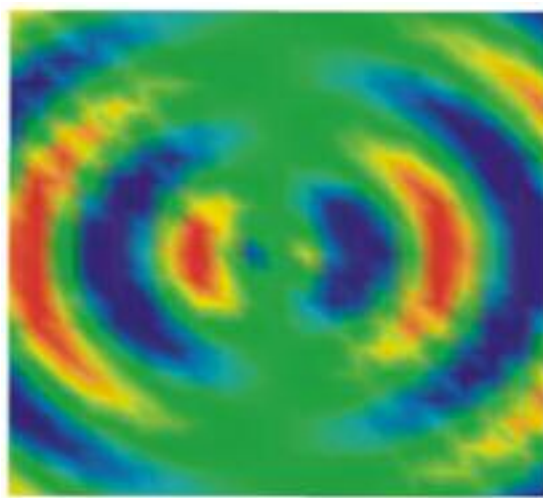
autor

ELEMENTY TEORII POLA

Pojęcie pola

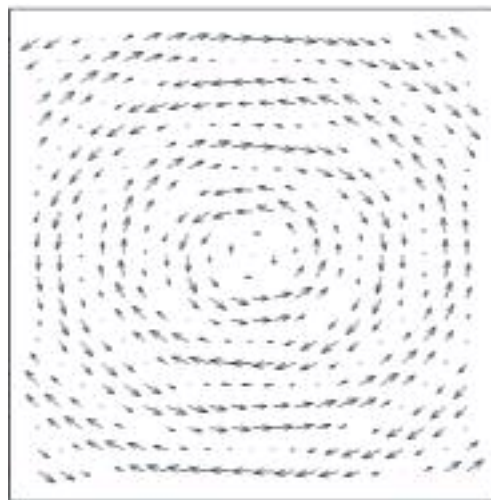
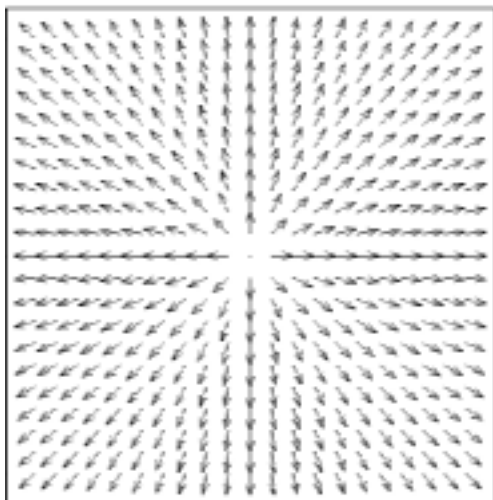
Pole w przestrzeni definiuje się przypisując każdemu punktowi tej przestrzeni określoną wielkość. Jeśli wielkość ta jest skalarą to zdefiniowane zostaje pole skalarne, jeśli wektorem – wektorowe. Termin „punkt” nie oznacza tu miejsca wyznaczonego z nieskończoną wielką dokładnością, ale, w sensie mechaniki klasycznej, pewien obszar przestrzeni zanedbywalnie mały w porównaniu z opisywanym za pomocą pola obszarem.

W tym sensie dobrym przykładem pola skalarnego jest mapa bitowa (jeśli każdemu kolorowi przypiszemy liczbę). Innym przykładem pola skalarnego może być rozkład temperatur na danym obszarze.



Bardzo dobrym przykładem pola skalarnego jest mapa fizyczna – każdy kolor odpowiada na niej liczbie oznaczającej wysokość nad poziomem morza.

Jako przykład pola wektorowego można podać pole prędkości cząstek gazu, pole sił grawitacji w otoczeniu masy itp.



Pokazane tu przykłady pól są oczywiście polami płaskimi (zdefiniowanymi na płaszczyźnie). Matematycznie można opisać jednak pola w dowolnej liczbie wymiarów (mniejsza o to czy można je sobie wyobrazić).

Przypisanie każdemu punktowi wielkości skalarnej lub wektorowej byłoby zadaniem niezwykle czasochłonnym. Zamiast tego pola określa się za pomocą potencjałów czyli funkcji opisujących pole. Jeśli tylko dana funkcja jest określona w całej przestrzeni to można oczywiście wyznaczyć wartość pola w dowolnym punkcie tej przestrzeni. W takim przypadku „punkt” rozumiany jest w sensie analitycznym tzn. jako konkretna wartość argumentów funkcji.

O istnieniu pola w przestrzeni można się przekonać obserwując określone efekty fizyczne. W przypadku pola elektrycznego obserwuje się wywieranie wpływu na ładunki elektryczne (obecność sił) oraz indukowanie ładunków elektrycznych na neutralnych powierzchniach ciał.

Pomiar siły F działającej na ładunek często prowadzi do wniosku, że siła ta w różnych punktach pola ma różną wartość i zwrot (jak np. na rys. B z poprzedniej strony) ale jest zawsze proporcjonalna do ładunku q . Daje to możliwość charakteryzowania pola elektrycznego za pomocą wyrażenia opisującego tę proporcjonalność:

$$F = Eq$$

Współczynnik proporcjonalności E nazywany jest natężeniem pola elektrycznego.

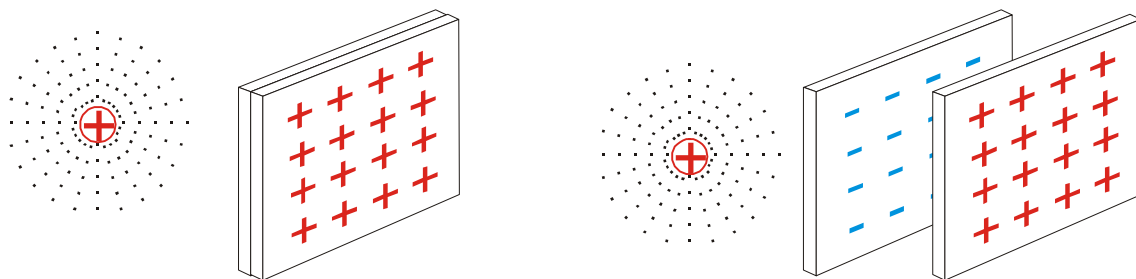
Przykłady liczbowe

Natężenie pola elektrycznego Ziemi przy jej powierzchni rzędu 100 V/m.

Natężenie pola elektrycznego, przy którym następuje elektryczne przebicie powietrza to ok. 30 kV/cm.

Dowolne pole wektorowe można scharakteryzować także przy pomocy tzw. linii wektorowych określanych w ten sposób, że styczna do każdego punktu tych linii pokrywa się z wektorem pola w tym punkcie.

Pomiaru ładunków elektrycznych indukowanych przez pole można dokonać za pomocą umieszczonych w polu płytek próbnych o małej powierzchni (płytek Mie).



Jeśli początkowo płytki się stykają to po rozsunięciu ich na każdej zgromadzi się ładunek – dodatni na jednej, a ujemny na drugiej. Ilość tego ładunku można zmierzyć np. za pomocą elektrometru.

Z doświadczeń wynika, że jeżeli wektor $\mathbf{E}$ jest do płytek równoległy to nie zgromadzi się na nich żaden ładunek, natomiast jeśli jest prostopadły to zgromadzony ładunek osiąga maksimum.

Fakt ten pozwala na wprowadzenie jeszcze jednej wielkości charakteryzującej pole elektryczne, a mianowicie wektora indukcji elektrycznej $\mathbf{D}$ o kierunku zgodnym z $\mathbf{E}$, zwrocie w kierunku płytki naładowanej ujemnie i długości równej powierzchniowej gęstości ładunku rozmieszczonego na płycie Mie.

$$|\mathbf{D}| = \frac{dq}{dS_{\perp}}$$

Jeśli wektory natężenia pola elektrycznego nie są prostopadłe do powierzchni płytek Mie, ale odchylone od normalnej do nich o kąt α to:

$$|\mathbf{D}| = \frac{dq}{dS \cos \alpha}$$

Linie wektora $\mathbf{D}$ są nazywane liniami sił pola elektrycznego.

Wielkość:

$$N = \int |\mathbf{D}| \cos \alpha dS$$

Nazywa się strumieniem pola elektrycznego przez powierzchnię dS . Jeśli powierzchnia jest zamknięta to:

$$N = \oint |\mathbf{D}| \cos \alpha dS.$$

Wektory $\mathbf{D}$ i $\mathbf{E}$ są do siebie w danym punkcie ośrodka proporcjonalne, a współczynnikiem proporcjonalności jest tzw. przenikalność elektryczna w tym punkcie ośrodka ϵ . Przyjęto:

$$\mathbf{D} = \epsilon \mathbf{E} \quad \text{w układzie jednostek CGS}$$

$$\mathbf{D} = \epsilon \epsilon_0 \mathbf{E} \quad \text{w układzie jednostek SI}$$

Wielkość ta charakteryzuje własności elektryczne ośrodka w danym punkcie. W szczególności w ośrodkach izotropowych i jednorodnych wartość ϵ w każdym punkcie jest stała.

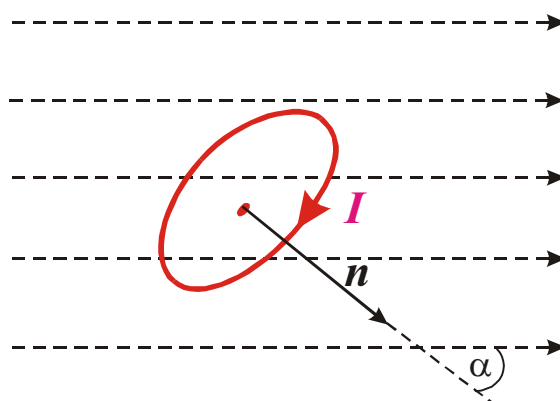
Ponieważ nie istnieją ciała o $\epsilon < 1$ przeto próżnię uznaje się często za początek skali dla ϵ i przyjmuje dla niej wielkość $\epsilon_0 = 1$.

$$\text{W jednostkach technicznych wartość } \epsilon_0 = \frac{10^{-9}}{36\pi} \left[\frac{\text{C}^2}{\text{Nm}^2} \right].$$

Stosuje się też tzw. absolutny układ jednostek, w którym $\epsilon_0 = \frac{1}{4\pi}$ (CGS)

Analogiczne wielkości do tutaj omówionych można zdefiniować dla wielu typów pól, w szczególności dla pola magnetycznego. Specyfika tego pola sprawia jednak, że mają one inny charakter. Źródłami pola magnetycznego nie są bowiem ładunki, a zgodnie z prezentowanym tu podejściem ruch tych ładunków czyli prądy. Z tego powodu

zamiast ładunków próbnych używa się w polu magnetycznym zamkniętych pętli z prądem stałym o bardzo małym promieniu.



Dla takiego obwodu próbnego definiuje się wielkość wektorową zwaną momentem magnetycznym. Ma on wartość:

$$\mathbf{M} = I S \mathbf{n}$$

gdzie S jest powierzchnią objętą przez pętlę (I i $\mathbf{n}$ jak na rysunku).

Moment siły skręcającej $\mathbf{N}$ działający na pętlę próbną wyraża się poprzez:

$$\mathbf{N} = B \mathbf{M} \sin \alpha,$$

gdzie B jest współczynnikiem proporcjonalności nazywanym indukcją magnetyczną.

Jeśli moment skręcający przedstawić jako wektor o kierunku zgodnym z osią obrotu i zwrotem wynikającym z reguły śruby prawoskrętnej to można zapisać:

$$\mathbf{N} = \mathbf{M} \times \mathbf{B}$$

Wzór ten wyraża fakt, że pole magnetyczne dąży do ustawienia obwodu z prądem w taki sposób, aby jego własny moment magnetyczny był zgodny z kierunkiem pola.

Siła z jaką działa pole magnetyczne na nieskończenie krótki odcinek przewodu z prądem określona jest przez prawo Ampera:

$$d\mathbf{F} = I(d\mathbf{l} \times \mathbf{B})$$

Siła działająca na przewód o skończonej długości będzie wynosić:

$$\mathbf{F} = I \int d\mathbf{l} \times \mathbf{B}.$$

Analogicznie do strumienia pola elektrycznego definiuje się strumień pola magnetycznego:

$$\Phi = \int \mathbf{B} \cos \alpha dS,$$

lub, jeśli powierzchnia jest zamknięta,

$$\Phi = \oint \mathbf{B} \cos \alpha dS.$$

Znając siłę działającą na przewód z prądem oraz strumień pola przez zadaną powierzchnię można następująco zdefiniować natężenie pola magnetycznego:

$$\mathbf{H} = \int \frac{1}{\Phi} d\mathbf{F}$$

Wektory $\mathbf{B}$ i $\mathbf{H}$ są do siebie proporcjonalne (podobnie jak $\mathbf{D}$ i $\mathbf{E}$), a współczynnikiem proporcjonalności jest przenikalność magnetyczna μ .

W fizycznym układzie jednostek przenikalność magnetyczna próżni $\mu_0=1$.

W układzie jednostek technicznych $\mu_0 = 4\pi \cdot 10^{-7} [\text{J/A}^2 \cdot \text{m}]$.

Pojęcia tu wprowadzone w oparciu o rozważania nad pętlą z prądem można także rozszerzyć na takie układy, dla których pojęcie obwodu z prądem nie ma sensu (magnesy trwałe, momenty magnetyczne cząstek i in.). Rozszerzenie taki jest uprawnione ze względu na tzw. równoważność prądów i magnesów. Jak się okazuje skutki ich działania są jednakowe z dokładnością do stałego czynnika równego $\mu\mu_0$.

Ruch ładunku w polu

Rozpocznijmy od elementarnej zasady wariacyjnej (zasady Hamiltona), która mówi, że dla ruchów rzeczywistych wariacja działania cząstki swobodnej (tutaj ładunku):

$$S = -\alpha \int_a^b ds \quad (1)$$

wynosi zero, co zapisuje się w formie: $\delta S = 0$.

α - stała charakteryzująca cząstkę (np. ładunek, masa itp.)

$\overline{ab}$ - droga między dwoma zjawiskami (położeniami cząstki) wzdłuż linii świata.

Dla zjawisk zachodzących w chwilach t_1 i t_2 można napisać:

$$S = \int_{t_1}^{t_2} L dt \quad (2)$$

gdzie L naz. lagranżianem, wyraża energię zużytą na ruch.

Z zasady względności czasu:

$$S = -\int_{t_1}^{t_2} \alpha c \sqrt{1 - \frac{v^2}{c^2}} dt. \quad (3)$$

Z (2) i (3) wynika zaś:

$$L = -\alpha c \sqrt{1 - \frac{v^2}{c^2}}. \quad (4)$$

Po rozwinięciu (4) w szereg potęgowy względem v/c i odrzuceniu wyrazów wyższych rzędów otrzymujemy:

$$L \approx \frac{\alpha v^2}{2c}. \quad (5)$$

Energia ruchu w mechanice klasycznej wyraża się przez:

$$L = \frac{mv^2}{2}, \quad (6)$$

zatem

$$\alpha = mc \quad (7)$$

Jeśli ładunek znajduje się w polu to w całce działania należy uwzględnić także oddziaływanie pola z ładunkiem (opis polowy!)

Wyraz uwzględniający to oddziaływanie ma postać:

$$\frac{e}{c} \int_a^b A_i dx_i \quad (8)$$

i charakteryzuje zarówno ładunek jak i samo pole, wektor A_i naz. potencjałem pola.

Na mocy (7) i po uwzględnieniu (8) całka działania (1) wyrazi się poprzez:

$$S = \int_a^b -mc ds + \frac{e}{c} A_k dx_k . \quad (9)$$

Wektor A_k ma w tym przypadku cztery składowe (czterowektor): trzy przestrzenne i jeden czasowy. Składowa czasowa jest urojona i ma postać :

$$A_4 = i\varphi , \quad (10)$$

φ – skalarny potencjał pola.

Rozpisując składowe otrzymamy:

$$S = \int_a^b -mc ds + \frac{e}{c} \mathbf{A} d\mathbf{r} - e\varphi dt . \quad (11)$$

Biorąc po uwagę, że $\mathbf{v} = \frac{d\mathbf{r}}{dt}$ równanie (11) można podać w postaci:

$$S = \int_{t_1}^{t_2} \underbrace{\left(-mc^2 \sqrt{1 - \frac{v^2}{c^2}} + \frac{e}{c} \mathbf{A} \mathbf{v} - e\varphi \right)}_L dt . \quad (12)$$

Równanie to opisuje obustronne oddziaływanie ładunku z polem.

Równanie (12) można przekształcić do postaci różniczkowej:

$$\frac{\partial}{\partial t} \frac{\partial \mathcal{L}}{\partial \mathbf{v}} = \frac{\partial \mathcal{L}}{\partial \mathbf{r}} \quad (13)$$

gdzie L jest wyrażeniem podcałkowym (12).

Równanie (13) nosi nazwę zwyczajnego równania Lagrange'a.

Rozpisując prawą stronę (13):

$$\frac{\partial \mathcal{L}}{\partial \mathbf{r}} \equiv \nabla L = \frac{e}{c} \text{grad} \mathbf{A} \mathbf{v} - e \text{grad} \varphi \quad (14)$$

oraz korzystając z tożsamości różniczkowej:

$$\text{grad} \mathbf{a} \mathbf{b} = (\mathbf{a} \nabla) \mathbf{b} + (\mathbf{b} \nabla) \mathbf{a} + \mathbf{b} \text{rot} \mathbf{a} + \mathbf{a} \text{rot} \mathbf{b} , \quad (15)$$

równanie Lagrange'a przybiera postać:

$$\frac{d}{dt} \left(\mathbf{p} + \frac{e}{c} \mathbf{A} \right) = \frac{e}{c} (\mathbf{v} \nabla) \mathbf{A} + \frac{e}{c} \mathbf{v} \text{rot} \mathbf{A} - e \text{grad} \varphi \quad (16)$$

gdzie $\mathbf{p}$ – jest pędem uogólnionym.

Różniczka zupełna $\frac{d\mathbf{A}}{dt}$ zawiera w istocie dwie składowe:

- zmianę potencjału wektorowego w czasie w ustalonym punkcie przestrzeni $\frac{\partial \mathbf{A}}{\partial t}$
- zmianę związaną z przejściem z ustalonego punktu w przestrzeni do drugiego punktu o wektor $d\mathbf{r}$. Z analizy wektorowej wynika, że ta składowa ma postać: $(d\mathbf{r} \nabla) \mathbf{A}$

zatem:

$$\frac{d\mathbf{A}}{dt} = \frac{\partial \mathbf{A}}{\partial t} + (\mathbf{v} \nabla) \mathbf{A}. \quad (17)$$

Na podstawie (16) i (17) można więc zapisać:

$$\frac{d\mathbf{p}}{dt} = \underbrace{-\frac{e}{c} \frac{\partial \mathbf{A}}{\partial t} - e \operatorname{grad} \varphi + \frac{e}{c} \mathbf{v} \operatorname{rot} \mathbf{A}}_{\text{siła Lorentza}}. \quad (18)$$

Pierwsza i druga składowa siły Lorentza nie zależą od prędkości, trzecia składowa zależy od prędkości i (z uwagi na rotację) jest do niej prostopadła.

Człony niezależne od prędkości zwykło się nazywać natężeniem pola elektrycznego $\mathbf{E}$, zaś współczynnik stojący przy prędkości w trzecim członie (dokładnie przy $\frac{\mathbf{v}}{c}$) natężeniem pola magnetycznego $\mathbf{H}$.

Przyjmując, że e jest ładunkiem jednostkowym mamy (def.):

$$\mathbf{E} = -\frac{1}{c} \frac{\partial \mathbf{A}}{\partial t} - \operatorname{grad} \varphi \quad (19)$$

$$\mathbf{H} = \operatorname{rot} \mathbf{A} \quad (20)$$

Jeśli $\mathbf{E} \neq 0$ i $\mathbf{H} = 0$ mówi się o polu elektrycznym.

Jeśli $\mathbf{E} = 0$ i $\mathbf{H} \neq 0$ mówi się o polu magnetycznym.

Ogólnie jeśli $\mathbf{E} \neq 0$ i $\mathbf{H} \neq 0$ mówi się o superpozycji pól elektrycznego i magnetycznego czyli o polu elektromagnetycznym.

W związku z (19) i (20) równanie ruchu ładunku w polu elektromagnetycznym można napisać w postaci:

$$\frac{d\mathbf{p}}{dt} = e\mathbf{E} + \frac{e}{c} \mathbf{v} \mathbf{H} \quad (21)$$

Ze związków (19) i (20) można łatwo uzyskać równania pól pozbywając się potencjałów. W tym celu poddaje się równanie (19) obustronnej rotacji i wstawia równanie (20).

$$\operatorname{rot} \mathbf{E} = -\frac{1}{c} \frac{\partial}{\partial t} \operatorname{rot} \mathbf{A} - \operatorname{rot} \operatorname{grad} \varphi \quad (22)$$

Ponieważ $\text{rot grad } \mathbf{U} = 0$ i $\mathbf{H} = \text{rot} \mathbf{A}$ otrzymujemy:

$$\boxed{\text{rot} \mathbf{E} = -\frac{1}{c} \frac{\partial \mathbf{H}}{\partial t}} \quad (23)$$

Działając zaś na równanie (20) obustronnie operatorem div i pamiętając, że $\text{div rot} \mathbf{U} = 0$ otrzymujemy:

$$\boxed{\text{div} \mathbf{H} = 0} \quad (24)$$

Równania (23) i (24) to pierwsza para równań Maxwella uzyskana w roku 1860. Drugą parę równań otrzymuje się podobnie, ale z uwzględnieniem faktu generacji pola magnetycznego na skutek przepływu prądów (konwekcyjnych lub przewodzenia).

Równania Maxwella często zapisuje się w równoważnej postaci:

	w układzie SI	w układzie CGS
(28)	$\nabla \times \mathbf{H} = \mathbf{j} + \frac{\partial \mathbf{D}}{\partial t}$	$\nabla \times \mathbf{H} = \frac{4\pi}{c} \mathbf{j} + \frac{1}{c} \frac{\partial \mathbf{D}}{\partial t}$
(29)		
(30)	$\nabla \times \mathbf{E} = -\frac{\partial \mathbf{B}}{\partial t}$	$\nabla \times \mathbf{E} = -\frac{1}{c} \frac{\partial \mathbf{B}}{\partial t}$
(31)		
	$\nabla \mathbf{D} = \rho$	$\nabla \mathbf{D} = \rho$
	$\nabla \mathbf{B} = 0$	$\nabla \mathbf{B} = 0$
	gdzie:	gdzie:
	$\mathbf{D} = \varepsilon \varepsilon_0 \mathbf{E}$	$\mathbf{D} = \varepsilon \mathbf{E}$
	$\mathbf{B} = \mu \mu_0 \mathbf{H}$	$\mathbf{B} = \mu \mathbf{H}$
	$\mathbf{j} = \sigma \mathbf{E}$	$\mathbf{j} = \sigma \mathbf{E}$

$\mathbf{D}$ – indukcja elektryczna [C/m²]

$\mathbf{B}$ - indukcja magnetyczna [Wb/m² = Vs/m²]

$\mathbf{j}$ – gęstość prądu [A /m²]

ρ - gęstość ładunku [C /m³]

σ - przewodność [A/Vm]

Działając operatorem div na (28) i kładąc równanie (30) otrzymuje się prawo zachowania ładunku w postaci:

$$\nabla \mathbf{j} = -\frac{\partial \rho}{\partial t} \quad (32)$$

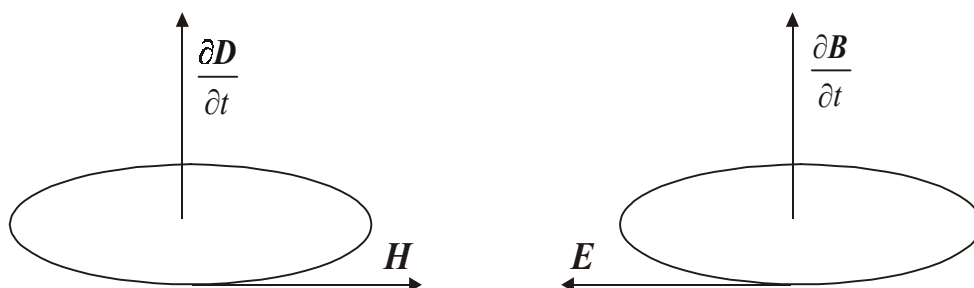
często dołączane do równań (28-31) jako uzupełnienie.

Równania Maxwella w próżni

W ośrodkach bez prądów przewodzenia ($\sigma=0$) oraz bez ładunków ($\rho=0$) równania Maxwella upraszczają się do postaci:

$$\nabla \times \mathbf{H} = \frac{\partial \mathbf{D}}{\partial t} \quad \nabla \times \mathbf{E} = -\frac{\partial \mathbf{B}}{\partial t}$$

$$\nabla \mathbf{D} = 0 \quad \boxed{\nabla \mathbf{B} = 0}$$



Ilustracja równań Maxwella

Pola statyczne – ładunki, magnesy i pętle z prądem są nieruchome - wektory $\mathbf{E}$ i $\mathbf{H}$ niezależne.

$\nabla \mathbf{D} = \rho$		$\nabla \mathbf{B} = 0$	
$\nabla \times \mathbf{E} = 0$		$\nabla \times \mathbf{H} = \mathbf{j}$	

Pola dynamiczne - ładunki, magnesy i pętle z prądem są ruchome - wektory $\mathbf{E}$ i $\mathbf{H}$ sprzężone.

$\nabla \mathbf{D} = \rho$ prawo Gaussa $\nabla \times \mathbf{H} = \mathbf{j} + \frac{\partial \mathbf{D}}{\partial t}$ prawo Ampera	$\nabla \mathbf{B} = 0$ brak ładunków magnetycznych $\nabla \times \mathbf{E} = -\frac{\partial \mathbf{B}}{\partial t}$ prawo Faradaya

Równania Maxwella w postaci całkowej

Można je uzyskać korzystając z twierdzenia Gaussa i Stokesa.

Twierdzenie Gaussa:

$$\int \text{div} \mathbf{H} dV = \oint \mathbf{H} d\mathbf{f} \quad (25)$$

Całka po prawej stronie równania (25) rozciągnięta jest na dowolną powierzchnię zamkniętą obejmującą obszar, po którym całkuje się lewą stronę; $d\mathbf{f}$ jest wektorem normalnym do powierzchni całkowania. Prawa strona równania (25) nazywana jest **strumieniem wektora pola** przez powierzchnię, po której się całkuje.

Ponieważ $\text{div} \mathbf{H} = 0$ to

$$\oint \mathbf{H} d\mathbf{f} = 0$$

Jak widać strumieniem wektora pola przez dowolną powierzchnię zamkniętą jest równy 0.

Twierdzenie Stokesa:

$$\int \text{rot} \mathbf{E} df = \oint \mathbf{E} d\mathbf{l} \quad , \quad (26)$$

gdzie po prawej stronie wykonujemy całkę po zamkniętym konturze, na którym rozpięta jest powierzchnia, po której całkujemy lewą stronę; $d\mathbf{l}$ element konturu całkowania.

Całkując obustronnie pierwsze równanie Maxwella i korzystając z powyższych twierdzeń otrzymujemy pierwszą parę równań Maxwella w postaci całkowej:

$\oint \mathbf{E} d\mathbf{l} = -\frac{1}{c} \frac{\partial}{\partial t} \int \mathbf{H} df$	(27)
$\underbrace{\oint \mathbf{E} d\mathbf{l}}$ cyrkulacja pola E czyli SEM konturu $\underbrace{-\frac{1}{c} \frac{\partial}{\partial t} \int \mathbf{H} df}$ strumień pola magnetycznego przez powierzchnię rozpiętą na konturze	

Równania falowe

Z równań Maxwella dla próżni można bezpośrednio uzyskać jednorodne równania falowe.

$$\nabla \times \mathbf{H} = \varepsilon \varepsilon_0 \frac{\partial \mathbf{E}}{\partial t} \quad (1)$$

$$\nabla \times \mathbf{E} = -\mu \mu_0 \frac{\partial \mathbf{H}}{\partial t} \quad (2)$$

Aby uzyskać jednorodne równanie falowe dla wektora $\mathbf{H}$ (lub dowolnej jego składowej), równanie (1) poddaje się obustronnej rotacji i w miejsce $\nabla \times \mathbf{E}$ wstawia prawą stronę równania (2).

Wedle ww. procedury uzyskamy:

$$\nabla \times (\nabla \times \mathbf{H}) = (\nabla \times \mathbf{E}) = \varepsilon \varepsilon_0 \frac{\partial}{\partial t} \left(-\mu \mu_0 \frac{\partial \mathbf{H}}{\partial t} \right) = -\varepsilon \varepsilon_0 \mu \mu_0 \frac{\partial^2 \mathbf{H}}{\partial t^2} \quad (3)$$

Na mocy tożsamości różniczkowej: $\nabla \times (\nabla \times \mathbf{A}) = \nabla (\nabla \cdot \mathbf{A}) - \nabla^2 \mathbf{A}$

mamy:

$$\nabla (\nabla \cdot \mathbf{H}) - \nabla^2 \mathbf{H} = -\varepsilon \varepsilon_0 \mu \mu_0 \frac{\partial^2 \mathbf{H}}{\partial t^2} \quad (4)$$

Ponieważ z równań Maxwella wiadomo, że $\nabla \mathbf{H} = 0$ to po przeniesieniu wyrazów na lewą stronę otrzymamy jednorodne równanie falowe dla wektora $\mathbf{H}$:

$$\nabla^2 \mathbf{H} - \varepsilon \varepsilon_0 \mu \mu_0 \frac{\partial^2 \mathbf{H}}{\partial t^2} = 0 \quad (5)$$

W analogiczny sposób można uzyskać jednorodne równanie falowe dla wektora $\mathbf{E}$ (lub dowolnej jego składowej).

Występujący w równaniu (5) operator $\square = \nabla^2 - \varepsilon \varepsilon_0 \mu \mu_0 \frac{\partial^2}{\partial t^2}$ ma postać operatora d'Alemberta.

(Ściślej, równanie (5) w trzech wymiarach nazywane jest równaniem Kirchhoffa, w dwóch wymiarach Poissone'a, zaś w jednym właśnie d'Alemberta.)

Po wprowadzeniu operatora, równanie falowe (5) przyjmuje postać:

$$\square \mathbf{H} = 0 \quad (6)$$

Rozwiązaniem równania (6) jest superpozycja fal propagujących się przeciwbieżnie.

Twierdzenie Poyntinga

Zapiszmy równania Maxwella dla wektorów rzeczywistych w przestrzeni jednorodnej i izotropowej:

$$\nabla \times \mathbf{H} = \mathbf{J} + \frac{\partial \mathbf{D}}{\partial t} \quad (1)$$

$$\nabla \times \mathbf{E} = -\frac{\partial \mathbf{B}}{\partial t} \quad (2)$$

Mnożąc obustronnie i skalarnie równanie (1) przez $\mathbf{E}$, i analogicznie drugie przez $\mathbf{H}$ oraz dodając je stronami otrzymamy:

$$\mathbf{E} \cdot \nabla \times \mathbf{H} - \mathbf{H} \cdot \nabla \times \mathbf{E} = \mathbf{E} \cdot \frac{\partial \mathbf{D}}{\partial t} + \mathbf{H} \cdot \frac{\partial \mathbf{B}}{\partial t} + \mathbf{J} \cdot \mathbf{E} \quad (3)$$

Korzystając z tożsamości różniczkowej: $\mathbf{A} \cdot \nabla \times \mathbf{B} - \mathbf{B} \cdot \nabla \times \mathbf{A} = -\nabla \cdot (\mathbf{A} \times \mathbf{B})$ otrzymamy:

$$\nabla \cdot (\mathbf{E} \times \mathbf{H}) + \mathbf{J} \cdot \mathbf{E} + \mathbf{E} \cdot \frac{\partial \mathbf{D}}{\partial t} + \mathbf{H} \cdot \frac{\partial \mathbf{B}}{\partial t} = 0, \quad (4)$$

gdzie składnik $\mathbf{J} \cdot \mathbf{E}$ wyraża gęstość mocy strat.

W ośrodku z ładunkami swobodnymi gęstość prądu $\mathbf{J}$ należy uzupełnić o gęstość prądów konwekcyjnych, wynikających z ruchu ładunków, wtedy:

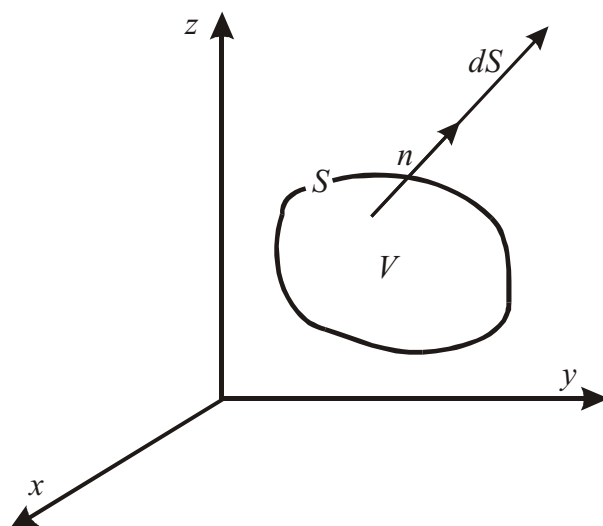
$$\nabla(\mathbf{E} \times \mathbf{H}) + \mathbf{E}(\sigma \mathbf{E} + \rho \mathbf{v}) + \mathbf{E} \frac{\partial \mathbf{D}}{\partial t} + \mathbf{H} \frac{\partial \mathbf{B}}{\partial t} = 0$$

Wektor $\mathbf{E} \times \mathbf{H}$ nazywa się wektorem Poyntinga.

Równanie (4) fizycznie oznacza, że w każdym punkcie przestrzeni suma rozbieżności wektora Poyntinga, gęstości mocy strat oraz szybkości zmian gęstości energii elektrycznej i magnetycznej jest zawsze równa 0. Twierdzenie to wyraża więc bilans mocy w każdym punkcie przestrzeni. Wynika z niego, że sumaryczna energia pola zawarta w całej przestrzeni jest zerowa.

Zależność ta powinna też obowiązywać w przypadku obszaru zamkniętego.

Rozpatrzmy pewną objętość przestrzeni ograniczoną dodatnio zorientowaną powierzchnią S .



Ponieważ bilans energii musi być spełniony dla każdego punktu tej objętości to całkując po niej wyrażenie (4) otrzymamy bilans energii dla całej objętości. Po zastosowaniu do pierwszego składnika twierdzenia Gaussa otrzymamy:

$$\oint_S (\mathbf{E} \times \mathbf{H}) dS + \iiint_V \mathbf{E} \mathbf{J} dV + \iiint_V \left(\mathbf{H} \frac{\partial \mathbf{B}}{\partial t} + \mathbf{E} \frac{\partial \mathbf{D}}{\partial t} \right) dV = 0 \quad (5)$$

Wyrażenie podcałkowe pierwszego składnika opisuje gęstość powierzchniową strumienia mocy $\boldsymbol{\rho} = \mathbf{E} \times \mathbf{H}$, drugi ilość mocy wydzielanej w objętości V , zaś trzeci zmianę energii elektromagnetycznej zmagazynowanej w objętości V w jednostce czasu. Z analizy wektorowej wynika, że wektor Poyntinga jest w rozpatrywanym ośrodku zawsze skierowany w zgodzie z kierunkiem propagacji fali.

Płaskie fale elektromagnetyczne

Rozpatrzmy jednorodne równanie falowe w postaci:

$$\nabla^2 A - \varepsilon\mu \frac{\partial^2 A}{\partial t^2} = 0 \quad \text{gdzie } A = E_x, E_y, E_z, H_x, H_y, H_z \quad (1)$$

Jeśli

$$\frac{\partial}{\partial x} = \frac{\partial}{\partial y} = 0 \quad \wedge \quad \frac{\partial}{\partial z} \neq 0 \quad (2)$$

To równanie (1) upraszcza się do tzw. równania fali płaskiej:

$$\begin{aligned} \nabla^2 E_x - \varepsilon\mu \frac{\partial^2 E_x}{\partial t^2} &= 0 \\ \nabla^2 H_x - \varepsilon\mu \frac{\partial^2 H_x}{\partial t^2} &= 0 \end{aligned} \quad (3)$$

Rozwiązaniem tego równania jest superpozycja dwóch funkcji (dwukrotnie różniczkowalnych) reprezentujących dwie fale o dowolnej zależności od czasu rozchodzące się w przeciwnych kierunkach, równoległe do osi z, z prędkością:

$$v = \frac{1}{\sqrt{\varepsilon\mu}}$$

$$\begin{aligned} E_x(z, t) &= F_1(z - vt) + F_2(z + vt) \\ H_x(z, t) &= F_3(z - vt) + F_4(z + vt) \end{aligned} \quad (4)$$

Uogólniając powyższe wyrażenia na dowolną składową otrzymamy:

$$\begin{aligned} E(z, t) &= F_1(z - vt) + F_2(z + vt) \\ H(z, t) &= F_3(z - vt) + F_4(z + vt) \end{aligned} \quad (5)$$

Uwzględniając (2) oraz drugą parę równań Maxwella dla próżni ($\nabla \varepsilon E = 0$, $\nabla \mu H = 0$) otrzymujemy:

$$\frac{\partial E_z}{\partial z} = 0 \quad \frac{\partial H_z}{\partial z} = 0, \quad (6)$$

z czego od razu wynika, że $E_z = \text{const}$ i $H_z = \text{const}$, co oznacza brak zmian w kierunku z.

Tego typu fala nazywa się falą poprzeczną (TEM – ang. transversal electric and magnetic).

Przedstawmy każdą składową zmienną poprzez pewną funkcję czasu i drogi:

$$\begin{aligned}E_x &= f_1(z - vt) \\E_y &= f_2(z - vt) \\H_x &= f_3(z - vt) \\H_y &= f_4(z - vt)\end{aligned}\tag{7}$$

Z I równania Maxwella dla próżni ($\text{rot}\mathbf{H} = \frac{\partial \mathbf{D}}{\partial t}$) otrzymamy m. in. równanie dla składowych x :

$$\text{rot}H_x = \frac{\partial H_z}{\partial y} - \frac{\partial H_y}{\partial z} = \varepsilon \frac{\partial E_x}{\partial t}\tag{8}$$

Na mocy (6), (7) i (8) otrzymamy:

$$\frac{\partial f_4(z - vt)}{\partial z} = -\varepsilon \frac{\partial f_1(z - vt)}{\partial t}$$

i dalej

$$f_4' \frac{\partial(z - vt)}{\partial z} = -\varepsilon f_1' \frac{\partial(z - vt)}{\partial t}$$

ostatecznie

$$f_4' = \varepsilon v f_1' .\tag{9}$$

Ponieważ

$$v = \frac{1}{\sqrt{\varepsilon\mu}} \text{ to } f_4' = \sqrt{\frac{\varepsilon}{\mu}} f_1' .\tag{10}$$

Po obustronnym całkowaniu (10) i pominięciu stałych całkowania otrzymuje się:

$$f_4 = \sqrt{\frac{\varepsilon}{\mu}} f_1\tag{11}$$

Z podstawienia (7) widać, że odpowiada to zależności:

$$H_y = \sqrt{\frac{\varepsilon}{\mu}} E_x\tag{12}$$

Postępując analogicznie dla składowych y $\left(\text{rot}H_y = \frac{\partial H_z}{\partial x} - \frac{\partial H_x}{\partial z} = \varepsilon \frac{\partial E_y}{\partial t} \right)$

otrzymamy:

$$H_x = -\sqrt{\frac{\varepsilon}{\mu}} E_y \quad (13)$$

W skrócie można zapisać:

$$\left| \frac{H_{\perp}}{E_{\perp}} \right| = \sqrt{\frac{\varepsilon}{\mu}} = Y \quad (14)$$

Wielkość Y nazwano admitancją falową lub właściwą bezstratnego ośrodka, w którym rozchodzi się fala. Odwrotnością tej wielkości jest charakterystyczna impedancja falowa ośrodka (rzadziej nazywana impedancją właściwą). Dla próżni, gdzie zmierzono:

$$\mu_0 = 4\pi \cdot 10^{-7} \approx 1,257 \cdot 10^{-6} \left[\frac{\text{H}}{\text{m}} \right] \quad \varepsilon_0 = \frac{10^{-9}}{36\pi} \approx 8,854 \cdot 10^{-12} \left[\frac{\text{F}}{\text{m}} \right]$$

przyjmuje ona wartość: $Z_0 = \sqrt{\frac{\mu_0}{\varepsilon_0}} = 120\pi \approx 377 \, \Omega$.

Równania Maxwella w postaci zespolonej

Jeśli w równaniach Maxwella wektory $\mathbf{E}$, $\mathbf{D}$, $\mathbf{H}$ i $\mathbf{B}$ są harmonicznymi funkcjami czasu. Wtedy dają się one przedstawić następująco:

$$\mathbf{E} = \text{Re} \left[\hat{\mathbf{E}} e^{j\omega t} \right] = \frac{1}{2} \left[\hat{\mathbf{E}} e^{j\omega t} + \hat{\mathbf{E}}^* e^{-j\omega t} \right]$$

$$\mathbf{H} = \text{Re} \left[\hat{\mathbf{H}} e^{j\omega t} \right] = \frac{1}{2} \left[\hat{\mathbf{H}} e^{j\omega t} + \hat{\mathbf{H}}^* e^{-j\omega t} \right]$$

$$\mathbf{D} = \text{Re} \left[\hat{\mathbf{D}} e^{j\omega t} \right] = \frac{1}{2} \left[\hat{\mathbf{D}} e^{j\omega t} + \hat{\mathbf{D}}^* e^{-j\omega t} \right]$$

$$\mathbf{B} = \text{Re} \left[\hat{\mathbf{B}} e^{j\omega t} \right] = \frac{1}{2} \left[\hat{\mathbf{B}} e^{j\omega t} + \hat{\mathbf{B}}^* e^{-j\omega t} \right]$$

Podstawiając te reprezentacje do równań Maxwella dla ośrodka liniowego, izotropowego i jednorodnego, pamiętając, że

$$\frac{\partial \hat{\mathbf{A}}}{\partial t} = \frac{1}{2} j\omega \left[\hat{\mathbf{A}} e^{j\omega t} + \hat{\mathbf{A}}^* e^{-j\omega t} \right],$$

otrzymamy (np. dla równania $\nabla \times \mathbf{E} = -\frac{\partial \mathbf{B}}{\partial t}$):

$$\nabla \times \left(\hat{\mathbf{E}} e^{j\omega t} + \hat{\mathbf{E}}^* e^{-j\omega t} \right) = -j\omega \left(\hat{\mathbf{B}} e^{j\omega t} + \hat{\mathbf{B}}^* e^{-j\omega t} \right)$$

Równanie to musi być spełnione w dowolnej chwili czasu.

W szczególności dla $t = 0$ mamy:

$$\nabla \times \left(\hat{\mathbf{E}} + \hat{\mathbf{E}}^* \right) = -j\omega \left(\hat{\mathbf{B}} + \hat{\mathbf{B}}^* \right),$$

$$\text{a dla } t = \frac{T}{4} = \frac{\pi}{2\omega}$$

$$\nabla \times \left(\hat{\mathbf{E}} - \hat{\mathbf{E}}^* \right) = -j\omega \left(\hat{\mathbf{B}} - \hat{\mathbf{B}}^* \right)$$

Dodając stronami równania dla $t = 0$ i $t = \frac{\pi}{2\omega}$ otrzymujemy:

$$\nabla \times \hat{\mathbf{E}} = -j\omega \hat{\mathbf{B}} \quad (1)$$

i analogicznie

$$\nabla \times \hat{\mathbf{H}} = j\omega \hat{\mathbf{D}} + \hat{\mathbf{J}}. \quad (2)$$

Oczywiście w dalszym ciągu:

$$\nabla \hat{\mathbf{D}} = \rho, \quad \hat{\mathbf{D}} = \varepsilon \hat{\mathbf{E}} \quad (3)$$

$$\nabla \hat{\mathbf{B}} = 0, \quad \hat{\mathbf{B}} = \mu \hat{\mathbf{H}} \quad (4)$$

$$\hat{\mathbf{J}} = \sigma \hat{\mathbf{E}}$$

Równania Maxwella w postaci zespolonej nie zależą od czasu

Równania falowe w postaci zespolonej

Wychodząc z równań Maxwella w postaci zespolonej:

$$\nabla \times \hat{\mathbf{E}} = -j\omega \hat{\mathbf{B}} \quad (1)$$

$$\nabla \times \hat{\mathbf{H}} = j\omega \hat{\mathbf{D}} + \hat{\mathbf{J}}, \quad \hat{\mathbf{J}} = \sigma \hat{\mathbf{E}} \quad (2)$$

$$\nabla \hat{\mathbf{D}} = \rho, \quad \hat{\mathbf{D}} = \varepsilon \hat{\mathbf{E}} \quad (3)$$

$$\nabla \hat{\mathbf{B}} = 0, \quad \hat{\mathbf{B}} = \mu \hat{\mathbf{H}} \quad (4)$$

Wstawiając (1) do (2) z uwzględnieniem (3) i (4) otrzymamy:

$$\nabla \times \left(\frac{1}{\omega\mu} \nabla \times \hat{\mathbf{E}} \right) = (\sigma + j\omega\varepsilon) \hat{\mathbf{E}},$$

a po przekształceniu: $\nabla \times \nabla \times \hat{\mathbf{E}} = -j\omega\mu(\sigma + j\omega\varepsilon) \hat{\mathbf{E}}$

Wiedząc, że

$\nabla \times \nabla \times \mathbf{A} = \nabla \nabla \cdot \mathbf{A} - \nabla^2 \mathbf{A}$ i przenosząc wyrażenia na lewą stronę otrzymamy:

$$\nabla^2 \hat{\mathbf{E}} - j\omega\mu(\sigma + j\omega\varepsilon) \hat{\mathbf{E}} = 0 \quad (5)$$

analogicznie $\nabla^2 \hat{\mathbf{H}} - j\omega\mu(\sigma + j\omega\varepsilon) \hat{\mathbf{H}} = 0 \quad (6)$

Wielkość stała $\gamma = \sqrt{j\omega\mu(\sigma + j\omega\varepsilon)}$ nazywana jest stałą propagacji. Po jej wprowadzeniu do równań (5) i (6) otrzymamy:

$$\nabla^2 \hat{E} - \gamma^2 \hat{E} = 0 \quad (7)$$

$$\nabla^2 \hat{H} - \gamma^2 \hat{H} = 0 \quad (8)$$

Są to eliptyczne równania różniczkowe cząstkowe noszące nazwę równań **Helmholtza**.

Część rzeczywista stałej propagacji nosi nazwę stałej tłumienia i wyraża zależnie od znaku logarytm naturalny ze względnego wzrostu lub zmniejszenia wartości amplitudy fali na jednostkę przebytej drogi przedstawiany w dB/m.

$$\operatorname{Re}(\gamma) = \alpha$$

Część urojona stałej propagacji nosi nazwę współczynnika fazowego i wyraża zmianę fazy fali na jednostkę przebytej drogi wyrażoną w rad/m.

$$\operatorname{Im}(\gamma) = \beta = \frac{2\pi}{\lambda}$$

Rozwiązanie równań Helmholtza poprzez rozdzielenie zmiennych

$$\nabla^2 \hat{E} - \gamma^2 \hat{E} = 0 \quad (1)$$

Załóżmy, że istnieje elementarne rozwiązanie tego równania w postaci (taka metoda będzie również zastosowana przy rozwiązywaniu równania falowego w przestrzeni ograniczonej):

$$\hat{E} = X(x) \cdot Y(y) \cdot Z(z) \quad (2)$$

Po podstawieniu rozwiązania (2) do (1) uzyskamy:

$$X'' \cdot Y \cdot Z + X \cdot Y'' \cdot Z + X \cdot Y \cdot Z'' - \gamma^2 X \cdot Y \cdot Z = 0 \quad (3)$$

Dzieląc obustronnie przez XYZ otrzymuje się:

$$\frac{X''}{X} + \frac{Y''}{Y} + \frac{Z''}{Z} - \gamma^2 = 0 \quad (4)$$

Po przeniesieniu jednego ze zmiennych wyrazów (np. $\frac{Y''}{Y}$) na prawą stronę równania uzyskujemy równanie, w którym lewa strona nie zależy od prawej:

$$\frac{X''}{X} + \frac{Z''}{Z} - \gamma^2 = -\frac{Y''}{Y}, \quad (5)$$

ponieważ strony są równe, każda z nich musi być równa pewnej stałej. Stąd (5) można rozdzielić na trzy niezależne równania:

$$\begin{aligned}X'' - \gamma_x^2 X &= 0 \\Y'' - \gamma_y^2 Y &= 0 \\Z'' - \gamma_z^2 Z &= 0\end{aligned}\tag{6}$$

Stałe $\gamma_x, \gamma_y, \gamma_z$ noszą nazwę stałych rozdziału, a z (4) wprost wynika, że suma ich kwadratów wynosi γ .

Elementarne rozwiązania równań (6) mają postać:

$$A(x) = A_a e^{\gamma_x x} + A_b e^{-\gamma_x x},\tag{7}$$

przy czym, w przypadku ogólnym, stałe całkowania A_a i A_b oraz stałe rozdziału mogą być zespolone, a jeśli tak to równanie (7) przedstawia superpozycję dwóch fal rozchodzących się w przeciwnych kierunkach.

Dyskusja

Zakładając w (7) $A_a = 0$ i $\text{Im}(\gamma_x) = 0$ otrzymamy w punktach $x = 0$ i $x = 1$ związek:

$$|A(1)| = |A(0)| e^{-\alpha_x}\tag{8}$$

stąd:

$$\alpha_x = \ln \frac{|A_b(0)|}{|A_b(1)|}\tag{9}$$

Zakładając natomiast w (7) $A_a = 0$ i $\text{Re}(\gamma_x) = 0$ otrzymamy:

$$|A_b| e^{j\varphi_1} = |A_b| e^{j\varphi_0} e^{j\beta_x}\tag{10}$$

stąd wynika, że:

$$\beta_x = |\varphi_1 - \varphi_0|\tag{11}$$

Prędkość zmian fazy każdej z fal w wyrażeniu (7) w danym kierunku nosi nazwę prędkości fazowej. Wynosi ona:

$$v_f = \frac{\omega}{\beta}.\tag{12}$$

Jeśli weźmiemy pod uwagę ich superpozycję to obwódnia zdudnienia przemieszczać się będzie z prędkością

$$v_g = \frac{d\omega}{d\beta}. \quad (13)$$

Prędkość ta nazywana jest prędkością grupową (prędkością rozchodzenia się grupy fal).

Pomiędzy tymi prędkościami zachodzi związek:

$$v_g = \frac{v_f}{1 - \frac{\omega}{v_f} \frac{dv_f}{d\omega}} \quad (14)$$

Przy braku dyspersji (tzn. braku zależności parametrów ośrodka od częstotliwości)

$$\frac{dv_f}{d\omega} = 0 \quad v_g = v_f$$

W ośrodkach o **dyspersji normalnej** prędkość fazowa maleje ze wzrostem częstotliwości:

$$\frac{dv_f}{d\omega} < 0, \quad v_g < v_f$$

zaś w ośrodkach o **dyspersji anormalnej** prędkość fazowa rośnie ze wzrostem częstotliwości:

$$\frac{dv_f}{d\omega} > 0, \quad v_g > v_f$$

Niejednorodne równanie falowe

$$\square \Psi = -4\pi \rho(r, t) \quad (1)$$

$\rho(r, t)$ gęstość ładunków.

Pojedyncze zaburzenie rozchodzące się od źródła w r_0 opisywane jest funkcją Greena $G(r, t; r_0, t_0)$.

Rozwiązanie niejednorodnego równania falowego dla dowolnej gęstości ρ jest superpozycją wszystkich funkcji Greena odpowiadającym wszystkim zaburzeniom. Jeśli fale znikają w nieskończoności (nie są zaburzone) to rozchodząca się fala kulista (funkcja Greena) jest rozwiązaniem równania:

$$\square G = -4\pi \delta(r - r_0) \delta(t - t_0) \quad (2)$$

Rozwiązanie to ma postać:

$$G(r, t, r_0, t_0) = \frac{\delta(t - t_0 - R/c)}{R} \quad (3)$$

gdzie $R = |r - r_0|$

Oznacza to, że fala kulista osiągnie odległość R tylko w ściśle wyznaczonej chwili $t = t_0 + R/c$. Współczynnik $1/R$ jest wynikiem obowiązywania zasady zachowania energii: powierzchnia kuli rośnie jak R^2 zatem gęstość energii maleje jak $1/R^2$, co oznacza, że amplituda musi maleć jak $1/R$ (gęstość energii jest proporcjonalna do kwadratu amplitudy).

Wobec powyższego, korzystając z zasady superpozycji można podać rozwiązanie niejednorodnego równania falowego w postaci:

$$\Psi(r, t) = \int d^3r_0 \cdot \int G(r, t, r_0, t_0) \rho(r_0, t_0) dt_0 \quad (4)$$

W równaniach dla potencjału elektrycznego ϕ wielkość ρ jest gęstością ładunku, zaś dla potencjału wektorowego A ρ jest odpowiednią składową $\mathbf{j}/c$.

Równania falowe dla potencjałów mają postać:

wektorowy $\square A = -4\pi \frac{\mathbf{j}}{c}$

skalarny $\square \phi = -4\pi \rho$

Ponieważ potencjały muszą znikać w nieskończoności musi być spełnione równanie (3) stąd:

$$A(r, t) = \int \frac{\frac{j}{c} \left(r_0, t - \frac{R}{c} \right)}{R} d(r_0)^3 \quad \phi(r, t) = \int \frac{\rho \left(r_0, t - \frac{R}{c} \right)}{R} d(r_0)^3$$

W jaki sposób przewodnik z prądem $j(r)\exp(-i\omega t)$ emituje fale?

Dla takiego układu potencjał wektorowy opisany jest wzorem:

$$A(r)e^{j\omega t} = \frac{1}{c} \int j(r_0) \frac{e^{ikR}}{R} d(r_0)^3 \quad k = \frac{\omega}{c} \quad (A)$$

Podobna zależność obowiązuje dla potencjału skalarnego.

Rozwińmy R i $1/R$ w szereg Taylora:

$$R = |r - r_0| = |r| - \frac{rr_0}{|r|} + \dots \quad \frac{1}{R} = \frac{1}{|r|} + \frac{rr_0}{r^3} + \frac{\frac{3(rr_0)^2}{r^5} - \frac{r_0^2}{r^3}}{2} + \dots$$

Rozwinięcia te są słuszne dla odległości r znacznie większych niż rozmiary anteny (r_0).

Po wstawieniu rozwinięć do równania (A) otrzymuje się:

$$cAe^{j\omega t} = \frac{e^{ikr}}{r} \left[\int j(r_0) d(r_0)^3 + \left(\frac{1}{r} - jk \right) \int j(r_0) e_r r_0 d(r_0)^3 \right] = \\ = ce^{j\omega t} (A_0 + A_1 + \dots)$$

gdzie e_r jest jednostkowym ładunkiem elektrycznym, zaś

$$rcA_0 = e^{j(kr - \omega t)} \int j(r_0) d(r_0)^3$$

jest tzw. wkładem dipolowym

Po skorzystaniu z równania ciągłości:

$$\text{div}(j) + \frac{d}{dt} \rho = \text{div}(j) - j\omega \rho = 0$$

całka ta daje się zapisać w formie:

$$-\int r_0 \text{div}(j(r_0)) d(r_0)^3 = -jck \int r_0 \rho(r_0) d(r_0)^3 = -jckP$$

P – elektryczny moment dipolowy

Jeśli ładunki (dodatnie i ujemne) rozmieszczone są w odległości d to:

$$|P| = qd$$

$$A_0 + A_1 + \dots = -jkP \frac{e^{j(kr - \omega t)}}{r} + \dots$$

Antena zatem promieniuje fale kuliste o amplitudzie proporcjonalnej do wektora P . Wysyłane promieniowanie jest więc jednakowo intensywne we wszystkich kierunkach.

Znajomość A_0 pozwala wyznaczyć pola elektryczne i magnetyczne:

$$B_0 = \text{rot} A_0 = k^2 (e_r \times P) \left(1 - \frac{1}{ikr} \right) \frac{e^{i(kr - \omega t)}}{r}$$

$$E_0 = \frac{j}{k} \text{rot} B_0 = \left[k^2 (e_r \times P) \times \frac{e_r}{r} + [3e_r - P] (r^{-3} - ikr^{-2}) \right] e^{i(kr - \omega t)}$$

Wyrażenia te znacznie się upraszczają dla strefy dalekiej tzn. wtedy gdy odległość od anteny jest znacznie większa od długości fali.

Dla dużych kr otrzymuje się:

$$B_0 = k^2 e_r \times \frac{Pe^{i(kr - \omega t)}}{r} \qquad E_0 = B_0 \times e_r$$

Relatywistyczna postać równań Maxwella

Wprowadźmy do równania d'Alemberta zmienną czasową jako czwarty wymiar:

$$x_4 = ict.$$

Jeśli skorzystamy z definicji iloczynu skalarnego: $x \cdot y = \sum_m x_m y_m$,

to operator d'Alemberta: $\square = \nabla^2 - c^{-2} \frac{\partial^2}{\partial t^2}$ przyjmie postać:

$$\square = \sum_m \frac{\partial^2}{\partial x_m^2}$$

Wprowadźmy asymetryczny tensor f_{mn} :

$$f_{mn} = \frac{\partial A_m}{\partial x_n} - \frac{\partial A_n}{\partial x_m} \quad m, n = 1 \dots 4$$

Z własności asymetrii $f_{mn} = -f_{nm}$ wynika, że tensor ten zawiera 6 niezależnych składowych. W odniesieniu do pól elektrycznego i magnetycznego będziemy mieli:

$$f_{mn} = \begin{bmatrix} 0 & B_z & -B_y & -iE_x \\ -B_z & 0 & B_x & -iE_y \\ B_y & -B_x & 0 & -iE_z \\ iE_x & iE_y & iE_z & 0 \end{bmatrix}$$

Korzystając z tego tensora można z równania d'Alemberta wrócić do równań Maxwella. Przyjmą one postać:

$$\frac{\partial f_{mn}}{\partial x_l} + \frac{\partial f_{nl}}{\partial x_m} + \frac{\partial f_{lm}}{\partial x_n} = 0,$$

lub korzystając z konwencji Einsteina (sumowania po powtarzających się wskaźnikach):

$$\partial_m f^{mn} = 0$$

W wyniku poddania tensora $\hat{f}$ przekształceniu Lorentza, czyli wymnożeniu jego składowych przez macierz:

$$L = \begin{bmatrix} \gamma & 0 & 0 & i\gamma v/c \\ 0 & 1 & 0 & 0 \\ 0 & 0 & 1 & 0 \\ -i\gamma v/c & 0 & 0 & \gamma \end{bmatrix}$$

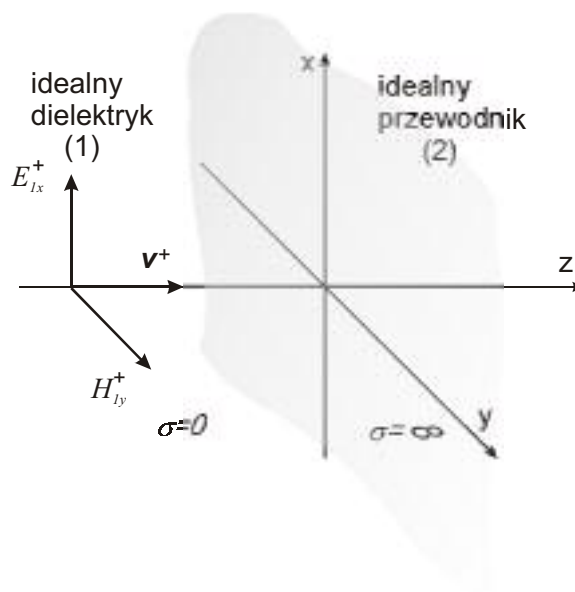
(co formalnie oznacza przetransformowanie składowych tego tensora do dowolnego układu współrzędnych) pole elektryczne transformuje się w magnetyczne i vice versa.

W ujęciu relatywistycznym, to czy mówimy o polu elektrycznym czy też magnetycznym zależy tylko od wyboru układu odniesienia.

Fakt ten jest powodem wielu analogii dostrzeganych przy analizie tych pól.

Odbicie fal elektromagnetycznych

Fala pada prostopadłe na idealny przewodnik



Dla wektora natężenia pola elektrycznego:

$$E_{1x}^+ = E_{1x}^+(0) e^{j\omega\sqrt{\mu_1\epsilon_1}z}$$

$$H_{1y}^+ = \sqrt{\frac{\epsilon_1}{\mu_1}} \cdot E_{1x}^+ = Y_{01} E_{1x}^+$$

Ponieważ fala elektromagnetyczna nie może się propagować w doskonałym przewodniku:

$$E_1(0) = E_1^+(0) + E_1^-(0) = 0$$

to

$$E_1^-(0) = -E_1^+(0) = -E_{1x}^+(0).$$

Stąd równanie fali odbitej ma postać:

$$E_{1x}^- = -E_{1x}^+(0) e^{j\omega\sqrt{\epsilon_1\mu_1}z}$$

Całkowite pole w odległości z od niejednorodności (graniczy rozdziału) będzie superpozycją fali padającej i odbitej:

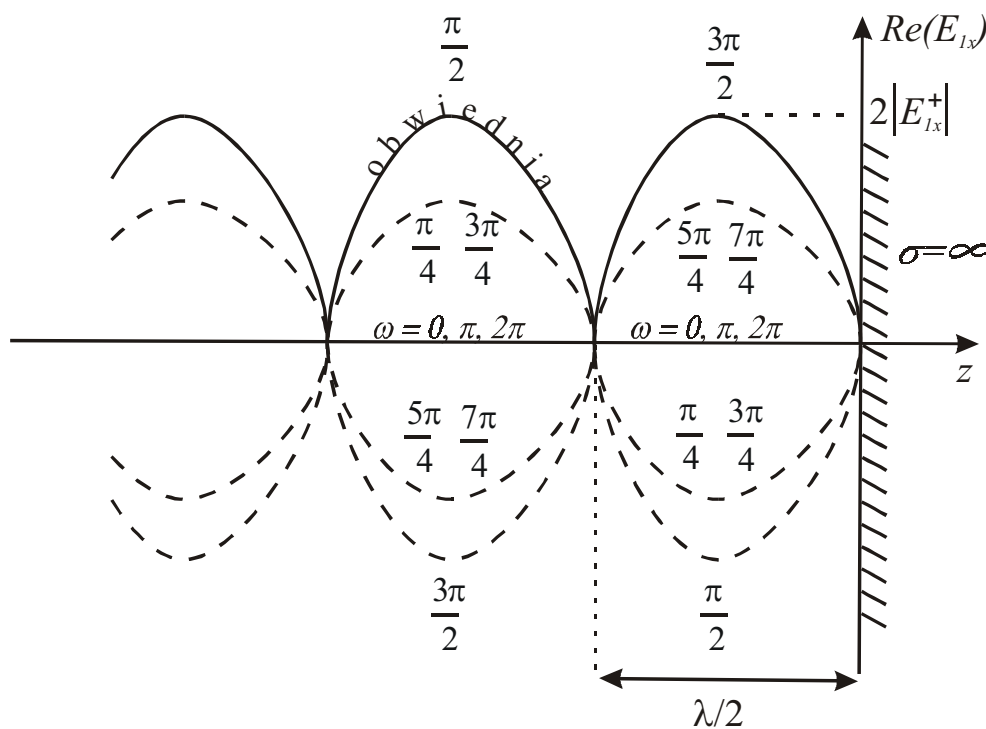
$$E_{1x} = E_{1x}^+ + E_{1x}^- = E_{1x}^+(0) \left(e^{-j\omega\sqrt{\epsilon_1\mu_1}z} - e^{j\omega\sqrt{\epsilon_1\mu_1}z} \right) = -2jE_{1x}^+(0) \sin \left(\omega\sqrt{\epsilon_1\mu_1}z \right).$$

Dla fazy początkowej φ_0 (przy $t = 0$ i $z = 0$):

$$E_{1x} = 2 \left| E_{1x}^+ \right| e^{j \left(\varphi_0 - \frac{\pi}{2} \right)} \sin \left(\omega \sqrt{\varepsilon_1 \mu_1} z \right)$$

Część rzeczywista tego pola ma postać:

$$\text{Re}(E_{1x}) = 2|E_{1x}^+| \sin(\omega\sqrt{\varepsilon_1\mu_1}z) \cos\left(\omega t + \varphi_0 - \frac{\pi}{2}\right)$$



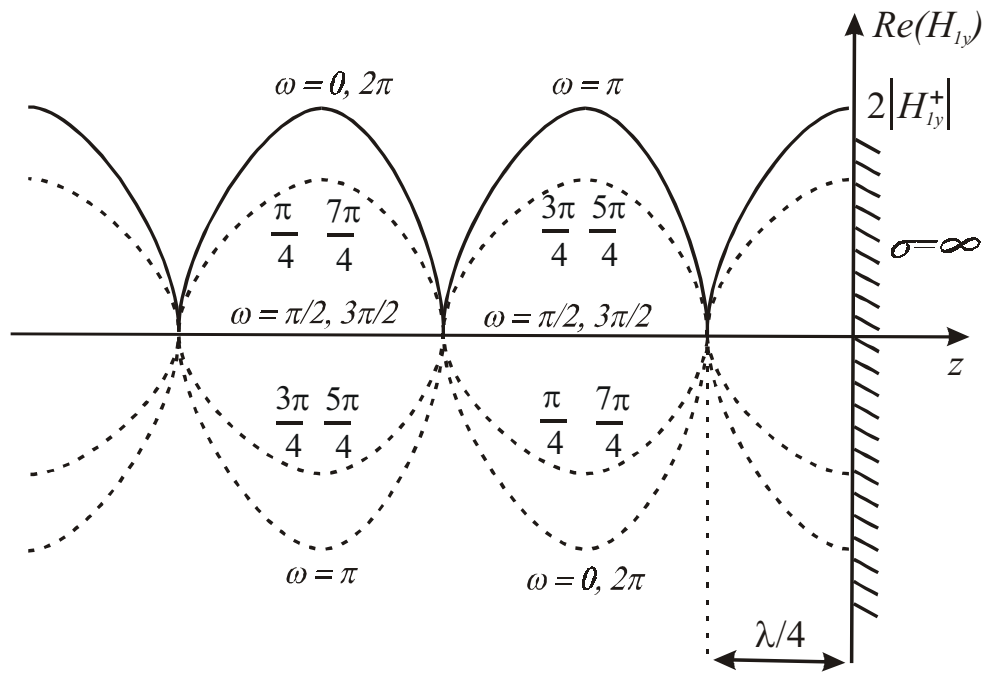
Zupełna fala stojąca ($\varphi_0 = 0$)

Analogicznie dla wektora natężenia pola magnetycznego:

$$H_{1y} = 2Y_{01}E_{1x}^+(0)\cos\omega\sqrt{\varepsilon_1\mu_1}z.$$

Część rzeczywista pola magnetycznego (z uwzględnieniem kąta fazowego) ma postać:

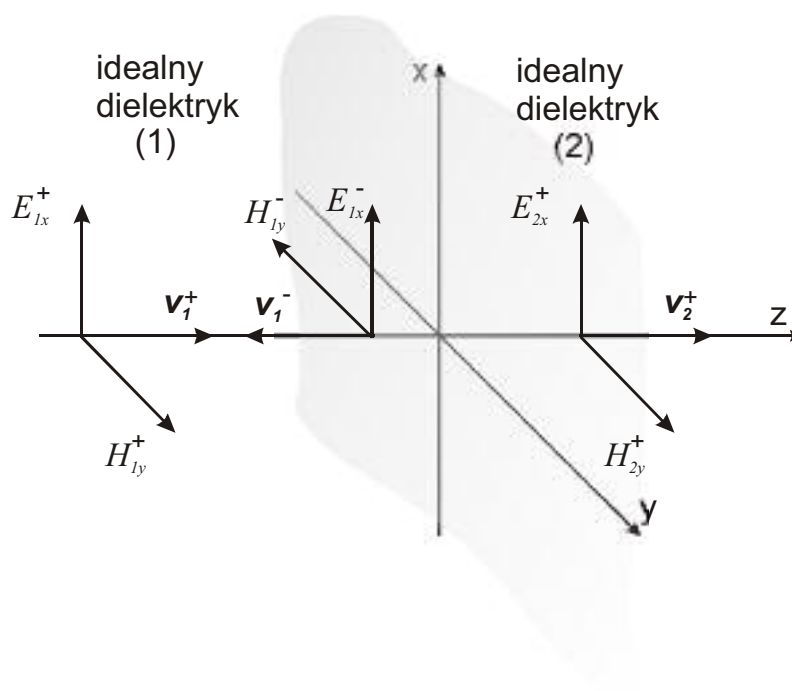
$$H_{1y} = 2Y_{01} \left| E_{1x}^+ \right| \cos \omega \sqrt{\epsilon_1 \mu_1} z \cdot \cos(\omega t + \varphi_0)$$



W wyniku odbicia fali elektromagnetycznej od powierzchni idealnego przewodnika:

- A – pola elektryczne i magnetyczne mają postać zupełnej fali stojącej,
- B – maksima (i minima) fali stojącej pola elektrycznego przesunięte są względem maksimów (i minimów) fali stojącej pola magnetycznego o ćwierć fali,
- C – w każdej chwili i w każdym punkcie wektory $\mathbf{E}$ i $\mathbf{H}$ są przesunięte w fazie o $\pi/2$ – fala taka nie niesie energii (nie jest falą par excellence).

Fala pada prostopadłe na idealny dielektryk.



Def. współczynników odbicia

$$\Gamma_e = \frac{E_{1x}^-}{E_{1x}^+}$$

$$\Gamma_h = \frac{H_{1y}^-}{H_{1y}^+}$$

w każdym punkcie $z \leq 0$ otrzymamy:

$$E_{1x} = E_{1x}^+ + E_{1x}^- = E_{1x}^+ (1 + \Gamma_e) \quad \text{ i } \quad H_{1y} = H_{1y}^+ + H_{1y}^- = H_{1y}^+ (1 + \Gamma_h) \quad (1)$$

a ponieważ $H_{1y}^+ = Y_{01} E_{1x}^+$ i $H_{1y}^- = -Y_{01} E_{1x}^-$ to:

$$H_{1y} = Y_{01} E_{1x}^+ - Y_{01} E_{1x}^- = Y_{01} (E_{1x}^+ - E_{1x}^-) = Y_{01} E_{1x}^+ (1 - \Gamma_e) = H_{1y}^+ (1 - \Gamma_e) = H_{1y}^+ (1 + \Gamma_h) \quad (2)$$

skąd wynika, że:

$$\Gamma_h = -\Gamma_e$$

Oznaczmy $\Gamma_e \equiv \Gamma$.

Warunki brzegowe (ciągłości) w punkcie $z = 0$ (na granicy rozdziału faz).

$$E_{1x}^+(0) + E_{1x}^-(0) = E_{2x}^+(0) \quad (3a)$$

$$H_{1y}^+(0) + H_{1y}^-(0) = H_{2y}^+(0) \quad (3b)$$

Uwzględniając (1) mamy:

$$E_{1x}^+(0)[1 + \Gamma(0)] = E_{2x}^+(0) \quad (4a)$$

$$H_{1y}^+(0)[1 - \Gamma(0)] = H_{2y}^+(0) \Leftrightarrow Y_{01}E_{1x}^+(0)[1 - \Gamma(0)] = Y_{02}E_{2x}^+(0) \quad (4b)$$

Stąd dzieląc stronami (4a) przez (4b) mamy:

$$\frac{1 + \Gamma(0)}{1 - \Gamma(0)} = \frac{Y_{01}}{Y_{02}},$$

stąd:

$$\Gamma(0) = \frac{Z_{02} - Z_{01}}{Z_{02} + Z_{01}}$$

dla $Z_{02} > Z_{01}$ $\Gamma(0) > 0$, a E_{1x}^+ i E_{1x}^- są w fazie

dla $Z_{02} < Z_{01}$ $\Gamma(0) < 0$, a E_{1x}^+ i E_{1x}^- są w przeciwfazie

dla $Y_{01} = Y_{02}$ $\Gamma(0) = 0$

Jeśli ośrodki nie są idealne to:

$$E_{1x}^+ = E_{1x}^+(0)e^{-\gamma_1 z} \quad E_{1x}^- = E_{1x}^-(0)e^{\gamma_1 z}$$

$$\Gamma = \Gamma(0)e^{2\gamma_1 z}, \text{ a dla ośrodków bezstratnych } \Gamma = \Gamma(0)e^{2\beta_1 z}$$

Def. współczynniki przenoszenia pola:

$$T_e = \frac{E_{2x}^+(0)}{E_{1x}^+(0)} \quad T_h = \frac{H_{2x}^+(0)}{H_{1x}^+(0)},$$

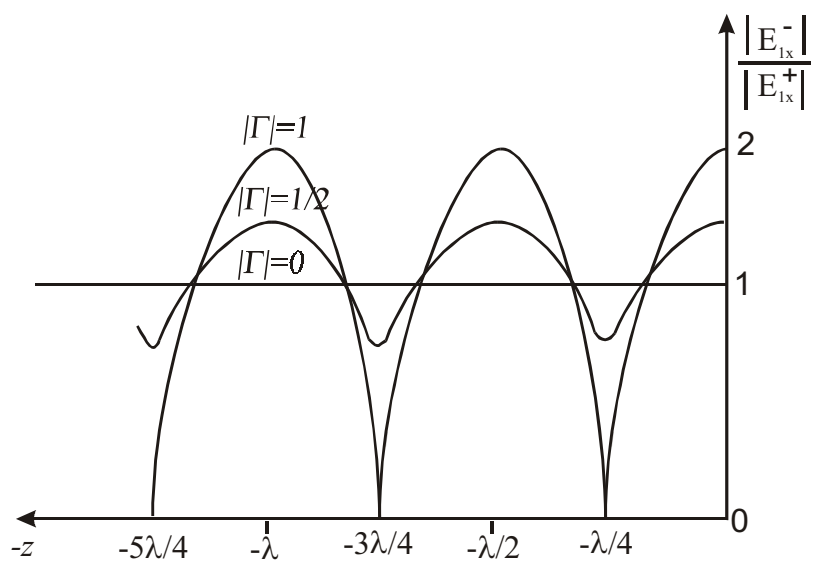
z def. wynika:

$$T_e = 1 + \Gamma(0) \quad T_h = 1 - \Gamma(0)$$

Równania pól dla fali przechodzącej (propagującej się w ośrodku 2) mają więc postać:

$$E_{2x}^+ = T_e E_{1x}^+(0)e^{-\gamma_2 z} \quad H_{2x}^+ = T_h H_{1x}^+(0)e^{-\gamma_2 z}$$

Na skutek tylko częściowego odbicia w ośrodku 1 tworzą się częściowe fale stojące:



Def. Współczynnik fali stojącej $WFS = \frac{|E_{1x}|_{MAX}}{|E_{1x}|_{MIN}} = \frac{|H_{1y}|_{MAX}}{|H_{1y}|_{MIN}} = \frac{1+|\Gamma|}{1-|\Gamma|}$.

Głębokość wnikania fal elektromagnetycznych.

Rozpatrzmy niemagnetyczny przewodnik o impedancji:

$$\text{def. } Z = j\omega\mu,$$

admitancji:

$$\text{def. } Y = \sigma + j\omega\varepsilon,$$

współczynnika propagacji:

$$\text{def. } \gamma = \sqrt{ZY} = \sqrt{j\omega\mu\sigma - \omega^2\varepsilon\mu}$$

oraz impedancji charakterystycznej:

$$Z_0 = \sqrt{\frac{Z}{Y}} \approx \sqrt{\frac{\omega\mu}{2\sigma}} \cdot (1 + j) = \sqrt{\frac{\omega\mu}{\sigma}} e^{j\frac{\pi}{4}}.$$

Jeśli (co jest słuszne dla przewodników) $\sigma \gg \omega\varepsilon$ to następuje bardzo szybkie tłumienie pola w przewodniku.

Grubość warstwy, w której amplituda pola spada do wartości 1/e nazywa się głębokością wnikania δ .

Wg powyższej definicji w odległości δ wartość amplitudy wyniesie:

$$A(\delta) = \frac{1}{e} A(0) = A(0) e^{-\alpha\delta},$$

z czego wynika, że
$$\delta = \frac{1}{\alpha} = \frac{1}{\text{Re}(\gamma)} = \sqrt{\frac{2}{\omega\mu\sigma}}$$

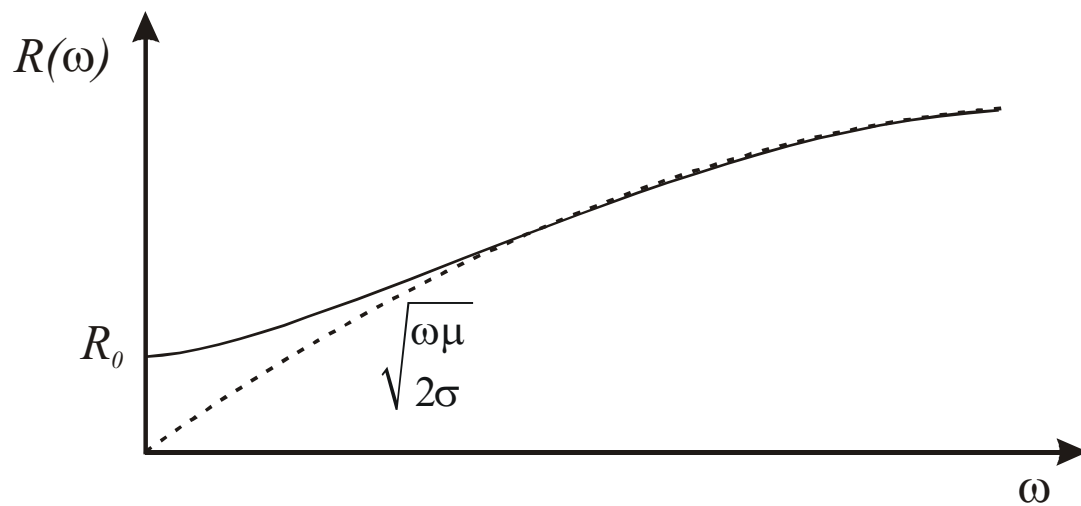
(dla dobrego przewodnika jest to ok. 0,02 mm)

Przytoczone zależności pozwalają także wyznaczyć zależność głębokości wnikania od długości fali:

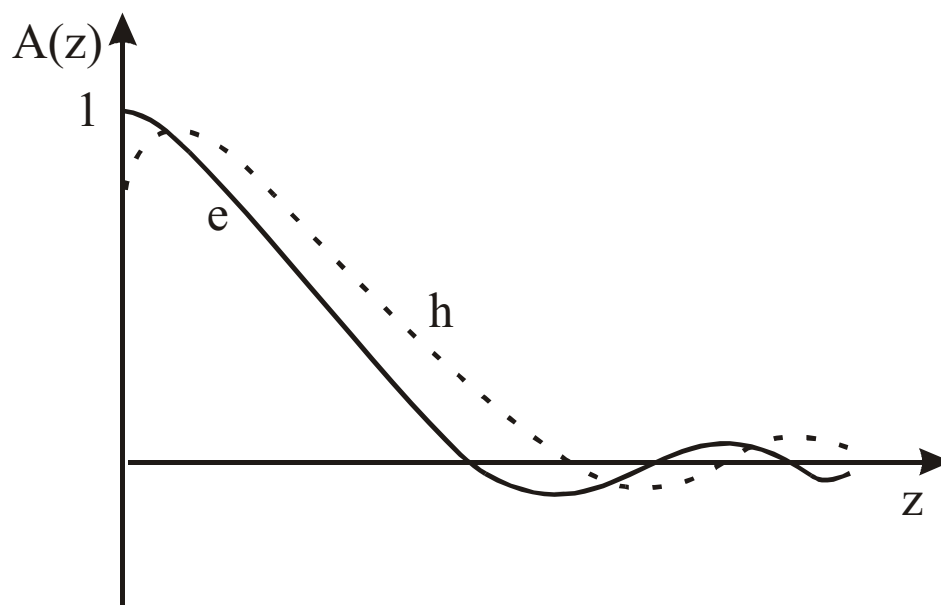
$$\lambda = 2\pi\delta$$

Zjawisko wnikania pola na ograniczoną głębokość nazywa się **zjawiskiem naskórkowym**.

Ze względu na zależność między polem magnetycznym, a gęstością prądu zjawisko to dotyczy także prądów w. cz. płynących w przewodnikach i jest odpowiedzialne za wzrost rezystancji przewodników.



Zmiana rezystancji przewodnika w funkcji częstotliwości.



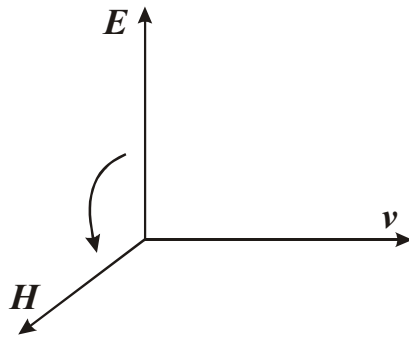
Unormowane charakterystyki wnikania fali elektromagnetycznej w przewodnik

Wybrane własności fal elektromagnetycznych

- Wektory $\mathbf{E}$ i $\mathbf{H}$ oraz ich wszystkie rzuty na osie kartezjańskie w ośrodku jednorodnym i izotropowym spełniają jednorodne równania falowe.
- Prędkość fal elektromagnetycznych wynosi:

$$c = \frac{1}{\sqrt{\epsilon\epsilon_0\mu\mu_0}}; \text{ w próżni } c = \frac{1}{\sqrt{\epsilon_0\mu_0}} \approx 3 \cdot 10^8 \left[\frac{m}{s} \right].$$

- Fale elektromagnetyczne są w próżni falami poprzecznymi tzn. leżą w płaszczyźnie prostopadłej do kierunku propagacji i są do siebie wzajemnie prostopadłe.
- W próżni wektory $\mathbf{E}$, $\mathbf{H}$ oraz $\mathbf{v}$ tworzą prawoskrętny układ wektorów:



- Dla dowolnej fali elektromagnetycznej wektory $\mathbf{E}$ i $\mathbf{H}$ drgają w fazie, ich długości są równe i wynoszą:

$$|\mathbf{E}| = |\mathbf{H}| = \sqrt{\epsilon\epsilon_0} \mathbf{E} = \sqrt{\mu\mu_0} \mathbf{H}$$

- Gęstość energii pola elektromagnetycznego w próżni wynosi:

$$\mathbf{W} = \frac{\epsilon\epsilon_0 \mathbf{E}^2}{2} + \frac{\mu\mu_0 \mathbf{H}^2}{2} = \frac{\sqrt{\epsilon\mu}}{c} \mathbf{E}\mathbf{H}$$

- Gęstość energii fali elektromagnetycznej w próżni wynosi:

$$\mathbf{W} = \epsilon\epsilon_0 \mathbf{A}^2 \sin^2(\omega t - kx), \quad k = \frac{2\pi}{\lambda} = \frac{\omega}{v}$$

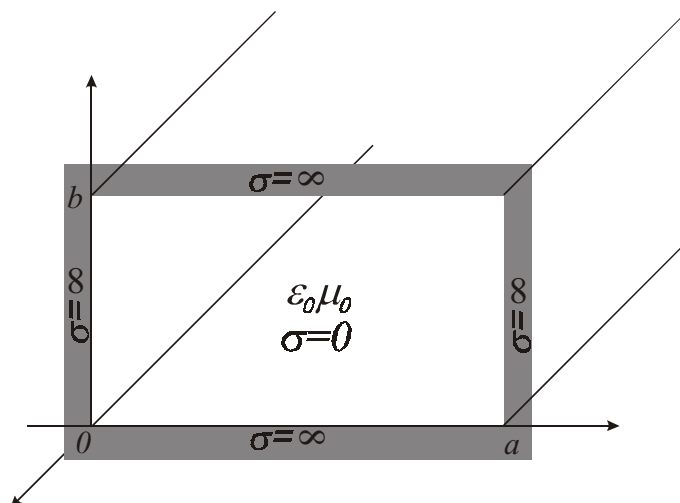
- Wektor gęstości strumienia energii fali elektromagnetycznej (wektor Poyntinga) w próżni wynosi:

$$\mathbf{P} = \mathbf{E} \times \mathbf{H}.$$

Wektor Poyntinga jest zgodny z kierunkiem propagacji.

TEORIA LINII TRANSMISYJNYCH

Falowody prostokątne



Dla wektorów zespolonych harmonicznie zmiennych równanie falowe przyjmuje postać Helmholtza:

$$\nabla^2 E_z - \gamma^2 E_z = 0 \quad (1)$$

$$\gamma = \sqrt{(R_1 + j\omega L_1)(G_1 + j\omega C_1)} = \alpha + j\beta, \quad \beta = \frac{2\pi}{\lambda}$$

współczynnik tłumienia
stała fazowa

przyjmując bezstratność falowodu $\alpha = 0$ otrzymamy:

$$\nabla^2 E_z + \beta^2 E_z = 0 \quad (2)$$

ROZWIĄZANIE

Zakładając rozwiązanie w postaci:

$$E_z = X(x)Y(y)Z(z) \quad (3)$$

otrzymuje się (jak przy rozwiązywaniu równania Helmholtza):

$$\frac{X''}{X} + \frac{Y''}{Y} + \frac{Z''}{Z} + \beta^2 = 0, \quad (4)$$

a po rozdzieleniu zmiennych:

$$X'' - \gamma_x^2 X = 0, \quad (5a)$$

$$Y'' - \gamma_x^2 Y = 0, \quad (5b)$$

$$Z'' - \gamma_x^2 Z = 0. \quad (5c)$$

Jednym z możliwych rozwiązań równań (5) jest superpozycja przeciwbieżnych fal.

$$E_z(x) = Ae^{\gamma_x x} + Be^{-\gamma_x x}. \quad (6)$$

Z w. b.

$$E_z(x)=0 \text{ dla } x=0$$

$$E_z(x)=0 \text{ dla } x=a$$

(znikanie rozwiązania w płaszczyznach ścianek falowodu) wynika, że składowa γ_x musi być zespolona zatem:

$$E_z(x) = Ae^{j\beta_x x} + Be^{-j\beta_x x} \quad (7)$$

Korzystając ze wzorów Eulera

$$e^{jz} = \cos z + i \sin z, \quad e^{-jz} = \cos z - i \sin z,$$

po przegrupowaniu wyrazów otrzymuje się:

$$E_z(x) = C \cos \beta_x x + D \sin \beta_x x \quad (8)$$

z 1. w.b. wynika, że $C = 0$, a $E_z(x) = D \sin \beta_x x$

2. w.b. będzie spełniony dla wszystkich β_x , dla których kąt $\beta_x a$ jest wielokrotnością π więc

$$\beta_x a = m\pi \quad m > 0 \text{ i } m \in \mathbb{C} \quad \text{skąd} \quad \beta_x = m \frac{\pi}{a}$$

przeprowadzając podobne rozumowanie dla funkcji Y , uzyskuje się

$$\beta_y = n \frac{\pi}{b}$$

ponieważ $E_z = 0$ dla $y = 0$ i $y = b$, $n > 0$ i $n \in \mathbb{C}$, n nie zależy od m .

Wartości β_x i β_y umożliwiają wyznaczenie rozkładu składowej E_z w płaszczyźnie $\perp$ do osi z .

$$[E_z(x, y)]_{m,n} = \cos\left(\frac{m\pi}{a}x\right)\cos\left(\frac{n\pi}{b}y\right)$$

wskaźniki m i n określają rodzaj fali typu E ozn. TE_{mn} .

Analogiczne rozważania dla składowej pola magnetycznego prowadzą do rozkładu:

$$[H_z(x, y)]_{m,n} = \cos\left(\frac{m\pi}{a}x\right)\cos\left(\frac{n\pi}{b}y\right)$$

wskaźniki m i n określają rodzaj fali typu H ozn. TM_{mn} .

Długość fali krytycznej - największa długość fali, dla której nie jest ona tłumiona. Długość ta wynosi:

$$\lambda_c = \frac{2}{\sqrt{\left(\frac{m}{a}\right)^2 + \left(\frac{n}{b}\right)^2}}$$

Podstawowym rodzajem fali jest taki rodzaj, dla którego $\lambda_c = \max$.

Dla pola E jest to TE_{10} , zaś dla TE_{10} $\lambda_c = 2a$ i nie zależy od b .

Aby nie dopuścić do pojawienia się w falowodzie innych rodzajów fal należy dobrać wymiar b . Zakładając możliwość pojawienia się sąsiednich dla TE_{10} rodzajów fal (czyli TE_{01} i TE_{20}) uzyskuje się:

$$\left(\frac{\pi}{b}\right)^2 = \left(\frac{2\pi}{a}\right)^2 \quad \text{stąd } a = 2b$$

Impedancja falowa dla takiego falowodu wynosi:

$$Z_0 = \frac{377}{\sqrt{1 - \left(\frac{\lambda}{\lambda_c}\right)^2}}$$

Falowody są ośrodkami o silnej dyspersji (zależności parametrów falowych w tym także prędkości fali od częstotliwości). W takich ośrodkach prędkość fazowa (prędkość zmian fazy) nie pokrywa się z prędkością grupową (prędkość przesuwania się obwiedni grupy fal - prędkość przenoszenia energii).

Przy zbliżaniu się długości fali w falowodzie do długości fali krytycznej, kierunek wektora prędkości grupowej zbliża się do kierunku prostopadłego do wektora prędkości fazowej, a jego wartość spada do zera, co oznacza, że fala o takiej długości nie może przenosić energii wzdłuż falowodu. Jednocześnie prędkość fazowa przy zbliżaniu się długości fali w falowodzie do długości fali krytycznej rośnie do nieskończoności.

Prędkość fazowa wynosi:

$$v_f = \frac{c}{\sqrt{1 - \left(\frac{\lambda}{\lambda_c}\right)^2}},$$

a grupowa:

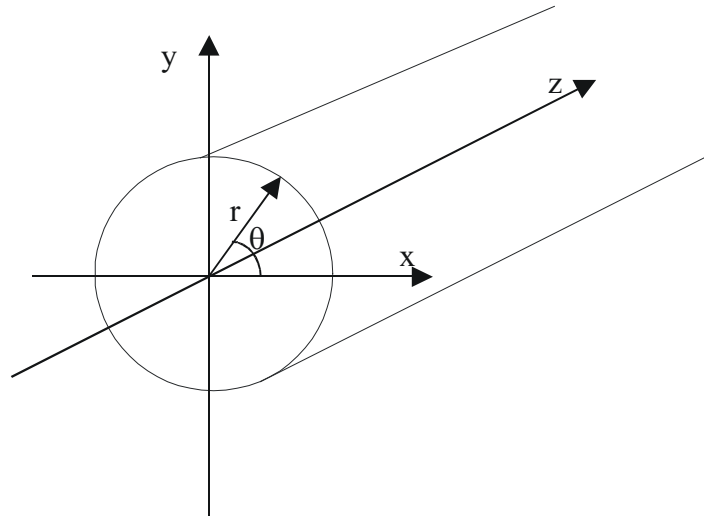
$$v_g = c \sqrt{1 - \left(\frac{\lambda}{\lambda_c}\right)^2}.$$

Przy czym: $v_f v_g = c^2$.

Fale w falowodzie dla określonej częstotliwości są dłuższe niż w wolnej przestrzeni:

$$\lambda_F = \frac{\lambda}{\sqrt{1 - \left(\frac{\lambda}{\lambda_c}\right)^2}}.$$

Falowody cylindryczne



Przyjmuje się, z oczywistego powodu, cylindryczny układ współrzędnych, którego związki z układem kartezjańskim są następujące.

$$x = r \cos \theta, \quad y = r \sin \theta, \quad z = z$$

$$r = \sqrt{x^2 + y^2}, \quad \theta = \arctg \frac{y}{x}$$

Podobnie jak dla falowodu prostokątnego w ściankach $\sigma = \infty$, zaś wewnątrz mamy $\sigma = 0$ oraz $\mu = \mu_0$ i $\varepsilon = \varepsilon_0$.

W cylindrycznym układzie odniesienia równanie falowe przyjmie postać:

$$\frac{1}{r} \frac{\partial}{\partial r} \left(r \frac{\partial E}{\partial r} \right) + \frac{1}{r^2} \frac{\partial^2 E}{\partial \theta^2} + \chi_e E = 0 \quad (1)$$

gdzie χ_e jest pewną stałą charakteryzującą pole

Rozdzielając zmienne $E = f(r)g(\theta)$ i mnożąc przez $\frac{r^2}{f(r)g(\theta)}$ otrzymujemy układ równań ze względu na f i g :

$$r^2 f''(r) + r f'(r) + \chi_e r^2 f(r) = n^2 f(r) \quad (2)$$

$$g''(\theta) + n^2 g(\theta) = 0 \quad (3)$$

Równanie (2) można sprowadzić do równania Bessela za pomocą funkcji $z(r\sqrt{\chi_e}) = f(r)$, wtedy:

$$r^2 \chi_e z''(r\sqrt{\chi_e}) + r\sqrt{\chi_e} z'(r\sqrt{\chi_e}) + (\chi_e r - n^2) z(r\sqrt{\chi_e}) = 0. \quad (4)$$

Zagadnienie tego typu rozwiązują funkcje Bessela.

$$\text{Dla } n \in \mathbb{C} \quad z = C_1 I_n \left(r \sqrt{\chi_e} \right) + C_2 Y_n \left(r \sqrt{\chi_e} \right) \quad (5)$$

I_n - funkcja Bessela I-ego rodzaju n -ego rzędu

Y_n - funkcja Bessela II-ego rodzaju n -ego rzędu

Ponieważ (z własności funkcji) $\lim_{r \rightarrow 0} Y_n \left(r \sqrt{\chi_e} \right) \rightarrow \infty$
można przyjąć

$$z \left(r \sqrt{\chi_e} \right) = I_n \left(r \sqrt{\chi_e} \right) \quad (6)$$

wobec tego rozwiązanie równania falowego ma postać:

$$E(r, \theta) = I_n \left(r \sqrt{\chi_e} \right) \cos n\theta \quad (7)$$

i spełnia w.b. $I_n \left(R \sqrt{\chi_e} \right) = 0$

Pierwiastki funkcji Bessela (lub jej pochodnych) ozn. α_{mne} oraz α_{mnh}

n - rząd funkcji Bessela

m - miejsca zerowego funkcji

Wartościom m i n odpowiadają różne rodzaje pola.

Rodzaje podstawowe $\alpha_{01e} = 2,4$ oraz $\alpha_{11h} = 1,84$

Ogólnie $R \sqrt{\chi} = \alpha_{mn} \rightarrow \chi = \frac{\alpha_{mn}^2}{R^2}$

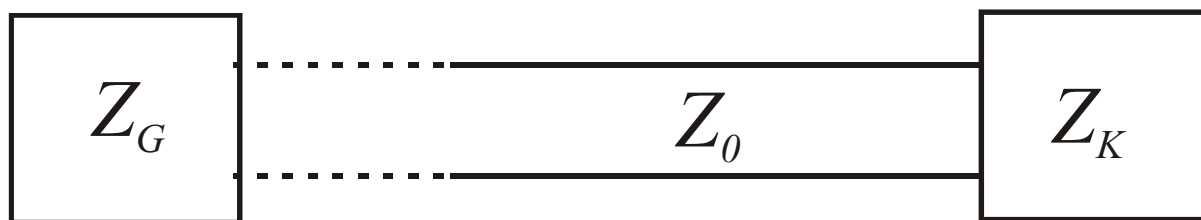
Krytyczna długość fali dla falowodu cylindrycznego wynosi:

$$\lambda_c = \frac{2\pi}{\sqrt{\chi}} = \frac{2\pi R}{\alpha_{mn}}$$

dla TM $\lambda_{Ce} = \frac{2\pi R}{2,4} = 2,61R$

dla TE $\lambda_{Ch} = \frac{2\pi R}{1,84} = 3,41R$

Miary dopasowania



Stan dopasowania linii do obciążenia oznacza, że energia fali elektromagnetycznej propagującej się w linii w całości przedostanie się do obciążenia. Warunkiem dopasowania jest równość impedancji $Z_0 = Z_K$. W każdym innym przypadku mówi się o niedopasowaniu.

Używa się dwóch miar dopasowania mających genezę w teorii odbicia fali elektromagnetycznej na granicy ośrodków:

1. Współczynnik odbicia

$$\Gamma = \frac{u_{odb}}{u_{pad}} = -\frac{i_{odb}}{i_{pad}} = \frac{Z_K - Z_0}{Z_K + Z_0}$$

$$\Gamma = |\Gamma| e^{j\theta} \quad \Gamma \in \langle -1; 1 \rangle$$

θ – zmiana fazy w miejscu odbicia

2. Współczynnik fali stojącej

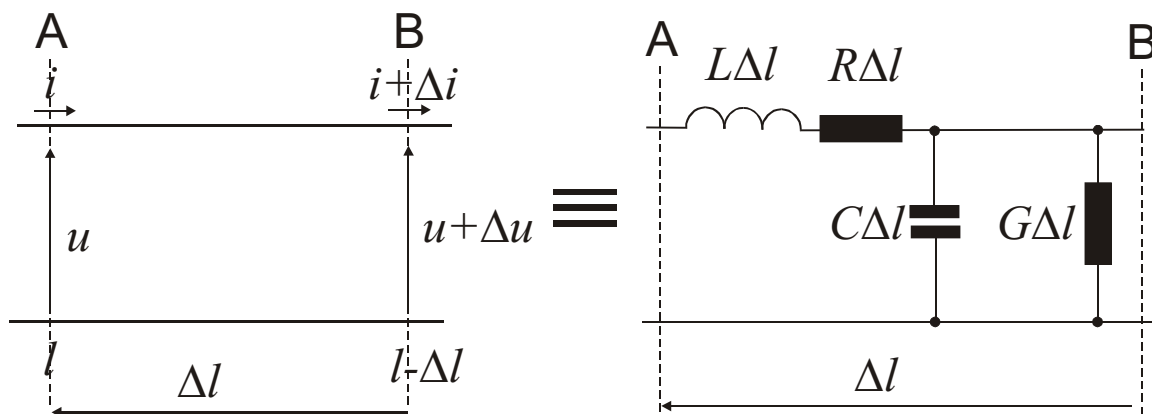
$$WFS = \frac{u_{\max}}{u_{\min}} = \frac{|u_{pad}| + |u_{odb}|}{|u_{pad}| - |u_{odb}|} = \frac{1 + |\Gamma|}{1 - |\Gamma|} \quad WFS \in (1; \infty)$$

Uwaga. W anglojęzycznej literaturze używa się oznaczenia SWR (standing wave ratio).

$$|\Gamma| = \frac{WFS - 1}{WFS + 1}$$

Transformacyjne własności linii przesyłowych

Odcinek linii transmisyjnej i jego schemat zastępczy z elementami skupionymi:



Definicje:

Impedancja charakterystyczna linii:

$$Z_0 = \sqrt{\frac{(R + j\omega L)}{(G + j\omega C)}}$$

gdzie R , L , G , C są tzw. parametrami jednostkowymi linii wyrażanymi w odpowiednich jednostkach na metr.

Stała propagacji:

$$\gamma = \sqrt{(R + j\omega L)(G + j\omega C)}$$

$$\operatorname{Re}(\gamma) = \alpha \quad \text{stała tłumienia [dB/m]}$$

$$\operatorname{Im}(\gamma) = \frac{2\pi}{\lambda} = \frac{\omega}{c} = \beta \quad \text{stała fazowa [rad/m]}$$

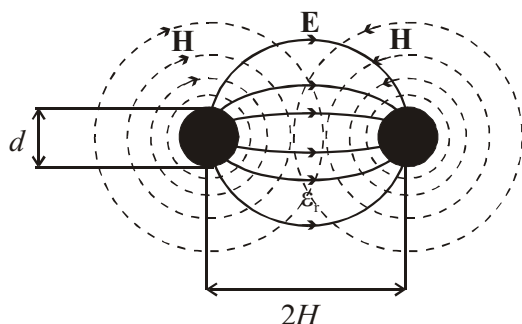
dla linii bezstratnej $R = 0$, $G = 0$ $\alpha = 0$.

Prędkość fali w linii bezstratnej:

$$v = \frac{1}{\sqrt{LC}}.$$

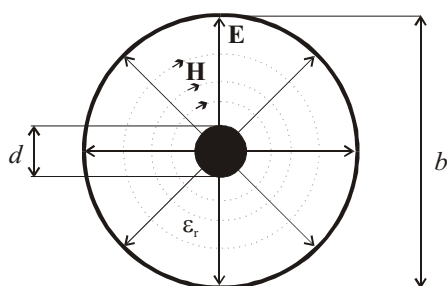
Przykłady przewodowych linii przesyłowych

linia dwuprzewodowa



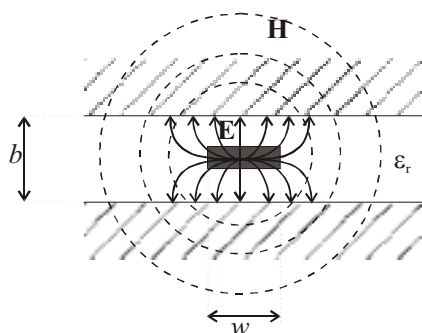
$$Z_0 = \frac{120}{\sqrt{\epsilon_r}} \ln \left(\frac{2H}{d} + \sqrt{\left(\frac{2H}{d} \right)^2 - 1} \right)$$

linia koncentryczna



$$Z_0 = \frac{60}{\sqrt{\epsilon_r}} \ln \frac{b}{d}$$

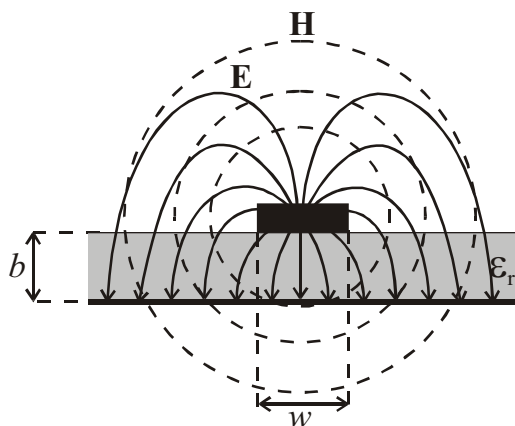
symetryczna linia paskowa (SLP)



$$Z_0 = \frac{60}{\sqrt{\epsilon_r}} \ln \frac{8b}{\pi w} \quad \text{dla } w \ll b$$

$$Z_0 = \frac{15\pi^2}{\sqrt{\epsilon_r}} \left(\frac{\pi w}{2b} + \ln 2 \right)^{-1} \quad \text{dla } w \gg b$$

niesymetryczna linia paskowa



$$Z_0 = 337 \frac{b}{w} \left\{ \sqrt{\epsilon_r} w \left[1 + 1,735 \epsilon_r^{-0,0724} \left(\frac{w}{b} \right)^{-0,386} \right] \right\}^{-1}$$

Równania telegrafistów

Korzystając z prawa Ohma dla odcinka linii Δl ([patrz schemat zastępczy](#)) otrzymamy:

$$\Delta u(l) = (R + j\omega L)\Delta l \cdot i(l)$$

$$\Delta i(l) = (G + j\omega C)\Delta l \cdot u(l)$$

Po przejściu do przyrostów nieskończenie małych mamy:

$$\frac{du(l)}{dl} = (R + j\omega L) \cdot i(l)$$

$$\frac{di(l)}{dl} = (G + j\omega C) \cdot u(l),$$

co po podzieleniu na obydwie strony pierwszego z równań operatorem różniczkowania po długości i wstawieniu drugiego równania oraz podzieleniu na obydwie strony drugiego z równań operatorem różniczkowania po długości i wstawieniu pierwszego równania, prowadzi w rezultacie do układu równań typu falowego opisujących zmiany napięć i prądów w linii transmisyjnej (tzw. równania telegrafistów):

$$\frac{d^2 u(l)}{dl^2} = (R + j\omega L)(G + j\omega C) \cdot u(l) = \gamma^2 \cdot u(l)$$

$$\frac{d^2 i(l)}{dl^2} = (R + j\omega L)(G + j\omega C) \cdot i(l) = \gamma^2 \cdot i(l)$$

Jednoznaczne rozwiązanie otrzymuje się dla ustalonego obciążenia końcowego:

$$Z_K = \frac{U_K}{I_K}, \quad u(l)\big|_{l=0} = U_K, \quad i(l)\big|_{l=0} = I_K$$

ma ono postać:

$$u(l) = U_K \cosh \gamma l + I_K Z_0 \sinh \gamma l$$

$$i(l) = I_K \cosh \gamma l + \frac{U_K}{Z_0} \sinh \gamma l$$

Zatem impedancja wejściowa w linii w odległości l od obciążenia wynosi:

$$Z_{we}(l) = \frac{u(l)}{i(l)} = Z_0 \frac{Z_K + Z_0 \tanh \gamma l}{Z_0 + Z_K \tanh \gamma l}$$

Dla linii bezstratnej zależność ta się upraszcza do:

$$Z_{we}(l) = Z_0 \frac{Z_K + jZ_0 \tan \beta l}{Z_0 + jZ_K \tan \beta l}$$

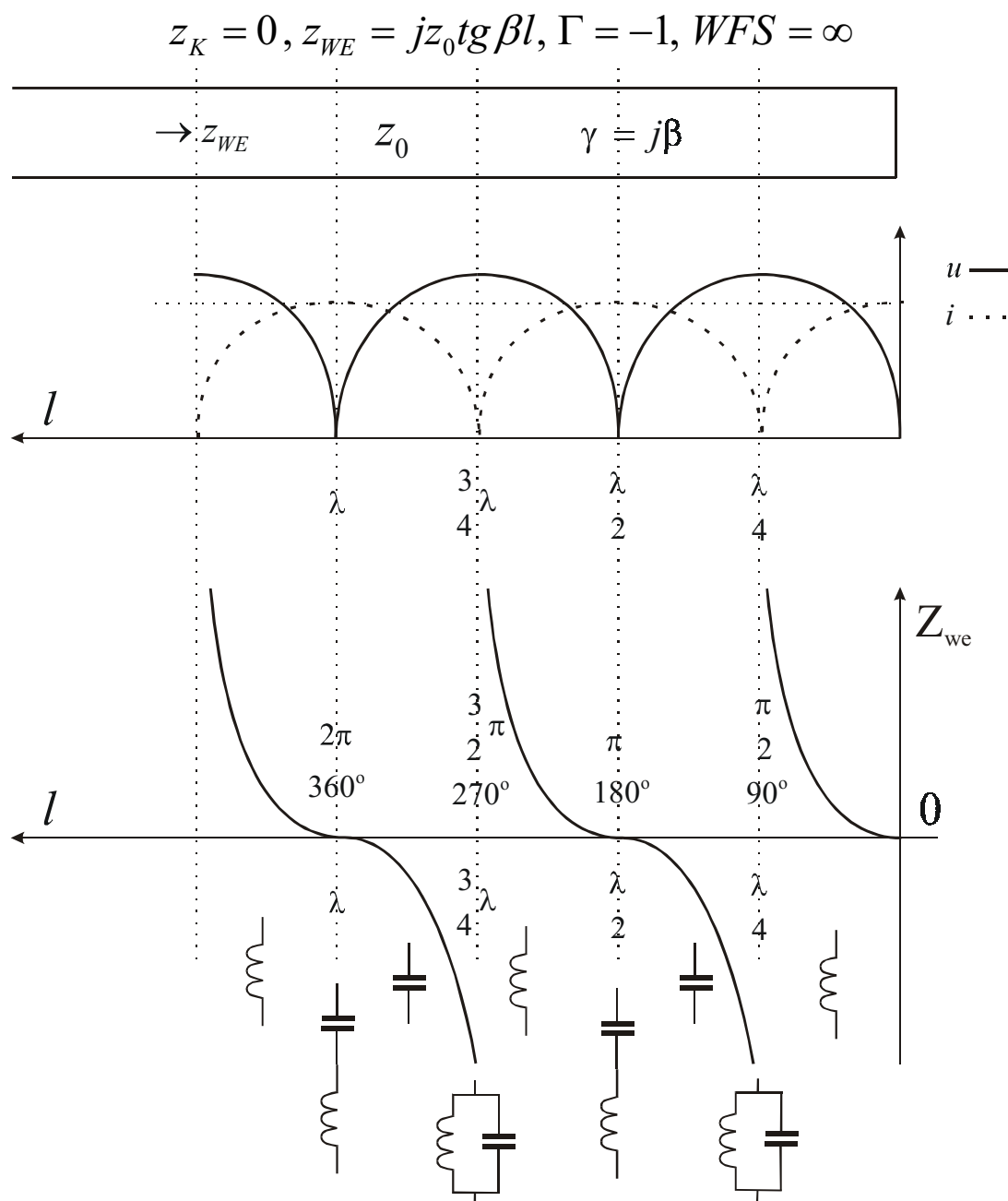
Z racji tego, że większość linii transmisyjnych można uznać z dobrym przybliżeniem za bezstratne, powyższa zależność ma kardynalne znaczenie przy obliczaniu parametrów wielu obwodów mikrofalowych działających w oparciu o teorię linii transmisyjnych.

Jak widać impedancja linii zmienia się wraz z odległością od obciążenia stąd mówi się (w sensie impedancyjnym) o transformacyjnych własnościach linii.

Trywialnym przypadkiem, jest obciążenie linii impedancją $Z_K = Z_0$ (idealne dopasowanie), dla którego w linii nic się nie zmienia i impedancja wejściowa na całej jej długości wynosi Z_0 . Można więc powiedzieć, że impedancja charakterystyczna linii to taka impedancja, że po obciążeniu nią linii prąd, napięcie, a co za tym idzie także impedancja wejściowa, utrzymują się wzdłuż linii na stałym poziomie.

Pozostałe przypadki omówione są poniżej.

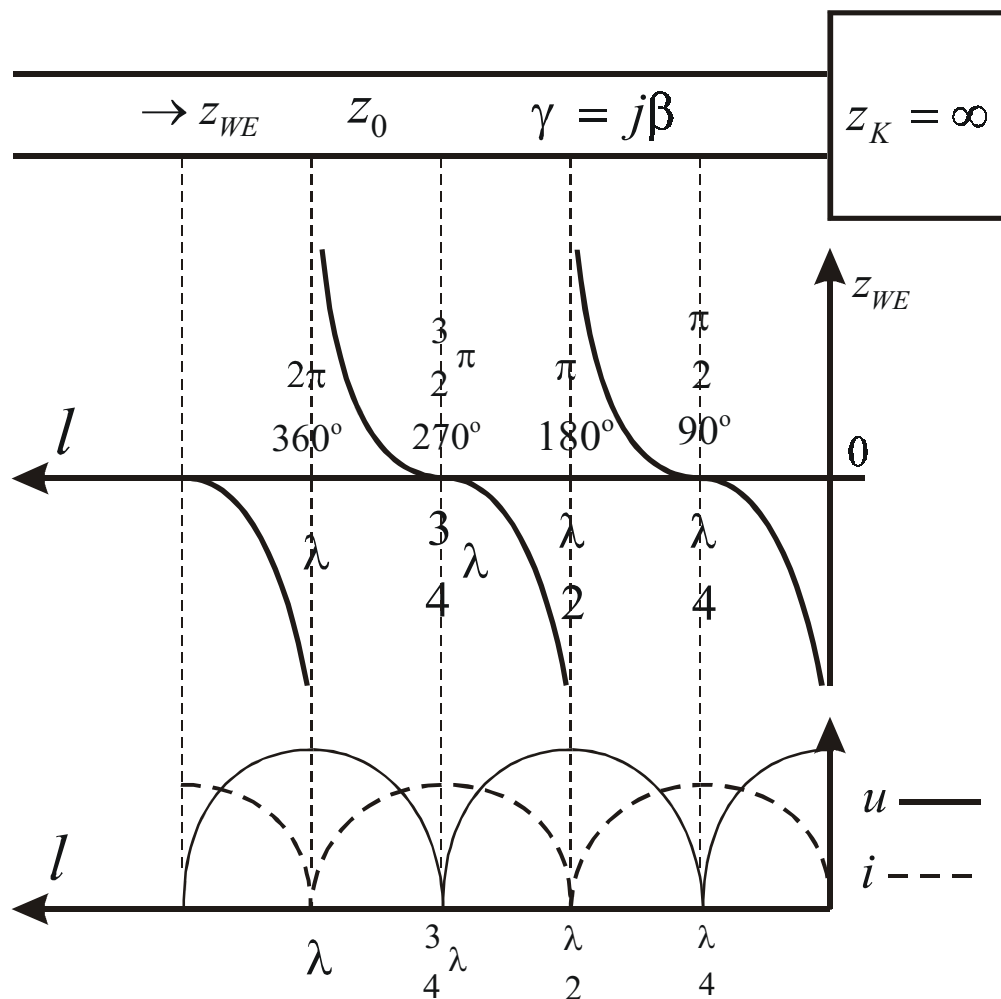
Linia zwarta na końcu



Jak widać zmienia się nie tylko impedancja ale także jej charakter. Jest tak nie tylko dla zwarcia, ale i dla dowolnego obciążenia.

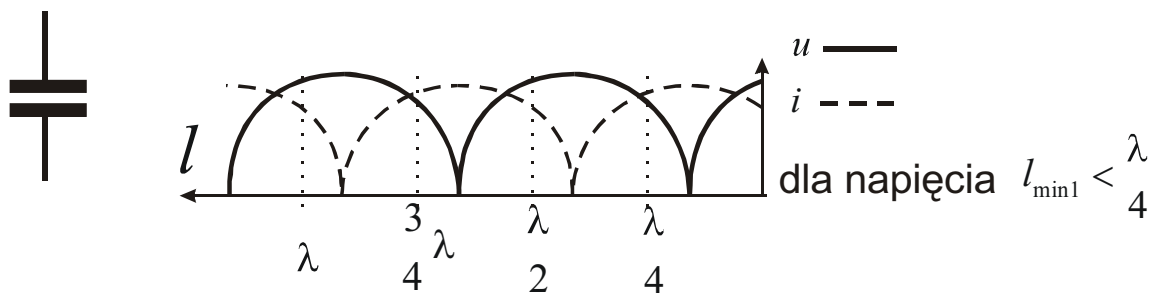
Linia rozwartą

$$z_K = \infty, z_{WE} = -jz_0 \operatorname{ctg} \beta l, \Gamma = 1, WFS = \infty$$



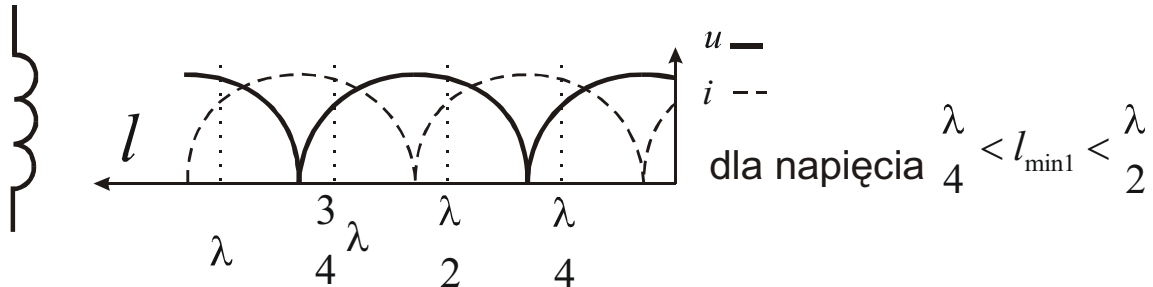
Linia obciążona pojemnościowo

$$z_K = -jX_{KC} = -j\frac{1}{\omega C}, z_{WE} = -jz_0 \frac{X_{KC} - z_0 \operatorname{tg} \beta l}{z_0 + X_{KC} \operatorname{tg} \beta}, \Gamma = 1, WFS = \infty$$



Linia obciążona indukcyjnie

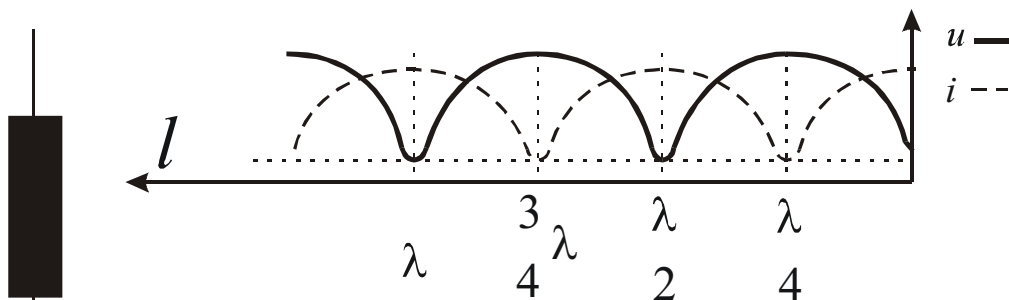
$$z_K = jX_{KL} = -j\omega L, z_{WE} = jz_0 \frac{X_{KL} + z_0 \operatorname{tg} \beta l}{z_0 - X_{KL} \operatorname{tg} \beta}, \Gamma = -1, WFS = \infty$$



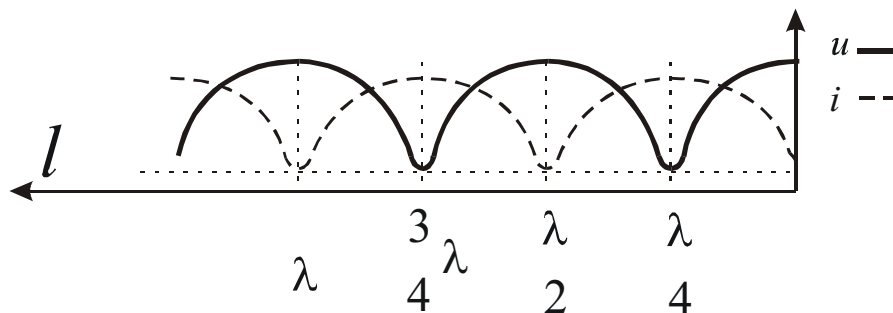
Dwa przypadki obciążenia rezystancyjnego

$$0 < z_K = R_K < \infty, z_{WE} \in \operatorname{Im}, U_{ODB} < U_{PAD}$$

$$0 < R_K < z_0, 1 > \Gamma > 0, 1 < WFS < \infty$$



$$z_0 < R_K < \infty, -1 < \Gamma < 0, 1 < WFS < \infty$$



Wykresy impedancji dla linii obciążonej rezystancyjnie (lub reaktancją z częścią rzeczywistą) nie będą przebiegały do ∞ , ale będą ograniczone wartością rezystancji.

Wykres Smitha.

Odwzorowanie homograficzne

- Funkcja homograficzna wiążąca ze sobą dwie zmienne zespolone w i z ma postać:

$$w = \frac{az + b}{cz + d} \quad ; \quad z \neq -\frac{d}{c}$$

przy czym a , b , c i d są stałymi zespolonymi.

Odwzorowaniem homograficznym nazywamy przyporządkowanie punktom na płaszczyźnie zespolonej Z punktów na płaszczyźnie zespolonej W , opisane funkcją homograficzną. Podstawowe własności odwzorowania homograficznego:

- ✓ odwzorowanie homograficzne $w(z)$ jest wzajemnie jednoznaczne,
- ✓ okrąg na płaszczyźnie Z transformuje się na okrąg na płaszczyźnie W (prosta jest szczególnym przypadkiem okręgu),
- ✓ zachowana zostaje ortogonalność okręgów.

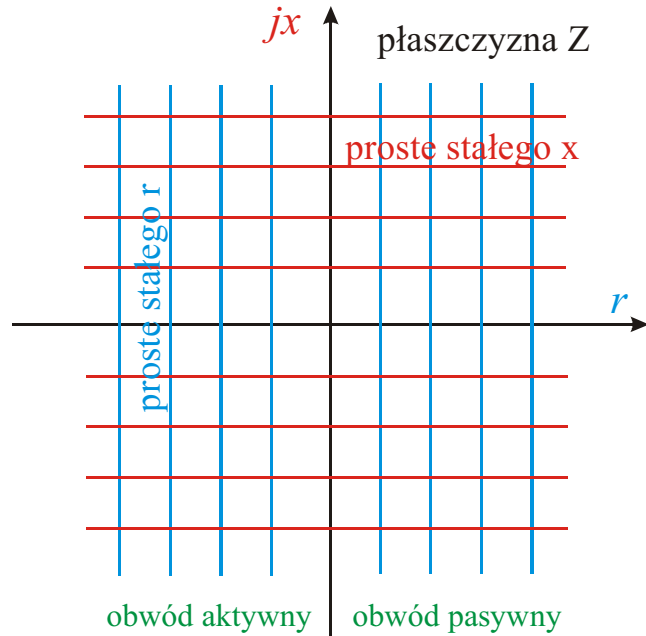
Znana już zależność:

$$\Gamma = \frac{Z_K - Z_0}{Z_K + Z_0} = \frac{\frac{Z_K}{Z_0} - 1}{\frac{Z_K}{Z_0} + 1} = \frac{z_{zn} - 1}{z_{zn} + 1}$$

to typowa funkcja homograficzna, z_{zn} – impedancja znormalizowana.

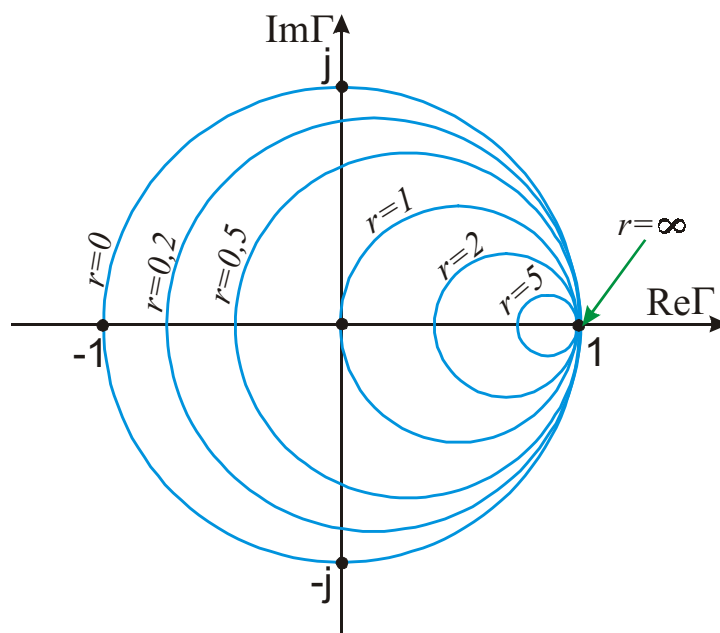
Konstrukcja wykresu Smitha

Wykres Smitha powstaje przez przetransformowanie siatki prostych $r = const$ i $x = const$ z części pasywnej płaszczyzny impedancji Z na płaszczyznę współczynnika odbicia Γ .

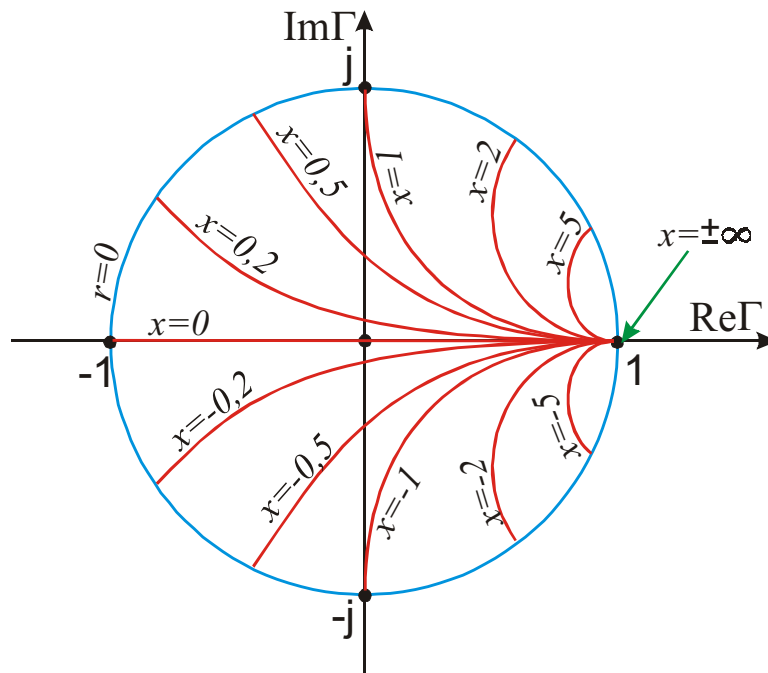


$$\Gamma = \frac{r + jx - 1}{r + jx + 1}$$

Proste $r = const$ na płaszczyźnie Z transformują się na płaszczyznę Γ jako okręgi o promieniach $1/(r+1)$ i środkach $[r/(r+1), 0]$.



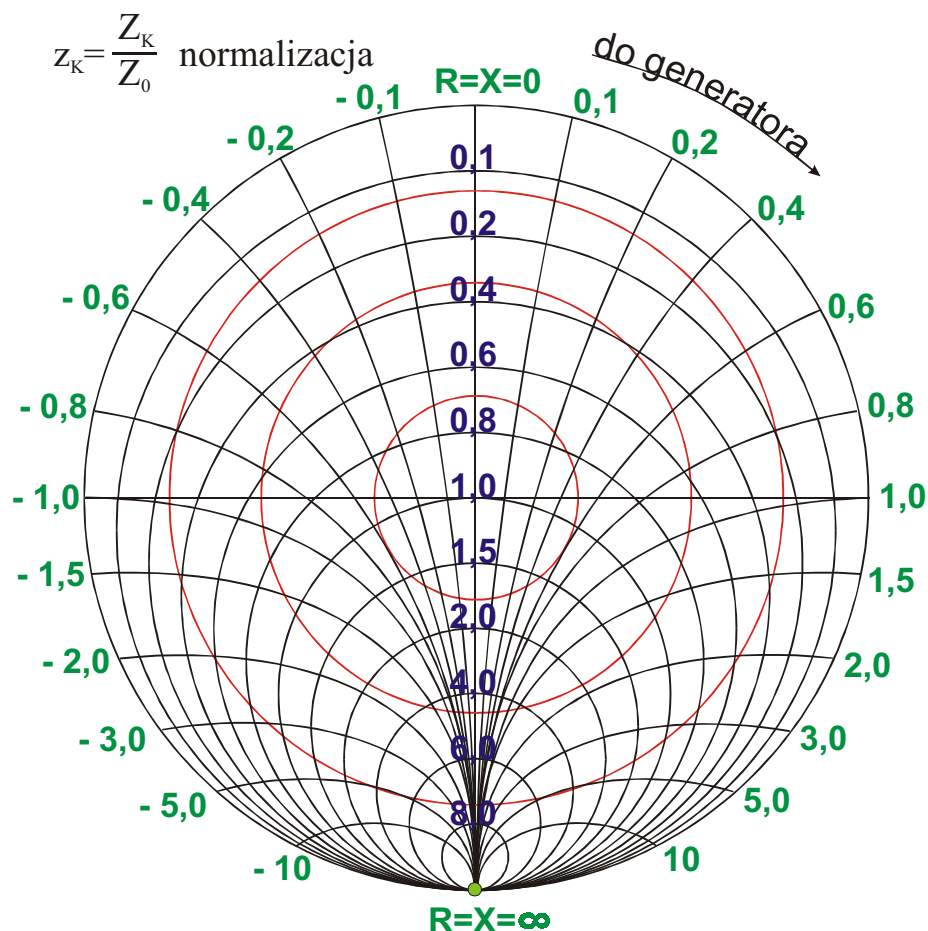
Proste $x = \text{const}$ transformują się na okręgi o promieniach $1/|x|$ i środkach leżących w punktach o współrzędnych $[1, 1/x]$.



- Obie rodziny okręgów są względem siebie ortogonalne.
- Jeżeli transformację ograniczyć do prawej (pasywnej) półpłaszczyzny $r \geq 0$ to otrzymuje się wykres Smitha.
- Można też przetransformować z płaszczyzny admitancji Y proste $g = \text{const}$ i $b = \text{const}$ na odpowiednie okręgi na płaszczyźnie Γ - otrzymuje się wtedy identyczną, siatkę współrzędnych, ale obróconą o 180° .
- Punkty prawej półpłaszczyzny Z transformują się do wnętrza okręgu o promieniu 1, punkty lewej półpłaszczyzny transformują się do zewnątrz okręgu.
- Przy graficznej prezentacji parametrów obwodów generacyjnych korzysta się z poszerzonego wykresu Smitha, zawierającego także siatkę współrzędnych poza okręgiem jednostkowym.

Po nałożeniu obydwu rodzin kół otrzymuje się wykres Smitha. Koła współśrodkowe to koła stałego $|\Gamma|$ (albo WFS).

UWAGA. Prezentowany rysunek jest tylko poglądowy i nie nadaje się do obliczeń

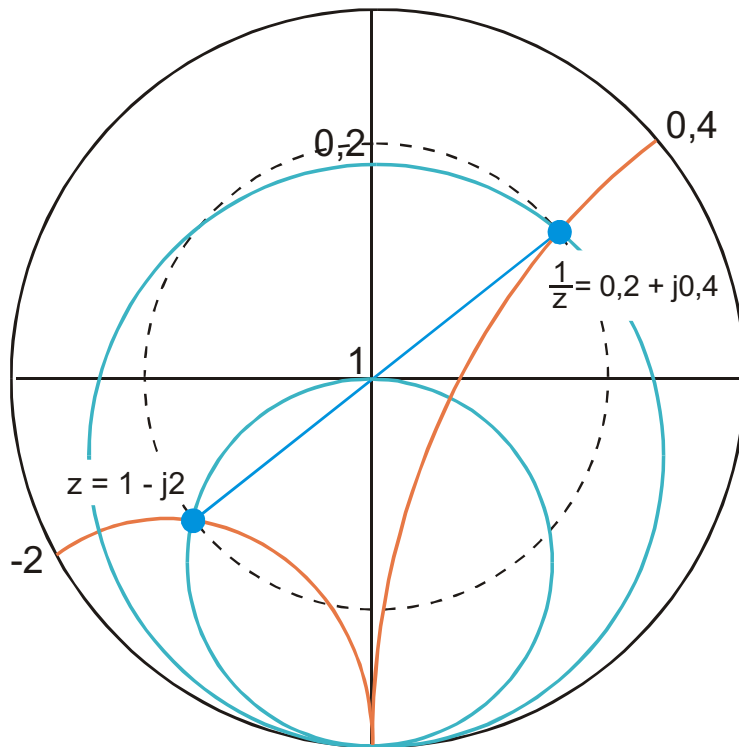


Aby skorzystać z wykresu Smitha często konieczne jest przeprowadzenie normalizacji liczby zespolonej. W przypadku rozważań impedancyjnych normalizacji dokonuje się dzieląc daną impedancję przez impedancję charakterystyczną linii. Impedancję znormalizowaną można nanieść na wykres – znajduje się ona w miejscu przecięcia się odpowiedniego koła stałej rezystancji z odpowiednim kołem stałej reaktancji (np. impedancja $1 - j2$ znajdzie się na przecięciu koła stałej rezystancji $r=1$ z kołem stałej reaktancji $jx = -2$).

Do najprostszych operacji, które można wykonać na wykresie Smitha jest odwracanie liczb zespolonych. Jak już wiadomo przetransformowanie płaszczyzny admitancji Y czyli rodziny prostych $g = \text{const}$ i $b = \text{const}$, na odpowiednie okręgi na płaszczyźnie Γ - daje wykres Smitha z identyczną siatką współrzędnych, ale obróconą o 180° względem wykresu otrzymanego poprzez przetransformowanie płaszczyzny impedancji.

Wynika z tego, że dwie dowolne liczby położone symetrycznie względem środka wykresu są swoimi odwrotnościami.

Po dokonaniu operacji na wykresie Smitha i znalezieniu szukanej wartości impedancji (ogólnie liczby zespolonej) w celu ustalenia jej faktycznej wartości (a nie znormalizowanej) należy dokonać denormalizacji tj. pomnożyć liczbę z wykresu Smitha przez tą samą, za pomocą której dokonana była normalizacja. W przypadku rozważań impedancyjnych jest to impedancja charakterystyczna linii.



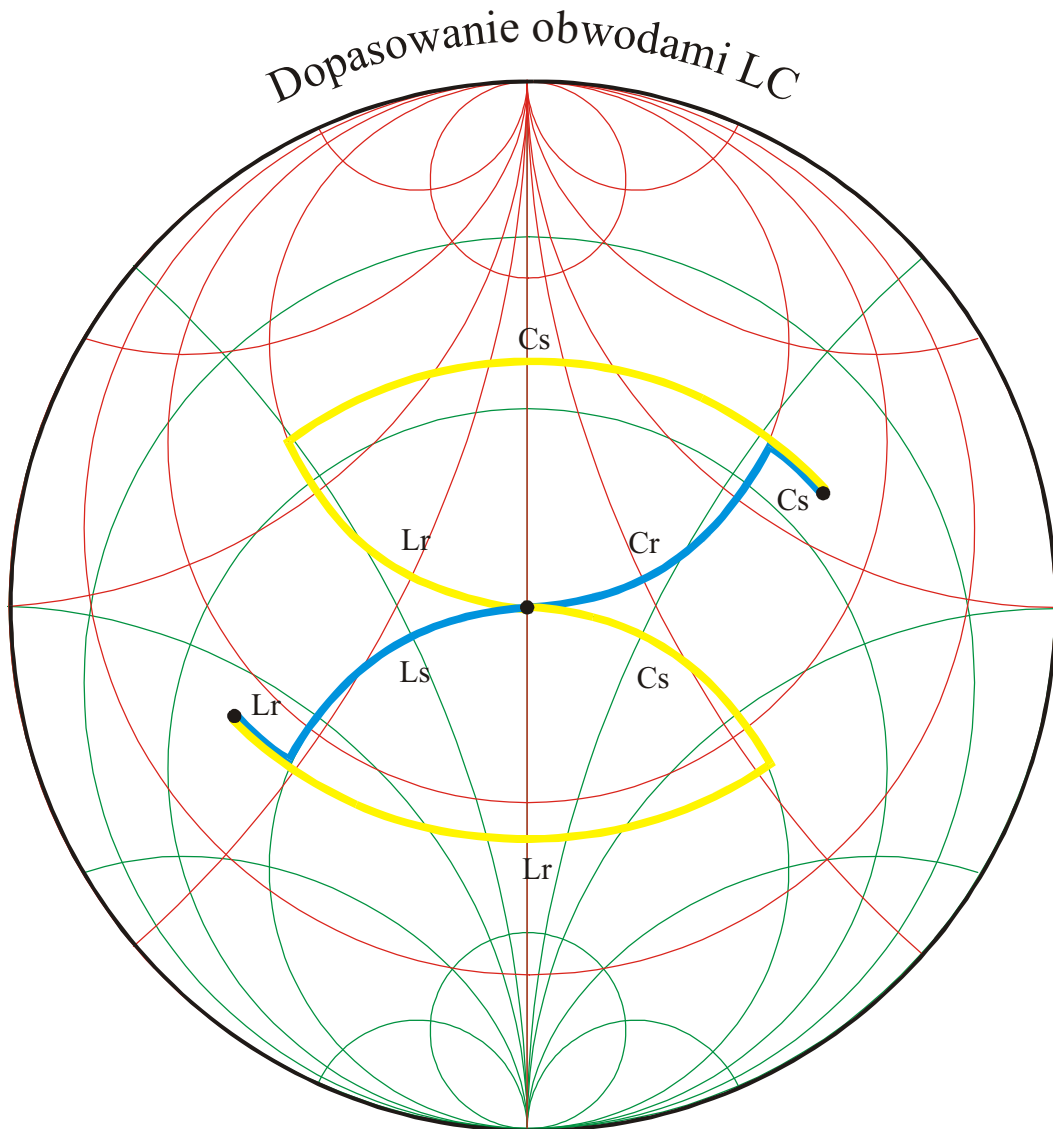
Wykres Smitha jest pod pewnym względem graficzną reprezentacją znanej już zależności na impedancję wejściową linii.

$$Z_{we}(l) = Z_0 \frac{Z_K + jZ_0 \operatorname{tg} \beta l}{Z_0 + jZ_K \operatorname{tg} \beta l}$$

Zaznaczoną na nim impedancję końcową (lub wejściową) linii można bowiem przetransformować (czyli przenieść po kole stałego $|\Gamma|$) w dowolny punkt linii odległy o l od impedancji końcowej (lub punktu o danej impedancji wejściowej) w stronę generatora lub obciążenia (zależnie od potrzeb). Do tego celu wystarczy pamiętać, że pełny obrót po kole stałego $|\Gamma|$ odpowiada odległości równej połowie długości fali.

Metody dopasowania impedancji

Skupione elementy reaktancyjne są stosowane w celu zniwelowania nadmiarowej reaktancji w miejscu włączenia obciążenia lub też w odpowiednim miejscu linii (tj. spełniającym warunek $Z_{we} = Z_0 + jX$).

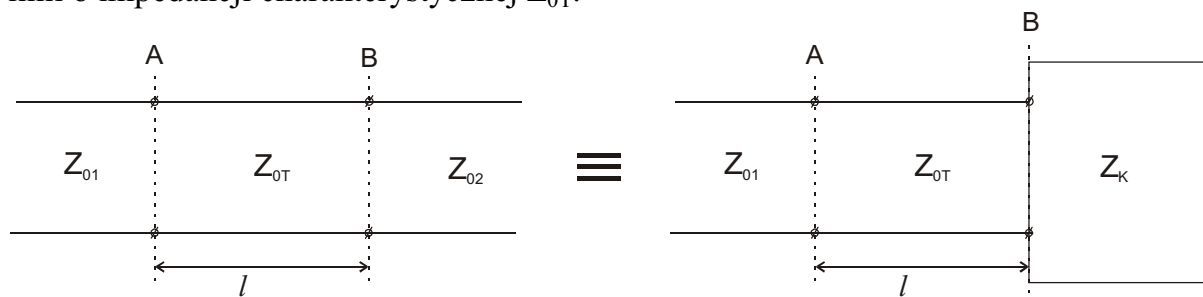


Reaktancyjne elementy szeregowe transformują impedancję po kołach stałego r , równoległe po kołach stałego g , przy czym indukcyjności zawsze w prawo, a pojemności zawsze w lewo.

Rezystancje szeregowe transformują impedancję po kołach stałego x , a równoległe po kołach stałego b zawsze w kierunku biegunów wykresu.

Transformator ćwierćfalowy

Należy dopasować dwie linie transmisyjne o impedancjach charakterystycznych Z_{01} i Z_{02} (rzeczywistych) za pomocą pewnego odcinka pewnej linii o impedancji charakterystycznej Z_{0T} .



Jeśli zadanie jest wykonalne to w płaszczyźnie A będzie:

$$Z_{01} = Z_T \frac{Z_{02} + jZ_T \operatorname{tg} \beta l}{Z_T + jZ_{02} \operatorname{tg} \beta l} \quad | : \tau \rightarrow \infty$$

$$Z_{01} = Z_T \frac{\frac{Z_{02}}{\tau \rightarrow \infty} + jZ_T \frac{\operatorname{tg} \beta l}{\tau \rightarrow \infty}}{\frac{Z_T}{\tau \rightarrow \infty} + jZ_{02} \frac{\operatorname{tg} \beta l}{\tau \rightarrow \infty}}$$

$$\operatorname{tg} \beta l \rightarrow \infty \quad \text{dla} \quad \beta l = \frac{\pi}{2} \quad \text{wtedy} \quad Z_{01} = Z_T \frac{jZ_T}{jZ_{02}} = \frac{Z_T^2}{Z_{02}}$$

zatem

$$Z_T = \sqrt{Z_{01} Z_{02}}$$

ponieważ relacje te zachodzą dla $\beta l = \frac{\pi}{2}$, a $\beta = \frac{2\pi}{\lambda}$ to wynika z tego, że $l = \frac{\lambda}{4}$.

Zadanie zatem rozwiązuje odcinek linii o impedancji charakterystycznej $Z_T = \sqrt{Z_{01} Z_{02}}$ i długości $l = \frac{\lambda}{4}$.

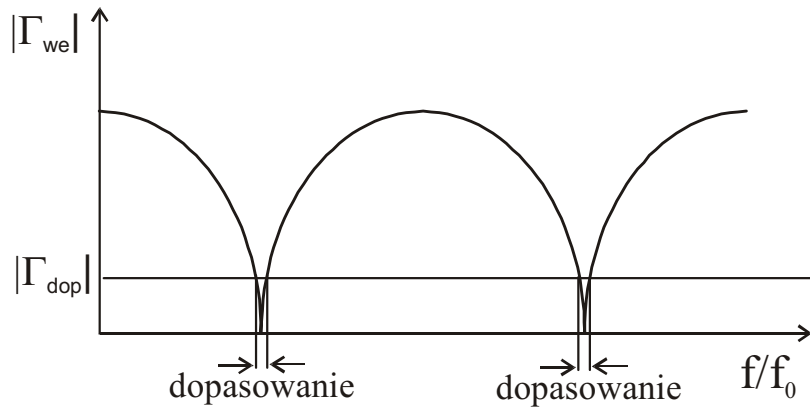
Dopasowanie to jest oczywiście wąskopasmowe (tylko dla fali o długości λ).

Z obydwu stron transformatora współczynnik odbicia wynosi:

$$\Gamma_{we} = \frac{Z_{02} - Z_{01}}{Z_{02} + Z_{01} + j2\sqrt{Z_{01} Z_{02}} \cdot \operatorname{tg} \beta l}$$

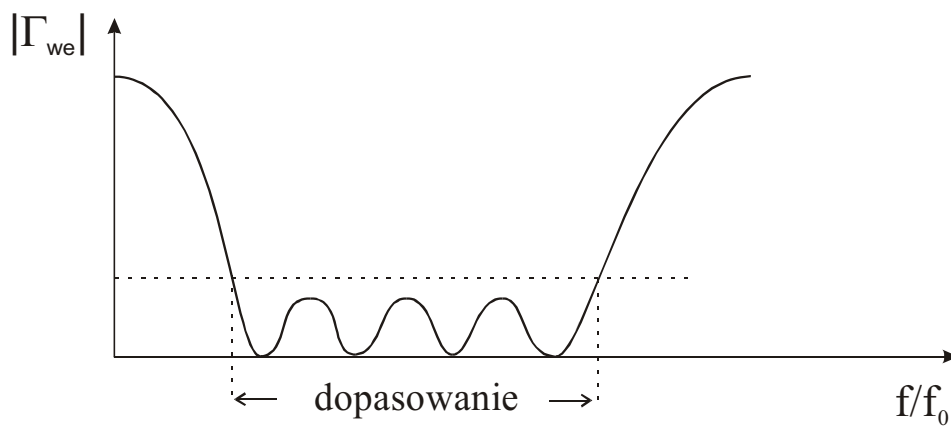
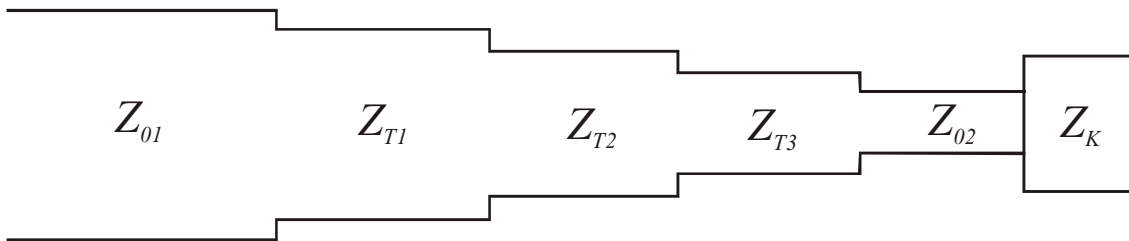
$$|\Gamma_{we}| = \frac{Z_{02} - Z_{01}}{\sqrt{(Z_{02} + Z_{01})^2 + 4Z_{01}Z_{02}tg^2\beta l}}$$

Jeśli założyć dopuszczalne niedopasowanie na niewielkim poziomie to transformator będzie realizować dopasowanie w bardzo wąskim zakresie.

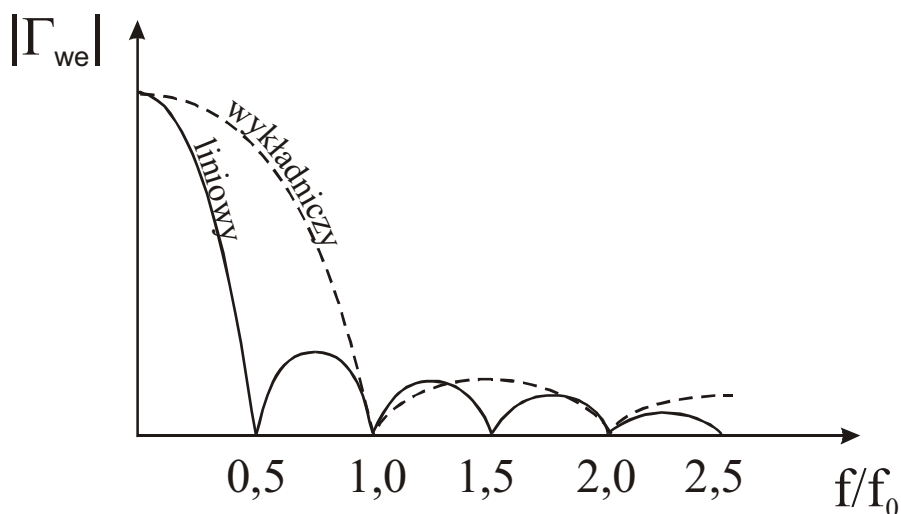
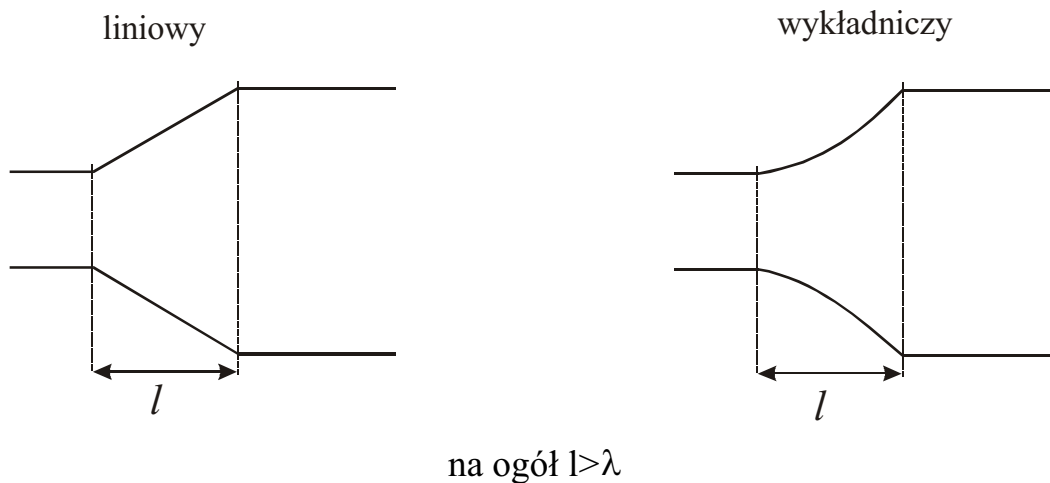


Łańcuch transformatorów

Szersze pasmo można uzyskać stosując łańcuch transformatorów.



Falowodowe realizacje transformatorów $\lambda/4$



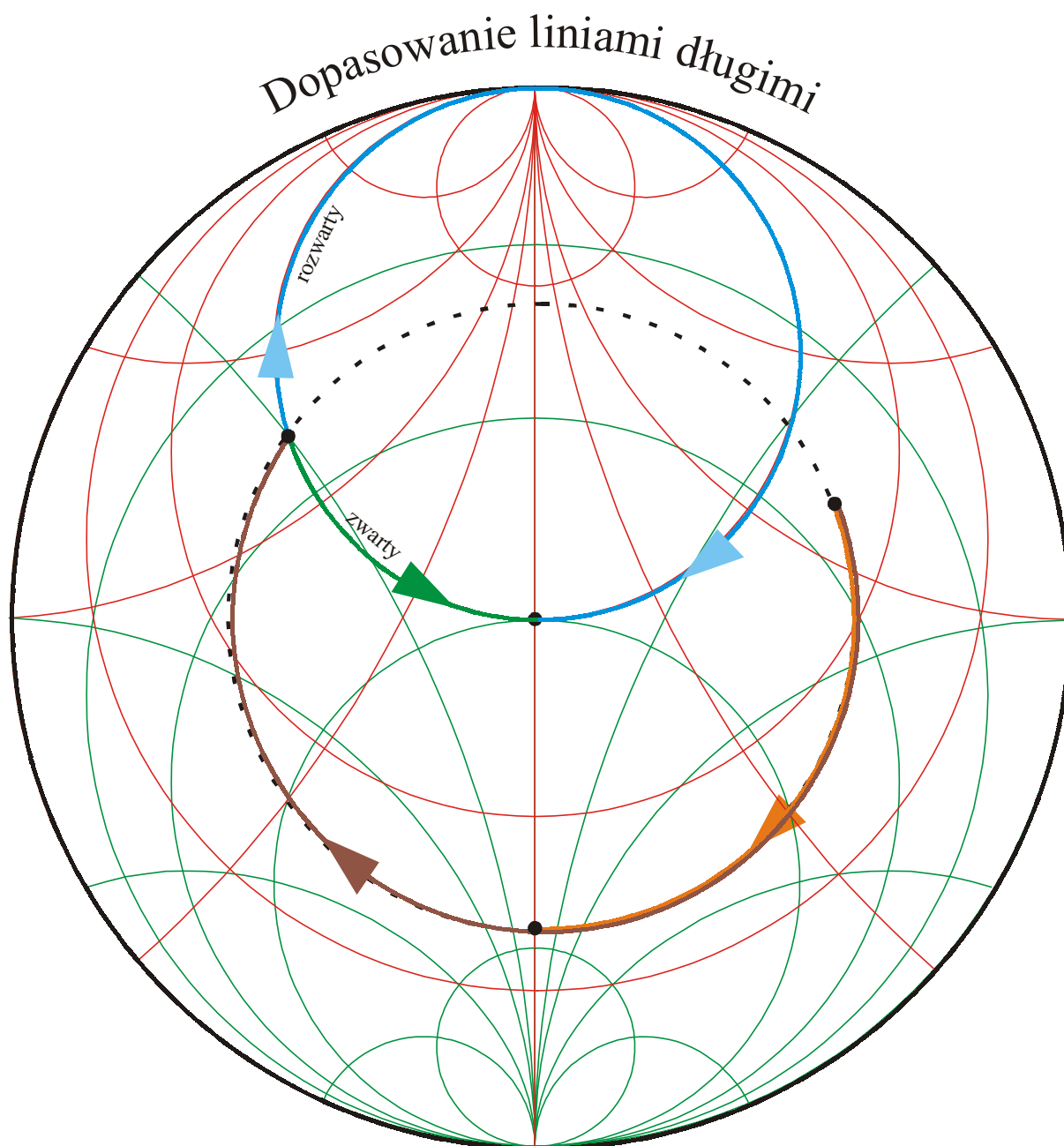
Stroik

Stroik jest odcinkiem linii przesyłowej zwartej lub rozwartej na końcu, równolegle dołączonej do linii dopasowywanej w miejscu gdzie $Z_{we} = Z_0 + jX$ z zadaniem skompensowania nadmiarowej reaktancji (podobnie jak reaktancje skupione). Parametry stroika można obliczyć analitycznie:

Odległość stroika od obciążenia l (miejsce jego włączenia) wyznacza się z równania $\text{Re}[Z_{we}(l)] = Z_0$. Jego długość l_s powinna niwelować $\text{Im}[Z_{we}(l)]$ zatem znaleźć można ją rozwiązując równanie:

$$-\text{Im}[Z_{we}(l)] = Z_0 \frac{Z_K + jZ_0 \text{tg}(\beta l_s)}{Z_0 + jZ_K \text{tg}(\beta l_s)}$$

Łatwiej znaleźć jest parametry stroika (i transformatora ćwierćfalowego) korzystając z wykresu Smitha. Poniżej prezentowane są dwie metody.

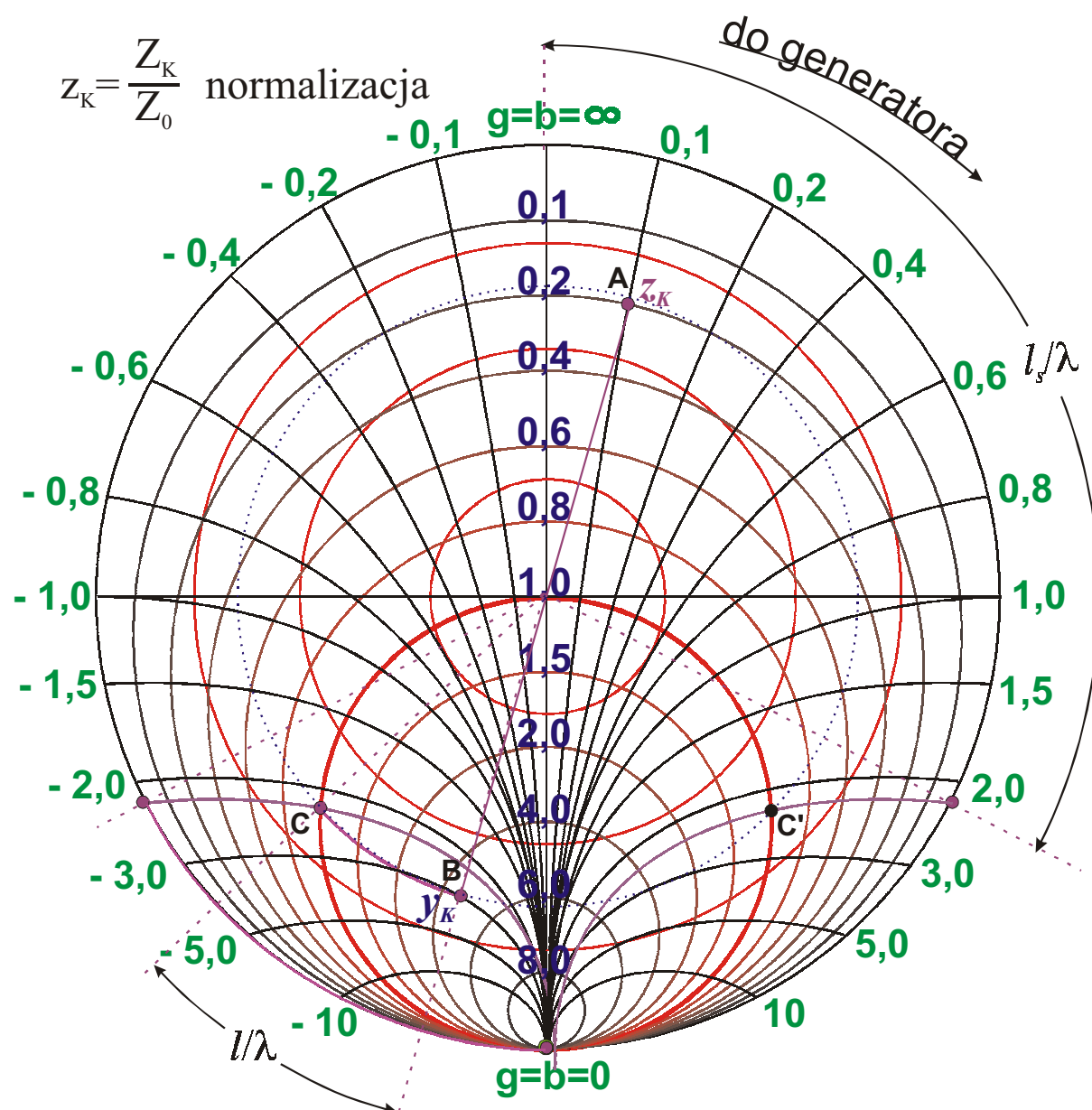


Transformator ćwierćfalowy $Z_{0T} = \sqrt{R_{WE} Z_0}$

Stroik zwarty na końcu

Stroik rozarty na końcu

Dopasowanie za pomocą stroika - inny sposób.

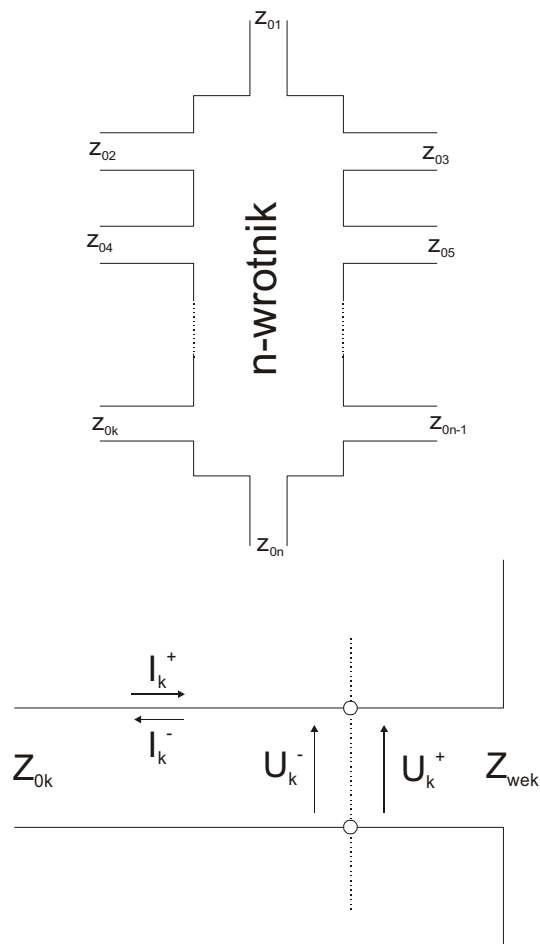


Po przejściu na rachunek admitancyjny (Z punktu A do B) należy przetransformować admitancję obciążenia do miejsca gdzie część rzeczywista jest równa 1 (czyli Z_0 po renormalizacji) – są dwa takie punkty C lub C' i należy wybrać jeden z nich. Susceptancja wejściowa w wybranym punkcie musi być skompensowana przez susceptancję wejściową dołączonego stroika (musi mieć ona znak przeciwny).

Długość stroika rozwartego wyznacza się więc transformując rozwarcie do wartości równej przeciwnej susceptancji (jak na rysunku dla punktu C), zaś stroika zwartego analogicznie transformując zwarcie. Łatwo zauważyć, że różnica długości między stroikami zwartym a rozwartym wyniesie $\lambda/4$.

Tu trzeba pamiętać, że na skutek przejścia na rachunek admitancyjny (odwrócenie wykresu o 180°) zwarcie zamieniło się z rozwarciem.

MACIERZ ROZPROSZENIA



Def.

Znormalizowane zespolone napięcie fali dobiegającej do k-tych wrót:

$$a_k = \frac{U_k^+}{\sqrt{Z_{0k}}}$$

Znormalizowane zespolone napięcie fali wybiegającej z k-tych wrót:

$$b_k = \frac{U_k^-}{\sqrt{Z_{0k}}}$$

Podobnie dla prądów:

$$a_k = I_k^+ \sqrt{Z_{0k}} \quad b_k = I_k^- \sqrt{Z_{0k}}$$

Całkowite napięcie i prąd w linii:

$$U_k = U_k^+ + U_k^- \quad I_k = I_k^+ + I_k^- ,$$

Całkowite znormalizowane napięcie i prąd w linii:

$$\frac{U_k}{\sqrt{Z_{0k}}} = a_k + b_k \quad I_k \sqrt{Z_{0k}} = I_k^+ \sqrt{Z_{0k}} - I_k^- \sqrt{Z_{0k}}$$

Współczynnik odbicia we wrotach k:

$$\Gamma_k = \frac{b_k}{a_k}$$

Unormowana impedancja wejściowa we wrotach k:

$$Z_{wek} = \frac{Z_{wek}}{Z_{0k}} = \frac{a_k + b_k}{a_k - b_k}$$

Dla wielowrotnika liniowego napięcie (prąd) fali wybiegającej z wrót k jest superpozycją fal dobiegających do pozostałych wrót:

$$\left. \begin{aligned} b_1 &= S_{11}a_1 + S_{12}a_2 + \dots + S_{1k}a_k + \dots + S_{1n}a_n \\ b_2 &= S_{21}a_1 + S_{22}a_2 + \dots + S_{2k}a_k + \dots + S_{2n}a_n \\ &\vdots \\ b_k &= S_{k1}a_1 + S_{k2}a_2 + \dots + S_{kk}a_k + \dots + S_{kn}a_n \\ b_n &= S_{n1}a_1 + S_{n2}a_2 + \dots + S_{nk}a_k + \dots + S_{nn}a_n \end{aligned} \right\}$$

$$\begin{aligned} &\quad \quad \quad \text{|||} \\ &\quad \quad \quad [\mathbf{b}] = [\mathbf{S}][\mathbf{a}] \end{aligned}$$

W przypadku ogólnym macierz $\mathbf{S}$ zawiera $l_w = n^2$ niezależnych wyrazów zespolonych. Jeśli na wrota 1 pada fala o amplitudzie a_1 , zaś pozostałe obciążone są impedancjami dopasowanymi (brak fal dobiegających do pozostałych wrót) to powyższy układ równań przyjmie postać:

$$\left. \begin{array}{l} b_1 = S_{11}a_1 \\ \vdots \\ b_k = S_{k1}a_1 \\ \vdots \\ b_n = S_{n1}a_1 \end{array} \right\} \Rightarrow \left. \begin{array}{l} S_{11} = \frac{b_1}{a_1} \\ \vdots \\ S_{k1} = \frac{b_k}{a_1} \\ \vdots \\ S_{n1} = \frac{b_n}{a_1} \end{array} \right\}$$

$S_{ii} = \Gamma_i$ współczynnik odbicia od i-tych wrót

S_{ij} współczynnik transmisji pomiędzy wrotami j-tymi a i-tymi

Szczególne przypadki

Wielowrotniki odwracalne

$$S_{ij} = S_{ji} \quad l_w = \frac{n}{2}(n+1)$$

Wielowrotniki symetryczne

a) pełnosymetryczne	niepełnosymetryczne
$S_{11} = S_{22} = \dots = S_{nn}$	$S_{ik} = S_{jk}$
$S_{ij} = \text{const}$	$S_{ki} = S_{kj}$
$i \neq j$	$S_{ii} = S_{jj}$
	$k \neq i \quad k \neq j$

Wielowrotniki bezstratne

$$\sum_{k=1}^n \frac{(U_k^+)^2}{Z_{ok}} = \sum_{k=1}^n \frac{(U_k^-)^2}{Z_{ok}}$$

(równość sum mocy na wejściowych i wyjściowych)

$$\sum_{k=1}^n \frac{(U_k^+)^2}{Z_{ok}} = \sum_{k=1}^n a_k^2 = a_1 a_1^* + a_2 a_2^* + \dots + a_n a_n^* = \mathbf{a} \mathbf{a}^{*T}$$

$$\sum_{k=1}^n \frac{(U_k^-)^2}{Z_{ok}} = \sum_{k=1}^n b_k^2 = b_1 b_1^* + b_2 b_2^* + \dots + b_n b_n^* = \mathbf{b} \mathbf{b}^{*T}$$

warunek bezstratności $\mathbf{a} \mathbf{a}^{*T} = \mathbf{b} \mathbf{b}^{*T} = \mathbf{a} \mathbf{S} \mathbf{S}^{*T} \mathbf{a}^{*T}$

po przeniesieniu macierzy na lewą stronę $\mathbf{a}^{*T} (\mathbf{E} - \mathbf{S}^{*T} \mathbf{S}) \mathbf{a} = 0$

równanie to jest spełnione dla $\mathbf{S}^{*T} = \mathbf{S}^{-1}$

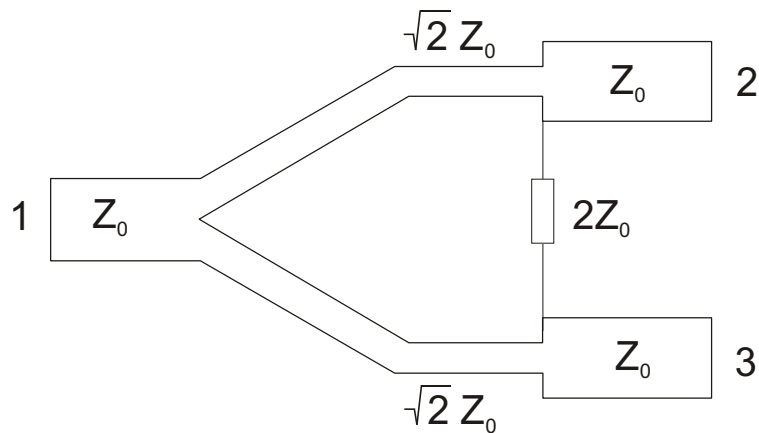
(dla wielowrotnika bezstratnego macierz $\mathbf{S}$ jest unitarna).

Wielowrotniki bezstratne i odwracalne

$$\sum_{k=1}^n S_{ki}^* S_{kl} = \begin{cases} 1 & \text{dla } i = l \\ 0 & \text{dla } i \neq l \end{cases}$$

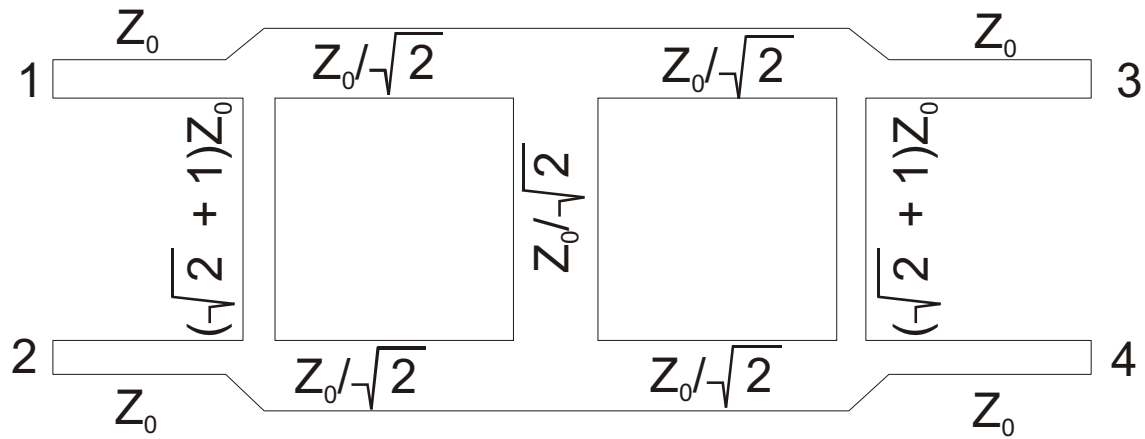
Przykłady:

Jednostopniowy dzielnik Wilkinsona



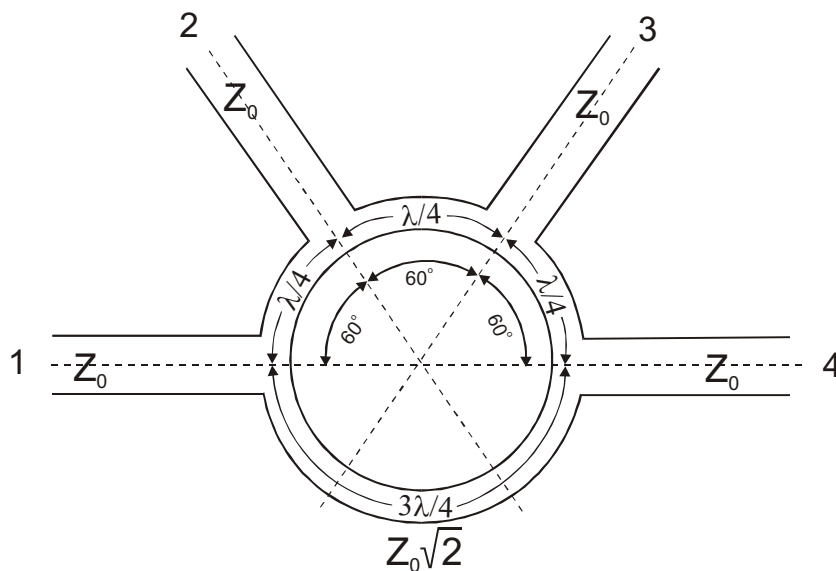
$$\mathbf{S} = \frac{1}{\sqrt{2}} \begin{bmatrix} 0 & -j & -j \\ -j & 0 & 0 \\ -j & 0 & 0 \end{bmatrix}$$

Dwustopniowy sprzęgacz gałęziowy



$$S = -\frac{1}{\sqrt{2}} \begin{bmatrix} 0 & 0 & 1 & j \\ 0 & 0 & j & 1 \\ 1 & j & 0 & 0 \\ j & 1 & 0 & 0 \end{bmatrix}$$

Przykład wyznaczania macierzy rozproszenia



Założenie: układ jest bezstratny i dopasowany od strony wszystkich wrót.

$$S_{11} = S_{22} = S_{33} = S_{44} = 0$$

Moc doprowadzona do wrót 1 dzieli się po równo na sąsiednie wrota 2 i 4 (z fazami przeciwnymi), wrota 3 są izolowane.

$$S_{12} = -\frac{1}{\sqrt{2}}, \quad S_{14} = \frac{1}{\sqrt{2}}, \quad S_{13} = 0$$

Analogicznie dla mocy doprowadzonej do wrót 2 (tym razem fazy zgodne)

$$S_{21} = -\frac{1}{\sqrt{2}}, \quad S_{23} = -\frac{1}{\sqrt{2}}, \quad S_{24} = 0$$

oraz wrót 3 i 4

$$S_{32} = -\frac{1}{\sqrt{2}}, \quad S_{34} = -\frac{1}{\sqrt{2}}, \quad S_{31} = 0$$

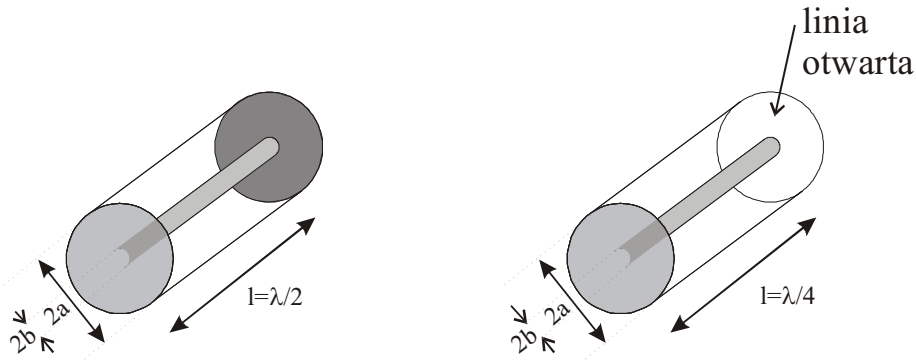
$$S_{41} = \frac{1}{\sqrt{2}}, \quad S_{43} = -\frac{1}{\sqrt{2}}, \quad S_{42} = 0$$

$$\mathbf{S} = \frac{1}{\sqrt{2}} \begin{bmatrix} 0 & -1 & 0 & 1 \\ -1 & 0 & -1 & 0 \\ 0 & -1 & 0 & -1 \\ 1 & 0 & -1 & 0 \end{bmatrix}$$

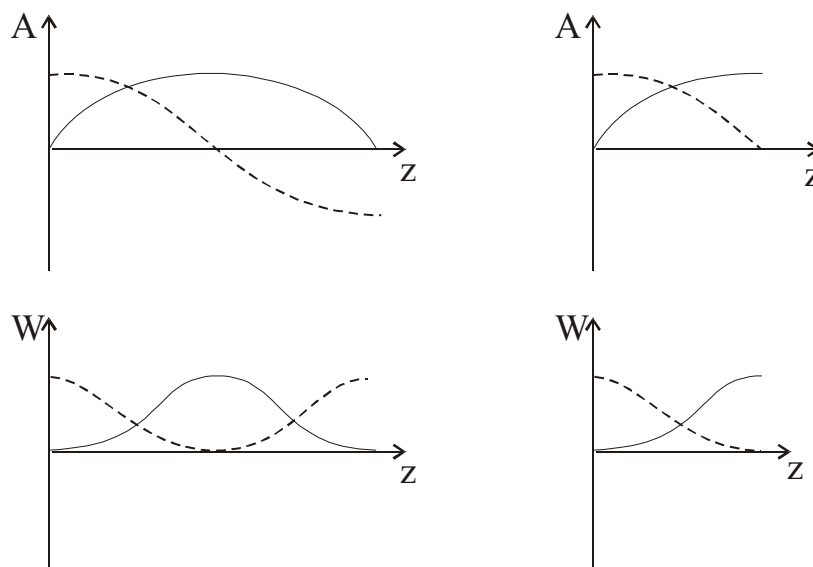
REZONATORY MIKROFALOWE

Rezonator mikrofalowy to odcinek linii transmisyjnej obustronnie lub jednostronnie zwarty (rozarty) na końcu. Wektory E i H wewnątrz takiego „pudełka” spełniają jednorodne równania falowe.

Przykłady rezonatorów

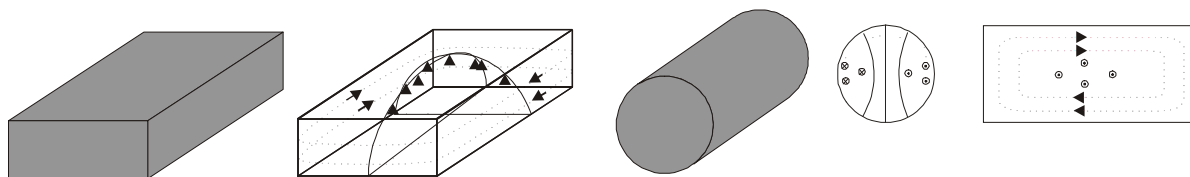


Rezonans dla długości fali $\lambda_{rez} = \frac{2l}{p}$ gdzie p oznacza ilość połówek fali mieszczących się wewnątrz rezonatora. Stosownie do tego rezonatory oznacza się TE_{mnp} lub TM_{mnp} . Dla przykładu rezonator o długości $\lambda/2$ będący fragmentem falowodu prostokątnego z podstawowym rodzajem pola TE oznaczany jest TE_{101} .

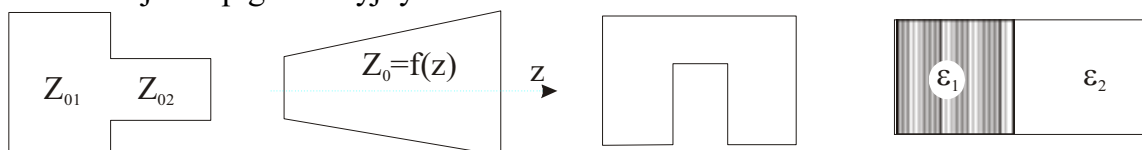


Rozkład amplitud A i energii W w falowodach półfalowym (po lewej) i ćwierćfalowym (po prawej).

Rezonatory falowodowe jednorodne $\lambda/2$ z podstawowym rodzajem pola



Rezonatory falowodowe niejednorodne geometrycznie i materiałowo umożliwiają uzyskanie złożonych charakterystyk rezonansowych i stosowane są najczęściej w konstrukcji lamp generacyjnych.



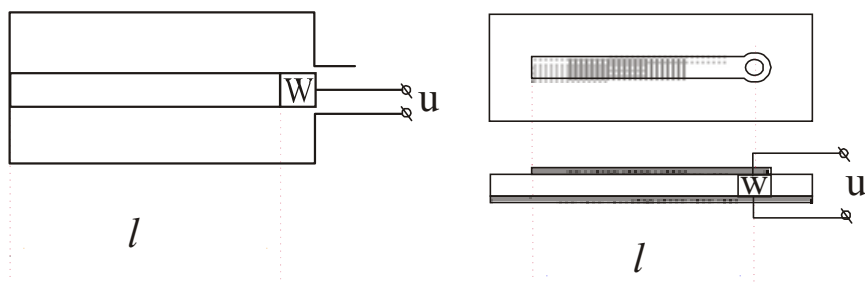
[Patrz też rezonatory z akustyczną falą powierzchniową.](#)

Rezonatory strojone

mechanicznie

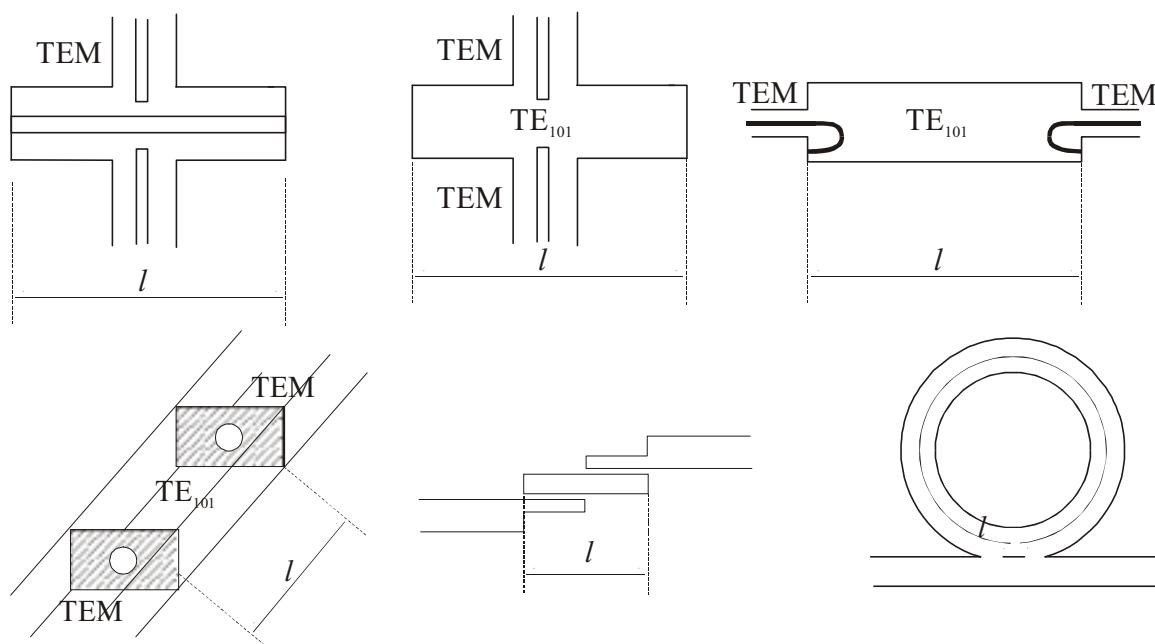
elektronicznie

Przykład rezonatora strojonego elektronicznie



Rezonator z diodą waraktorową zmieniającą „elektryczną” długość paska

Sposoby sprzęgania rezonatorów z liniami przesyłowymi



Dobroć rezonatorów mikrofalowych

Po odłączeniu źródła zasilania energia zmagazynowana wewnątrz rezonatora i w jego obwodach zewnętrznych ulega dyssypacji zgodnie z zależnością:

$$W(t) = W_0 \exp\left(-\frac{\omega_0 t}{Q}\right) \quad (1)$$

W_0 – energia początkowa

ω_0 – pulsacja rezonansowa

Q – dobroć

Moc tracona w jednostce czasu wyrazi się poprzez:

$$P = -\frac{dW}{dt} \quad (2)$$

Z (1) i (2) wynika, że:

$$Q = \omega_0 \frac{W_0}{P} \quad (3)$$

Zależnie od tego, jaką część obwodu opisują powyższe wyrażenia można mówić o trzech rodzajach dobroci:

wewnętrznej Q_0 związanej z mocą traconą wewnątrz rezonatora

zewnętrznej Q_Z związanej z mocą traconą w obwodach zewnętrznych

całkowitej Q_L związanej z sumą ww. traconych mocy

$$Q_0 = \omega_0 \frac{W_0}{P_R}, \quad Q_Z = \omega_0 \frac{W_0}{P_Z}, \quad Q_L = \omega_0 \frac{W_0}{P_R + P_Z}, \quad \frac{1}{Q_L} = \frac{1}{Q_0} + \frac{1}{Q_Z}$$

Def. współczynnika sprzężenia: $\beta = \frac{Q_0}{Q_Z}$

$\beta < 1$ – sprzężenie podkrytyczne

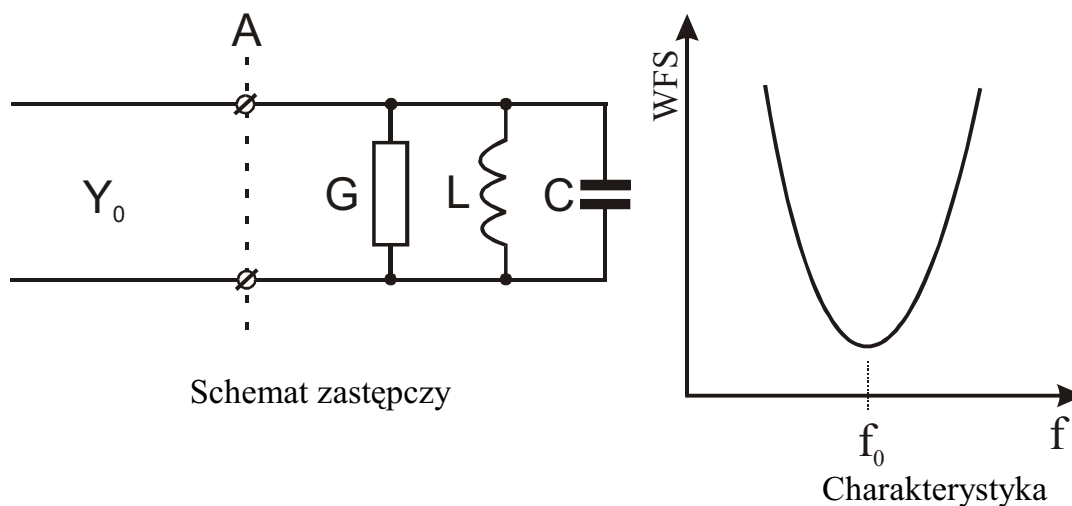
$\beta = 1$ – sprzężenie krytyczne

$\beta > 1$ – sprzężenie nadkrytyczne

Def. współczynnika rozstrojenia: $\alpha = Q_L \left(\frac{\omega}{\omega_0} - \frac{\omega_0}{\omega} \right)$

Trzy rodzaje sprzężeń rezonatorów mikrofalowych

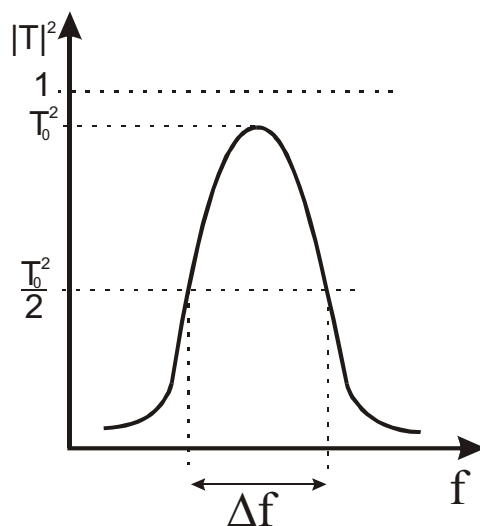
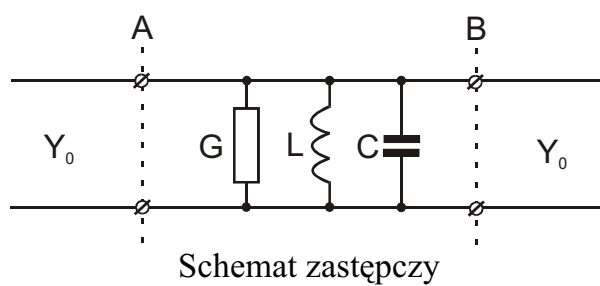
odbiciowe



$$Q_L = \frac{\omega_0 C}{Y_0 + G}, \quad Q_Z = \frac{\omega_0 C}{Y_0}, \quad \beta = \frac{Y_0}{G},$$

$$\Gamma = \frac{\Gamma_0 - j\alpha}{1 + j\alpha}, \quad \Gamma_0 = \frac{\beta - 1}{\beta + 1}$$

transmisyjne



Charakterystyka

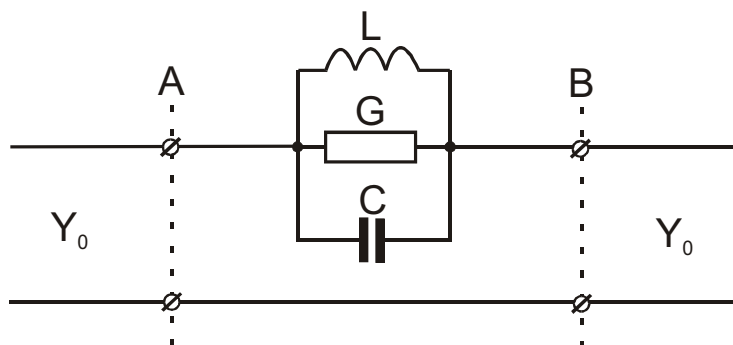
$$Q_0 = \frac{\omega_0 C}{G}, \quad Q_L = \frac{\omega_0 C}{G + 2Y_0} = \frac{Q_0}{1 + 2\beta},$$

$$\Gamma = \frac{\Gamma_0 - j\alpha}{1 + j\alpha}, \quad \Gamma_0 = \frac{-1}{1 + 2\beta}$$

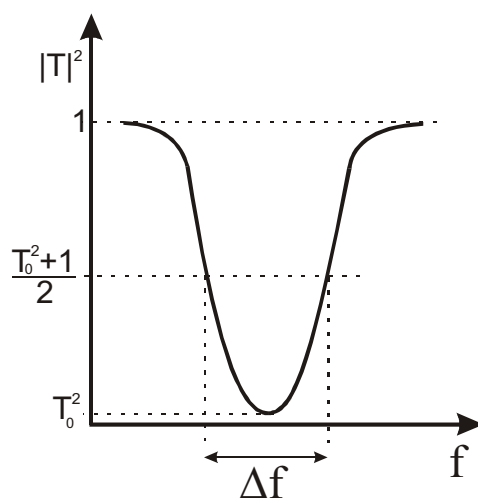
Transmisja

$$T = \frac{T_0}{1 + j\alpha}, \quad \text{w rezonansie } T_0 = \frac{2\beta}{1 + 2\beta}, \quad S = \begin{bmatrix} \Gamma & T \\ T & \Gamma \end{bmatrix}$$

reakcyjne



Schemat zastępczy



Charakterystyka

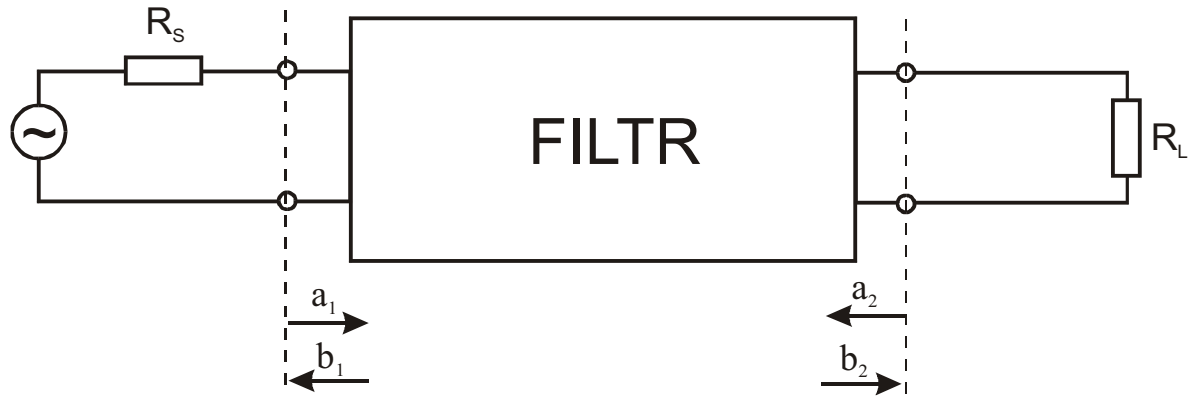
$$\beta = \frac{Y_0}{2G}, \quad Q_0 = \frac{\omega_0 C}{G}, \quad Q_L = \frac{\omega_0 C}{G + Y_0/2} = \frac{Q_0}{1 + \beta}$$

$$\Gamma = \frac{\Gamma_0 + j\alpha}{1 + j\alpha}, \quad \Gamma_0 = \frac{1}{1 + \beta},$$

Reakcja

$$R = \frac{R_0}{1 + \beta}, \quad \text{w rezonansie } R_0 = \frac{\beta}{1 + \beta}, \quad R + T = 1$$

FILTRY MIKROFALOWE



PARAMETRY

straty wnoszone (in. wtrąceniowe – ang. insertion loss):

$$IL = -10 \log \frac{P_L}{P_{SA}} = -10 \log \frac{|b_2|^2}{|a_1|^2} = -10 \log |S_{21}|^2 \text{ [dB]}$$

P_{SA} – moc dysponowana generatorem

P_L – moc czynna wydzielana na obciążeniu

straty odbiciowe (ang. return loss):

$$RL = -10 \log \frac{P_R}{P_{SA}} = -10 \log \left(\frac{WFS - 1}{WFS + 1} \right)^2 = -10 \log |\Gamma|^2 \text{ [dB]}$$

P_R – moc odbita od wejścia filtru

przesunięcie fazy (ang. phase shift):

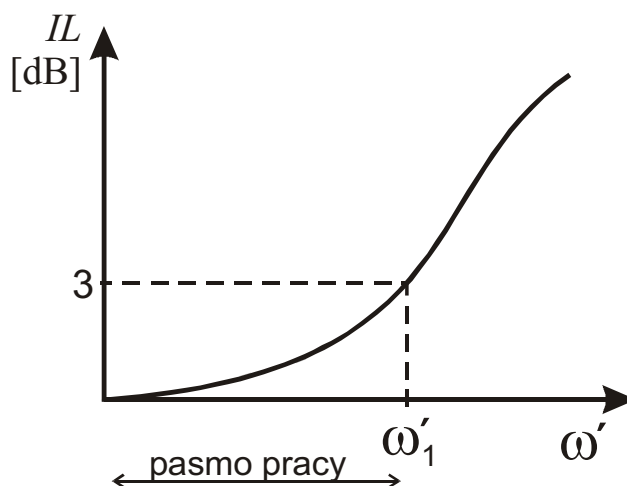
$$\Phi_T = \arg \{ S_{21} \}$$

opóźnienie grupowe (ang. group delay):

$$\tau_D = \frac{d\Phi_T}{d\omega} = \frac{1}{2\pi} \frac{d\Phi_T}{df}$$

Typowe charakterystyki (na przykładzie filtra dolnoprzepustowego)

charakterystyka płaska

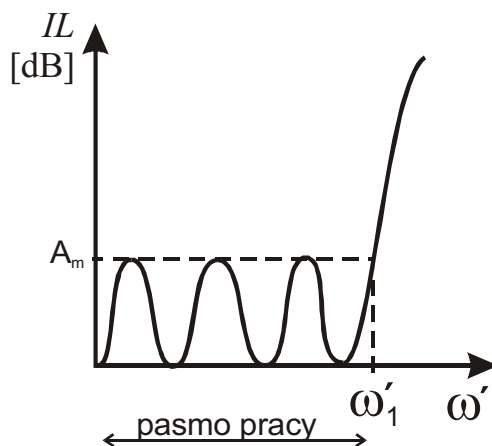


$$IL = 10 \log(1 + \omega'^{2n}) \text{ [dB]}$$

n - rząd filtra

ω' – pulsacja unormowana

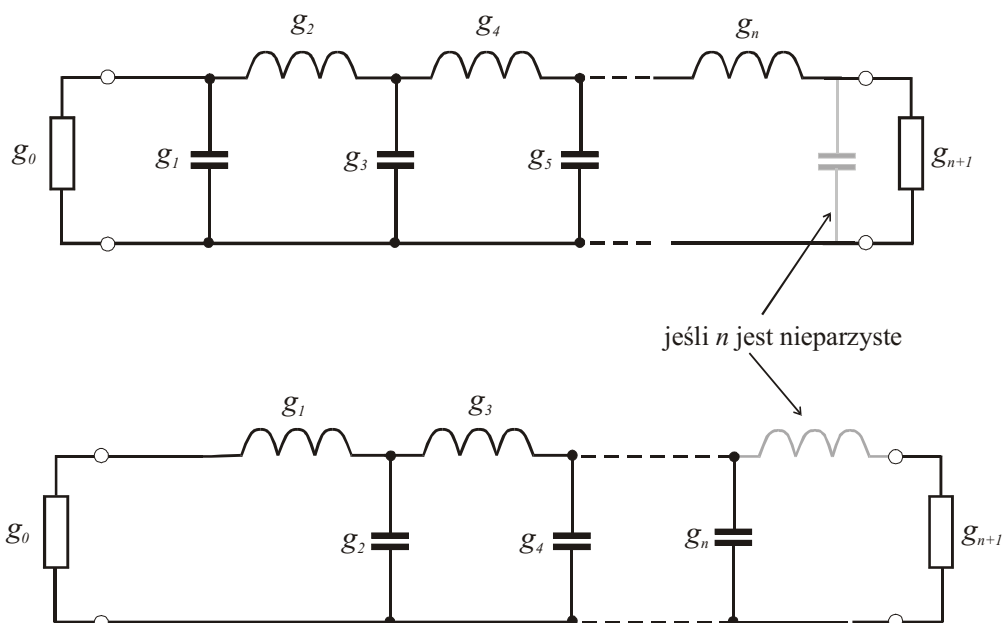
charakterystyka równomiernie falista



$$\text{dla } |\omega'| \leq 1 \quad IL = 10 \log[1 + (10^{A_m/10} - 1) \cos^2(n \arccos \omega')] \text{ [dB]}$$

$$\text{dla } |\omega'| \geq 1 \quad IL = 10 \log[1 + (10^{A_m/10} - 1) \cosh^2(n \operatorname{arc} \cosh \omega')] \text{ [dB]}$$

Wyznaczanie elementów filtrów



Elementy g można obliczyć z następujących zależności:
dla $g_0 = 1$ i $\omega' = 1$ (3 dB pulsacja graniczna)

charakterystyka płaska

$$g_k = 2 \sin \left[\frac{(2k-1)\pi}{2n} \right], \quad k = 1, 2, 3 \dots n$$

$$g_{n+1} = 1 \text{ dla wszystkich } n$$

charakterystyka równomiernie falista

$$g_1 = \frac{2a_1}{\gamma}, \quad \gamma = \sinh \frac{\beta}{2n}, \quad \beta = \ln \left(\operatorname{ctgh} \frac{A_m}{17,37} \right)$$

$$g_k = \frac{4a_{k-1}a_k}{b_{k-1}g_{k-1}} \quad k = 2, 3 \dots n$$

$$a_k = \sin \left[\frac{(2k-1)\pi}{2n} \right], \quad k = 1, 2 \dots n \quad b_k = \gamma^2 + \sin^2 \frac{k\pi}{n}, \quad k = 1, 2 \dots n$$

$$g_{n+1} = \operatorname{tgh} \frac{\beta}{4} \text{ dla } n \text{ parzystych} \quad g_{n+1} = 1 \text{ dla } n \text{ nieparzystych}$$

Znajomość parametrów g_n pozwala określić rzeczywiste wartości odpowiadających im reaktancji.

Transformacja częstotliwości

Dla filtrów dolnoprzepustowych: $\omega' = \frac{\omega}{\omega_G}$,

gdzie ω_G jest górną częstotliwością graniczną.

Dla filtrów górnoprzepustowych: $\omega' = -\frac{\omega_D}{\omega}$,

gdzie ω_D jest dolną częstotliwością graniczną.

Dla filtrów środkowoprzepustowych:

$$\omega' = \frac{\omega_0}{\omega_D - \omega_G} \left(\frac{\omega}{\omega_0} - \frac{\omega_0}{\omega} \right).$$

Szerokość pasma wynosi: $\Delta\omega_0 = \sqrt{\omega_D \omega_G}$

Dla filtrów środkowozaporowych:

$$\omega' = \frac{\omega_D - \omega_G}{\omega_0 \left(\frac{\omega}{\omega_0} - \frac{\omega_0}{\omega} \right)}.$$

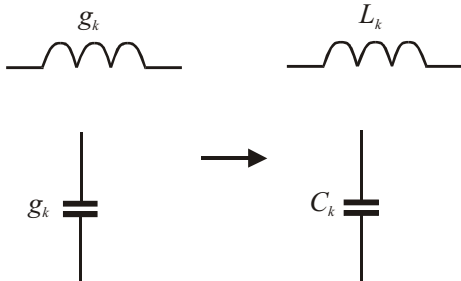
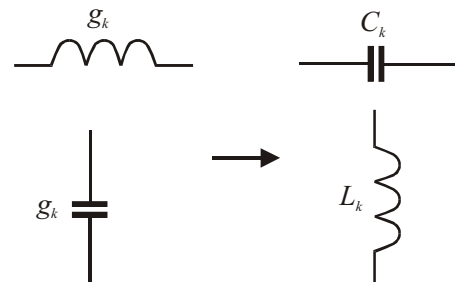
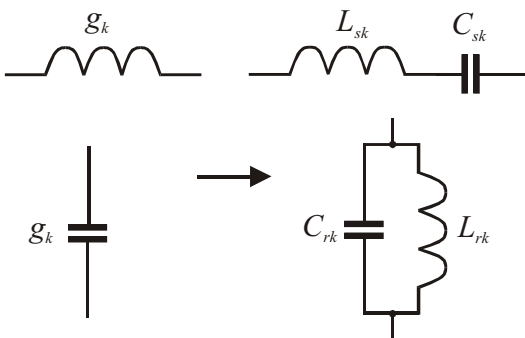
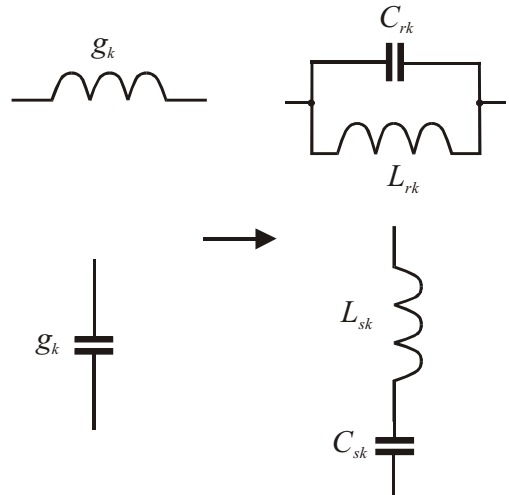
Szerokość pasma wynosi: $\Delta\omega_0 = \sqrt{\omega_D \omega_G}$

Skalowanie wartości elementów

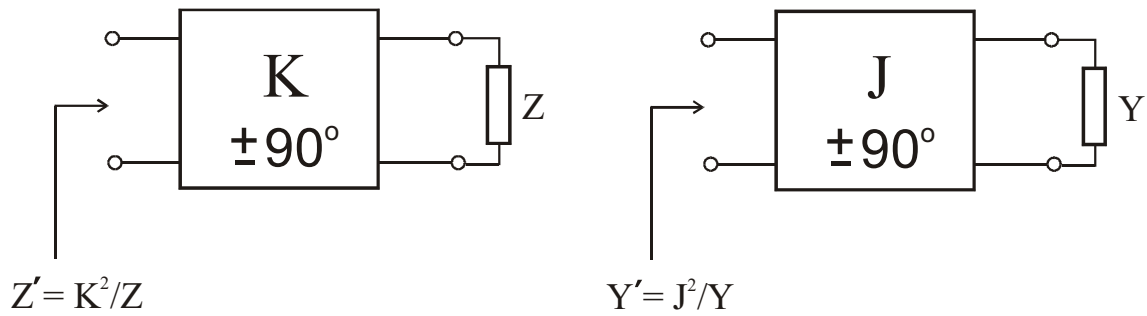
Realizuje się poprzez $R_0 = R/r$ krotną zmianę wartości parametrów obliczonych.

$r = g_0$, R - rzeczywista impedancja generatora.

Transformacja częstotliwości i skalowanie impedancji

rodzaje transformacji Dolnoprzepustowego filtra prototypowego		transformacja elementów	wartości reaktancji rzeczywistych
	FDP		$L_k = R_0 \frac{g_k}{\omega_G}$ $C_k = \frac{g_k}{R_0 \omega_G}$
	FGP		$C_k = \frac{1}{R_0 \omega_D g_k}$ $L_k = \frac{R_0}{\omega_D g_k}$
	FPP		$L_{sk} = R_0 \frac{g_k}{\omega_G - \omega_D}$ $C_{sk} = \frac{\omega_G - \omega_D}{R_0 \omega_0^2 g_k}$ $C_{rk} = \frac{g_k}{R_0 (\omega_G - \omega_D)}$ $L_{sk} = R_0 \frac{\omega_G - \omega_D}{\omega_0^2 g_k}$
	FPZ		$C_{rk} = \frac{1}{R_0 (\omega_G - \omega_D) g_k}$ $L_{rk} = R_0 \frac{g_k (\omega_G - \omega_D)}{\omega_0^2}$ $C_{sk} = \frac{g_k (\omega_G - \omega_D)}{R_0 \omega_0^2}$ $L_{sk} = \frac{R_0}{g_k (\omega_G - \omega_D)}$

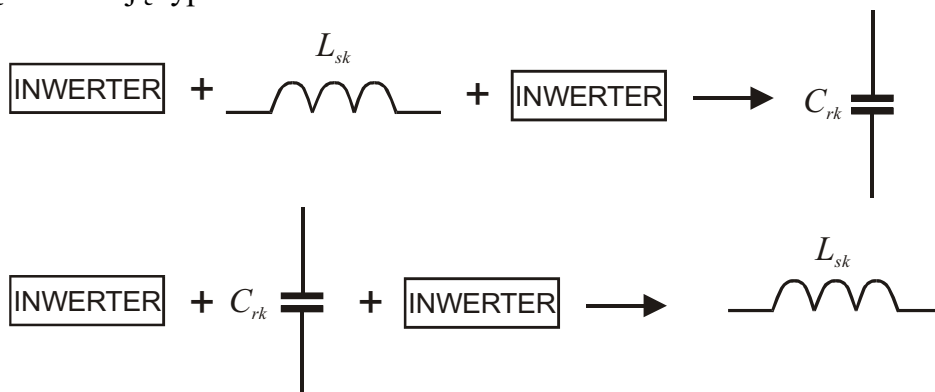
Inwertery immitancji



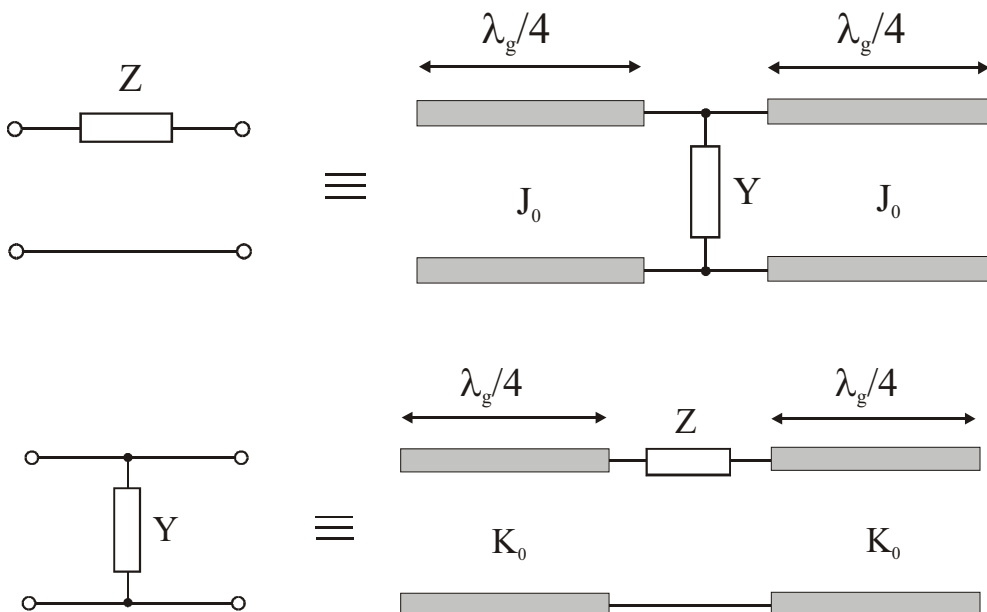
K – współczynnik inwersji impedancji

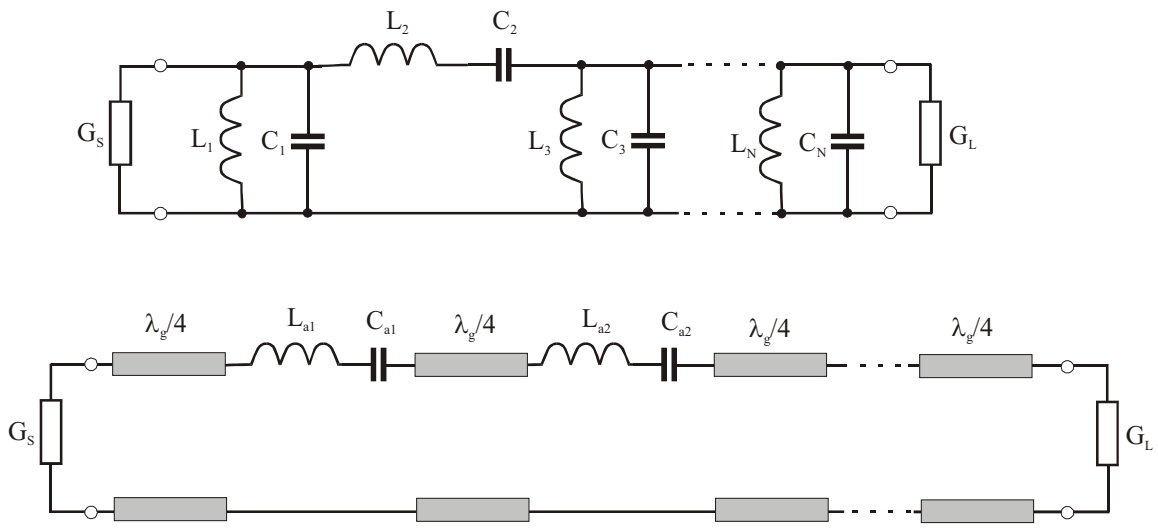
J – współczynnik inwersji admitancji

Realizują konwersję typu:

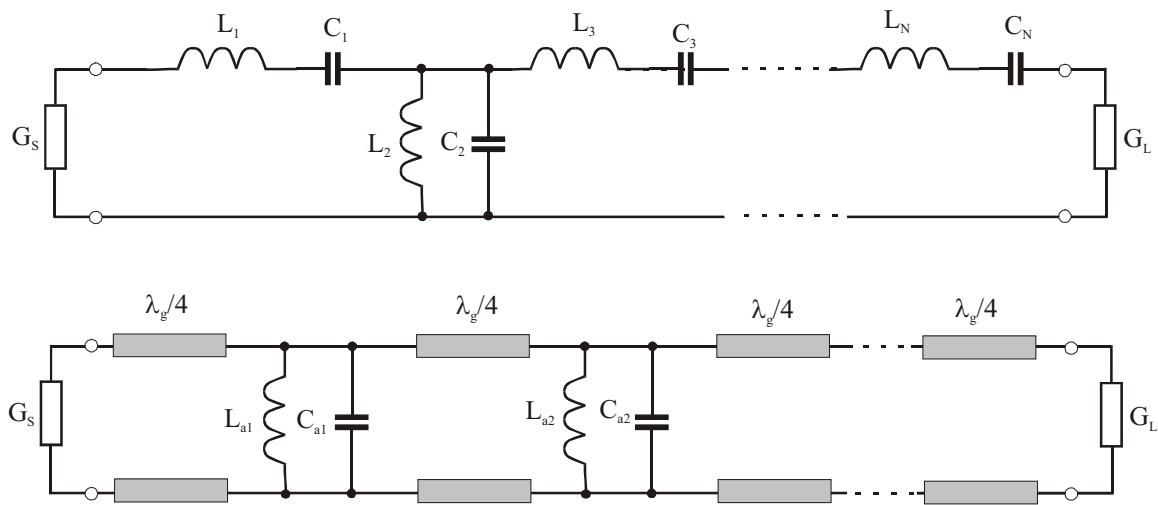


Dzięki inwerterom można realizować filtry z elementów reaktancyjnych tylko jednego rodzaju, co w praktyce daje możliwość konstruowania filtrów mikrofalowych o stałych rozłożonych.





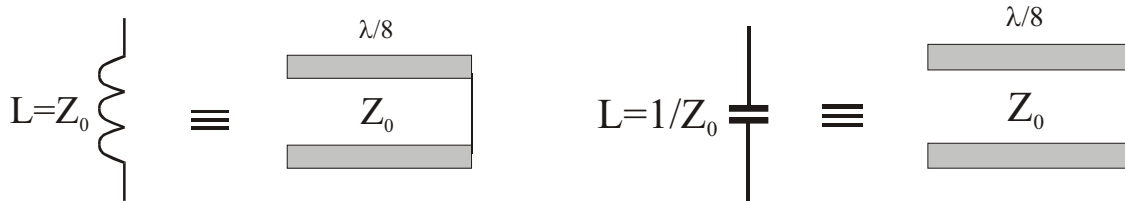
Realizacja filtru na elementach szeregowych



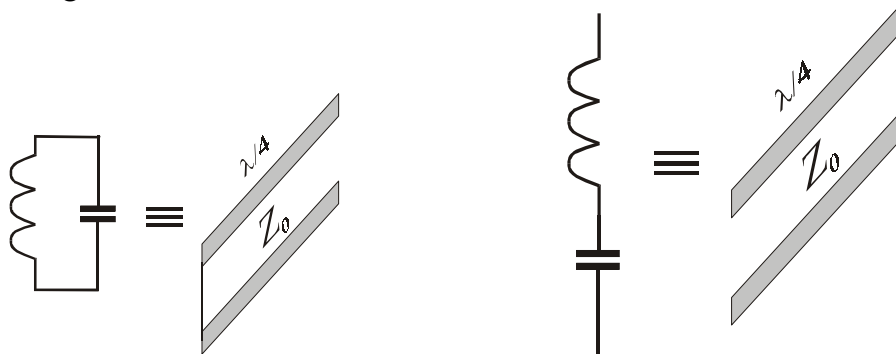
Realizacja filtru na elementach równoległych

Przykłady realizacji reaktancji w zakresie mikrofalowym

linie szeregowie

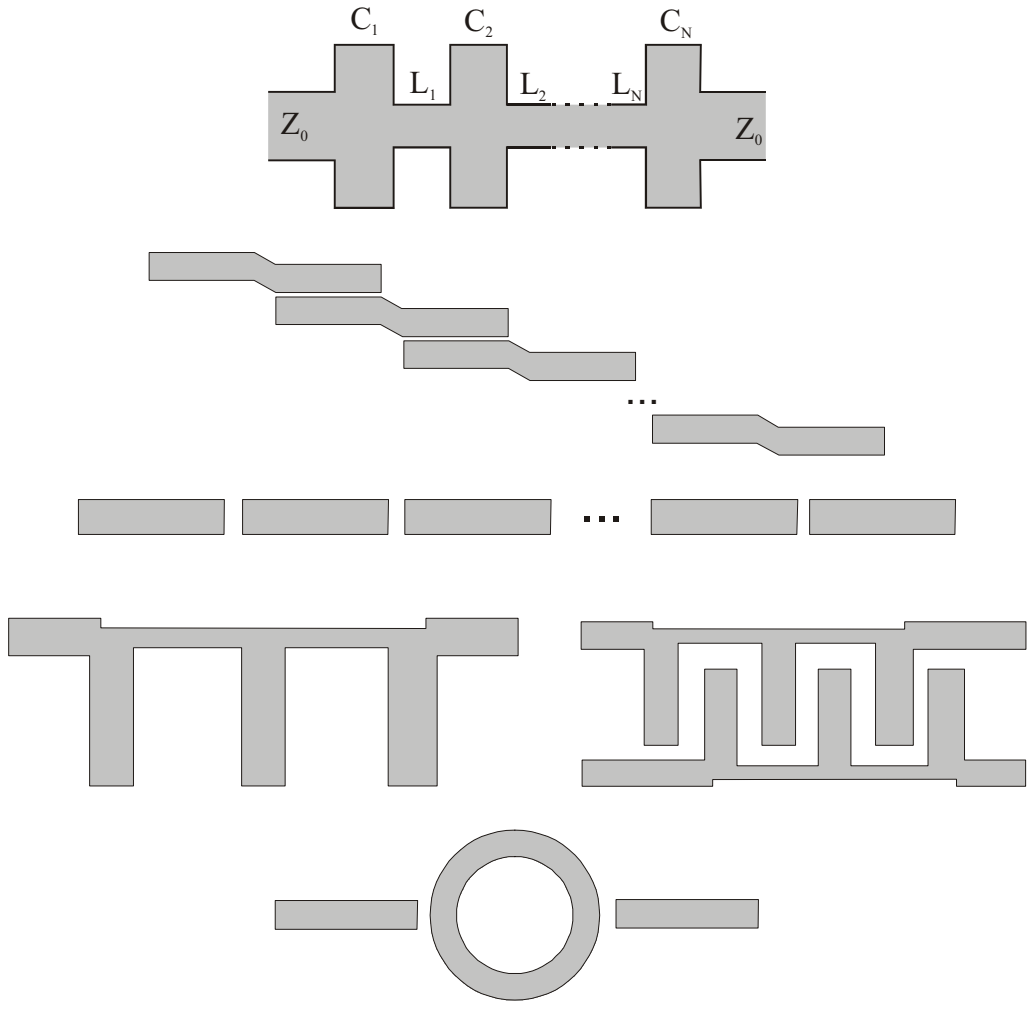
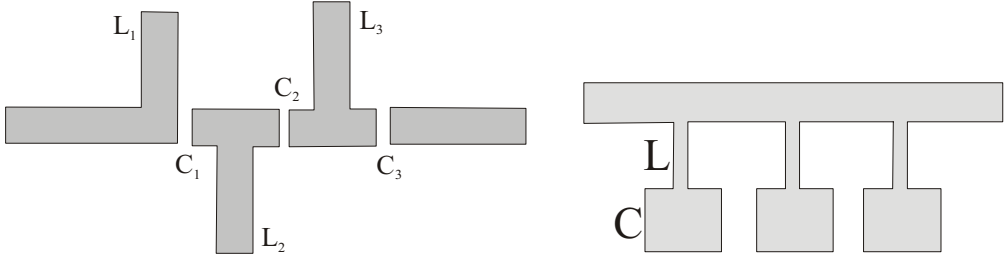


linie równoległe

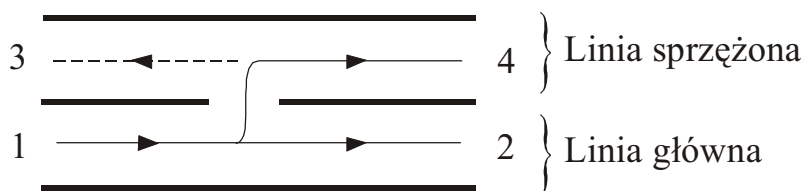


Patrz także $\Rightarrow$ [własności linii transmisyjnych](#) oraz [niejednorodności linii przesyłowych](#).

Przykłady praktycznych realizacji filtrów

środkowoprzepustowe	
środkowozaporowe	

SPRZĘGACZE KIERUNKOWE



Parametry

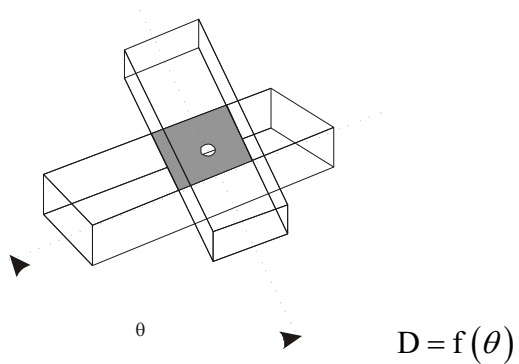
Sprzężenie $C = 10 \log \frac{P_1}{P_4} \text{ [dB]}$

Kierunkowość $D = 10 \log \frac{P_4}{P_3} \text{ [dB]}$

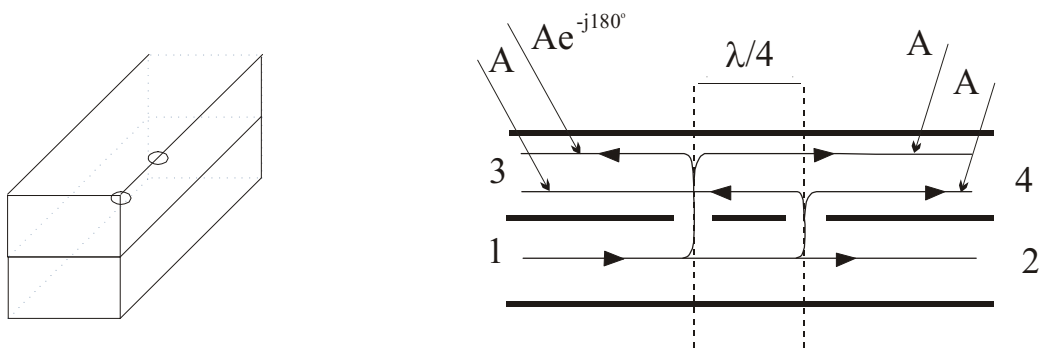
Izolacja $I = 10 \log \frac{P_1}{P_3} \text{ [dB]}$

Przykłady

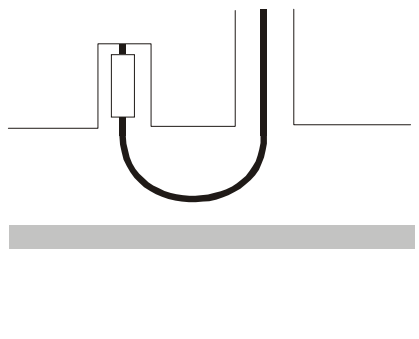
Sprzęgacz jednootworowy typu Bethe



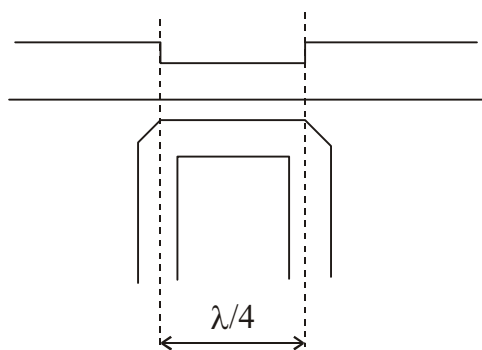
Sprzęgacz dwuotworowy



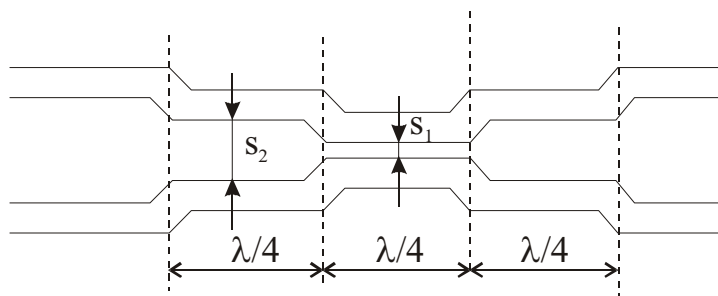
Sprzęgacz na linii współosiowej



Sprzęgacze na NLP



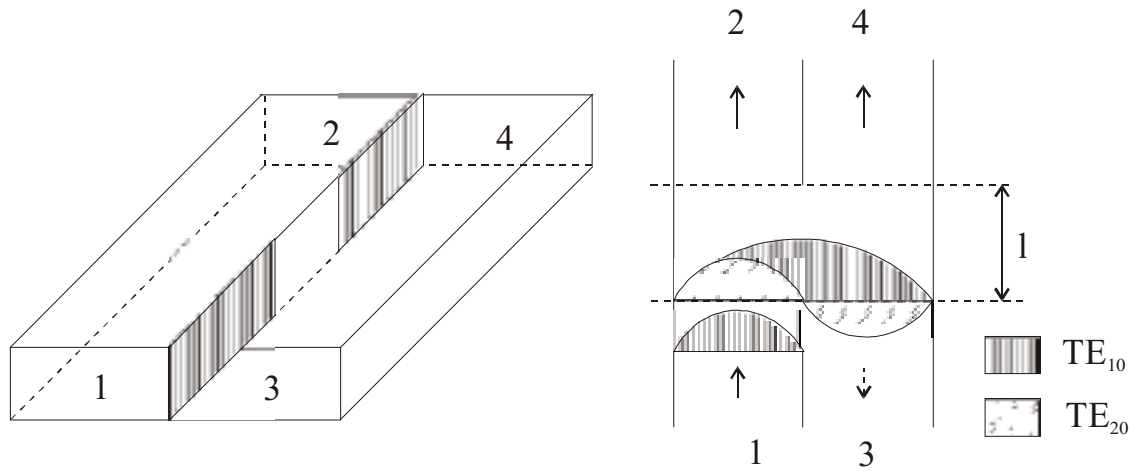
ćwierćfalowy



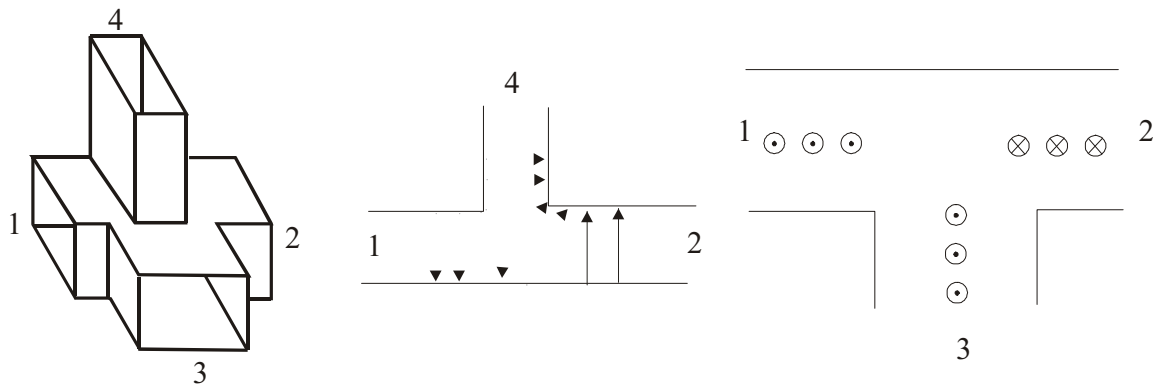
3/4 falowy

Sprzęgacze 3 dB

Sprzęgacz o krótkiej szczelinie



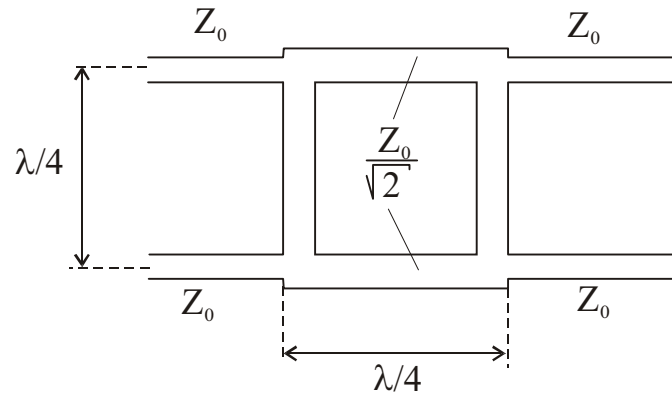
Układ "magiczne T"



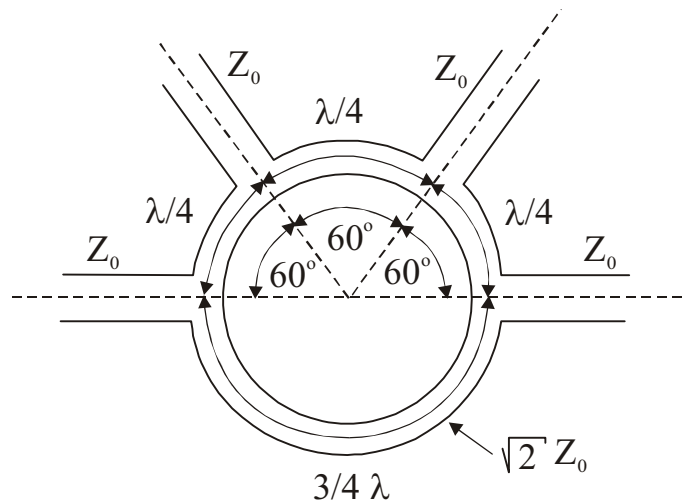
Sprzęgacze 3 dB w technice NLP

Działanie poniższych układów (pokazanych już w rozdziale dotyczącym macierzy rozproszenia) odbywa się na zasadzie konstruktywnej bądź destruktywnej interferencji fal w poszczególnych wrotach. Przy ich analizie wystarczy pamiętać, że na dystansie $\lambda/4$ faza fali przesuwa się o $\pi/2$ czyli o 90° .

Równoramienny układ hybrydowy

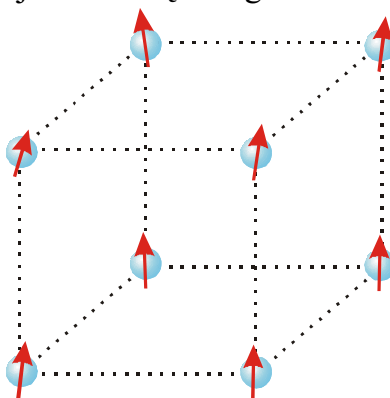


Pierścień hybrydowy



ELEMENTY FERRYTOWE

Materiałami najczęściej spotykanymi w przyrodzie są antyferromagnetyki, w których spiny (momenty pędu) elektronowe ustawione są tak, że wypadkowy moment magnetyczny jest równy zero. Zdarzają się jednak materiały, w których spiny elektronowe ustawione są zgodnie, a co za tym idzie wszystkie lokalne momenty magnetyczne mają zbliżoną wielkość i ten sam średni kierunek. W związku z tym ciało takie wykazuje niezerowy zewnętrzny moment magnetyczny reagujący na pole magnetyczne. Pole to może być zarówno pochodzenia wewnętrznego jak i zewnętrznego.



Materiały wykazujące takie cechy nazwane zostały ferromagnetykami. Własności takie posiadają np. mieszaniny (spieki) tlenku żelaza Fe_2O_3 z tlenkami Zn, Ni, Mn, Mg lub metali ziem rzadkich zwane ferrytami.

Ich przenikalność magnetyczna jest dość skomplikowana i reprezentowana zwykle przez zestaw liczb (tensor):

$$\mathbf{B} = \begin{bmatrix} \mu & jk & 0 \\ -jk & \mu & 0 \\ 0 & 0 & \mu_0 \end{bmatrix} \mathbf{H}.$$

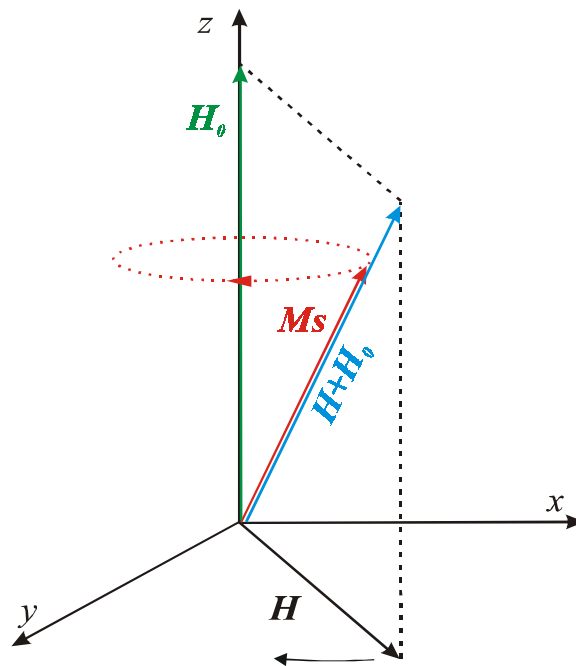
Charakteryzują się one stałą przenikalnością elektryczną $\varepsilon = 10 \dots 20$ i dużą rezystywnością $10 \dots 10^5 \Omega\text{m}$.

Wybrane zjawiska fizyczne zachodzące w ferrytach

Rezonans ferromagnetyczny

Spiny elektronowe pod wpływem pola magnetycznego H_0 (dowolnego pochodzenia) wykonują ruch precesyjny i taki też ruch wykonuje generowany przez nie wypadkowy moment magnetyczny M_s .

Jeśli przez takie ciało przechodzić będzie fala elektromagnetyczna z wektorem H rotującym w płaszczyźnie xy w kierunku zgodnym z kierunkiem precesji to wypadkowy wektor $H+H_0$ pokryje się z M_s i także będzie podlegał precesji. Fakt ten uzewnętrznia się pochłanianiem energii fali elektromagnetycznej.



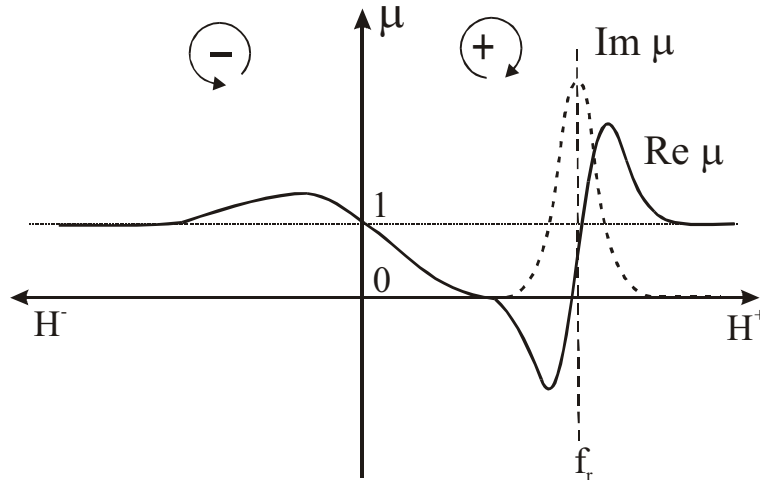
Pochłanianie to będzie najsilniejsze przy odpowiednim (widocznym na rys.) dobraniu wartości H_0 , H oraz częstotliwości wirowania H (równej częstotliwości precesji). Stan taki nazywa się rezonansem ferromagnetycznym. Ponieważ precesja zależy od H_0 to przybliżeniu można przyjąć, że częstotliwość rezonansowa wynosi:

$$f_r \approx 35 \cdot 10^3 H_0$$

Zjawiska te zachodzą tylko dla kierunku wirowania H zgodnego z kierunkiem precesji. W przeciwnym wypadku ferromagnetyk zachowuje się jak zwykły dielektryk. Kierunek rotacji H zależy od kierunku propagacji fali, jeśli zatem

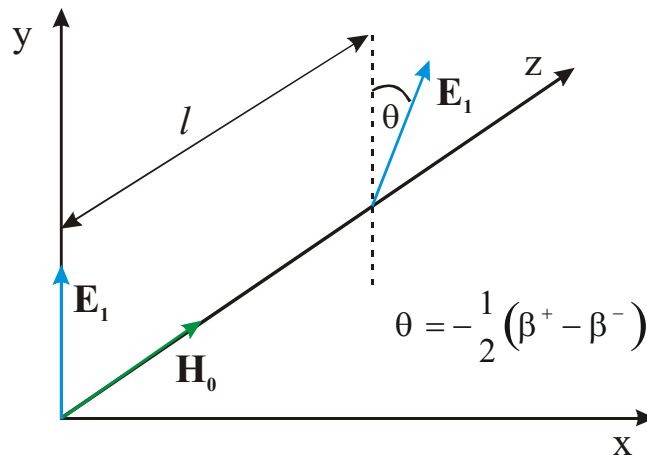
dla określonego kierunku propagacji odnotujemy rezonans, to dla kierunku przeciwnego zjawisko to nie będzie zachodzić.

W rezonansie i jego pobliżu przenikalność magnetyczna ferrytu jest zespolona, a jej część urojona odpowiada za pochłanianie energii fali elektromagnetycznej.



Rotacja Faradaya

Zjawisko to polega na skręcaniu płaszczyzny polaryzacji fali o pewien kąt po przejściu przez materiał ferrytowy poddany działaniu w pola magnetycznego (w kierunku jego linii sił).



Każdą falę elektromagnetyczną można rozłożyć na dwie przeciwnie rotujące składowe. Skręcenie płaszczyzny polaryzacji jest skutkiem różnych prędkości tych składowych ($\mu^+ \neq \mu^-$ czyli $c^+ \neq c^-$)

Kąt skręcenia można określić z zależności:

$$\theta = V l H_0$$

gdzie V jest stałą Verdetą zależną od rodzaju ferrytu, l długością ferrytu, zaś H_0 natężeniem zewnętrznego pola magnetycznego.

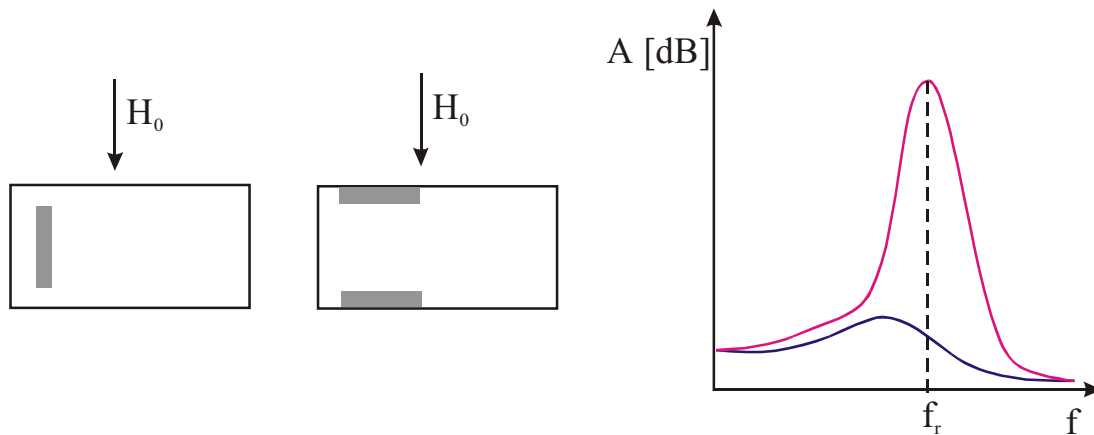
Kąt θ nie zależy od kierunku propagacji fali, co oznacza, że zwielokrotnia się po wielokrotnym przejściu fali w jedną i przeciwną stronę. Dodatnia dla większości ferrytów wartość V świadczy o tym, że patrząc wzdłuż kierunku propagacji fali elektromagnetycznej, zgodnie ze zwrotem linii sił zewnętrznego pola magnetycznego obserwuje się skręcenie płaszczyzny polaryzacji zgodne z ruchem wskazówek zegara.

Omówione tu pokrótce zjawiska można wykorzystać do konstrukcji mikrofalowych elementów ferrytowych.

Izolatory ferrytowe

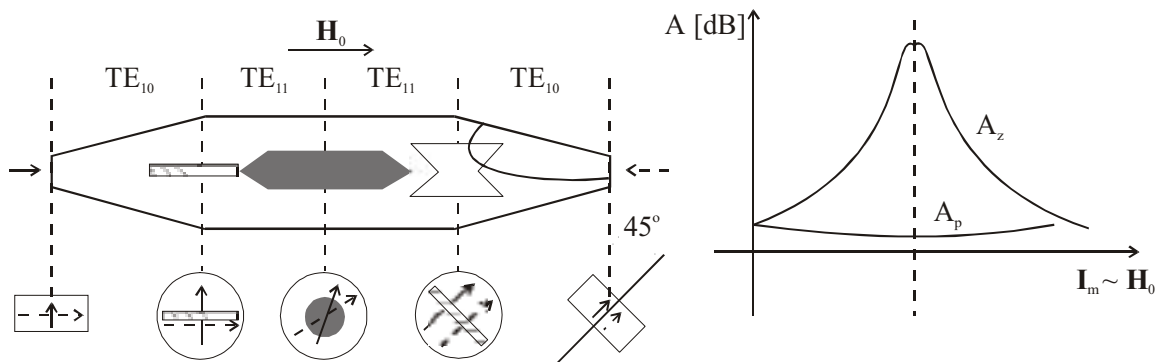
Izolator jest dwuwrotnikiem nieodwracalnym przepuszczającym falę w jednym (przepustowym), a tłumiącym w przeciwnym (zaporowym) kierunku. Umożliwia on separację falową układów mikrofalowych.

Przykłady izolatorów wykorzystujących rezonans magnetyczny



Zasada działania izolatorów wykorzystujących rezonans magnetyczny bazuje na fakcie silnego tłumienia fali w okolicy częstotliwości rezonansowej. Zjawisko to następuje tylko dla kierunku wirowania wektora H zgodnego z kierunkiem precesji, a zatem tylko dla określonego kierunku przepływu fali.

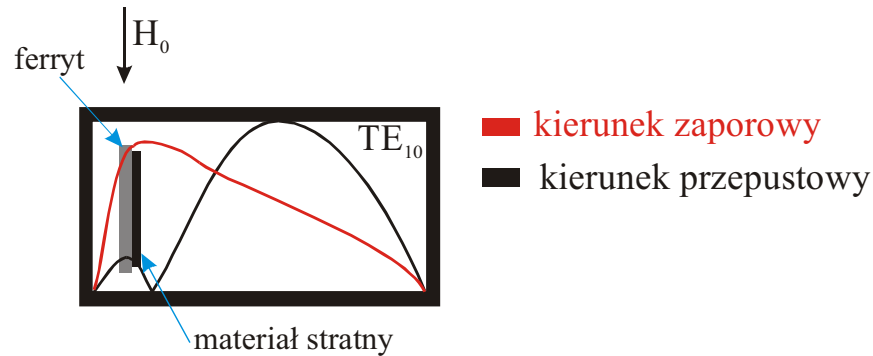
Izolator z efektem rotacji Faradaya i jego charakterystyka



Izolator ten składa się z trzech odcinków falowodów, każdy z podstawowym rodzajem pola. Widoczna w centrum odcinka falowodu cylindrycznego wkładka ferrytowa oraz stałe pole magnetyczne są tak dobrane, aby fala po pokonaniu wkładki doznała skręcenia swojej płaszczyzny polaryzacji o kąt równy 45° . Fala pochodząca z lewych wrót (zgodnie z rys.) po przejściu przez obszar z wkładką ferrytową dozna skręcenia płaszczyzny polaryzacji o kąt 45° i przedostanie się następnie do wrót prawych, które są skręcone względem lewych właśnie o 45° . Fala będzie więc mogła propagować się dalej. W kierunku przeciwnym do 45° -ego skręcenia falowodu prawego (względem lewego) doda się skręcenie płaszczyzny polaryzacji wprowadzane przez wkładkę ferrytową. Z tego powodu płaszczyzna polaryzacji będzie we wrotach lewych usytuowana poprzecznie do wymaganej dla dalszej propagacji fali, która w tym wypadku nie jest już możliwa. Fala zatem zostaje w całości odbita, a widoczne po obu stronach wkładki ferrytowej płytki mają za zadanie całkowite pochłonięcie jej energii.

Ze względu na wymagany kierunek stałego pola magnetycznego (wzdłuż wkładki ferrytowej) pole to uzyskuje się zwykle za pomocą elektromagnesu umieszczonego na wkładce. Kierunek i wartość izolacji w takim układzie będą zależeć od kierunku i wartości prądu w elektromagnesie. Ze względu na tę zależność można mówić o prądzie optymalnym, przy którym parametry izolatora z efektem rotacji Faradaya są najlepsze.

Izolator z przemieszczeniem pola



Fakt różnej przenikalności magnetycznej dla przeciwnych kierunków wirowania wektora H został wykorzystany w izolatorze z przemieszczeniem pola. Bliżej jednej z krótszych ścianek odcinka falowodu prostokątnego z podstawowym rodzajem pola TE_{10} umieszczona jest wkładka ferrytowa z naniesionym nań materiałem stratnym (płytką rezystywną pochłaniająca promieniowanie elektromagnetyczne).

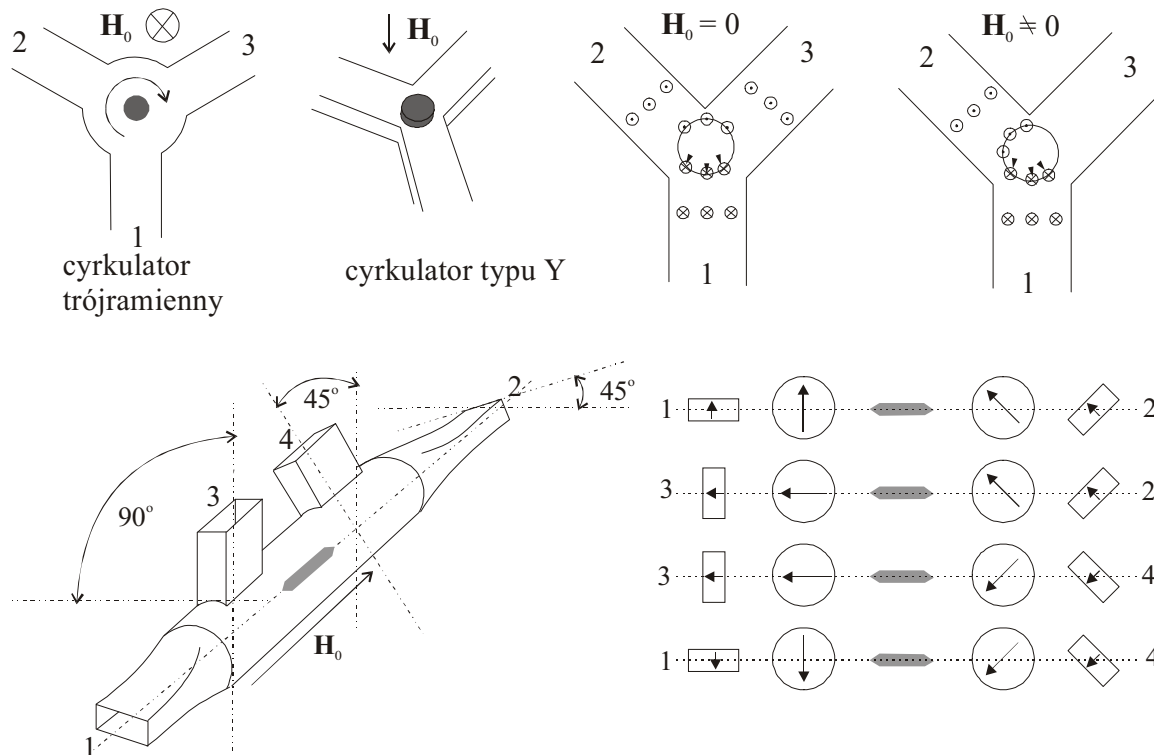
W jednym kierunku przepływu fali wkładka zachowuje się jak przeciętny dielektryk o $\mu \approx 1$. Rozkład pola magnetycznego, a co za tym idzie i elektrycznego ulega w tym kierunku tylko niewielkiemu zniekształceniu (czerwona linia na rys). W tym kierunku maksimum natężenia pola elektrycznego przypada na wkładkę ferrytową, na której umieszczona jest płytka rezystywna pochłaniająca energię fali. Jest to kierunek zaporowy, w którym fala elektromagnetyczna jest silnie tłumiona. W przeciwnym zaś kierunku propagacji fali (czarna linia na rys.) przenikalność magnetyczna μ osiąga duże wartości i rozkład pola zniekształca się tak, że jego minimum przypada na płytkę rezystywną i fala nie jest tłumiona.

Ponieważ nie bazuje się tu na rezonansie, a pracuje na zakresie krzywej $\mu(H^+)$, dla której wartości μ są duże, uzyskiwane pasmo pracy izolatora z przemieszczeniem pola jest zdecydowanie szersze od pasma pracy izolatorów z rezonansem magnetycznym.

Cyrkulatory

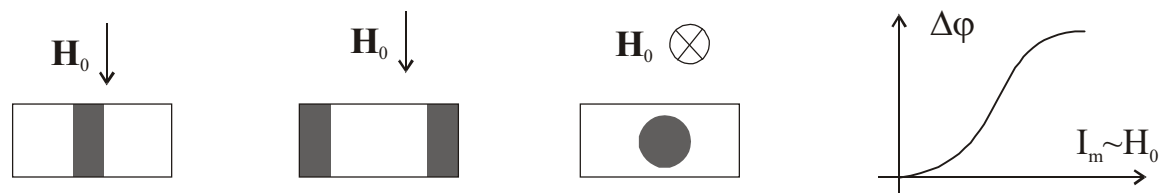
Są to najczęściej trójwrotniki symetryczne, w których moc cyркуluje między wrotami 1-2 zaś 3 są izolowane, 2-3, 1 są izolowane itd.

Przykłady konstrukcji i zasada działania



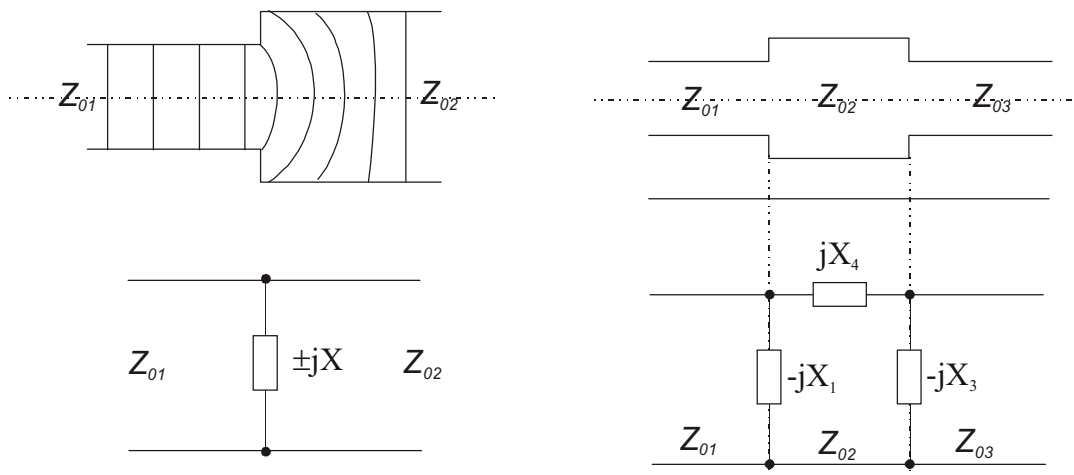
Cyrkulator z efektem rotacji Faradaya

Ferrytowe przesuwniki fazy (bazujące na dużej wartości μ)



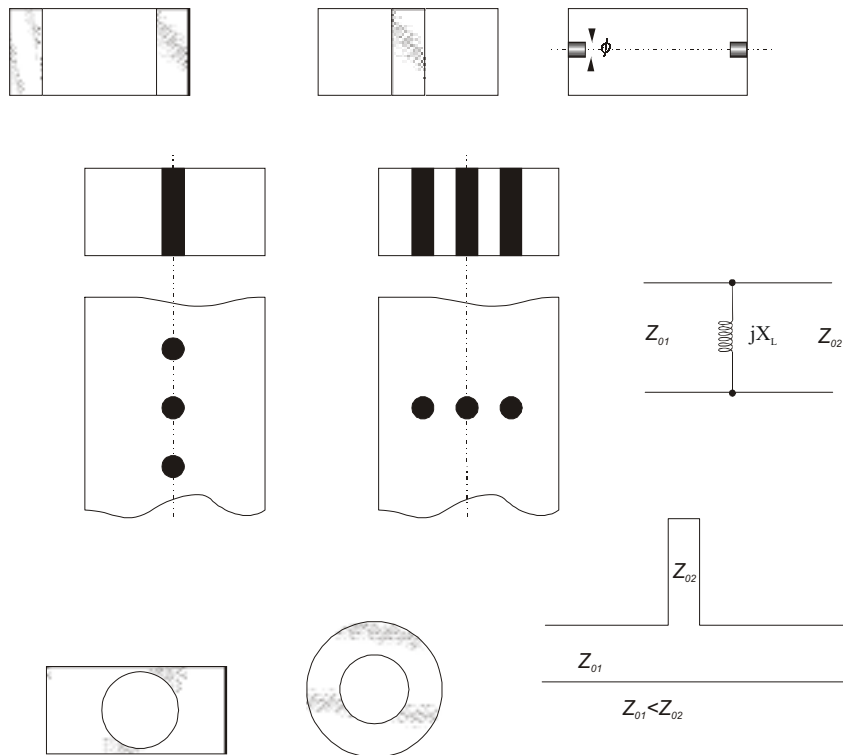
NIECIĄGŁOŚCI W LINIACH TRANSMISYJNYCH

Jakiegolwiek zniekształcenie linii transmisyjnej materiałowe bądź geometryczne powoduje powstanie pewnej reaktancji.

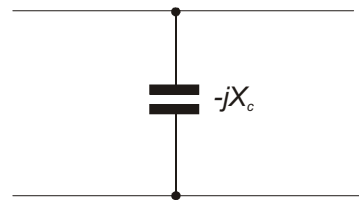
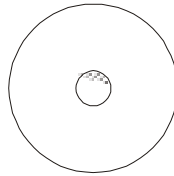
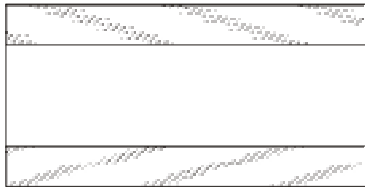


Nieciągłości o charakterze indukcyjnym

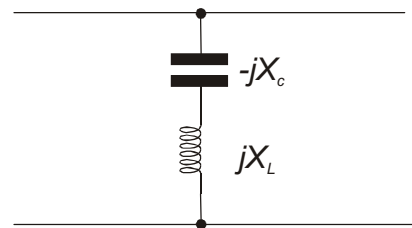
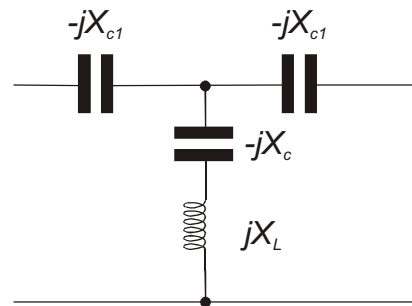
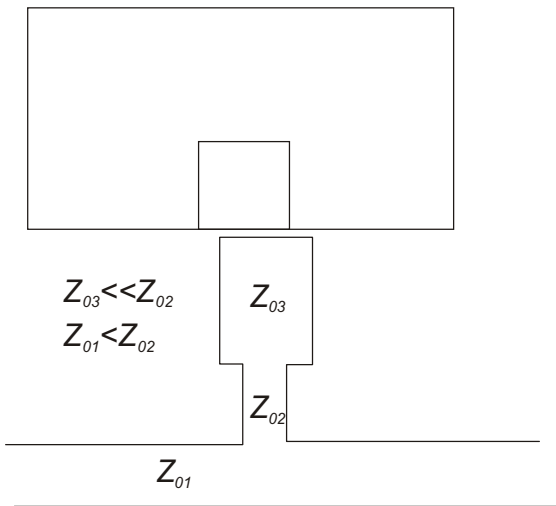
Czarnym kolorem zaznaczono elementy metalowe, szarym wkładki dielektryczne.



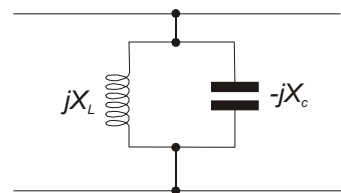
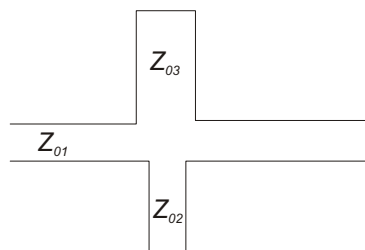
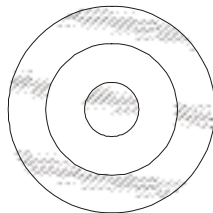
Nieciągłości o charakterze pojemnościowym



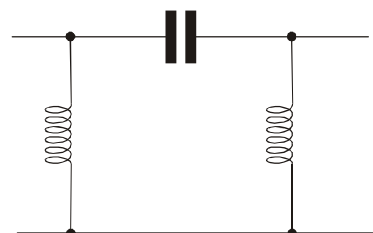
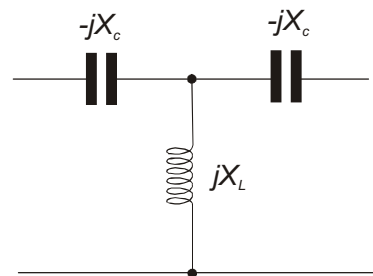
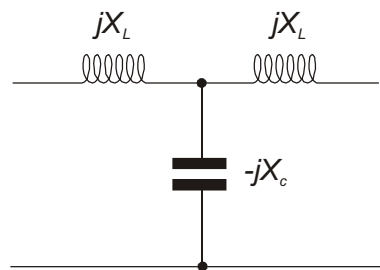
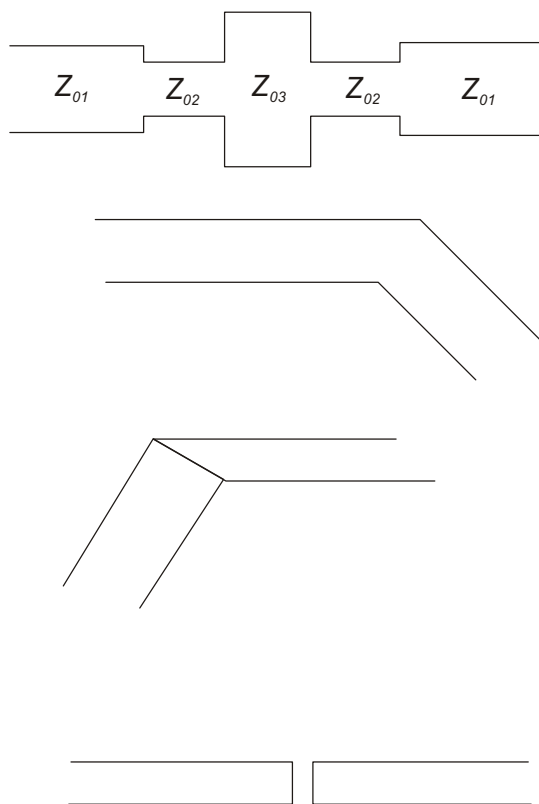
Nieciągłości o charakterze rezonansu szeregowego



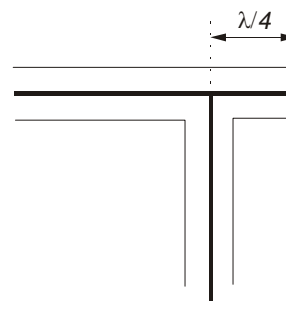
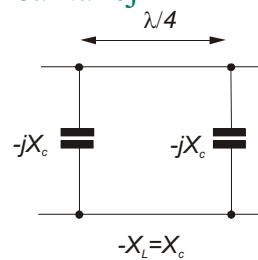
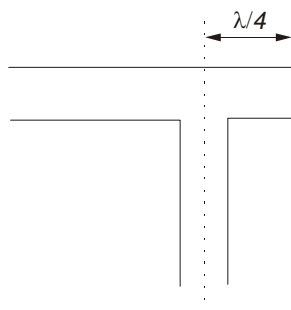
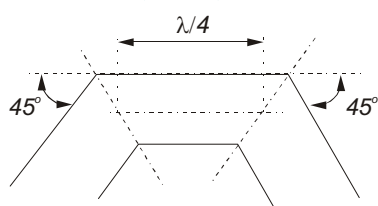
Nieciągłości o charakterze rezonansu równoległego



Nieciągłości o charakterze filtrów



Zaginanie linii przesyłowych z kompensacją reaktancji

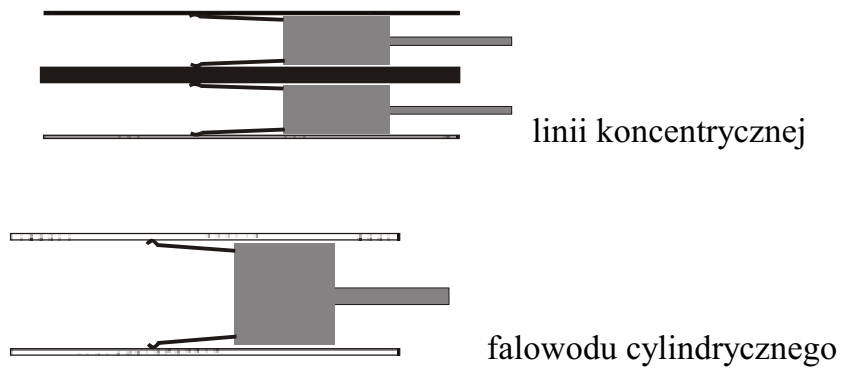


WYBRANE ELEMENTY PASYWNE TORU MIKROFALOWEGO

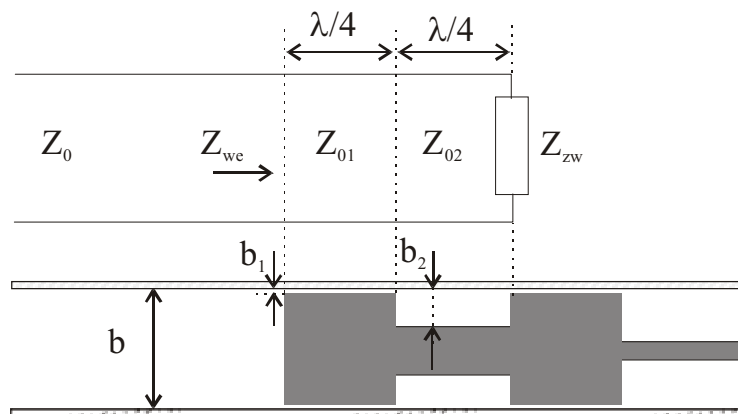
Przesuwnik fazy bazujący na dużej wartości ϵ



Zwieraki nastawne (umożliwiają strojenie długości linii)



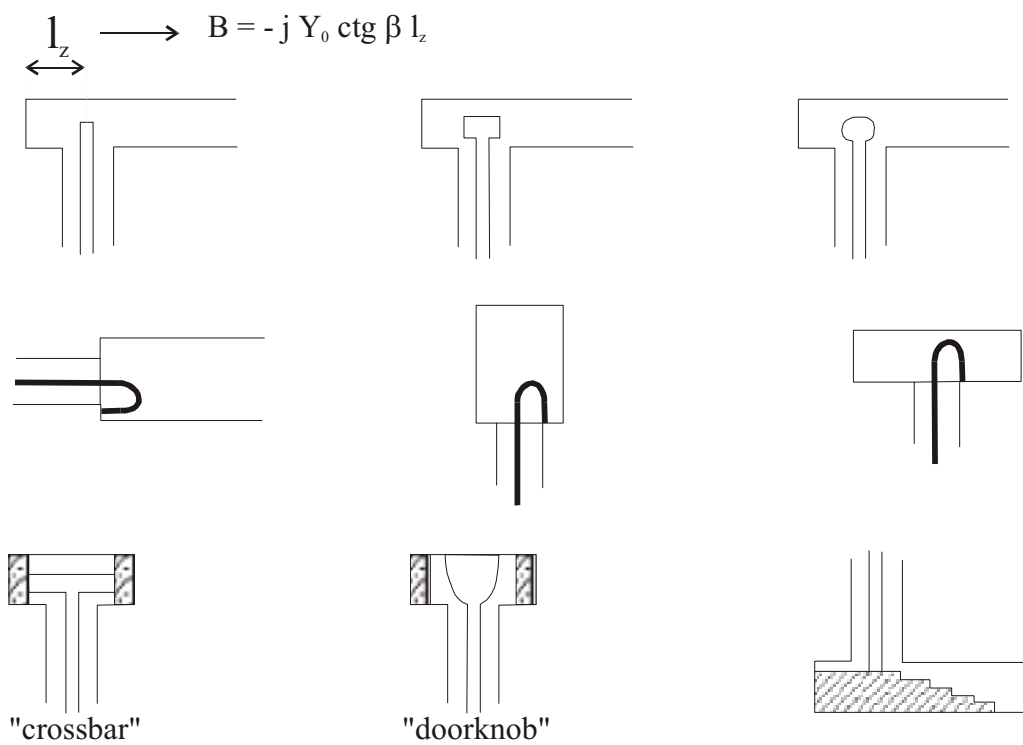
Zwierak z dwoma transformatorem ćwierćfalowymi



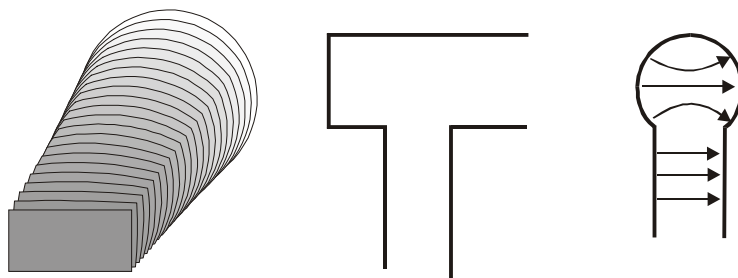
$$Z_{we} = \left(\frac{b_1}{b_2} \right)^2 Z_{zw} = \left(\frac{Z_{01}}{Z_{02}} \right)^2 Z_{zw}$$

Przykłady przejść między liniami różnego typu

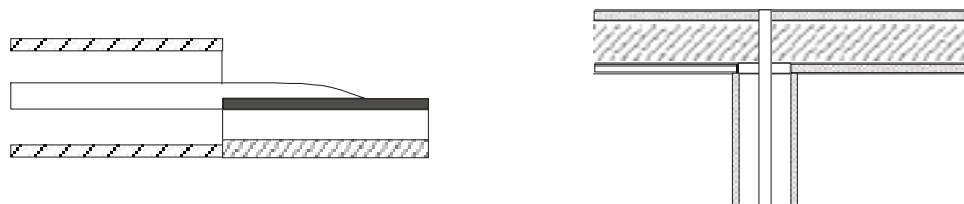
Przejścia linia koncentryczna – falowód



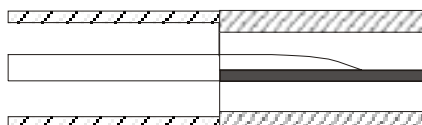
Przejścia falowód – falowód



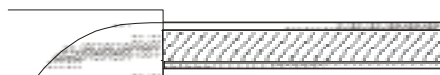
Przejścia linia koncentryczna –NLP



Przejście linia koncentryczna –SLP

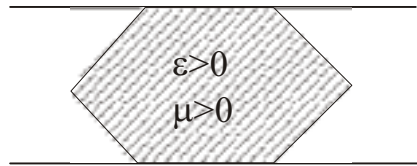


Przejście falowód – NLP

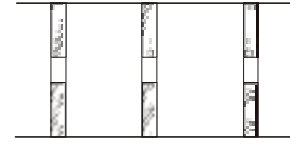
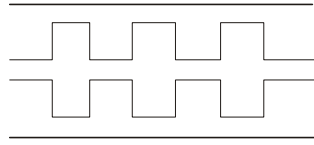
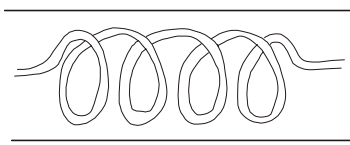


Linie opóźniające

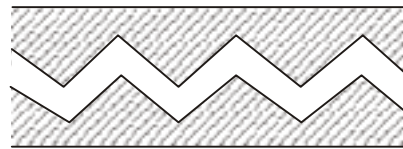
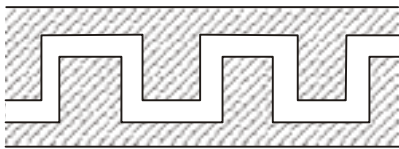
- jednorodne (bazujące na materiałach o dużym ε lub μ)



- periodyczne (niejednorodne) – „elektrycznie wydłużone”



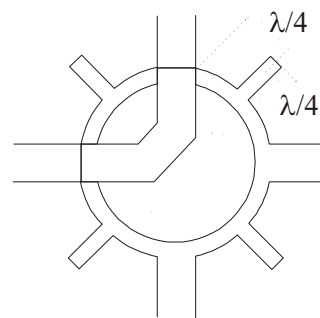
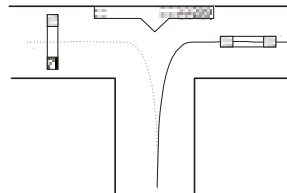
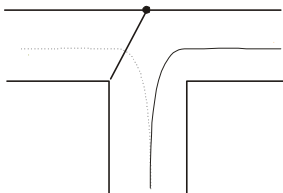
- zwinięte „fizycznie wydłużone”



- akustyczne ([patrz ostatni rozdział skryptu](#))

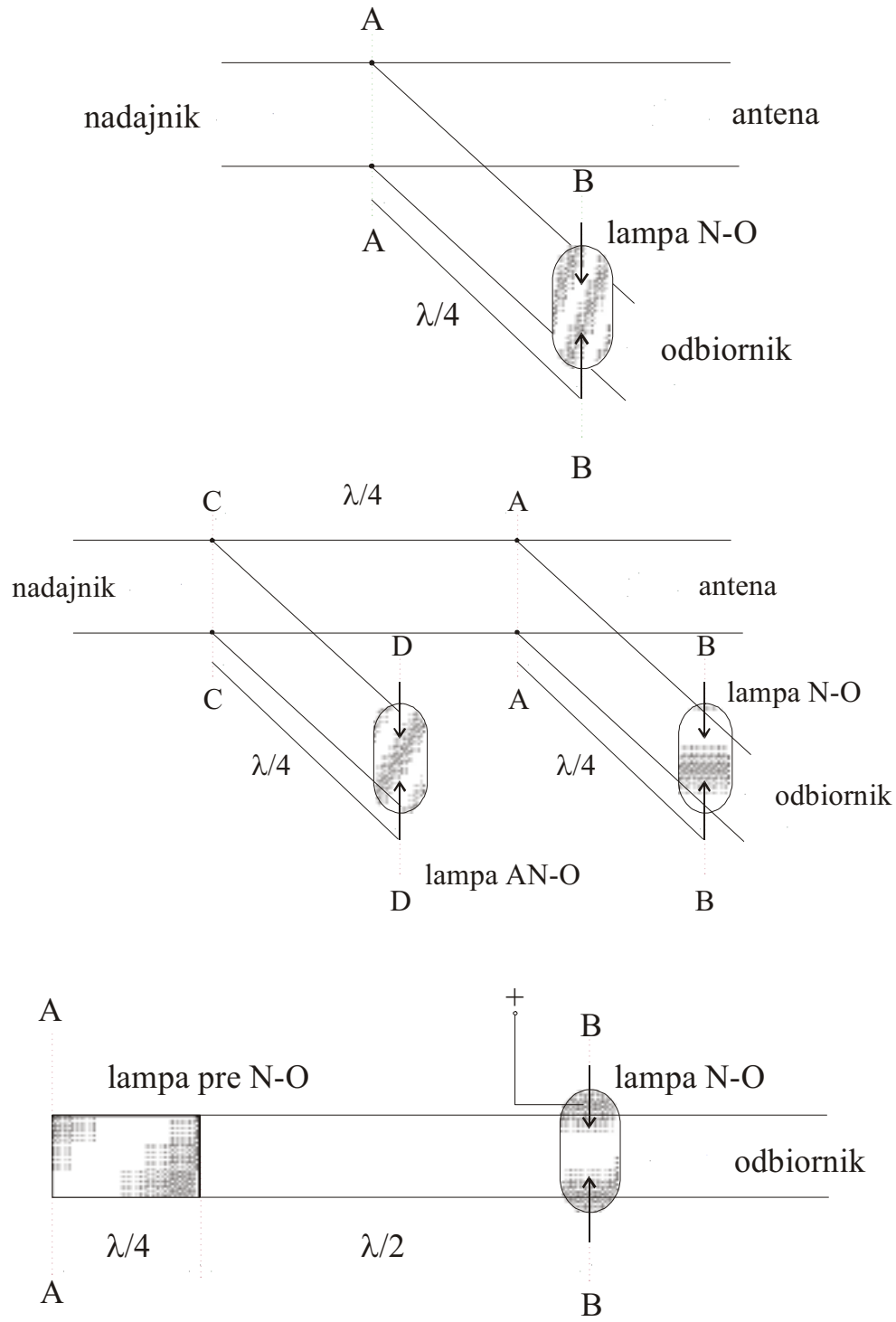
Przełączniki nadawczo odbiorcze

Mechaniczne



Elektryczne

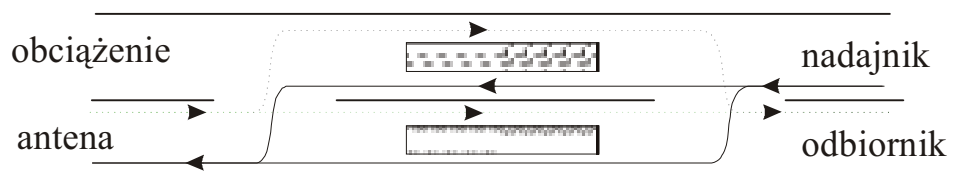
Przełączniki z lampami neonowymi



Przełącznik z wkładką ferrytową

 wkładka ferrytowa $\theta \begin{cases} 0^\circ \text{ dla } H_o = 0 \\ 180^\circ \text{ dla } H_o > 0 \end{cases}$

 wkładka dielektryczna θ



ELEMENTY AKUSTOELEKTRONIKI MIKROFALOWEJ

Układy akustoelektroniczne znajdują zastosowanie w technice mikrofalowej dzięki temu, że możliwe jest wygenerowanie fal mechanicznych (powierzchniowych lub przypowierzchniowych) o częstotliwościach dochodzących nawet do kilkunastu GHz.

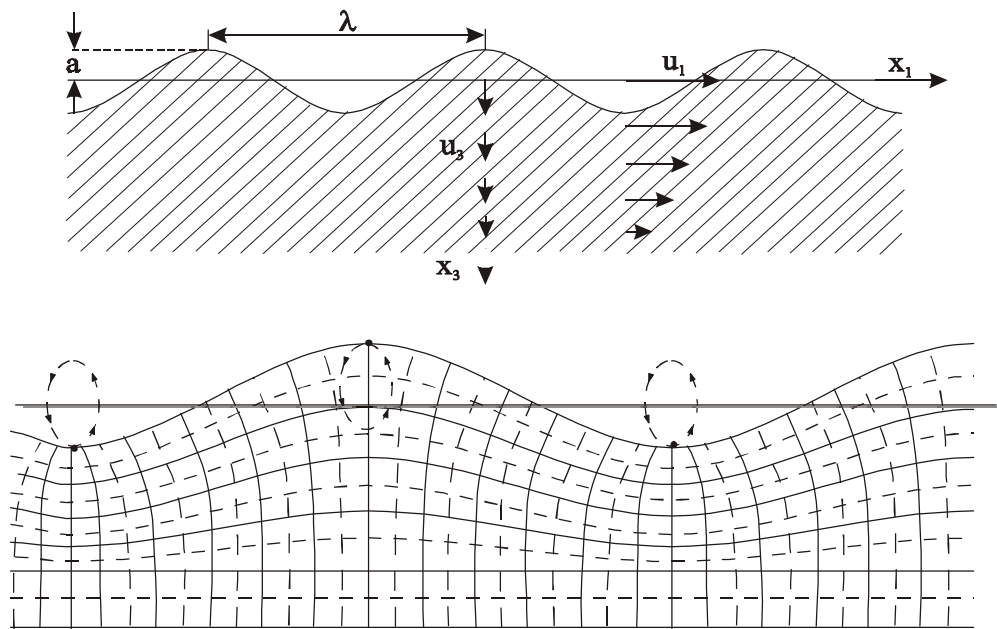
Prędkość propagacji fal akustycznych jest 10^5 rzędów niższa od elektromagnetycznych, co sprawia, że ich długości są rzędu μm . Fakt ten umożliwia konstruowanie unikalnych układów mikrofalowych o rozmiarach mikroskopowych. Należą do nich m. in.:

- filtry pasmowoprzepustowe,
- banki filtrów,
- linie opóźniające,
- przesuwniki fazy,
- sprzęgacze kierunkowe,
- konwoluty (układy realizujące splot),
- korelatory,
- transformatory Fouriera,
- transformatory Hilberta,
- transformatory Falkowe,
- czujniki m.in.:
 - stężenia gazów,
 - przyspieszenia,
 - wilgotności,
 - natężenia pola elektrycznego,
 - siły,
- minianteny,
- układy zdalnej identyfikacji,
- i wiele innych interesujących układów.

Fala Rayleigha

Najszerze techniczne zastosowania znalazła fala Rayleigha (FR) (fala tego samego typu przenosi energię trzęsień ziemi).

W ośrodku sprężystym ma ona dwie składowe przemieszczeń mechanicznych: poprzeczną u_1 i podłużną u_3 .



Chwilowe odkształcenia powierzchni wywołane FR. Punkt materialny na powierzchni, po której przechodzi taka fala zatacza elipsy.

Generacja i detekcja FR odbywa się za pomocą specjalnych przetworników.

W przypadku ośrodków piezoelektrycznych powszechnie stosowany jest przetwornik międzypalczasty (IDT – ang. interdigital transducer).

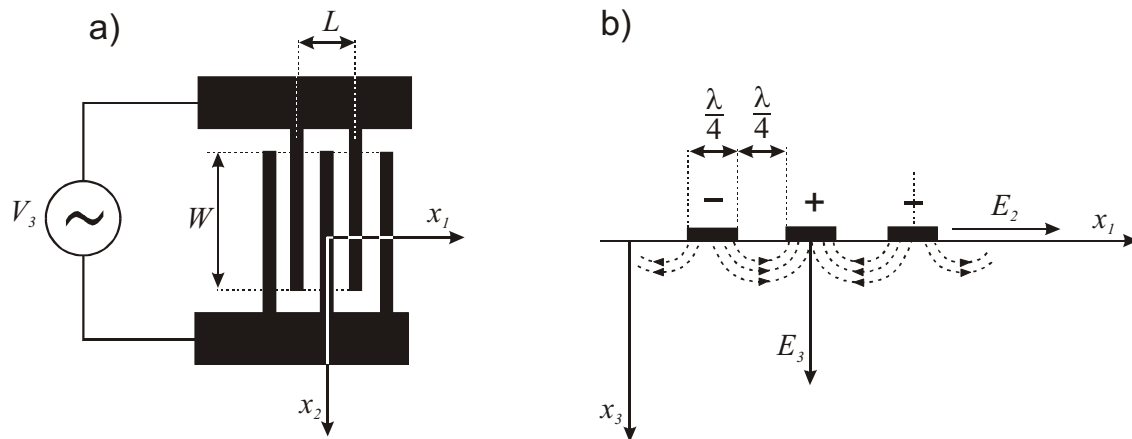
Najczęściej stosowanymi podłożami AFP są LiNbO_3 - materiał o dużym współczynniku sprzężenia piezoelektrycznego, oraz SiO_2 posiadający cięcie stabilne temperaturowo (ST).

Zasada działania przetworników IDT

W kryształach piezoelektrycznych na skutek prostego efektu piezoelektrycznego mechanicznym odkształceniom towarzyszy potencjał elektryczny φ . Składowe elektryczne mają postać:

$$E_1 = -\partial\varphi/\partial x_1 \text{ i } E_3 = -\partial\varphi/\partial x_3$$

Takie właśnie składowe generują przetworniki międzypalczaste stanowiące układ zachodzących na siebie elektrod metalowych leżących na powierzchni piezoelektryka.



Prosty przetwornik międzypalczasty
 a) układ elektrod b) rozkład pola elektrycznego

Napięcie elektryczne V_3 , przyłożone między elektrody przetworników, wytworzy składowe pola elektrycznego E_1 i E_3 . Jeśli napięcie to jest harmonicznie zmienne $V_3 = V_0 e^{j\omega t}$, to wzbudzone periodycznie zaburzenia mechaniczne będą rozchodzić się w postaci FR w obydwu kierunkach osi x_1 . Odształcenie mechaniczne wzbudzone np. przez składową pola elektrycznego E_3 pod pierwszą elektrodą, po przejściu w obszar drugiej elektrody będzie dodawać się (lub odejmować) do odształcenia wywołującego tą samą składową pola w obszarze drugiej elektrody. Jeśli odległość między sąsiednimi elektrodami L jest równa połowie długości fali, faza odształcenia przy przejściu drogi o długości $\lambda/2$ zmieni się o π . Jednocześnie o tę samą wartość zmieni się faza napięcia V_3 , co oznacza sumowanie odształceń generowanych przez kolejne elektrody. Amplituda fali powierzchniowej będzie więc liniowo narastać zarówno w kierunku $+x_1$, jak i $-x_1$, na odcinku równym iloczynowi NL , gdzie N jest liczbą par elektrod przetwornika. Prosty przetwornik międzypalczasty jest więc przetwornikiem dwukierunkowym. Jeśli apertura elektrod W jest znacznie większa od długości fali λ (zwykle jest to kilkadziesiąt λ), to przetwornik międzypalczasty można traktować jako układ dwuwymiarowy. Propagująca się FR niesie ze sobą ładunek elektryczny, który zbierany jest przez elektrody, zbudowanego na tej samej zasadzie co przetwornik nadawczy, przetwornika odbiorczego dołączonego do obciążenia.

Główną rolę przy wzbudzaniu i odbiorze fal powierzchniowych odgrywa składowa pola E_3 , prostopadła do kierunku propagacji. Z faktu zależności geometrii elektrod od λ wynikają własności filtracyjne IDT.

Filtry z FR posiadają wiele zalet, do których należy zaliczyć m.in.:

- małą tłumienność (FR jest falą płaską)
- łatwość kształtowania pożądanego przebiegu charakterystyk zarówno amplitudowej jak i fazowej (poprzez odpowiednie kształtowanie geometrii przetwornika międzypalczystego)
- łatwość wykonywania (technologia planarna)
- małe wymiary, tym mniejsze im częstotliwość podstawowa na której pracuje filtr jest większa.

Pewną niedogodność stanowi ograniczenie częstotliwości pracy zarówno od góry jak i od dołu. Od góry ono dyktowane możliwościami technologicznymi wytwarzania wysmukłych elektrod o grubościach rzędu pojedynczych μm , od dołu zaś możliwościami technologii uzyskiwania odpowiednio dużych kryształów piezoelektrycznych.

Jeśli częstotliwościową charakterystykę filtra określimy $G(\omega)$, a $H(\omega)$ jest widmem sygnału to sygnał wyjściowy po filtracji ma postać:

$$h(t) = \int H G e^{-j\omega t} d\omega = \int h(\tau) g(t - \tau) d\tau,$$

gdzie $h(t)$ jest sygnałem, a $g(t)$ odpowiedzią impulsową filtra.

Filtry z FR realizują powyższą funkcję, a ich projektowanie polega na realizacji pożądaney odpowiedzi impulsowej. Jest to zasadnicza różnica w porównaniu z filtrami typu RLC, przy projektowaniu których odpowiedź impulsowa nie jest brana pod uwagę.

Głównym elementem filtra z FR jest układ dwóch przetworników, zamieniających sygnał elektryczny na FR i odwrotnie.

Mała prędkość tej fali umożliwia dostęp do niej i ewentualną obróbkę w trakcie jej propagacji. Pozwala to m.in. na pobieranie próbek sygnału, czyli realizację tzw. filtracji transwersalnej. Stanowi ona podstawę działania podzespołów z AFP.

Powierzchniowe własności piezoelektryków

Przyłożone do powierzchni styczne pole elektryczne powoduje powstanie od strony piezoelektryka indukcji normalnej do powierzchni. Elektryczna admitancja powierzchniowa piezoelektryka ma postać:

$$Y(r) = D_{\perp}/E_{\parallel} = j Y_r \frac{r^2 - k_{\infty}^2}{r^2 - k_0^2} \text{sign}(r)$$

gdzie:

$D_{\perp}/E_{\parallel}$ - amplitudy indukcji normalnej i pola stycznego,

Y_r - wolnozmienna funkcja r , zależna od materiału,

k_{∞} , k_0 – wektory falowe dla powierzchni odpowiednio rozwartej i zwartej elektrycznie,

r – wektor falowy

$$\text{sign}(r) = \begin{cases} 1 & \text{dla } r \geq 0 \\ -1 & \text{dla } r < 0 \end{cases}$$

Dla r z obszaru (k_{∞}, k_0) dobrym przybliżeniem $Y(r)$ jest:

$$Y(r) = \varepsilon_o \varepsilon_r = \varepsilon_o \sqrt{\varepsilon_{11} \varepsilon_{22} + \varepsilon_{12}^2}$$

gdzie: ε_r jest tzw. efektywną stałą dielektryczną, ε_{ii} składowymi tensora przenikalności elektrycznej.

W obszarze $r < k_{\infty}$ wielkość Y_r ma niewielką część urojoną, wynikającą ze strat wskutek wypromieniowania AFP w głąb podłoża (transformacji AFP do fal objętościowych).

Istotną różnicę między dielektrykami a piezoelektrykami stanowi występowanie wektorów falowych k_{∞} i k_0 .

Ze składową harmoniczną pola E na powierzchni piezoelektryka związana jest gęstość mocy fali powierzchniowej:

$$P \sim [E_{\parallel}/(r - k_0)]^2$$

Można zdefiniować amplitudę fali powierzchniowej a jako:

$$a \sim E_{\parallel} (r - k_0)$$

przy czym $P = \frac{1}{2} a a^*$.

Powyższe rozważania upraszczają teorię przetworników międzypalczastych do problemów znanych z teorii pola. Związane z wektorami falowymi k_0 , k_{∞} prędkości fali v_0 , v_{∞} mogą istotnie różnić się dla różnych kierunków propagacji, na tej samej powierzchni. Fakt ten powoduje dodatkowe efekty w przypadku generacji wiązki fali o skończonej szerokości (aperturze).

Synteza podzespołów z AFP

Model funkcji δ

Pozwala on na syntezę struktur generujących i detekujących AFP z uwzględnieniem tylko podstawowych zjawisk fizycznych.

Pole E wymuszone przez układ elektrod w przybliżeniu jest odcinkiem sinusoidy, którego widmo ma postać $\text{sinc}x$. Zakłada się, że elektrody nie wpływają (słabo wpływają) na prędkość fali powierzchniowej: $k_\infty = \omega/v$

Składowa widma dla $r = k_\infty$ odpowiada amplitudzie generowanej fali powierzchniowej. Ponieważ w funkcji częstotliwości sygnału ω widmo pola nie zmienia się, a zmienia się k_∞ , to w funkcji częstotliwości amplituda fali powierzchniowej $a(\omega)$ będzie odwzorowywać $E_{||}(r)$:

$$a(\omega) \sim E_{||}(\omega/v)$$

Przyłożenie napięcia impulsowego $V(\omega) = 1$ do przetwornika spowoduje powstanie fali powierzchniowej o postaci:

$$a(x, t) \propto \int E_{||} \left(\frac{\omega}{v} \right) e^{j\omega \left(t - \frac{x}{v} \right)} d\omega = E_{||}(x - vt),$$

czyli fali o kształcie $E_{||}(x)$. Zatem akustyczna odpowiedź przetwornika na impuls elektryczny jest odtworzeniem kształtu przetwornika. Lepsze od poprzedniego przedstawienie widmowe pola $E(r)$ daje pole w postaci funkcji δ -Diraca. Wtedy:

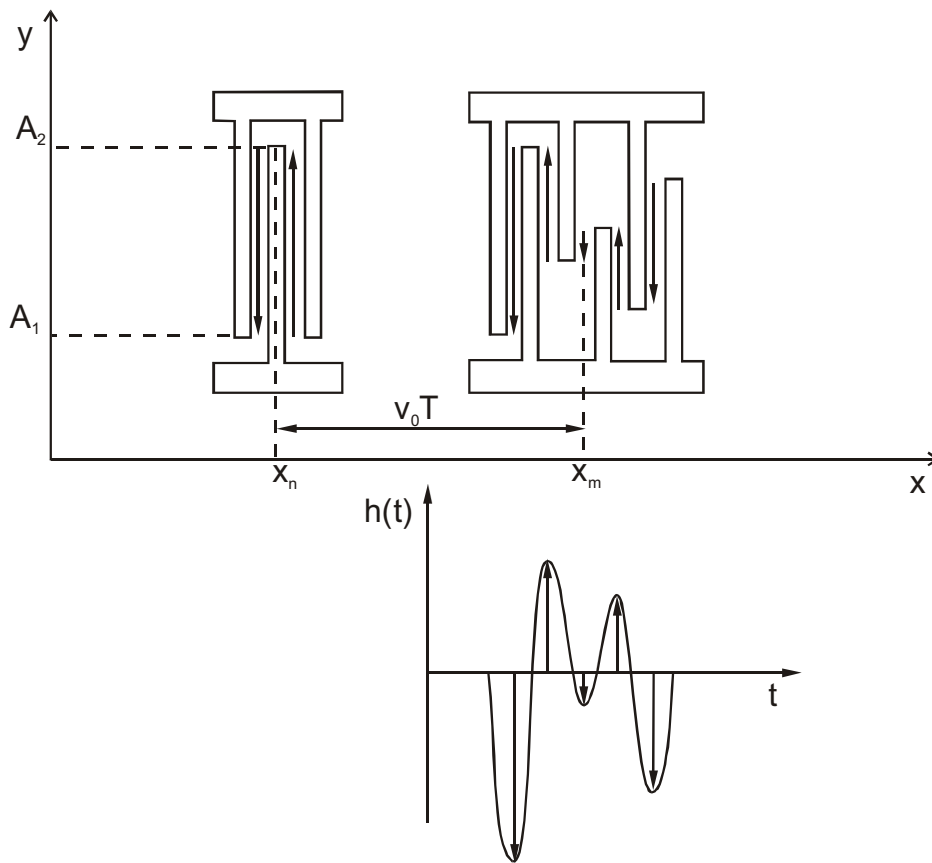
$$a(\omega) \approx \sum E_n e^{-j\omega x_n / v}$$

E_n - amplituda impulsowego pola na powierzchni, położonego w miejscach x_n . Pasma przenoszenia takiego filtru zawęża się wraz ze wzrostem ilości elektrod.

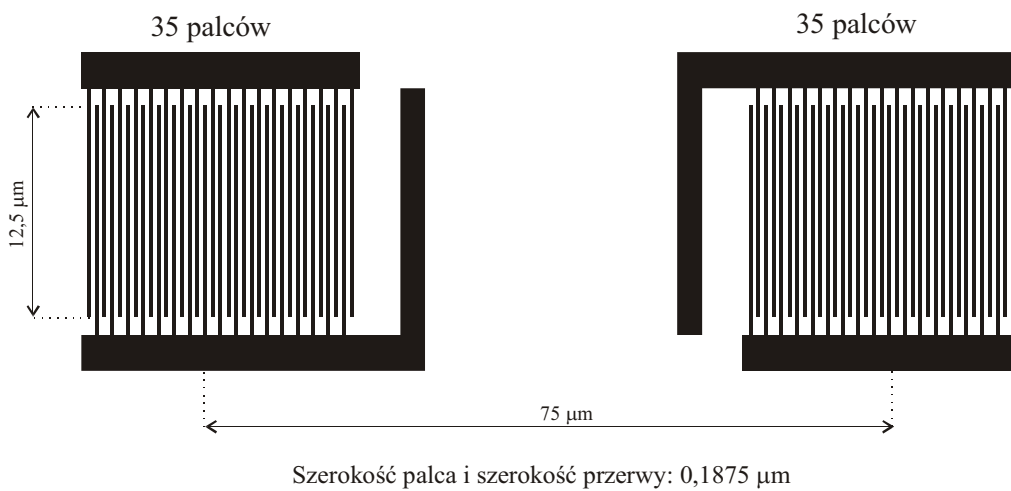
Odpowiedź impulsowa dwóch współpracujących przetworników jest splotem odpowiedzi impulsowych każdego z nich z osobna, dzięki czemu parametry sygnałów można łatwo kształtować przez odpowiedni dobór charakterystyk współpracujących przetworników, co sprowadza się z kolei do odpowiedniego ukształtowania geometrii projektowanego urządzenia. Długość obrabianych impulsów jest w zasadzie ograniczona jedynie rozmiarami podłoża, chociaż należy zaznaczyć, że wymiary przetworników maleją proporcjonalnie do wzrostu częstotliwości i dla dostatecznie dużych ich wartości takie ograniczenie w zasadzie nie ma istotnego znaczenia.

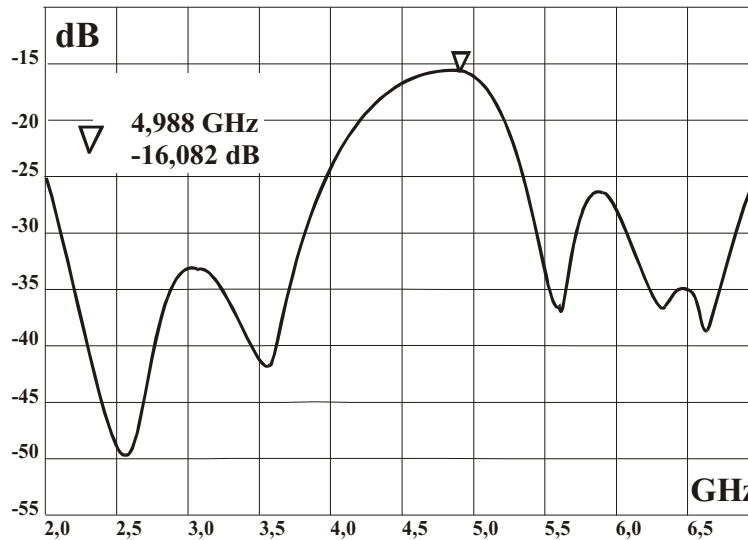
Ta prosta teoria umożliwia obliczenie geometrii filtru bez uwzględniania efektów dodatkowych (efektów II rzędu) związanych m. in. z ograniczoną ilością elektrod w przetworniku, odbiciami międzyelektrodowymi, odbiciami od przetworników, dyfrakcją fali, itd.

Przykład kształtowania charakterystyki impulsowej filtru.



Przykład geometrii filtru i jego charakterystyka amplitudowa





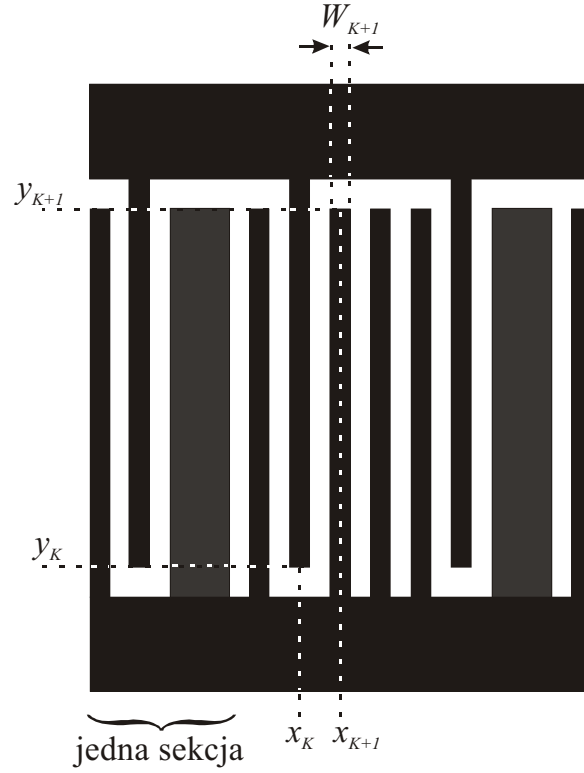
Filtr taki jest w istocie także linią opóźniającą o opóźnieniu zależnym od odległości przetworników (AFP jest 10^5 razy wolniejsza od elektromagnetycznej).

Przetworniki jednokierunkowe

Zasadniczą, w większości zastosowań, wadą prostych przetworników międzypalczastych jest ich dwukierunkowość. Z racji tego, że najczęściej wykorzystywany jest tylko jeden kierunek promieniowania nie można osiągnąć w wielu układach z AFP, zbudowanych w oparciu o tradycyjne przetworniki, strat niższych niż 3 dB.

Aby temu zaradzić wynaleziono całą gamę interesujących konstrukcji, które co prawda pracowały jednokierunkowo, ale niestety uzyskanie tej właściwości było okupione innymi niedogodnościami (najczęściej natury technicznej lub technologicznej).

Pokazany fragment przetwornika pracuje jednokierunkowo. Jedna jego sekcja składa się z dwóch palców o szerokości $\lambda/8$ doczepionych do przeciwnych szyn zbiorczych (pierwszy do uziemionej) i jednego reflektora $3\lambda/8$. Taki niesymetryczny układ elektrod powoduje, że amplitudy przemieszczeń generowanych w kierunku dodatnim osi interferują konstruktywnie, zaś w kierunku ujemnym destruktywnie.



Charakterystyka częstotliwościowa takiego przetwornika wyraża się następująco:

$$A(f) = \sum_{k=1}^{N-1} (y_{K+1} - y_K) \exp \left(-j2\pi f \frac{x_K + x_{K+1}}{2v_{m/2}} \right),$$

gdzie: N – ilość palców, $v_{m/2}$ – prędkość AFP na powierzchni połowicznie metalizowanej.

Położenia elektrod są następujące:

$$x_K = \frac{kv_{m/2}}{4f_0}, \quad K = 1 \dots N$$

$$y_K = y_{K-1} + WT_K,$$

$$WT_K = 0, 1, -1, 0, 1, -1 \dots,$$

$$K = 1 \dots N.$$

Rezonatory z AFP

Konstrukcja ich bazuje na możliwości uzyskania pełnego odbicia AFP od periodycznych zaburzeń przez wykorzystanie efektu sprzężenia fali postępującej i wstecznej, propagujących się pod periodycznie nałożonymi elektrodami.

Istnieje kilka metod realizacji takich zaburzeń powierzchni, należą do nich:

- wykorzystanie niejednorodności elektrycznej układu elektrod,
- zastosowanie periodycznych rowków na powierzchni podłoża o głębokości $0,01 \div 0,03$ długości odbijanej fali,
- naruszenie struktury materiałowej podłoża poprzez wtrącenie obcego materiału,
- nałożenie na powierzchnię dostatecznie ciężkich elektrod.

Pasmo pracy struktury odbijającej leży wokół częstotliwości środkowej, gdzie spełniony jest warunek synchronizmu fali i układu elektrod: $2k \approx K$

gdzie: k jest wektorem falowym fali powierzchniowej, $K = 2\pi/L$, L - odległość między niejednorodnościami.

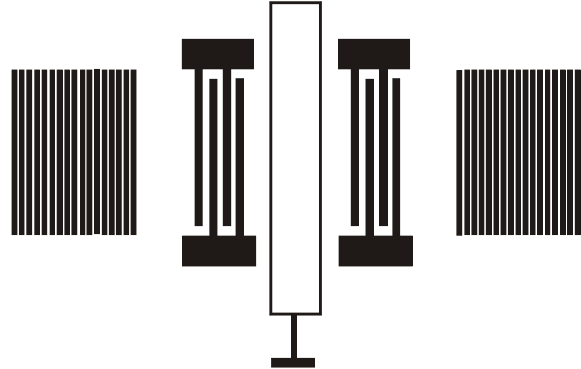
Struktury odbijające służą do konstrukcji rezonatorów typu Fabry-Perota.

Zmieniając położenie luster i przetworników można otrzymać trzy typy rezonatorów:

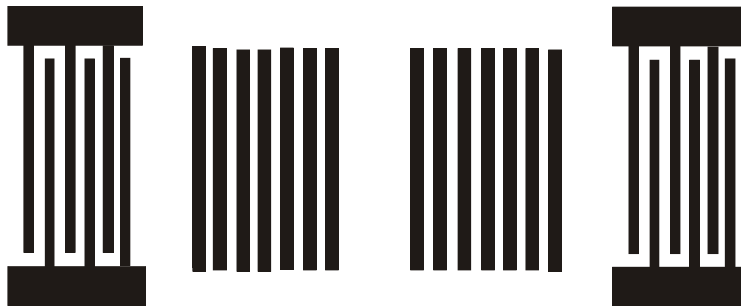
- rezonatory z jednym przetwornikiem między strukturami odbijającymi. Efekt rezonansowy polega na zmianie impedancji Z przetwornika,



- rezonator z dwoma przetwornikami między strukturami odbijającymi. Efekt rezonansowy polega na zmianie współczynnika transmisji H między przetwornikami. Znajdujący się między przetwornikami metalowy ekran elektrycznie separuje przetworniki.



- rezonator z dwoma przetwornikami na zewnątrz struktur odbijających, które położone są w pewnym odstępie L_c , zwanym obszarem rezonansu. Efekt rezonansowy polega tu również na zmianie transmisji między przetwornikami.



Rezonatory takie mogą osiągać bardzo duże dobroci rzędu kilkudziesięciu tysięcy. Małe straty transmisyjne (kilka dB) umożliwiają szeregowe łączenie rezonatorów (aby np. zwiększyć dobroć), a także konstrukcję wielobiegunowych rezonansowych filtrów pasmowych.

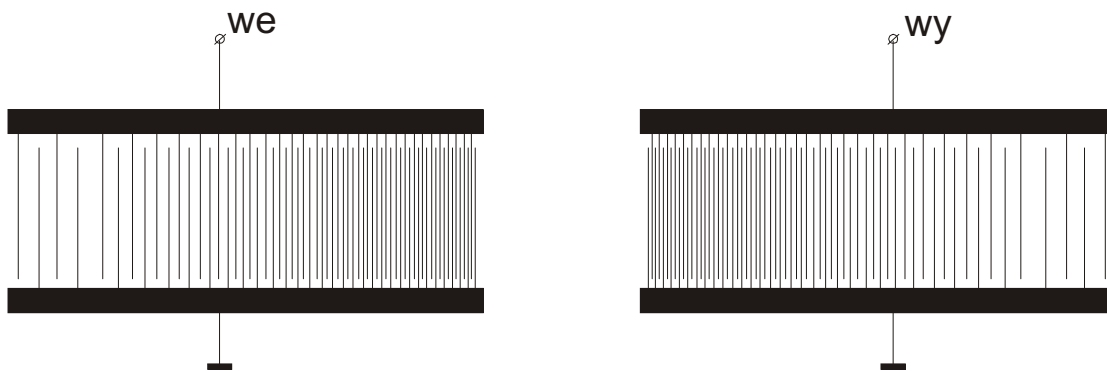
Czynnikami pogarszającymi parametry rezonatorów z AFP są m. in.:

- zły wybór odległości struktur, powodujący, że warunek fazy jest spełniany dla częstotliwości, dla której Γ nie jest maksymalne,
- tłumienie fali powierzchniowej,
- dyfrakcyjne „wyciekanie” fali poza obszar rezonatora,
- niepełne odbicie – transmisja fali poza struktury odbijające.

Linie dyspersyjne

Jednym z bardziej użytecznych, szczególnie w radiolokacji, zastosowań przyrządów z AFP są filtry służące do generacji oraz kompresji sygnałów z liniową modulacją częstotliwości (LMC), nazywane dyspersyjnymi liniami opóźniającymi.

Przykład filtra z LMC

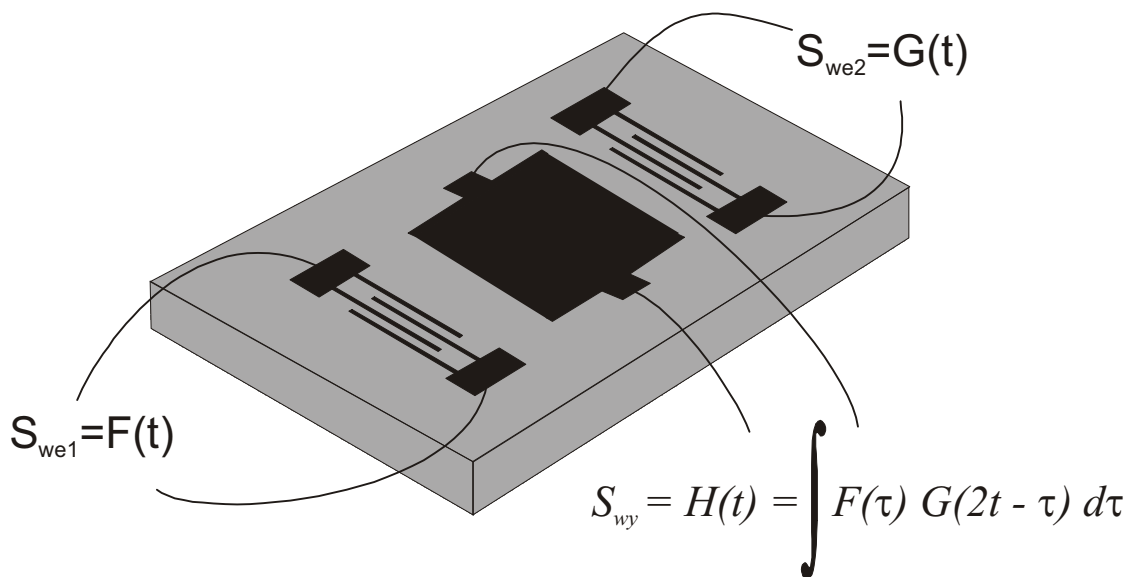


Odpowiedni zagęszczający się układ elektrod powoduje, że charakterystyka częstotliwościowa przetwornika wejściowego odtwarza widmo sygnału z modulacją częstotliwości. Sygnał ten ulega następnie kompresji, przechodząc w postaci fali powierzchniowej pod odpowiednio dobranym, rozrzedzającym się, układem elektrod przetwornika wyjściowego. Obydwie operacje, tj. generacji i kompresji, są tu wykonywane w sposób naturalny i w czasie rzeczywistym, co czyni tą metodę bardzo efektywną.

Innym znanym zastosowaniem linii dyspersyjnych są analizatory widma, w których wygenerowany sygnał akustyczny zostaje, dzięki dyspersji, rozdzielony w dziedzinie czasu na częstotliwości składowe.

Konwoluty

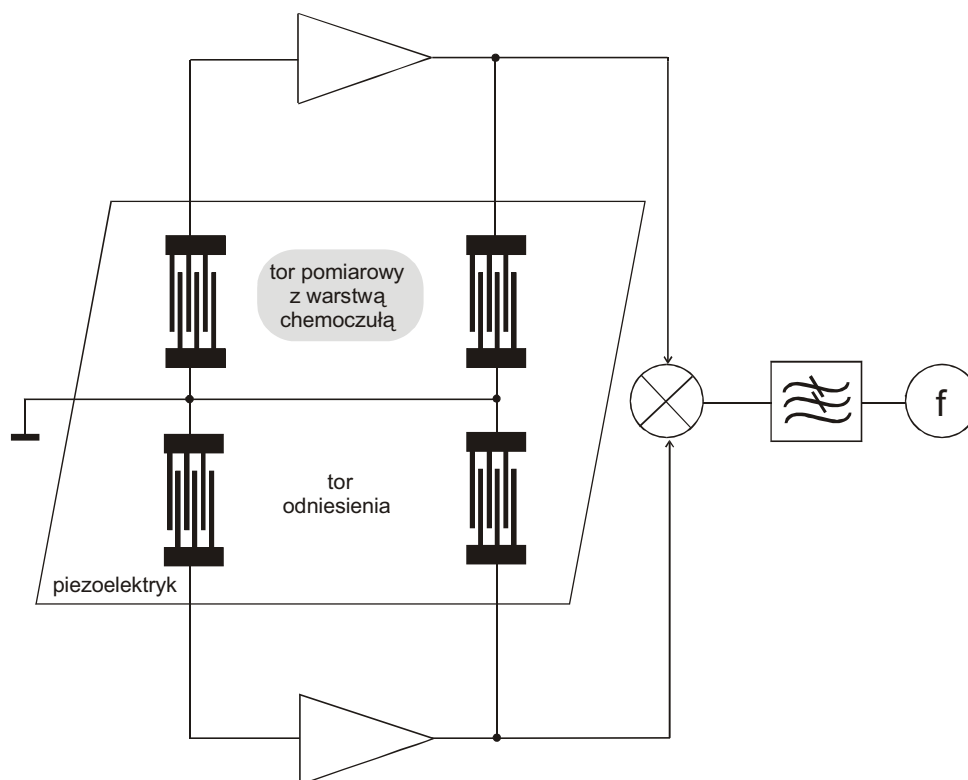
Prostą konsekwencją zasady współdziałania układu dwóch przetworników międzypalczastych jest możliwość szybkiego obliczania splotu, realizowana przez konwolutor (nazywany też konwolwerem). W przypadku tego układu obydwa współpracujące przetworniki są przetwornikami wejściowymi, a sygnał wyjściowy, w wersji z FR, jest pobierany za pomocą płaskiej elektrody leżącej na powierzchni w obszarze położonym pomiędzy przetwornikami



Przyrząd ten wykonuje tzw. splot zdegenerowany (w zwykłej całce splotowej drugi czynnik podcałkowy ma postać $G(t-\tau)$) degeneracja jest rezultatem względnego podwojenia prędkości fal propagujących się w przeciwnych kierunkach. Możliwa jest jednak, po dodaniu do układu dodatkowych struktur opóźniających, konstrukcja konwolutora wykonującego zwykły splot.

Czujniki

Do bardzo interesujących zastosowań przyrządów z AFP należą też czujniki wielkości nieelektrycznych. Zasada ich działania polega na wykorzystaniu zjawiska zmiany prędkości fali akustycznej wywołanego zmianami stanu powierzchni piezoelektryka, po którym się ta fala propaguje. Realizacja czujnika stężenia gazu polega na pokryciu odpowiednią substancją chemoczułą obszaru propagacji fali. Nałożona warstwa reagując na obecność gazu oddziałuje na powierzchnię, wywołując efekt zmiany jej obciążenia masą czy też jej przewodnictwa elektrycznego. Efekty te z kolei wpływają na prędkość AFP, której zmiany stosunkowo łatwo zmierzyć.



Układ do pomiaru częstotliwości różnicowej

Przedstawiony układ składa się z dwóch linii opóźniających zamykających pętle sprzężenia zwrotnego wzmacniaczy. Przy spełnieniu warunków generacji tworzą one dwa niezależne generatory o częstotliwościach zależnych od prędkości propagacji AFP, przy czym w jednym z tych generatorów prędkość ta zależy od zewnętrznego czynnika mierzonego (na który reaguje nałożona warstwa). Różnicowa częstotliwość generowana przez taki układ będzie więc zależać od stężenia czynnika mierzonego, a parametr ten można mierzyć stosunkowo łatwo i z dużą dokładnością.

Czujniki ciśnienia, przyspieszenia itp. (ogólnie czujniki naprężeń) reagują na zmianę stanu makronaprężenia powierzchni podłoża AFP (przełożoną na zmianę prędkości AFP) wywołaną przyłożonymi doń mierzonymi siłami zewnętrznymi.

LITERATURA ŹRÓDŁOWA

J. A. Dobrowolski, Technika wielkich częstotliwości, Oficyna Wydawnicza Politechniki Warszawskiej, 2001.

R. Litwin, Teoria pola elektromagnetycznego, WNT 1968.

J. S. Dobosz Podstawy techniki mikrofalowej, skrypt WAT S-48327.

L. Landau, E. Lifszic, Teoria pola, PWN 1958.

F. Lachowicz, Podstawy teorii pola elektromagnetycznego, skrypt Politechniki Łódzkiej 1994.

D. J. Griffiths, Podstawy elektrodynamiki, PWN 2001.

A. W. Bikadze, Równania fizyki matematycznej, PWN 1984.

D. Stauffer, H. E. Stanley, Od Newtona do Mandelbrota. Wstęp do fizyki teoretycznej, WNT 1996.

G. Kino, Akusticeskije wolny, ustrojstwa, wizualizacija i analogowaja obrabotka signalow, Moskwa, Mir, 1990.

H. Matthews Surface wave filters, J. Wiley & Sons, Inc. 1977.

IEEE Transactions on Sonics and Ultrasonics, SU-21, 1, 1974.