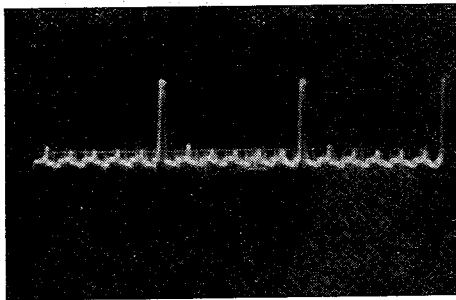


Stosunek sygnału użytecznego do zakłócenia jest określony przez stosunek sumy oporności obwodu z rdzeniem przełączającym się do sumy oporności obwodu z rdzeniem nie przełączającym się.

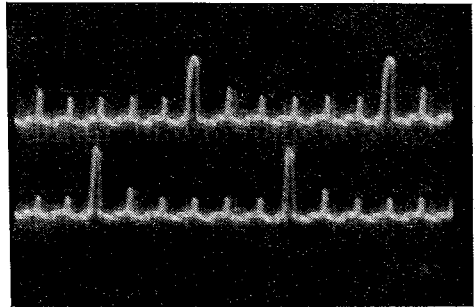
Dla ustalonej i dostatecznie małej wartości oporności wewnętrznej generatora prostokątnego napięcia przełączającego, ze wzrostem ilości zwojów uzyskujemy polepszenie się stosunku sygnału do zakłócenia, ponieważ impedancja zastępcza rdzenia jest proporcjonalna do kwadratu ilości zwojów.

Oczywiście polepszenie stosunku sygnału do zakłócenia tą metodą możliwe jest do tych granic, przy których wpływ impedancji zastępczej nasyczonego rdzenia (na którą składa się również indukcyjność własna uzwojenia) nie ograniczy wzrostu stosunków wspomnianych oporności.

Stosunek sygnału do zakłócenia można znacznie polepszyć stosując wejściowe i wyjściowe ograniczniki prądu, czyli tzw. „studnie”, oraz trze-



Rys. 14. Kształt impulsów wyjściowych na wyjściu przełącznika pozytorowego



Rys. 15. Impulsy wyjściowe na dwu różnych wyjściach przełącznika pozytorowego

cie uzwojenie polaryzujące o niewielkiej ilości zwojów, przez które przepuszcza się prąd stały o określonym kierunku. Użyteczność w tym układzie studni wejściowych i wyjściowych oraz trzeciego uzwojenia opisuje praca [5].

Zastosowanie studni powoduje, że negatorowy układ lub pozytorowy jest bardziej skomplikowany, co jest cechą ujemną tego rozwiązania.

Na fotografiach zamieszczonych obok pokazano kształt impulsów wyjściowych na różnych wyjściach dwunastokanałowego przełącznika pozytorowego pracującego w układzie z rys. 10.

Układ przełącznika zbudowanego na transformatorach nasyconych (zob. rys. 12) przenosi jedyńki prawidłowo. Jednak w porównaniu z przełącznikami zbudowanymi na negatorach i pozytorach daje stosunkowo duże napięcie zakłócające (stosunek sygnału do zakłócenia 4 : 1). Jest to spowodowane oddziaływaniem obwodów wejściowych rdzeni przełączających się na obwody wyjściowe rdzeni nieprzełączanych.

7. WNIOSKI

W wyniku dokonanej pracy uzyskano materiały i dane umożliwiające realizację przełącznika kanałowego dla telefonicznej krotnicy czasowej, opartego na wykorzystaniu rdzeni magnetycznych o prostokątnej pętli histerezy.

Spośród przebadanych układów najbardziej odpowiednim wydaje się być przełącznik zbudowany z szeregu pozytorów. Zalety tej konstrukcji są m.in. następujące:

1. Ilość wzmacniaczy magnetycznych (uzwojonych rdzeni) niezbędnych do realizacji przełącznika jest równa ilości potrzebnych wyjść — zalety tej nie posiada przełącznik zbudowany z szeregu negatorów.

2. Przełącznik zbudowany z szeregu transformatorów nasycanych wymaga wprowadzić do realizacji takiej samej ilości rdzeni jak przełącznik pozytorowy, jednak stosunek sygnału do zakłócenia jest dla tego układu gorszy niż dla układu pozytorowego i negatorowego.

3. Zasilanie układu pozytorowego (również negatorowego) jest bardziej korzystne niż układu na transformatorach nasycanych dlatego, że w tym pierwszym możliwe jest równoległe łączenie (w stosunku do zasilania) wzmacniaczy magnetycznych.

Praca niniejsza dotyczy w zasadzie konstrukcji magnetycznego przełącznika kanałowego przeznaczonego do wykorzystania w telefonicznej krotnicy czasowej. Mimo to również inne zespoły teletransmisyjnych systemów impulsowych mogą być zrealizowane przy pomocy układów magnetycznych. Dlatego też wydaje się, że przeprowadzone w niniejszej pracy rozważania poparte badaniami eksperymentalnymi mogą stanowić interesujący materiał do badań nad możliwością magnetycznej realizacji takich układów.

Politechnika Warszawska

Katedra Urządzeń Teletransmisyjnych i Telegraficznych

WYKAZ LITERATURY

1. Góral A.: Magnetyczne przekaźniki bezstykowe. Wydawnictwo M.O.N., Warszawa 1960.
2. Istvanffy E.: Materiały magnetyczne i ich zastosowanie. P.W.T., Warszawa 1956.
3. Materiały na obrady Krajowej Konferencji Ferrytowej: Ferryty o prostokątnej pętli histerezy. Warszawa, maj 1959.
4. Jaworski W.: Zastosowanie ferrytów w programowanych maszynach cyfrowych. — *Przegląd Telekomunikacyjny* nr 1, 1960.
5. Zienkiewicz-Kulińska E.: Negator — szeregowy wzmacniacz magnetyczny. Praca dyplomowa wykonana pod kierunkiem prof. F. Błockiego P.A.N., Z.A.M., 1959.
6. Kaposi I. F.: Rectangulator Hysteresis-Loop Magnetic Cores as Switching Elements. — *Electronic Engineering*, May 1959, str. 278—283.

7. Tellefsen H., Alessio S.: Application of Magnetic Cores to Industrial Controls. — I.R.E. National Convention Record. 1958, part 6, str. 283—293.
8. Meyerhoff A. J.: Digital Applications of Magnetic Devices. J. Wiley, 1960.
9. Opracowanie wewnętrzne Katedry Teletransmisyjnych Urządzeń Przewodowych: Telefonía wielokrotna oparta na kryterium podziału w czasie z modulacją kodowaną.
10. Bielański W. F., Szamajew Ju. M.: Rasczot elektricheskich cepiej so-dierzaszczich sierdieczniki s priamougolnoj pietloj histierezisa. — Awtomatika i Tielemechanika nr 8, 1960, str. 1138—1192.
11. Błocki F.: Zasada działania telefonicznego urządzenia wielokrotnego I-12. Referat na konferencji naukowej „Transmisja Danych”, Warszawa 1962.
12. Bieńkowska T.: Przełącznik kanałowy. Komunikat na konferencji naukowej „Transmisja Danych”, Warszawa 1962.

W. T. BIENKOWSKA

MAGNETIC CHANNEL SWITCH FOR TELEPHONE TIME MULTIPLIER

Summary

The paper discusses the rules of operating and designing magnetic switch systems built on toroidal cores with rectangular magnetic hysteresis loops. It discusses transducer amplifiers built on those cores which are switched by the current as well as by voltage. The designed and constructed switch permits the switching of twelve channels with the commutation frequency equal to 48 kHz. The optimum construction of this switch, built from the chain of positors, allows to obtain at output the relation of the signal to disturbance such as 6 : 1. The circuit may be employed as a channel switch of a telephone time multiplier of the twelve-fold system.

W. T. BIENKOWSKA

COMMUTATEUR MAGNÉTIQUE DE CANAUX DE SYSTÈME TÉLÉPHONIQUE À RÉPARTITION TEMPORAIRE

Résumé

On a présenté les principes d'action des commutateurs magnétiques avec des noyaux toroïdals à cycle d'hystérésis rectangulaire ainsi que les principes de projeté de tels commutateurs. On a analysé des amplificateurs magnétiques avec des noyaux toroïdals commutés par les courant ou par la tension. Le commutateur projeté et construit permet de commuter douze canaux à fréquence de commutation égale 48 KHz. La construction de ce commutateur, constituée d'une chaîne des éléments, permet d'obtenir à la sortie le rapport de signal utile sur brouilleur — 6 : 1. Le schéma peut être employé comme un commutateur de canaux à répartition temporaire dans le système téléphonique multiple de douze canaux.

В. Т. БЕНЬКОВСКА

МАГНИТНЫЙ ПЕРЕКЛЮЧАТЕЛЬ КАНАЛОВ ДЛЯ ТЕЛЕФОННОГО КОММУТАТОРА С РАЗДЕЛЕНИЕМ ВО ВРЕМЕНИ

Резюме

В работе описано принципы действия и проектирования магнитных переключающих схем, построенных на тороидальных стержнях с прямоугольной петлей гистерезиса. Описано магнитные усилители построенные на таких сердечниках, переключаемых так током как и напряжением. Запроектированный и построенный переключатель позволяет переключать двенадцать каналов с частотой коммутирования равной 48 кГц. Оптимальная конструкция этого переключателя, состоящая из цепи элементов, позволяет получить на выходе соотношение сигнала к помехам 6:1. Схема может быть использована в качестве переключателя каналов телефонного коммутатора с разделением во времени, двенадцатиканальной системы.

KRZYSZTOF GRABOWSKI

Charakterystyki pompowanego złącza $p - n$

Rękopis dostarczono 23.6.1964

W pracy przedstawiono matematyczną reprezentację nieliniowej zależności dynamicznej elastancji złącza $p - n$ od sterującego je napięcia lub prądu oraz wykreślono wyniki obliczeń numerycznych współczynników głębokości modulacji elastancji i zawartości harmonicznych elastancji w zależności od głębokości modulacji napięcia lub prądu na złączu.

1. WSTĘP

Parametrem, dzięki zmianom którego zachodzi wzmacnianie, przemiana lub generacja częstotliwości w diodowych układach parametrycznych, jest pojemność złącza $p - n$, spolaryzowanego w kierunku nieprzewodzenia i sterowanego (pompowanego) dużą, pomocniczą mocą mikrofalową. Parametry samego złącza determinują najmniejszą, teoretycznie możliwą do osiągnięcia temperaturę szumów układu [11], podczas gdy pasytywne elementy samej oprawki, w której umieszczone jest złącze, określają największą wstęgę przenoszenia [4].

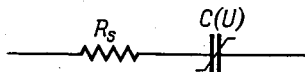
Nieliniowa zależność ładunku od napięcia na złączu $p - n$ w przypadku dużych, okresowo zmiennych sygnałów pompujących daje w wyniku okresowo zmienną w czasie funkcję reaktancji złącza, o złożonym na ogół przebiegu, zależnym od bardzo wielu czynników [1], [3], [6], [10] ÷ [14]. W praktyce przy analizie układów parametrycznych operujących sygnałami o amplitudach o kilka rzędów wielkości mniejszych w stosunku do amplitud sygnałów pompujących, wygodnie jest założyć, że dla sygnałów wzmacnianych lub podlegających przemianie reaktancja złącza, aczkolwiek zmienia się w czasie, jest liniowa. Zachodzi wówczas potrzeba:

- 1) znalezienia przebiegu w czasie funkcji reaktancji złącza,
- 2) rozłożenia tej funkcji na szereg Fouriera w okresie pulsacji pompującej.

Znajomość współczynników Fouriera reaktancji złącza jest czynnikiem koniecznym przy analizie i projektowaniu diodowych układów parametrycznych.

2. CHARAKTERYSTYKI ZŁĄCZA $p - n$ PRZY BRAKU GENERATORA POMPUJĄCEGO

Typowe złącze $p - n$ spolaryzowane zaporowo i pompowane dużym sygnałem daje się przedstawić w zakresie częstotliwości mikrofalowych [1] za pomocą szeregowo¹⁾ połączonych: stałej liniowej oporności R_s i nieliniowej pojemności C (rys. 1). Pojemność złącza jest funkcją napięcia pompy i zmiany jej zależą od szerokości warstwy ładunku przestrzenne-



Rys. 1. Schemat zastępczy złącza $p - n$ w zakresie mikrofal

go, która powiększa się lub maleje w zależności od tego, czy napięcie na złączu rośnie, czy maleje w kierunku nieprzewodzenia.

Specyfika układów parametrycznych narzuca potrzebę operowania pojęciem tzw. dynamicznej²⁾ elastancji złącza [5]

$$s(u) = \frac{\partial u}{\partial q}, \quad (1)$$

gdzie ∂q oznacza przyrost ładunku spowodowany przyrostem napięcia ∂u na złączu.

Korzystając z wyprowadzenia Shockleya [12] mamy dla złącza $p - n$ następującą ogólną zależność nieliniową

$$s(u) = s(0) \cdot \left[1 + \frac{u}{\Phi} \right]^\gamma \quad (2)$$

dla $-\Phi < u < U_{prz}$,

gdzie:

$$s(0) = s(u = 0),$$

Φ — potencjał dyfuzyjny złącza,

γ — wykładnik potęgi zależny od rodzaju złącza,

u — napięcie zewnętrzne doprowadzone do złącza i liczone jako dodatnie w kierunku nieprzewodzenia,

U_{prz} — napięcie przebicia złącza.

Często wygodnie jest [11] wprowadzić do równania (2) zamiast $s(0)$ wartość maksymalną elastancji możliwą do uzyskania w danym złączu

$$S_{mx} = s(u = U_{prz}).$$

¹⁾ Warto zauważyć, że istnienie liniowej niezależnej od częstotliwości oporności złącza R_s połączonej szeregowo z nieliniową elastancją $S(V)$ sugeruje celowość wyboru szeregowych, tzn. „oczkowych” metod analizy mikrofalowych diodowych układów parametrycznych [2], [11].

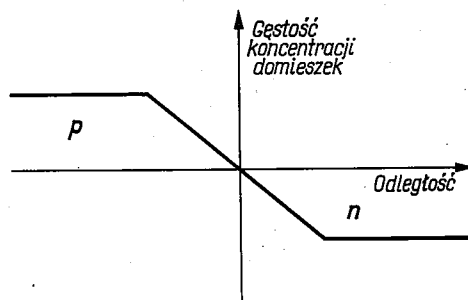
²⁾ Zwanej także elastancją przyrostową lub różniczkową.

Otrzymujemy wówczas zamiast (2)

$$s(u) = S_{mx} \left[\frac{u + \Phi}{U_{prz} + \Phi} \right]^\gamma \quad (3)$$

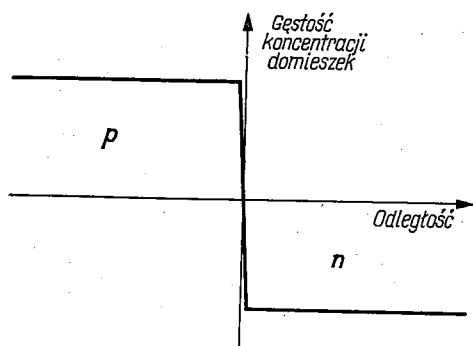
dla — $\Phi < u < U_{prz}$.

Jak już wspomniano, wykładnik potęgi γ zależy od rodzaju złącza, a ściślej od zaprogramowanego w danym złączu rozkładu gęstości domieszek. W praktyce najczęściej spotykane są dwa rozkłady: a) rozkład liniowy, w którym przejście od materiału p do materiału n odbywa się w sposób jednostajny, liniowy (rys. 2) — taki rozkład charakteryzuje się



Rys. 2. Rozkład koncentracji domieszek w złączu liniowym

wykładnikiem potęgi $\gamma = 1/3$, a złącze takie nosi nazwę liniowego; b) rozkład skokowy, w którym przejście od materiału p do materiału n odbywa się w sposób nagły (rys. 3) — taki rozkład charakteryzuje się wykładnikiem potęgi $\gamma = 1/2$, a złącze takie nosi nazwę stromego.

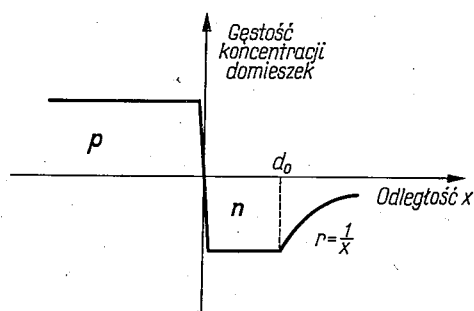


Rys. 3. Rozkład koncentracji domieszek w złączu stromym

W rzeczywistości możliwe są do uzyskania także inne rozkłady koncentracji domieszek, z których szczególnie interesujący wydaje się rozkład zaproponowany przez Marinosa [9] (rys. 4), charakteryzujący się wykładnikiem potęgi $\gamma = 1$. Cechą szczególną takiego złącza jest to, że koncentracja domieszek w materiale n jest stała do odległości d_0 , a dalej maleje odwrotnie proporcjonalnie do odległości.

Związki (2) lub (3) przedstawiają zależność elastancji złącza od sterującego je napięcia. W praktyce jednak częściej zachodzi potrzeba zna-

jomości zależności elastancji złącza od sterującego je prądu lub ładunku. Zależności te można otrzymać przez scałkowanie (3) wraz z wykorzystaniem definicji (1). Wynikającą przy takiej operacji stałą całkowania mo-



Rys. 4. Rozkład koncentracji domieszek w złączu Marinosa

żna wyznaczyć z warunku granicznego zanikania ładunku ($q = 0$) przy zwarciu złącza ($u = 0$).

W przypadku gdy $\gamma \neq 1$ wspomniane wyżej operacje prowadzą do prostego wyniku

$$s(q) = S_{mx} \left(\frac{q + Q_\Phi}{Q_{prz} + Q_\Phi} \right)^{\frac{\gamma}{1-\gamma}} \quad (4)$$

dla — $Q_\Phi < q < Q_{prz}$,

gdzie:

$$Q_\Phi = \frac{\Phi}{[1-\gamma]s(0)} = \frac{(U_{prz} + \Phi)^\gamma}{S_{mx}(1-\gamma)} \Phi^{1-\gamma} \quad (5)$$

jest ładunkiem dyfuzyjnym na złączu, gdy napięcie na nim wynosi — Φ ,

$$Q_{prz} = Q_\Phi \left[\left(\frac{U_{prz} + \Phi}{\Phi} \right)^{1-\gamma} - 1 \right] \quad (6)$$

jest ładunkiem powodującym przebicie złącza.

Warto zauważyć, że wykładnik potęgi $\frac{\gamma}{1-\gamma}$ w równaniu (4) dla złącza stromego przybiera wartość 1, a dla złącza liniowego 1/2. A zatem zależność elastancji od napięcia w przypadku złącza Marinosa jest identyczna jak zależność elastancji od ładunku w przypadku złącza stromego oraz zależność elastancji od napięcia na złączu stromym jest taka sama jak zależność elastancji od ładunku na złączu liniowym.

Z równania (4) widać także, że gdy $\gamma = 1$, wyrażenie to staje się nieokreślone i zachodzi potrzeba osobnego potraktowania tego przypadku. Podstawiając zatem w (3) $\gamma = 1$ i korzystając z (1), po scałkowaniu i wykorzystaniu wspomnianego wcześniej warunku granicznego otrzymujemy następujący związek dla złącza Marinosa:

$$q = \frac{U_{prz} + \Phi}{S_{mx}} \ln \frac{u + \Phi}{\Phi} \quad (7)$$

Następnie oznaczając

$$Q_e = \frac{U_{prz} + \Phi}{S_{mx}}, \quad (8)$$

gdzie Q_e jest ładunkiem na złączu, gdy napięcie doprowadzone do złącza wynosi

$$U_e = \frac{1-e}{e}\Phi, \quad (9)$$

oraz

$$Q_{prz} = \frac{U_{prz} + \Phi}{S_{mx}} \ln \frac{U_{prz} + \Phi}{\Phi} = Q_e \ln \frac{U_{prz} + \Phi}{\Phi} \quad (10)$$

i odejmując stronami (10) od (7) otrzymujemy

$$q - Q_{prz} = Q_e \ln \frac{u + \Phi}{U_{prz} - \Phi}. \quad (11)$$

Stąd ostatecznie mamy dla złącza Marinosa

$$s(q) = S_{mx} \exp \left[\frac{q - Q_{prz}}{Q_e} \right] \quad (12)$$

dla $q < Q_{prz}$.

3. CHARAKTERYSTYKI POMPOWANEGO ZŁĄCZA $p-n$

Zupełnie ogólną analizę sieci elektrycznej z nieliniowym operatorem impedancji, dostosowaną do przypadku złącza $p-n$ pompowanego dużym sygnałem można wyprowadzić z pracy [7] jako uproszczenie tej ostatniej na przypadek jednego tylko sygnału. Jednakże dokładność rozwiązania okupuje się skomplikowaną jego budową w postaci szeregu. Dlatego też w dalszym ciągu pracy zostaną przyjęte idealizowane warunki napięciowego i prądowego pompowania złącza wraz ze wskazaniem ich związku z warunkami spotykanymi w praktyce.

Rozważmy najpierw przypadek złącz, dla których $\gamma \neq 1$.

Przyjmijmy, że złącze $p-n$ charakteryzuje się małą w stosunku do swej reaktancji szeregową opornością strat R_s oraz że generator sterujący złącze ma małą w stosunku do reaktancji diody oporność wewnętrzną i może być uważany za idealny generator napięciowy¹⁾ napięcia sinusoidalnego o stałej pulsacji ω i amplitudzie U_M .

¹⁾ Należy dodać, że ww. warunki są w praktyce możliwe do spełnienia jedynie w zakresie fal krótkich i dolnych częstotliwości pasma ultrakrótkofalowego i są tu przytoczone głównie dla porównania z innymi warunkami oraz dlatego, że w wielu pracach dotyczących analizy układów parametrycznych warunki te są chętnie przyjmowane.

Możemy zatem przyjąć, że w stanie ustalonym przy odpowiednim doborze czasu początkowego, na złączu $p - n$ panuje napięcie

$$u(t) = U_0 - 2U_m \cos \omega t = U_0 - Ue^{j\omega t} - U^*e^{-j\omega t}. \quad (13)$$

Wówczas zgodnie z (3)

$$s(t) = S_{mx} \left(\frac{U_0 + \Phi}{U_{prz} + \Phi} \right)^\gamma \cdot \left(1 - \frac{2U_m}{U_0 + \Phi} \cos \omega t \right)^\gamma = s(U_0)(1 - 2a_u \cos \omega t)^\gamma, \quad (14)$$

$$\text{gdzie } a_u = \frac{U_m}{U_0 + \Phi} \leq \frac{1}{2} \quad (15)$$

można interpretować jako współczynnik głębokości modulacji napięcia na złączu, oraz

$$s(U_0) = S_{mx} \left(\frac{U_0 + \Phi}{U_{prz} + \Phi} \right)^\gamma \quad (16)$$

jest elastancją złącza odpowiadającą wartości stałego napięcia polaryzującego U_0 .

Przyjmijmy następnie inny przypadek, a mianowicie że generator sterujący złącze ma dużą w stosunku do reaktancji diody oporność wewnętrzną i może być uważany za idealny generator¹⁾ prądu sinusoidalnego o stałej pulsacji ω i amplitudzie I_M , której odpowiada amplituda zmian ładunku chwilowego na złączu Q_M .

Możemy zatem przyjąć, że w stanie ustalonym chwilowy ładunek na złączu jest opisany zależnością

$$q(t) = Q_0 - 2Q_M \cos \omega t = Q_0 - Qe^{j\omega t} - Q^*e^{-j\omega t}. \quad (17)$$

Wówczas zgodnie z (4)

$$\begin{aligned} s(t) &= S_{mx} \left(\frac{Q_0 + Q_\phi}{Q_{prz} + Q_\phi} \right)^{\frac{\gamma}{1-\gamma}} \cdot \left(1 - \frac{2Q_M}{Q_0 + Q_\phi} \cos \omega t \right)^{\frac{\gamma}{1-\gamma}} = \\ &= s(Q_0)(1 - 2a_q \cos \omega t)^{\frac{\gamma}{1-\gamma}}, \end{aligned} \quad (18)$$

gdzie

$$a_q = \frac{Q_M}{Q_0 + Q_\phi} \leq \frac{1}{2}$$

można interpretować jako współczynnik głębokości modulacji ładunku na złączu, oraz

$$s(Q_0) = S_{mx} \left(\frac{Q_0 + Q_\phi}{Q_{prz} + Q_\phi} \right)^{\frac{\gamma}{1-\gamma}} \quad (19)$$

¹⁾ W praktyce, w zakresie częstotliwości mikrofalowych, przy obecnie osiągalnych diodach ww. warunki dają się spełnić dość łatwo, dlatego też rozważania odnoszące się do pompowania prądowego należy uważać za bardziej zbliżone do warunków rzeczywistych.

jest elastancją złącza odpowiadającą wartości stałego ładunku polaryzującego Q_0 .

Z porównania prawych stron (14) i (18) wynika, że ich charakter jest identyczny. Spostrzeżenie to pozwala na ujednolicenie dalszej analizy obu przypadków sterowania złącza napięciem lub prądem przy poszukiwaniu odpowiednich harmoniczných elastancji.

Podtrzymując w dalszym ciągu warunek istnienia stanu ustalonego w obwodzie pompowania złącza możemy zależności (14) i (18) rozłożyć na szereg Fouriera w okresie pulsacji podstawowej

$$s(t) = \sum_{n=-\infty}^{\infty} S_n e^{jn\omega t}, \quad (20)$$

gdzie amplitudy zespolone S_n będą zależne nie tylko od indeksu sumowania n , lecz od sterującego złącze napięcia lub ładunku.

Definiując pojęcie głębokości modulacji elastancji M_{un} i M_{qn} odpowiednio dla przypadku sterowania napięciowego i prądowego

$$S_n = M_{un}s(U_0) = M_{qn}s(Q_0) \quad (21)$$

poszukujemy ostatecznie współczynników rozkładu na szereg Fouriera funkcji, która dla obu rozpatrywanych przypadków ma tę samą postać ogólną

$$m(t) = (1 - 2a \cos \omega t)^\alpha = \sum_{n=-\infty}^{\infty} M_n e^{jn\omega t}. \quad (22)$$

A zatem

$$M_n = \frac{1}{\pi} \int_0^\pi (1 - 2a \cos \omega t)^\alpha \cos(n\omega t) d(\omega t). \quad (23)$$

Podstawiając w (23) $\omega t = \Theta - \pi$ oraz korzystając z własności [15], [6] funkcji dołącznych Legendre'a $P_\alpha^n(x)$:

$$\int_0^\pi (x + \sqrt{x^2 - 1} \cos \Theta)^\alpha \cos(n\Theta) d\Theta = \frac{\pi \cdot P_\alpha^n(x)}{[a+1][a+2] \dots [a+n]}, \quad (24)$$

zdefiniowanych dla argumentu $x > 1$ przez Hobsona i stabelaryzowanych przez Lowana [8], otrzymujemy ostatecznie ogólne wyrażenie dla współczynników M_n :

$$M_n = \frac{[-1]^n [1 - 4a^2]^{\alpha/2} P_\alpha^n(1 - 4a^2)^{-1/2}}{[a+1][a+2] \dots [a+n]} \quad (25)$$

dla $\alpha < \frac{1}{2}$.

W przypadku braku¹⁾ odpowiednich tablic funkcji dołącznych Legendre'a dla indeksów α ułamkowych i argumentów $x > 1$ można skorzystać z ich związku z funkcjami hypergeometrycznymi $F(u, v; p; z)$:

$$F_{\alpha}^n(1-4a^2) = \frac{\Gamma(n+a+1) \left(\frac{4a^2}{1-4a^2} \right)^{n/2}}{2^n \cdot n! \Gamma(\alpha-n+1)} F \left[n-\alpha, n+a+1; n+1; \frac{1-(1-4a)^{-\frac{1}{2}}}{2} \right], \quad (26)$$

gdzie $\Gamma(m)$ jest funkcją gamma, zwaną także całką Eulera drugiego rodzaju.

W przypadku szczególnym pełnego występowania prądowego lub napięciowego na złączu, tj. gdy $\alpha = \frac{1}{2}$ wyrażenie (25) redukuje się do postaci

$$M_n = \frac{[-1]^n}{2} \frac{\Gamma(1+2\alpha)}{\Gamma(1+\alpha+n) \cdot \Gamma(1+\alpha-n)}. \quad (27)$$

Rozważmy teraz złącze Marisona, dla którego $\gamma = 1$.

Pompowanie napięciowe, zgodnie z (3), (13), (14) i (21) prowadzi do prostego wniosku, że $M_{u0} = 1$ oraz $M_{u1} = a$, zaś $M_{un} = 0$ dla $n \geq 2$.

W przypadku pompowania prądowego, zgodnie z (12) i (17) otrzymujemy dla złącza Marisona:

$$s(t) = S_{mx} \exp \left(\frac{Q_0 - Q_{prz}}{Q_e} \right) \cdot \exp \left(\frac{2Q_M}{Q_e} \cos \omega t \right) = s(Q_0) \exp \left(\frac{2Q_M}{Q_e} \cos \omega t \right). \quad (28)$$

Wprowadzając dla tego przypadku definicję współczynnika głębokości modulacji ładunku na złączu

$$b = \frac{Q_M}{Q_0 - Q_{prz}} \quad (29)$$

otrzymujemy ostatecznie

$$s(t) = s(Q_0) \left[\exp \left(\frac{Q_0 - Q_{prz}}{Q_e} \right) \right]^{2b \cos \omega t}. \quad (30)$$

Warto zauważyć, że identyczna podstawa potęgi w nawiasie kwadratowym występuje także w wyrażeniu na $s(Q_0)$, a mianowicie

$$s(Q_0) = S_{mx} \exp \left(\frac{Q_0 - Q_{prz}}{Q_e} \right). \quad (31)$$

Korzystając z wcześniejszych oznaczeń (7), (8), (10) ostatnią zależność można zapisać prościej

$$s(Q_0) = S_{mx} \frac{U_0 + \Phi}{U_{prz} + \Phi}, \quad (32)$$

¹⁾ Zgodnie z najlepszą wiedzą autora ww. tablice funkcji dołącznych Legendre'a dla indeksów α ułamkowych i zmiennej $x > 1$ są nieosiągalne obecnie w Polsce.

gdzie U_0 jest napięciem na złączu odpowiadającym stałemu, polaryzującemu ładunkowi Q_0 . Stąd widać, że podstawa potęgi $2b \cos x$, która jest zależna od obranego punktu pracy na charakterystyce złącza, może zawierać się w granicach $< 0, 1$). Oznaczając parametr zależny od wyboru punktu pracy przez PKT

$$PKT = \exp\left(\frac{Q_0 - Q_{prz}}{Q_e}\right) = \frac{U_0 + \Phi}{U_{prz} + \Phi} \quad (33)$$

otrzymujemy ostatecznie wzór do wyliczenia współczynników M_{qn} dla złącza Marinosa pompowanego prądowo

$$M_{qn} = \frac{1}{\pi} \int_0^\pi (PKT)^{2b \cos x} (\cos nx) dx, \quad (34)$$

gdzie:

$$\begin{aligned} PKT &\in < 0, 1 \rangle, \\ b &\in < 0, 1/2 \rangle. \end{aligned}$$

Ostatnia całka nie daje się prosto obliczyć, lecz może być wyznaczona przy pomocy jej związków z funkcjami Bessella dla urojonych argumentów i z funkcjami gamma, których tablice są łatwo dostępne.

4. NUMERYCZNE WYZNACZANIE WSPÓŁCZYNNIKÓW GŁĘBOKOŚCI MODULACJI ELASTACJI

Jak wspomniano wcześniej, całki oznaczone (23) nie są możliwe do bezpośredniego obliczenia, lecz wymagają stosowania przybliżonych metod obliczania przy użyciu tablic funkcji dołączonych Legendre'a lub dla pewnych szczególnych przypadków — tablic funkcji eliptycznych, hypergeometrycznych i tablic funkcji gamma. Podobnie całka (34) nie daje się wprost obliczyć, lecz wymaga korzystania z tablic funkcji Bessella dla argumentów urojonych oraz z tablic funkcji gamma. Korzystanie z tablic i wykonywanie dodatkowych operacji, jak sumowanie szeregów, jest dość kłopotliwe, dlatego do wyznaczenia współczynników M_{nu} i M_n posłużono się polską maszyną cyfrową ZAM2 β zainstalowaną w Ośrodku Maszyn Matematycznych w Politechnice Gdańskiej.

Obliczenia całek (23) i (34) dokonano przy zastosowaniu metody całkowania numerycznego Simpsona. Metoda ta jest opracowana dla maszyny ZAM2 β jako podprogram biblioteczny pod nazwą Całka Simpsona FUN 82. Drugim podprogramem była odpowiednia funkcja podcałkowa (23) lub (34). Obliczeń dokonano w skali binarnej 8, przy kroku całkowania $\pi/100$ wprowadzonym do podprogramu całkowania w skali binarnej 2. Otrzymano wyniki z dokładnością nie gorszą niż do piątej cyfry znaczącej¹⁾.

¹⁾ Forma niniejszej publikacji nie pozwala na przytoczenie otrzymanych tablic. Wyniki przedstawione wykreślnie uniemożliwiają zaprezentowanie pełnej dokładności uzyskanej w obliczeniach.

W przypadku obliczenia całki (23) parametrami były rodzaje złącza: liniowe $\gamma = \frac{1}{3}$ lub strome $\gamma = 1/2$ oraz numer harmonicznej. W przypadku obliczania całki (34) oprócz numeru harmonicznej jako parametr przyjęto wielkość PKT , charakteryzującą punkt pracy złącza zgodnie z (33). Wyniki obliczono w funkcji parametru głębokości modulacji prądu lub napięcia na złączu.

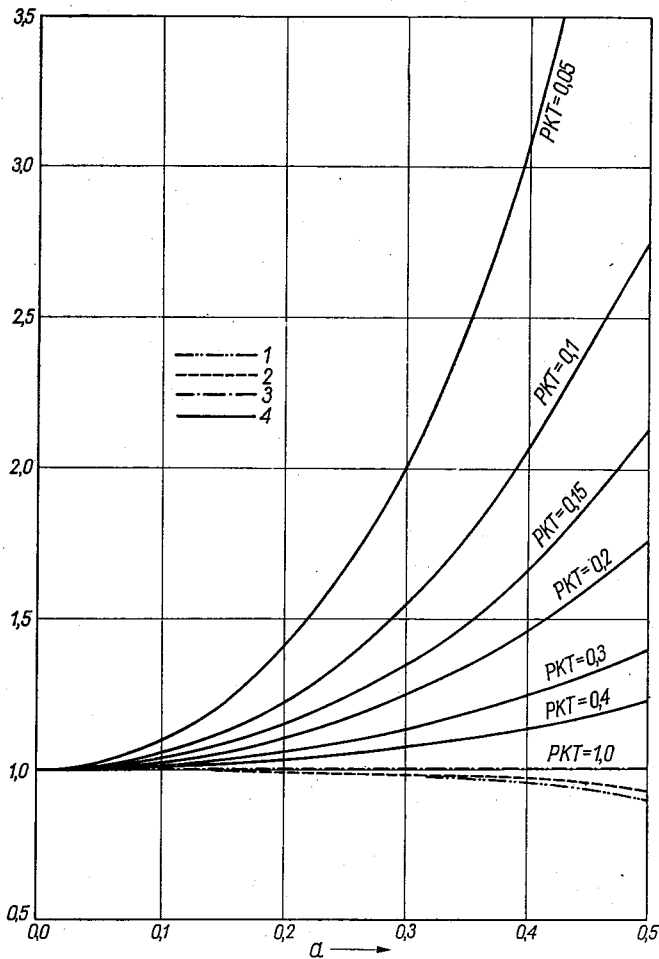
Oprócz wyznaczenia współczynników głębokości modulacji elastancji w zależności od rodzaju złącza i rodzaju pompowania (rys. 5, 6, 7, 8) obliczono także procentową zawartość drugiej i trzeciej harmonicznej w stosunku do pierwszej harmonicznej w funkcji parametru a (rys. 9, 10) oraz zawartość dalszych harmonicznych do dziesiątej włącznie dla $a = 0,5$ (tablica 1). Znajomość tych wielkości ma duże znaczenie przy projektowaniu wieloczęstotliwościowych wzmacniaczy i mieszaczy parametrycznych, a także przy projektowaniu generatorów harmonicznych z pomocniczymi obwodami jałowymi. W projektach takich najczęściej zaniedbuje się wyższe harmoniczne elastancji, podczas gdy, jak to widać z rysunków 9 i 10, zawartość np. drugiej harmonicznej w pewnych warunkach może przekraczać 50%, zaś zawartość trzeciej harmonicznej może dochodzić prawie do 25%, co wyklucza możliwość zaniedbania tak dużych wielkości.

5. DYSKUSJA OTRZYMANYCH WYNIKÓW OBLICZEŃ

Z dotychczasowych rozważań i przeprowadzonych obliczeń zilustrowanych wykresami na rysunkach 5, 6, 7, 8, 9 i 10 oraz na tablicy 1 wynikają następujące wnioski:

- a) Pompowanie prądowe złącza stromego nie zmienia składowej stałej elastancji w stosunku do jej wartości przy ładunku polaryzującym bez pompowania. Współczynnik głębokości modulacji dla pierwszej harmonicznej elastancji równa się współczynnikowi głębokości modulacji a_q ładunku na złączu, zaś amplitudy pozostałych harmonicznych elastancji są identycznie równe zeru.
- b) Pompowanie napięciowe¹⁾ złącza Marisona także nie zmienia składowej stałej elastancji w stosunku do jej wartości przy napięciu polaryzującym bez pompowania. Współczynnik głębokości modulacji pierwszej harmonicznej elastancji równa się współczynnikowi głębokości modulacji a_u napięcia na złączu, zaś amplitudy pozostałych harmonicznych elastancji są identycznie równe zeru.
- c) Złącza strome pompowane napięciowo¹⁾ oraz złącza liniowe pompowane prądowo wykazują takie same zależności elastancji od parametru a . W obu przypadkach zachodzi identyczna zależność składowej stałej elastancji od a , która maleje o ok. 7%, gdy a zmienia się od 0 do 0,5, zaś amplitudy wszystkich harmonicznych elastancji są różne od zera i wzrastają ze wzrostem a .

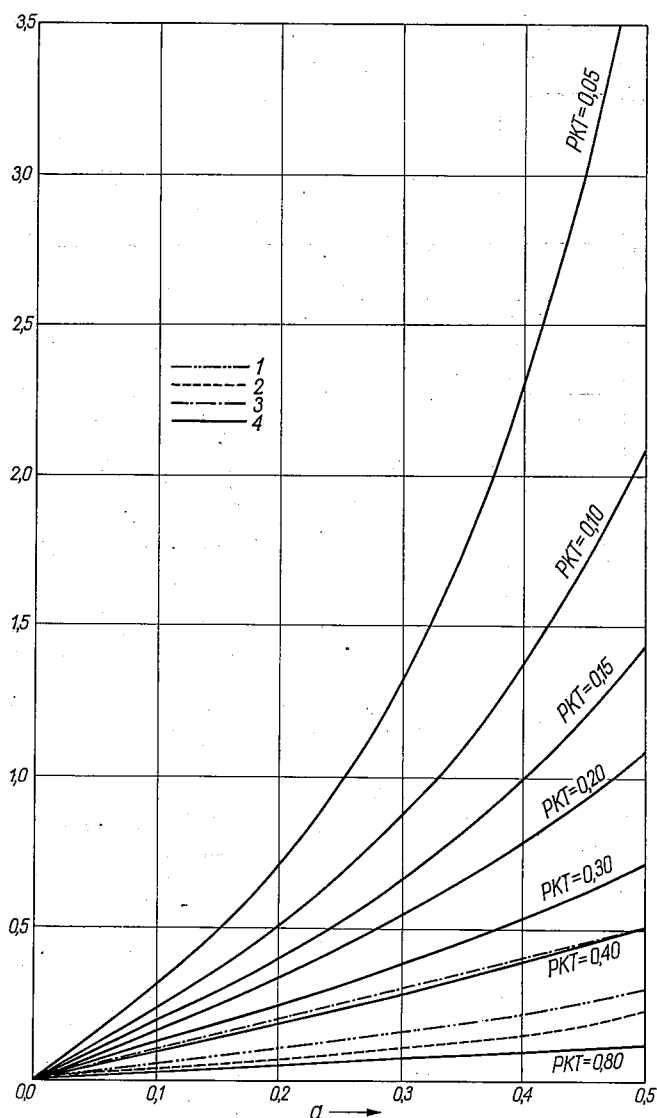
¹⁾ Porównaj uwagę ze strony 409.



Rys. 5. Wartość średnia elastancji

$$M_0 = \frac{S_0}{S(Q_0)} = \frac{S_0}{S(U_0)}; PKT = \exp \left[\frac{Q_0 - Q_{prz}}{Q_e} \right] = \frac{U_0 + \Phi}{U_{prz} + \Phi}; Q_e = \frac{U_{prz} + \Phi}{S_{mx}}$$

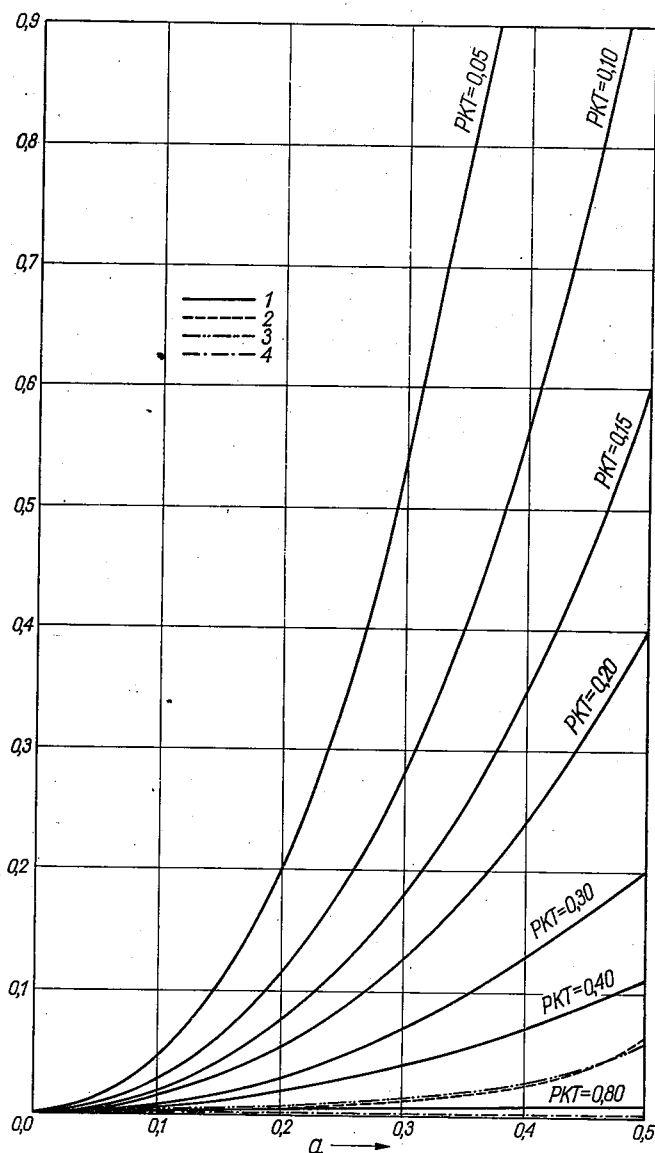
1 — złącze strome pompowane napięciowo oraz złącze liniowe pompowane prądowo, 2 — złącze liniowe pompowane napięciowo, 3 — złącze Marinosa pompowane napięciowo oraz złącze strome pompowane prądowo, 4 — złącze Marinosa pompowane prądowo



Rys. 6. Współczynnik głębokości modulacji pierwszej harmonicznej elastancji

$$M_1 = \frac{S_1}{S(Q_0)} = \frac{S_1}{S(U_0)}; PKT = \exp \left[\frac{Q_0 - Q_{prz}}{Q_e} \right] = \frac{U_0 + \Phi}{U_{prz} + \Phi}; Q_e = \frac{U_{prz} + \Phi}{S_{mx}}$$

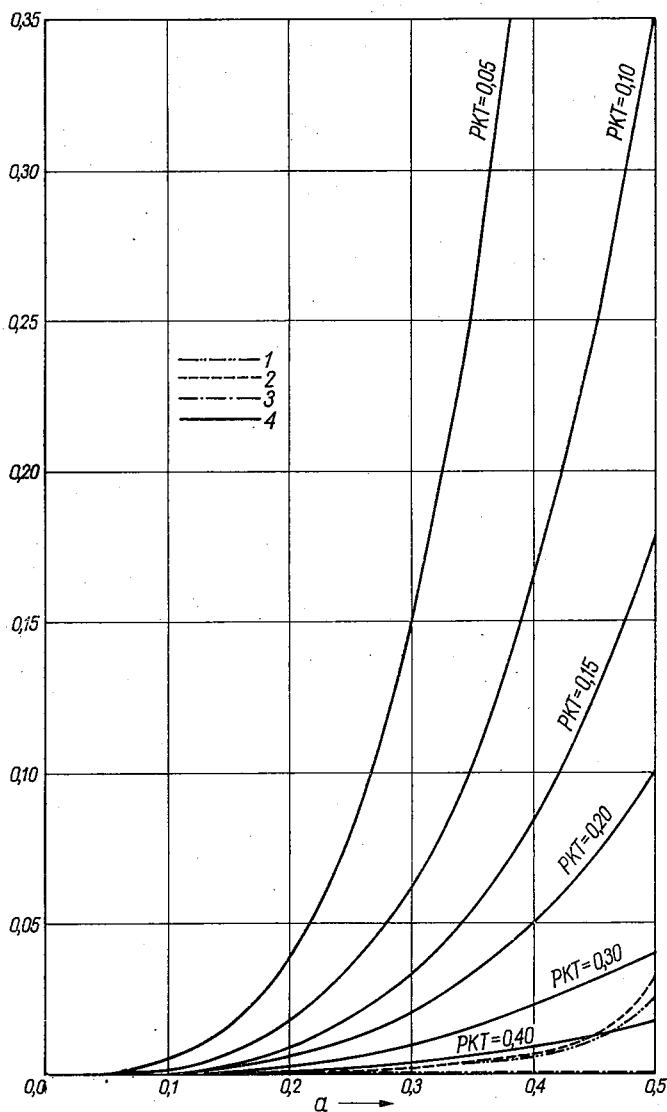
1 — złącze strome pompowane napięciowo oraz złącze liniowe pompowane prądowo, 2 — złącze liniowe pompowane napięciowo, 3 — złącze Marinosa pompowane napięciowo oraz złącze strome pompowane prądowo, 4 — złącze Marinosa pompowane prądowo



Rys. 7. Współczynnik głębokości modulacji drugiej harmonicznej elastancji

$$M_2 = \frac{S_2}{S(Q_0)} = \frac{S_2}{S(U_0)}; PKT = \exp\left[\frac{Q_0 - Q_{prz}}{Q_e}\right] = \frac{U_0 + \Phi}{U_{prz} + \Phi}; Q_e = \frac{U_{prz} + \Phi}{S_{mx}}$$

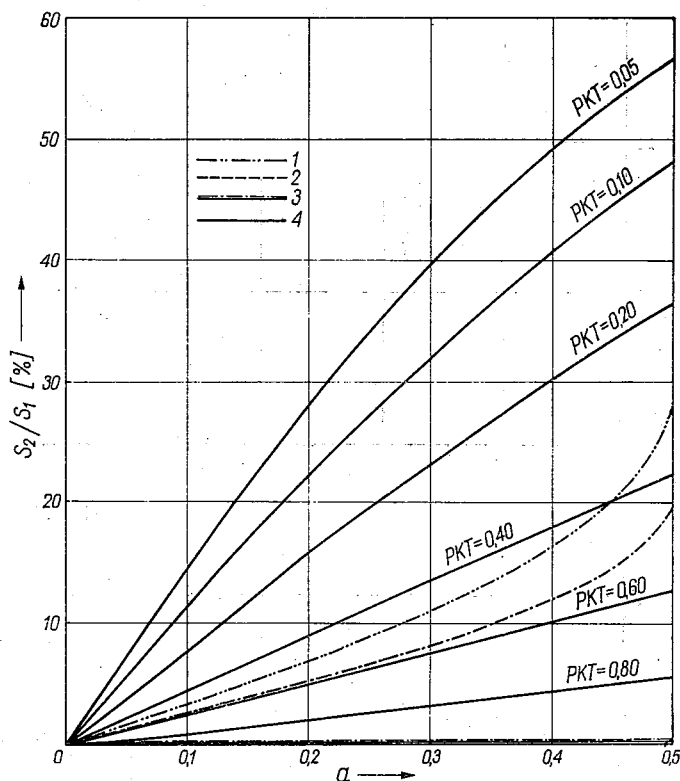
1 — złącze Marinosa pompowane prądowo, 2 — złącze liniowe pompowane napięciowo, 3 — złącze strome pompowane napięciowo oraz złącze liniowe pompowane prądowo, 4 — złącze Marinosa pompowane napięciowo oraz złącze strome pompowane prądowo



Rys. 8. Współczynnik głębokości modulacji trzeciej harmonicznej elastancji

$$M_3 = \frac{S_3}{S(Q_0)} = \frac{S_3}{S(U_0)}, \quad PKT = \exp \left[\frac{Q_0 - Q_{prz}}{Q_e} \right] = \frac{U_0 + \Phi}{U_{prz} + \Phi}; \quad Q_e = \frac{U_{prz} + \Phi}{S_{mx}}$$

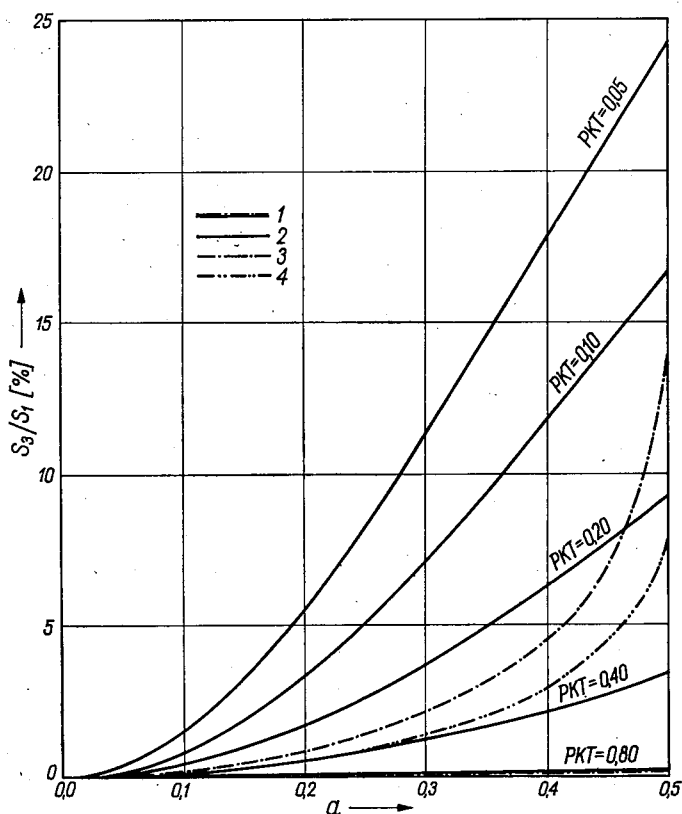
1 — złącze strome pompowane napięciowo oraz złącze liniowe pompowane prądowo, 2 — złącze liniowe pompowane napięciowo, 3 — złącze Marinsosa pompowane napięciowo oraz złącze strome pompowane prądowo, 4 — złącze Marinsosa pompowane prądowo



Rys. 9. Procentowa zawartość drugiej harmonicznej elastancji

$$PKT = \exp \left[\frac{Q_0 - Q_{prz}}{Q_e} \right] = \frac{U_0 + \Phi}{U_{prz} + \Phi}; \quad Q_e = \frac{U_{prz} + \Phi}{S_{mx}}$$

1 — złącze strome pompowane napięciowo oraz złącze liniowe pompowane prądowo, 2 — złącze liniowe pompowane napięciowo, 3 — złącze Marinosy pompowane napięciowo oraz złącze strome pompowane prądowo, 4 — złącze Marinosy pompowane prądowo



Rys. 10. Procentowa zawartość trzeciej harmonicznej elastancji

$$PKT = \exp \left[\frac{Q_0 - Q_{prz}}{Q_e} \right] = \frac{U_0 + \Phi}{U_{prz} + \Phi}; \quad Q_e = \frac{U_{prz} + \Phi}{S_{mx}}$$

1 — złącze Marinosa pompowane napięciowo oraz złącze strome pompowane prądowo, 2 — złącze Marinosa pompowane prądowo, 3 — złącze liniowe pompowane napięciowo, 4 — złącze strome pompowane napięciowo oraz złącze liniowe pompowane prądowo

T a b l i c a 1

Procentowa zawartość poszczególnych harmonicznych elastancji w stosunku do pierwszej harmonicznej przy parametrze $\alpha = 0,5$

[illegible]

- d) Złącze liniowe pompowane napięciowo¹⁾ charakteryzuje się także zależnością składowej stałej elastancji od współczynnika głębokości modulacji a napięcia na złączu. Składowa stała maleje o ok. 10%, gdy a zmienia się od 0 do 0,5. Amplitudy wszystkich harmonicznych elastancji są różne od zera i wzrastają ze wzrostem a .
- e) Pompowanie prądowe złącza Marinosa zmienia bardzo istotnie składową stałą elastancji. Np. przy parametrze $PKT = 0,05$ składowa stała rośnie ok. 5 razy, gdy a_q zmienia się od 0 do 0,5. Amplitudy poszczególnych harmonicznych rosną ze wzrostem a_q . Np. dla $PKT = 0,05$ amplituda pierwszej harmonicznej wzrasta od wartości 0 do wartości ok. czterokrotnie większej od wartości elastancji odpowiadającej ładunkowi polaryzującemu bez pompowania i do wartości równej 80% wypadkowej wartości średniej przy pompowaniu²⁾.

Procentowa zawartość drugiej i trzeciej harmonicznej elastancji w stosunku do pierwszej harmonicznej jest także największa dla złącza Marinosa przy małym parametrze PKT . Np. przy $PKT = 0,05$ i dla $\alpha = 1$ $S_2/S_1 = 56,75\%$, zaś $S_3/S_1 = 24,23\%$ ³⁾.

W miarę wzrostu PKT , tj. gdy stały ładunek polaryzujący złącze zbliża się do ładunku odpowiadającego przebicciu złącza, wspomniane wyżej parametry ulegają zmniejszeniu i w granicy, gdy $PKT = 1$, złącze takie całkowicie zatracą swoje właściwości, bowiem wówczas $M_{q0} = 1$ oraz $M_{qn} = 0$ dla $n > 0$ niezależnie od a .

Złącze Marinosa pompowane prądowo jest bardziej efektywne od wszystkich innych dyskutowanych złącz, jeśli chodzi o otrzymanie pierwszej, drugiej i ewentualnie trzeciej harmonicznej elastancji, natomiast dalsze harmoniczne w tym złączu i przy tym pompowaniu są bardzo małe (tablica 1) i dla dużych indeksów harmonicznych (np. dla szóstej i dalszych) są o jeden lub kilka rzędów wielkości mniejsze w stosunku do odpowiednich harmonicznych elastancji w złączu liniowym pompowanym prądowo lub napięciowo lub w złączu stromym pompowanym napięciowo.

6. DYSKUSJA ZASTOSOWAŃ PRAKTYCZNYCH ROZWAŻANYCH ZŁĄCZ

Z dotychczasowych rozważań wynikają następujące uwagi praktyczne:

- a) W układach parametrycznych typu diodowych wzmacniaczy i mieszaczy wieloczęstotliwościowych trzeba się liczyć z zasadniczym wpływem wyższych harmonicznych elastancji na pracę tych układów. Może to pociągnąć za sobą konieczność: 1° — rozszerzenia analizy teoretycznej układu na przypadek uwzględnienia tych harmonicznych (co w większości

¹⁾ Porównaj uwagę ze strony 409.

²⁾ Ostatnia liczba przy pompowaniu złącz omawianych w punktach a) i b) wynosi 50%, a złącz omawianych w punktach c) i d) wynosi odpowiednio ok. 30% i 24%.

³⁾ Odpowiednie liczby dla złącz dyskutowanych w punkcie c) są 20,00% i 8,75%, zaś w punkcie d) — 28,58% i 14,29%.

układów praktycznych nie zostało dotąd zrobione) oraz 2° — wykonania odpowiedniej konstrukcji wykorzystującej w sposób optymalny własności bardziej dogłębnie przeanalizowanego układu.

b) Do pompowania na podharmonicznej częstotliwości pompującej diod pojemnościowych pracujących w układach wzmacniaczy i mieszaczy parametrycznych nadaje się szczególnie dobrze złącze Marinosa pompowane prądowo, pod warunkiem, że indeks podharmonicznej nie przekracza trzech.

c) Zastosowanie złącza Marinosa w waraktorowych powielaczach częstotliwości może dać duże korzyści przy podwajaniu i potrajaniu. Powielanie częstotliwości z indeksem multiplikacji większym niż trzy przy zastosowaniu tego złącza jest niecelowe lub wręcz błędne.

d) Zależność składowej średniej S_0 elastancji złącza od mocy pompującej złącze jest inna dla złącz stromych i liniowych niż dla złącza Marinosa pompowanego prądowo. W przypadku tego ostatniego elastancja średnia rośnie bardzo silnie w miarę wzrostu mocy doprowadzonej do złącza (rys. 1), podczas gdy w złączu stromym i liniowym elastancja ta jest bądź to stała, bądź też maleje bardzo nieznacznie (o kilka procent). Wspomnianą wyżej własność złącza Marinosa, objawiającą się silnym rozstrajaniem obwodów rezonansowych w praktycznych układach wzmacniaczy, mieszaczy i powielaczy parametrycznych, należy uznać za poważną wadę utrudniającą konstrukcję i eksploatację.

Autor pragnie podziękować Władzom Politechniki Gdańskiej za umożliwienie skorzystania z maszyny elektronowej ZAM 2 Beta przy wykonywaniu obliczeń do niniejszej pracy oraz pracownikom Ośrodka Maszyn Matematycznych P.G. za koleżeńską pomoc przy programowaniu i liczeniu na maszynie.

*Politechnika Gdańska
Katedra Radiotechniki Odbiorczej*

WYKAZ LITERATURY

1. Eng S. T.: Characterization of Microwave Variable Capacitance Diodes. — I.R.E. Trans. on Microwave Theory and Techn., vol. MTT-9, str. 11—22, January 1961.
2. Grabowski K.: Uwagi o małosygnałowej analizie liniowej sieci elektrycznej z reaktancją okresowo zmienną w czasie. — Arch. Elektrotechniki (w druku).
3. Helgesson A. L.: Nonsymmetrical Properties of Nonlinear Elements in Low- and High-Impedance Circuits. — Proc. I.R.E., vol. 49, str. 1569, October 1961.
4. Humphreys B. L.: Characteristics of Broadband Parametric Amplifiers Using Filter Networks. — Proc. I.E.E., vol. 111, str. 264—274, February 1964.
5. Hyde F. J., Tucker D. G.: Parametric Amplifiers: Static and Dynamic Inductance and Capacitance and their Significance in the Non Linear and Time Varying Applications. — J. of. B.I.R.E., vol. 25, str. 353—355, April 1963.

6. Leeson D. B.: Capacitance and Charge Coefficients for Varactor Diodes. — Proc.I.R.E., vol. 50, str. 1854, August 1962.
7. Lenkowski J.: Power Series Development of Nonlinear Immitance Operator in the Presence of two External Signals. III Konferenz über Nichtlineare Schwingungen, Berlin 25—30 Mai 1964.
8. Lowan A.: Tables of Associated Legendre Functions. National Bureau of Standards, Columbia Univ. Press, New York 1945.
9. Marinos P. N.: A Note on Varactor Diodes. — Proc. I.E.E.E., vol. 51, str. 388—9, February 1963.
10. Penfield P., jr.: Fourier Coefficients of Power-Law Devices. — Journal of the Franklin Institute, vol. 273, str. 107—122, February 1962.
11. Penfield P., jr., Rafuse R. P.: Varactor Applications. — The M.I.T. Press, Cambridge, Mass. 1962.
12. Shockley W.: The Theory of p-n Junctions in Semiconductors and p-n Junction Transistors. — Bell Syst. Techn. J., vol. 28 str. 435—489, July 1949.
13. Thomson G.H.B.: Dependence of Parametric Element Nonlinearity on Tuning Circuit. — Proc. I.R.E., vol. 50, str. 1845—1846, August 1962.
14. Uhler A., jr.: The Potential of Semiconductor Diodes in High Frequency Communications. — Proc. I.R.E., vol. 46, str. 1099—1115, June 1958.
15. Whittaker E. T., Watson G. N.: A Course of Modern Analysis. Cambridge Univ. Press, London 1958 (czwarte wydanie).

K. GRABOWSKI

CHARACTERIZATION OF PUMPED $p - n$ JUNCTIONS

Summary

A mathematical representation of the dependence of the $p - n$ junction dynamic elastance on pumping voltage or current is presented in the paper. Results of the numerical computations of the depth of the elastance modulation versus the depth of voltage or charge modulation on the junction are plotted. Practical usefulness of the results carried out is indicated.

K. GRABOWSKI

LES CARACTÉRISTIQUES D'UNE JONCTION $p - n$ POMPÉE

Resumé

Dans l'article on a donné l'analyse mathématique de la dépendance non linéaire entre l'élastance dynamique d'une jonction $p - n$ et le courant ou la tension de pompe. On a présenté les résultats du calcul numérique des coefficients de la profondeur de modulation ainsi que des harmoniques de l'élastance en fonction de la profondeur de modulation du courant ou de la tension appliqués à la jonction. On a discuté aussi l'utilité pratique des résultats obtenus.

K. GRABOWSKI

DIE KENNLINIEN EINES GEPUMPTEN ÜBERGANGS $p-n$

Zusammenfassung

In diesem Aufsatz wurde die mathematische Repräsentation der non-linearen Abhängigkeit einer dynamischen Elastanz des Übergangs $p-n$ von der steuernden Spannung bzw. von dem steuernden Strom dargestellt. Es wurden hier weiter die numerischen Berechnungsergebnisse von den Koeffizienten der Modulationstiefe sowie des Inhalts der Harmonischen der Elastanz, abhängig von der Modulationstiefe der Spannung bzw. des Stromes am Übergang dargestellt. Man hat auch die praktische Nutzlichkeit der erhaltenen Ergebnisse erörtert.

K. ГРАБОВСКИ

ХАРАКТЕРИСТИКИ НАКАЧИВАЕМОГО $p-n$ ПЕРЕХОДА

Резюме

В статье дано математическое представление нелинейной зависимости величины S , обратной к динамической ёмкости $p-n$ перехода полупроводникового диода, от возбуждающего напряжения или тока. Приведены, в виде графиков, результаты нумерических вычислений коэффициентов глубины модуляции величины S и процентного содержания ее гармоник в зависимости от глубины модуляции напряжения или заряда на переходе. Полученные результаты обсуждено с точки зрения их практического использования.

RYSZARD GRZEGORZ STRUŻAK

Rezonanse w ceramicznych kondensatorach rurkowych¹⁾

Rękopis dostarczono 16.12.1963

Częstotliwość rezonansu wewnętrznego jest jednym z zasadniczych parametrów określających wartość użytkową kondensatorów przy wielkiej częstotliwości. W niniejszej pracy autor rozpatruje ceramiczne kondensatory rurkowe mające dwie końcówki. Omawia układy zastępcze kondensatorów uwzględniające wpływ umiejscowienia końcówek, wpływ naskórkowości, wpływ sprzężenia końcówek z elektrodami oraz wpływ wymiarów geometrycznych kondensatora i własności dielektryku. Wpływ indukcyjności końcówek jest rozpatrzony oddzielnie. W oparciu o analizę tych układów autor podaje wzory na częstotliwość rezonansu wewnętrznego. Wykazuje przy tym, że ta częstotliwość w dużym stopniu zależy od umiejscowienia końcówek. W dotychczasowej literaturze taka zależność nie jest brana pod uwagę. Następnie autor omawia optymalne umiejscowienie końcówek, zapewniające maksymalną częstotliwość rezonansową kondensatora.

W dalszej części pracy autor podaje opis i wyniki doświadczalnych badań częstotliwości rezonansowych ponad 20 kondensatorów, a także metodę i wyniki porównania teorii z doświadczeniem.

Autor potwierdził eksperymentalnie, że bez zmiany materiału, kształtu i wymiarów kondensatora, jedynie przez zmianę punktów dołączenia końcówek do elektrod, częstotliwość rezonansu wewnętrznego może być zmieniana w szerokich granicach: na przykład od 57 MHz do 190 MHz. Istnieje zatem nowa możliwość znacznego powiększenia wartości użytkowej omawianych kondensatorów bez dodatkowych kosztów i istotnych zmian technologii.

1. WSTĘP

Obecnie w 200 lat po odkryciu i w 100 lat od chwili pierwszego zastosowania technicznego [55] kondensatory elektryczne są jednymi z podstawowych podzespołów i wchodzi w skład rozmaitych urządzeń elektrycznych. Zarówno w Polsce, jak i w innych krajach są one produkowane w skali masowej, a produkcja ich z roku na rok wzrasta. Dla przykładu warto podać, że w Stanach Zjednoczonych A.P. w ciągu zaledwie połowy 1958 roku wyprodukowano ich ok. 588 milionów sztuk (z tego prawie połowę kondensatorów ceramicznych stałych niskiego napięcia), a wartość rocznej ich produkcji wyniosła w 1961 r. przeszło 300 milionów dolarów [8], [9].

1) Artykuł nagrodzony na konkursie prac naukowych zorganizowanym w 1963 r. przez Polskie Towarzystwo Elektrotechniki Teoretycznej i Stosowanej, Oddział we Wrocławiu.

W związku z tak masową produkcją kondensatorów jest zrozumiałe duże zainteresowanie nimi. Stale są patentowane nowe rozwiązania konstrukcyjne i technologiczne, pozwalające uzyskiwać nowe typy kondensatorów, doskonalsze pod pewnymi względami lub tańsze od dotychczas znanych. W prasie technicznej ukazują się ciągle publikacje poświęcone różnym zagadnieniom związanym z produkcją, technologią czy też zastosowaniami kondensatorów. Dla przykładu, wykaz literatury podany w monografii [7] zawiera 488 pozycji, przy czym nie obejmuje on wielu publikacji ogłoszonych w języku rosyjskim.

Stosunkowo znaczna ilość kondensatorów jest przeznaczona do pracy przy wielkiej częstotliwości. Wchodzą tutaj w rachubę kondensatory stosowane zarówno w obwodach rezonansowych, jak też w obwodach zasilania w urządzeniach elektronicznych wielkiej częstotliwości. Jak wiadomo, takie kondensatory powinny zachowywać określone własności w całym zakresie częstotliwości istotnych dla działania urządzenia. Znajomość tych własności oraz czynników, od których one zależą, jest konieczna nie tylko do prawidłowego konstruowania samych kondensatorów, ale również do celowego konstruowania urządzeń, w których takie kondensatory mają pracować.

Normalne wykorzystywanie kondensatora może odbywać się tylko przy częstotliwościach mniejszych od jego częstotliwości rezonansowej. Powyżej tej częstotliwości kondensator zachowuje się podobnie jak dławik. Dla stabilnej pracy urządzenia wymaga się najczęściej [46], aby maksymalna częstotliwość pracy kondensatora była co najmniej 2—3 razy mniejsza od jego częstotliwości rezonansowej. Wynika stąd, że ta częstotliwość jest jednym z podstawowych parametrów kondensatora, tak samo istotnym jak np. pojemność czy napięcie znamionowe. Fakt ten podkreślało niejednokrotnie wielu autorów; wypada w tym miejscu przypomnieć, że jednym z pierwszych (1930 r.) był Polak — Wilhelm Rotkiewicz [29]. Fakt ten znalazł również odbicie w wymaganiach wielu norm i warunków technicznych na kondensatory (por. np. [17], [18], [19]).

Niniejsza praca jest oparta na dysertacji [40]. Opisano w niej próby wykrycia i zbadania najważniejszych czynników określających maksymalną, możliwą do osiągnięcia częstotliwość rezonansową, tj. częstotliwość rezonansu wewnętrznego ceramicznych kondensatorów rurkowych o dwóch końcówkach. Została ona podzielona na dwie części.

W pierwszej części (teoretycznej), na podstawie rozpatrzenia jakościowego podstawowych zjawisk fizycznych zachodzących w rzeczywistym kondensatorze, ustalono kilka modeli teoretycznych równoważnych kondensatorowi rzeczywistemu w różnych zakresach częstotliwości. Następnie przeprowadzono analizę tych modeli.

Jednym z istotnych elementów tej części pracy jest analiza modelu zaproponowanego po raz pierwszy przez Ryżkę [30]. Ta część pracy opiera się głównie na artykułach [36], [37], [38]. Rozważania teoretyczne są do-

prowadzone do prostych zależności analitycznych, umożliwiających obliczenie częstotliwości rezonansowej kondensatora (lub jego indukcyjności) na podstawie jego pojemności, wymiarów i wielkości związanej z umiejscowieniem końcówek.

W drugiej części pracy (eksperymentalnej) podano szczegółowy opis metody pomiarowej i wyniki pomiarów częstotliwości rezonansowej szeregu kondensatorów. Oprócz tego podano w niej opis metody i wyniki doświadczalnego sprawdzenia wniosków wynikających z rozważań teoretycznych przedstawionych w części pierwszej.

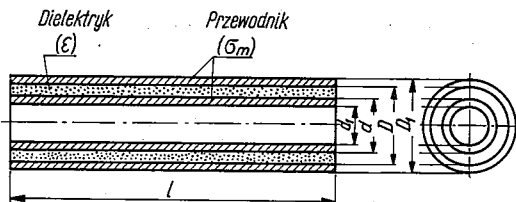
2. CZĘŚĆ TEORETYCZNA

2.1. Obiekt badań

Zasięg niniejszej pracy jest ograniczony do rozpatrzenia niskonapięciowych kondensatorów stałych małej mocy, zawierających rurkę ceramiczną pokrytą obustronnie warstwami srebra (lub innego materiału przewodzącego), stanowiącymi elektrody (rys. 1).

Takie kondensatory rurkowe mogą mieć różną ilość rozmaicie ukształtowanych końcówek. Przykłady kilku wariantów przedstawiono na ry-

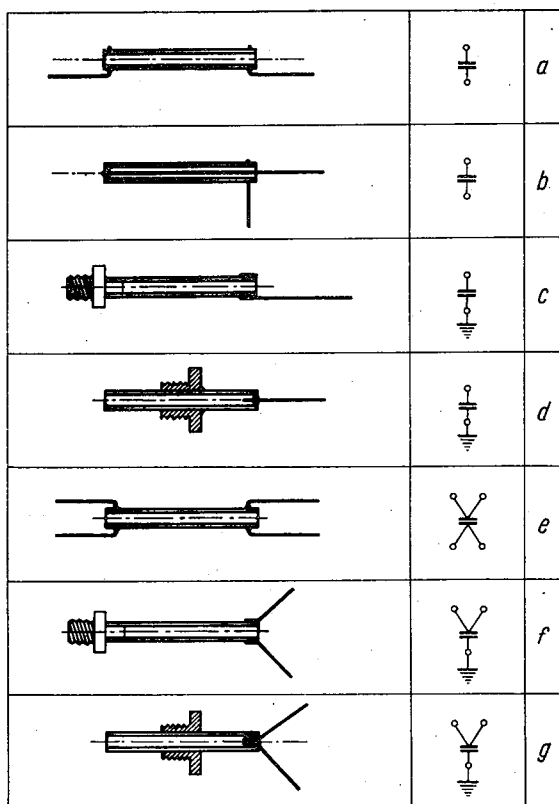
Rys. 1. Zasadnicze wymiary ceramicznego kondensatora rurkowego omawianego w artykule



sunku 2. Niniejsza praca dotyczy tylko takich kondensatorów, które mają końcówki dołączone do każdej elektrody w jednym tylko miejscu, tj. takich, które przy skracaniu długości końcówek (w granicy) można zastąpić dwójnikiem elektrycznym. Wynika stąd, że artykuł nie dotyczy przepustowych kondensatorów rurkowych, które wymagają odmiennej analizy teoretycznej. Własności takich ceramicznych (rurkowych) kondensatorów przepustowych przy wielkiej częstotliwości omówione są m.in. w pracach [31] i [44]. Ponadto nie będziemy także rozpatrywać kondensatorów z jedną końcówką umieszczoną wewnątrz rurki ceramicznej (typ „b” z rys. 2), ponieważ od razu można przewidzieć, że z taką konstrukcją jest związana duża indukcyjność kondensatora, większa niż kondensatorów pozostałych typów, uwidocznionych na rys. 2. Istotnie, przy takiej konstrukcji długość końcówki „wewnętrznej” nie może być mniejsza od długości rurki ceramicznej, zatem indukcyjność końcówek nie może być tu nigdy równa, ani nawet bliska zeru. Stąd, w odróżnieniu od pozostałych kon-

strukcji, nie można przybliżyć się tutaj do maksymalnej częstotliwości rezonansowej kondensatora, wynikającej z indukcyjności samych tylko elektrod.

Rurka ceramiczna kondensatora ma kształt wydrążonego walca kołowego. Długość jej zawiera się w granicach od kilku milimetrów do kilkudziesięciu milimetrów. Średnica zewnętrzna wynosi od kilku milimetrów do kilkunastu milimetrów, a grubość ścianek — od kilku dziesiątych



Rys. 2. Różne warianty konstrukcyjne ceramicznych kondensatorów rurkowych i ich schematy

części milimetra do kilku milimetrów. Na zewnętrzną i wewnętrzną powierzchnię rurki naniesione są warstwy srebra o grubości rzędu kilku (2—10) mikronów (por. np. [45], str. 65 i dalsze, oraz [55], str. 89). W praktyce dokładne określenie grubości elektrod nie jest sprawą tak łatwą jak mogłoby się wydawać. Na skutek przenikania cząsteczek srebra w głąb dielektryka, oraz w wyniku reakcji chemicznej na styku srebra i dielektryka powstaje często półprzewodząca warstwa, utrudniająca ustalenie wyraźnej granicy między przewodnikiem i dielektrykiem (por. np. [55], str. 70).

Mając to na uwadze, w dalszym ciągu pod grubością elektrod i grubością dielektryka będziemy rozumieć pewne zastępcze wielkości efektywne.

2.2. Przyczyny rezonansów w rurkowych kondensatorach ceramicznych

Przy rozpatrywaniu rezonansów w omawianych kondensatorach należy wziąć pod uwagę:

- 1) wpływ indukcyjności końcówek,
- 2) wpływ odbić poprzecznej fali elektromagnetycznej (TEM) w dielektryku kondensatora (efekt „linii długiej”),
- 3) wpływ umiejscowienia końcówek,
- 4) wpływ odbić fal elektromagnetycznych wyższych rzędów w dielektryku kondensatora (efekt falowodu),
- 5) ewentualnie wpływ rezonansów molekularnych w dielektryku kondensatora, piezoelektryczność i inne zjawiska.

2.2.1. Wpływ indukcyjności końcówek

Jeżeli pominąć wzajemne oddziaływanie prądów płynących w końcówkach kondensatora i jego elektrodach, to wpływ indukcyjności końcówek daje się łatwo oszacować.

Weźmy najpierw pod uwagę kondensator o dwóch końcówkach (rys. 1, typy a — d), przedstawiony schematycznie na rysunku 3a. Łatwo wykazać, że dla takiego kondensatora częstotliwość rezonansowa jest tym większa, im mniejsze są indukcyjności obu jego końcówek (L_1 , L_2) oraz im większa jest indukcyjność wzajemna między nimi (M). Istotnie, jeżeli przez Z_k oznaczyć impedancję wypadkową kondensatora (z końcówkami), a przez Z jego impedancję „wewnętrzną” (długość końcówek równa zeru), to jak wynika z rys. 3a

$$Z_k = j\omega[(L_1 + L_2) - 2M_{12}] + Z.$$

Widać stąd natychmiast, że wpływ końcówek może być znacznie zmniejszony przez uczynienie

$$[(L_1 + L_2) - 2M_{12}] \approx 0$$

na przykład na drodze zbliżenia i równoległego ułożenia obu końcówek (por. [23], str. 59).

W przypadku kondensatora o 3 lub 4 końcówkach (typy e — g z rys. 2) należy operować impedancją przejściową (wzajemną) kondensatora, określoną jako stosunek napięcia biegu jałowego „na wyjściu” kondensatora do prądu „wejściowego” (por. np. [33]). Można wykazać, że dla kondensatora o 3 końcówkach impedancja ta wynosi:

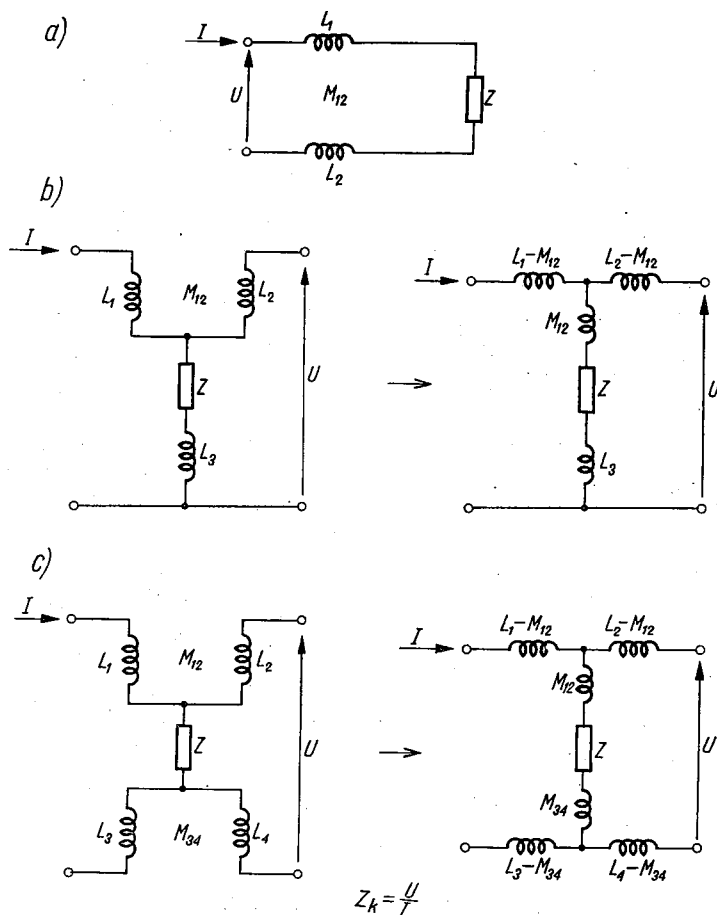
$$Z_k \approx j\omega(L_3 + M_{12}) + Z,$$

przy czym oznaczenia, analogiczne do poprzednich, podane są na rysunku 3b. Zatem przy takiej konstrukcji kondensatora jest istotne zapewnienie małej indukcyjności wzajemnej między końcówkami symetrycznymi i małej indukcyjności końcówki niesymetrycznej (L_3). Indukcyjności własne końcówek symetrycznych nie odgrywają tutaj roli.

Dla kondensatora o 4 końcówkach (rys. 3c)

$$Z_k \approx j\omega(M_{12} + M_{34}) + Z.$$

W takim kondensatorze jest ważne zachowanie małych indukcyjności wzajemnych między symetrycznymi końcówkami, natomiast indukcyjności własne tych końcówek nie są szkodliwe.



Rys. 3. Układy zastępcze kondensatorów z dwoma (a), trzema (b) i czterema (c) końcówkami

Powyższe elementarne rozważania odnoszą się do przypadku, kiedy można pominąć wzajemne oddziaływanie prądów w elektrodach i końcówkach kondensatora. Tylko w takim przypadku skracanie i wydłużanie

końcówek nie ma wpływu na impedancję „wewnętrzną” kondensatora (Z). Ze względu na olbrzymią ilość różnych możliwych układów geometrycznych wzajemnego ułożenia końcówek i elektrod ogólne ujęcie tego wpływu jest trudne i nie wydaje się zbyt celowe; pewien przypadek szczególny dyskutowany jest w jednym z dalszych punktów.

2.2.2. Propagacja poprzecznej fali elektromagnetycznej (TEM)

Biorąc pod uwagę niewielkie rozmiary geometryczne omawianych kondensatorów można by przypuszczać, że zachowują się one jak elementy o stałych skupionych. Rzeczywiście, jeżeli przyjąć, że układ ma parametry rozłożone, jeżeli jego rozmiary przekraczają $\frac{1}{10}$ część długości fali, to częstotliwości, przy których rozkład prądów i napięć na zewnętrznych powierzchniach kondensatora jest w przybliżeniu stały, ograniczone są wartością:

$$f \leq \frac{3 \cdot 10^8}{10l}.$$

Jeżeli w powyższym wzorze podstawić długość kondensatora l w metrach, to częstotliwość f wyrażona będzie w Hz. Dla kondensatora o długości $l = 3$ cm otrzymuje się z prawej strony powyższego wzoru częstotliwość równą 1000 MHz.

Jednocześnie jednak w dielektryku kondensatora następuje znaczne skrócenie fali. Jeżeli np. względna przenikalność elektryczna ceramiki wynosi 2500, to długość fali w dielektryku wynosi zaledwie 2% długości fali w powietrzu i kondensator jest jak gdyby 50-krotnie dłuższy. Zakres częstotliwości, dla których rozkład napięć na wewnętrznych (przylegających do dielektryku) powierzchniach można uważać w przybliżeniu za stały, ograniczony jest częstotliwością.

$$f \leq \frac{3 \cdot 10^8}{10\sqrt{\epsilon}l},$$

gdzie:

- ϵ — względna przenikalność elektryczna dielektryku,
- l — czynna długość kondensatora w metrach,
- f — częstotliwość w hercach.

Jeżeli w powyższym wzorze podstawić jak poprzednio $\epsilon = 2500$ i $l = 3$ cm, to teraz otrzyma się po prawej stronie wzoru zamiast 1000 — 20 MHz.

Zachodzi tutaj pozorny paradoks: ten sam kondensator może być jednocześnie elektrycznie „krótki” i „długi”, w zależności od tego, czy jego „wymiar elektryczny” są określane z zewnątrz, w powietrzu, czy od wewnątrz, w dielektryku. (Ta szczególna właściwość zostanie wykorzystana w dalszych częściach pracy).

Wynika stąd, że poprawny układ zastępczy kondensatora, odpowiedni dla szerokiego zakresu częstotliwości, powinien zawierać elementy o parametrach rozłożonych. Dla małych częstotliwości lub w wąskim pasmie częstotliwości możliwe jest utworzenie uproszczonego układu równoważnego, zawierającego skupione pojemności, oporności i indukcyjności.

Elektrody kondensatora tworzą układ współosiowych przewodników przedzielonych dielektrykiem. W przestrzeni między tymi przewodnikami może istnieć pole elektromagnetyczne, w szczególności w postaci fali poprzecznej (TEM). Jeżeli długość kondensatora jest znacznie większa od jego średnicy, a jego obwód znacznie mniejszy od długości fali w dielektryku, to pole TEM jest rozkładem dominującym. Wskutek skończonej długości toru elektrycznego utworzonego przez elektrody kondensatora i na skutek niejednorodności na jego końcach powstają odbicia, w wyniku czego w rurce ceramicznej istnieje fala stojąca. Częstotliwość, przy której w kondensatorze układa się ćwiartka tej fali, wynosi

$$f_0 = \frac{75 \cdot 10^6}{\sqrt{\epsilon l}}, \quad (1a)$$

gdzie:

ϵ — względna przenikalność elektryczna ceramiki,

l — czynna długość kondensatora (w metrach),

f_0 — częstotliwość odpowiadająca ćwiartce fali (w Hz).

Dalej pokażemy, że częstotliwość ta jest bezpośrednio związana z częstotliwością rezonansową kondensatora.

Z zależności (1a) wynika, że dla wszystkich kondensatorów wykonanych z takiego samego materiału ceramicznego (ϵ) i o takiej samej długości geometrycznej, częstotliwość f_0 jest taka sama, niezależnie od ich pojemności i pozostałych wymiarów.

Zależność (1a) może być przedstawiona także w innych, równoważnych postaciach. Stosując elementarne przekształcenia można otrzymać

$$f_0 = \frac{1}{4Z_0 C}. \quad (1b)$$

Tutaj częstotliwość f_0 wyrażona jest w Hz, pojemność kondensatora (określona przy małej częstotliwości) C — w faradach, a oporność falowa toru utworzonego przez elektrody kondensatora — w omach.

Ze związku (1b) widać, że wszystkie kondensatory o takiej samej oporności falowej (Z_0) i jednakowej pojemności (C) mają jednakową częstotliwość f_0 , niezależnie od ich długości, średnic i materiału, z którego są wykonane.

Stosując dalsze przekształcenia wzorów (1a) i (1b) można uzyskać jeszcze jedną postać równoważną wzoru (1a)

$$f_0 = \frac{560}{\sqrt{Cl \cdot \ln \frac{D}{d}}}, \quad (1c)$$

gdzie — podobnie jak poprzednio — częstotliwość f_0 wyrażona jest w Hz, pojemność C — w F, a długość czynna kondensatora l — w metrach. D i d oznaczają tutaj zewnętrzną i wewnętrzną średnicę rurki ceramicznej, odpowiednio, i powinny być wyrażone w jednakowych jednostkach.

Zależność (1c) wskazuje, że dla wszystkich kondensatorów o takich samych pojemnościach, długościach czynnych i stosunkach średnic rurki, częstotliwość f_0 jest taka sama, niezależnie od materiału ceramicznego i średnic rurki kondensatora.

Jeżeli do wzoru (1c) podstawić wartości liczbowe: pojemność $C = 520$ pF, długość czynną kondensatora $l = 62,4$ mm oraz stosunek średnic $\frac{D}{d} = 1,095$, to otrzyma się $f_0 = 325$ MHz.

2.2.3. Wpływ umiejscowienia końcówek

Produkowane obecnie kondensatory z reguły mają końcówki umieszczone na obu końcach rurki, dołączone do elektrod w punktach najbardziej od siebie oddalonych (rys. 2a, 2c, 2e, 2f). Nie jest to jednak jedyny możliwy sposób umiejscowienia końcówek.

Elektroda zewnętrzna kondensatora jest dostępna na całej swojej powierzchni zewnętrznej i jedna z końcówek może być do niej dołączona w dowolnym miejscu. Inaczej jest z końcówką drugą, która — ze względów technologicznych — musi być dołączona do jednego ze skrajnych punktów elektrody wewnętrznej.

Przy małych częstotliwościach, przy których prądy płyną w przybliżeniu równomiernie w całym przekroju elektrod, rozróżnianie ich powierzchni zewnętrznych i wewnętrznych jest nieuzasadnione. W miarę wzrostu częstotliwości coraz silniejszy staje się wpływ naskórkowości i przy odpowiednio wielkiej częstotliwości prądy płyną wyłącznie po powierzchniach elektrod. W tym przypadku szybkozmienne pole elektromagnetyczne nie wnika w głąb przewodników i rozróżnianie zewnętrznych i wewnętrznych powierzchni elektrod jest niezbędne do prawidłowego opisu procesów fizycznych zachodzących w rzeczywistym kondensatorze.

Jeżeli skuteczną grubość elektrod oznaczyć przez δ , to proste rozważania doprowadzają nas do wniosku, że przy częstotliwościach spełniających nierówność

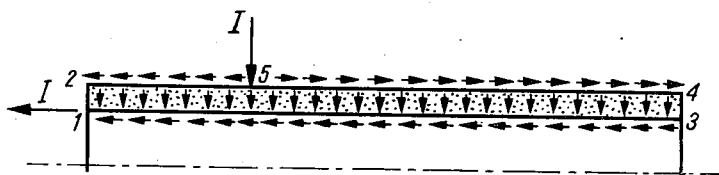
$$f \gg \frac{4 \cdot 10^{-3}}{\delta^2}$$

powierzchnie zewnętrzne i wewnętrzne elektrod kondensatora można uważać za „odizolowane” od siebie.

Dla elektrod najcieńszych (2 mikrony) otrzymuje się po prawej stronie powyższej zależności częstotliwość 1000 MHz, natomiast dla grubych

(10 mikronów) — 40 MHz. Zatem kondensator pracujący w szerokim zakresie częstotliwości, w pewnej części tego zakresu może zachowywać się tak jakby miał elektrody „cienkie” (w porównaniu z głębokością wnika-
nia), a w innej — jakby miał elektrody „grube” (w porównaniu z głębokością wnika-
nia).

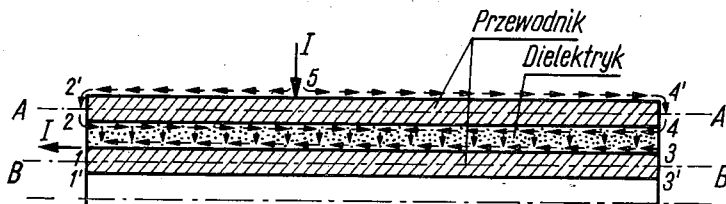
Jak wiadomo, w warunkach założonej symetrii geometrycznej i przy istnieniu tylko poprzecznej fali elektromagnetycznej układ charakteryzuje się pełną osiową symetrią elektryczną. Wewnętrzne powierzchnie



Rys. 4. Rozpływ prądów w kondensatorze o „cienkich” elektrodach (obraz uproszczony)

elektrod, przylegające bezpośrednio do dielektryku, można w przybliżeniu przedstawić w postaci odcinka jednorodnego współosiowego toru elektrycznego.

Biorąc pod uwagę częstotliwości na tyle małe, że elektrody kondensatora są cienkie w porównaniu z głębokością wnika-
nia, i pomijając wpływ ewentualnego sprzężenia końcówek z tymi elektrodami, jak również



Rys. 5. Rozpływ prądów w kondensatorze o „grubych” elektrodach (obraz uproszczony)

abstrahując chwilowo od ich indukcyjności (wpływ tej indukcyjności można, jak wiadomo, uwzględnić oddzielnie), uzyskujemy obraz kondensatora przedstawiony schematycznie na rysunku 4. Na rysunku tym pokazano w przekroju, w sposób idealizowany, kondensator ceramiczny rurkowy o cienkich elektrodach.

Prąd I płynący przez kondensator, wpływając do elektrody zewnętrznej w punkcie 5, rozdziela się na dwie składowe, płynące w przeciwnych kierunkach. Dalej przepływa on przez rozłożone pojemności cząstkowe do elektrody wewnętrznej kondensatora, z której wypływa w punkcie 1. Zatem rozpływ prądów w zewnętrznej elektrodzie istotnie zależy od umiejscowienia punktu 5, w którym dołączona jest końcówka kondensatora.

Przy odpowiednio wielkiej częstotliwości elektrody kondensatora stają się grube. Przy pominięciu końcówek kondensatora i ich sprzężenia z elektrodami otrzymujemy uproszczony obraz, przedstawiony schematycznie na rysunku 5. Jak zaznaczono na tym rysunku, prąd I wpływa do zewnętrznej elektrody kondensatora w punkcie 5. Następnie rozdziela się na dwie składowe płynące po zewnętrznej powierzchni tej elektrody w kierunkach przeciwnych, dopływając do punktów skrajnych 2 i 4. W tych miejscach prądy powierzchniowe opuszczają zewnętrzną powierzchnię zewnętrznej elektrody kondensatora i przenoszą się na wewnętrzną powierzchnię tejże elektrody. Na powierzchni wewnętrznej prądy powierzchniowe wypływają ze skrajnych krawędzi elektrody zewnętrznej naprzeciw siebie. Dalej przez rozłożone pojemności cząstkowe kondensatora prądy te przenoszą się na zewnętrzną powierzchnię (tj. przylegającą do ceramiki) elektrody wewnętrznej i wypływają z niej w punkcie 1 w postaci wypadkowego sumarycznego prądu I .

Ponieważ z założenia pole elektromagnetyczne wewnątrz elektrod kondensatora nie istnieje, to przecinając je umyślnymi współosiowymi powierzchniami cylindrycznymi, których ślady na rysunku 5 oznaczono literami $A-A$ i $B-B$, można uzyskać niezależnie od siebie układy elektryczne, połączone ze sobą jedynie w punktach $2'2$, $4'4$, $1'1$ oraz $3'3$ (rys. 5 i rys. 8).

Jeżeli punkt 5 pokrywa się z punktem $4'$, to prąd wpływający do punktu 4 jest dużo większy od prądu wpływającego do punktu 2 i odwrotnie — jeżeli punkt 5 pokrywa się z punktem $2'$, to prąd przepływający przez punkt 2 jest dużo większy od prądu przepływającego przez punkt 4. Naturalnie powyższe wnioski są prawdziwe przy założeniu, że impedancja powierzchniowa zewnętrznej elektrody jest większa od zera. Gdyby była ona dokładnie równa zeru, tj. punkty 2 i 4 byłyby zwarte ze sobą, to oczywiście przesuwanie punktu 5 nie miałoby żadnego wpływu na rozpływ prądów w kondensatorze.

Z powyższych rozważań wynika, że odcinek toru elektrycznego, o którym wcześniej wspomniano, utworzony przez przylegające do ceramiki powierzchnie elektrod może być „zasilany” w sposób szczególny, nie brany zazwyczaj pod uwagę w „klasycznej” teorii torów elektrycznych. Przy tym sposób ten zależy od umiejscowienia punktów dołączenia końcówek do elektrod kondensatora. Sprawa ta jest przedmiotem bardziej szczegółowej analizy, podanej w rozdz. 2.3.

2.2.4. Propagacja fal wyższych rzędów

Oprócz fali podstawowej typu TEM, w układzie przewodników i dielektryków, jaki stanowi kondensator rurkowy, mogą w zasadzie istnieć pola elektromagnetyczne wyższych rzędów. Dla uproszczenia, przyjmie-

my tutaj, że elektrody kondensatora są grube. W tym przypadku możemy rozróżnić trzy odrębne układy:

— układ utworzony przez zewnętrzną powierzchnię zewnętrznej elektrody kondensatora, umieszczoną w powietrzu,

— układ utworzony przez powierzchnie elektrod przylegające do ceramiki i zawartą między nimi warstwę dielektryku o dużej przenikalności dielektrycznej,

— układ utworzony przez wewnętrzną powierzchnię wewnętrznej elektrody kondensatora.

Pierwszy z nich może w zasadzie prowadzić energię wielkiej częstotliwości w postaci fali powierzchniowej, tak jak np. falowód drutowy (Goubeau) (por. np. [23] str. 188 — 191), a ostatni może zachowywać się jak normalny falowód kołowy. Ponieważ jednak znajdują się one w powietrzu i — jak wykazano wcześniej — ich wymiary „elektryczne” są małe, przeto propagację falowodową można z nich praktycznie zaniedbać, przynajmniej w zakresie częstotliwości użytkowych kondensatora. Jeżeli nawet w pobliżu niejednorodności powstaną pola wyższych rzędów, to ulegną one tłumieniu. Jak wiadomo, tłumienie pola w falowodzie przy częstotliwościach mniejszych od krytycznej wynosi ([23] str. 166):

$$A = \frac{2\pi l}{\lambda_k} \sqrt{1 - \left(\frac{\lambda_k}{\lambda}\right)^2}, \quad (2)$$

gdzie:

λ_k — krytyczna długość fali w falowodzie,

λ — długość fali w próżni,

l — długość falowodu (= długość czynna kondensatora),

A — wypadkowe tłumienie falowodu w neperach.

Dla pola typu H_{11} w falowodzie kołowym krytyczna długość fali wynosi

$$\lambda_k \approx 1,71d_1, \quad (3)$$

gdzie d_1 oznacza wewnętrzną średnicę falowodu (elektrody wewnętrznej kondensatora).

Odpowiadająca tej długości fali częstotliwość w przypadku $d_1 = 2$ cm wynosi około 9 tysięcy megaherców.

Odmienne przedstawia się sprawa w układzie utworzonym przez powierzchnie elektrod przylegające do ceramiki. Na skutek obecności dielektryku o dość dużej na ogół przenikalności dielektrycznej i „zagęszczenia” pola elektromagnetycznego fala krytyczna może być znacznie dłuższa, niż wynika to z wymiarów geometrycznych. Jak wiadomo ([23] str. 170), rozkładem pola, który może powstać przy stosunkowo najniższych częstotliwościach, jest w tym układzie przewodników rozkład typu H_{11} , dla którego długość fali krytycznej λ_k wynosi:

$$\lambda_k \approx \frac{\pi}{2} \sqrt{\varepsilon(D+d)}, \quad (4)$$

gdzie ε oznacza względną przenikalność dielektryczną materiału ceramicznego, a pozostałe oznaczenia — jak na rysunku 2. Tej długości fali odpowiada częstotliwość krytyczna

$$f_k = \frac{1}{\pi} 6 \cdot 10^8 \frac{1}{\sqrt{\varepsilon(D+d)}}. \quad (5)$$

Jeżeli w powyższym wzorze wymiary średnic (D i d) zostaną podstawione w metrach, to częstotliwość krytyczna (f_k) będzie wyrażona w hercach.

Dla przykładu: jeżeli przenikalność dielektryczna ceramiki wynosi 2500, a suma średnic rurki jest równa 2 cm, to dla takiego kondensatora częstotliwość krytyczna jest około 191 MHz.

Przy częstotliwościach mniejszych od krytycznej energia podlega tłumieniu określonymu zależnościami (2) i (4). Przy częstotliwościach większych od krytycznej — w zasadzie — może odbywać się transmisja energii. Wykażemy jednak, że w praktyce zdarza się to rzadko.

W tym celu podzielimy najpierw stronami zależności (1a) i (5). Uzyskamy wówczas

$$\frac{f_k}{f_0} \approx 2,54 \frac{l}{D+d}. \quad (6)$$

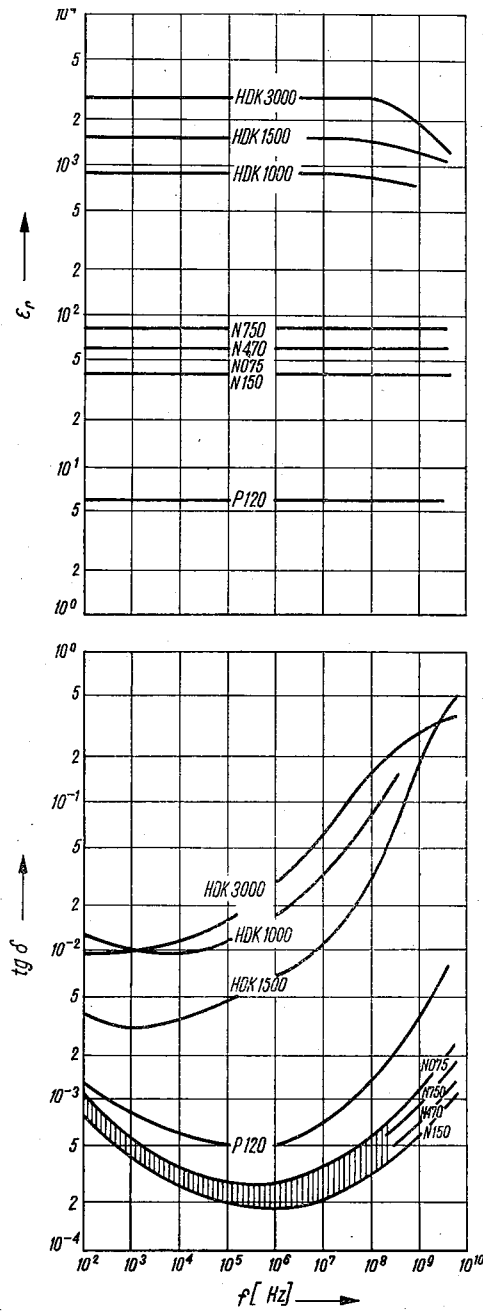
Wynika stąd, że w przypadkach, kiedy zachowany jest warunek

$$\frac{l}{D+d} > 0,394, \quad (7)$$

zawsze jest $f_k > f_0$. W omawianych kondensatorach warunek (7) z reguły jest spełniony. Jak to zostanie wykazane dalej, częstotliwość f_0 jest teoretyczną górną granicą częstotliwości rezonansowej kondensatora: granicą, do której można się w zasadzie dowolnie zbliżyć, lecz której w praktyce nie można przekroczyć. Ponieważ częstotliwość rezonansowa jest najwyższą częstotliwością użytkową kondensatora, to pominięcie w dalszych rozważaniach efektów propagacji falowodowej w dielektryku kondensatora wydaje się dostatecznie uzasadnione.

2.2.5. Rezonanse molekularne, piezoelektryczne i inne zjawiska

W zależności od składu chemicznego, stanu fizycznego, obróbki, temperatury i innych czynników, w materiale ceramicznym mogą w zasadzie zachodzić rozmaite zjawiska, jak na przykład rezonans molekularny [31], rezonans elektromechaniczny, spowodowany elektrostrykcją (efekt piezoelektryczny), i inne (por. np. [49]). Ponadto niektóre materiały ceramiczne mogą wykazywać znaczne własności nieliniowe i ich parametry charakterystyczne zależą wówczas od wartości natężenia pola elektrycznego, w którym są one umieszczone [49].



Rys. 6. Zależność przenikalności dielektrycznej (ϵ_r) i stratności ($\text{tg } \delta$) od częstotliwości dla typowych materiałów kondensatorowych (wg [4])

Jednakże w kondensatorach stosowanych do celów wspomnianych wcześniej, specjalnie dobierany jest taki materiał ceramiczny, dla którego efekty nieliniowe, podobnie jak i efekty piezoelektryczności w normalnych warunkach pracy są praktycznie pomijalne ([49], [55]).

Pozostałe zjawiska — mówiąc najogólniej — przejawiają się na zewnątrz w postaci mniej lub bardziej gwałtownych zmian przenikalności dielektrycznej i stratności materiału ceramicznego w funkcji częstotliwości. Na szczęście jednak w praktyce sprawa ta przedstawia się znacznie prościej. Jak wykazały eksperymentalne badania typowych materiałów ceramicznych, stosowanych w omawianych kondensatorach, ich przenikalność dielektryczna jest w przybliżeniu stała w bardzo szerokim zakresie częstotliwości [4], [16], [31], [52], [54]. Stratność ulega większym zmianom, ale i dla niej dla większości materiałów ceramicznych można w pierwszym przybliżeniu przyjąć, że zmiany te są niewielkie. Dla przykładu na rysunku 6 zamieszczono wyniki pomiarów Cirklera i Kuny'ego [4].

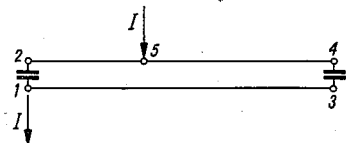
W związku z tym w dalszych rozdziałach pracy — podobnie jak i w pracach dotychczas opublikowanych — pominiemy efekty nieliniowości i przyjmiemy, że parametry charakterystyczne ceramiki nie zależą od częstotliwości.

2.3. Modele teoretyczne kondensatora rurkowego

2.3.1. Kondensator o cienkich elektrodach

W przypadku gdy grubość elektrod jest dużo mniejsza od głębokości wnikania pola elektromagnetycznego, można przyjąć, że prądy w elektrodach kondensatora płyną równomiernie w całym ich przekroju. Wówczas,

Rys. 7. Model teoretyczny kondensatora o „cienkich” elektrodach. Końcówki kondensatora pominięto (por. rys. 4)



przy założeniu symetrii osiowej, kondensator można przedstawić ogólnie jak na rysunku 4, lub równoważnym mu rysunku 7. Na rysunkach tych punkty 1, 2, 3 i 4 oznaczają krawędzie elektrod kondensatora. Punkty dołączenia końcówek oznaczono cyframi 1 i 5.

Jeżeli punkt 5 pokrywa się z punktem 2, to impedancja kondensatora jest w przybliżeniu równa impedancji wejściowej rozwartego odcinka toru 1-2-3-4. Jeżeli natomiast punkt 5 pokrywa się z punktem 4, to wówczas impedancja kondensatora jest zbliżona do impedancji skośnika [27], utworzonego z tego odcinka toru. Przypadek ogólny, kiedy punkt 5 znajduje się w dowolnym miejscu między punktami 2 i 4 wymaga specjalnej ana-

lizi. Ponieważ w torze przedstawionym na rys. 7 suma prądów płynących w obu jego przewodach w jednym przekroju poprzecznym może być różna od zera (równa dokładnie prądowi I), przeto taka analiza musiała zostać poprzedzona wyznaczeniem prądów i napięć w przewodach toru w ogólnym przypadku dowolnego sposobu zasilania toru i dowolnej jego symetrii. Sprawom tym poświęcone są prace [36] i [37].

W artykule [37] wykazano, że impedancja dwójnika z rys. 7 (między punktami 1 i 5) dla toru współosiowego ($\sigma = 1$) wynosi:

$$Z = Z_0[\text{cth}\Gamma + b\mathcal{K}\Gamma], \quad (8)$$

gdzie:

Z_0 — impedancja charakterystyczna toru współosiowego,

γ — tamowość charakterystyczna,

$\mathcal{K}$ — wsp. sprzężenia przewodów toru współosiowego,

l — czynna długość kondensatora,

l_1 — odległość między punktami dołączenia końcówek

oraz

$$\Gamma = \gamma l; \quad b = \frac{l_1}{l}. \quad (9)$$

Jeżeli straty kondensatora są pomijalne, to tamowność jest liczbą urojoną

$$\Gamma = jB$$

oraz

$$\frac{Z}{+jZ_0} = -\text{ctg}B + b\mathcal{K}B. \quad (10a)$$

Wprowadzimy dla wygody częstotliwość znormalizowaną

$$x = \frac{f}{f_0}, \quad (11)$$

gdzie częstotliwość odniesienia f_0 jest określona zależnościami (1a) — (1c).

Korzystając z wzorów na impedancję charakterystyczną toru współosiowego i pojemność kondensatora, łatwo jest wykazać, że zachodzi równość

$$B = \frac{\omega}{v}l = \omega Z_0 C, \quad (12a)$$

gdzie C oznacza pojemność kondensatora, v — prędkość fali elektromagnetycznej w dielektryku, $\omega = 2\pi f$ (wszystkie wymiary w jednostkach podstawowych). Uwzględniając zależności (1b) i (11) otrzymujemy z (12a) po prostych przekształceniach

$$B = x \frac{\pi}{2}. \quad (12b)$$

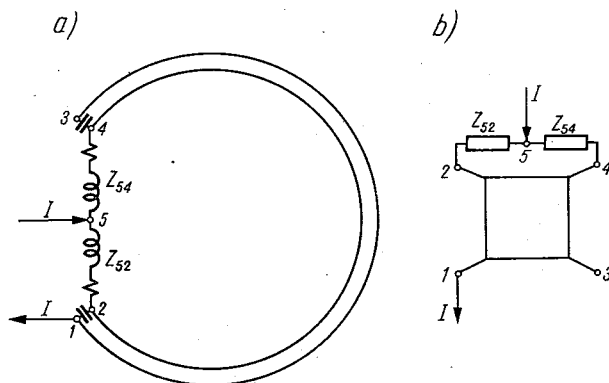
Zatem zależność (10) można teraz przedstawić w poniższej, równoważnej postaci

$$\frac{Z}{jZ_0} = -\operatorname{ctg}\left(x\frac{\pi}{2}\right) + b\frac{\pi}{2}x\chi. \quad (10b)$$

Z powyższych związków wynika, że impedancja kondensatora jest równa impedancji wejściowej odcinka toru rozwartego, połączonego w szereg z pewną impedancją (indukcyjnością), jak przedstawiono na rysunku 13. Należy podkreślić, że ta indukcyjność nie ma nic wspólnego z indukcyjnością końcówek. Wynika ona wyłącznie z rozptywu prądów w zewnętrznej elektrodzie kondensatora i istnieje nawet wówczas, gdy długość końcówek jest równa zeru.

2.3.2. Kondensator o grubych elektrodach

Jeżeli grubość elektrod kondensatora jest znacznie większa od głębokości wnikania pola elektromagnetycznego, to można przyjąć, że prądy płynące na ich powierzchniach zewnętrznych i wewnętrznych nie oddziałują na siebie. W tym przypadku — jak na to zwrócił uwagę Ryżko



Rys. 8. Model teoretyczny kondensatora o „grubych” elektrodach (a) i jego układ równoważny (b). Końcówki kondensatora pominięto (por. rys. 5)

[30] — obraz kondensatora przedstawiony powyżej nie jest prawdziwy. Dla tego zakresu częstotliwości Ryżko [30] proponuje układ przedstawiony schematycznie na rysunku 8. Rysunek ten odpowiada rysunkowi 5. Podobnie jak na poprzednich rysunkach tak i tutaj krawędzie elektrod kondensatora oznaczono cyframi 1, 2, 3, 4, a punkty dołączenia końcówek — 1 i 5. Przyjmujemy przy tym, że punkty 1 i 1', 2 i 2', 3 i 3', oraz 4 i 4' pokrywają się. Wtedy zgodnie z wcześniejszymi rozważaniami zewnętrzną powierzchnię elektrody zewnętrznej można przedstawić w postaci dwóch impedancji Z_{52} i Z_{54} załączonych odpowiednio między punktami 2 i 5 oraz 5 i 4. Wewnętrzne, przylegające do dielektryku powierzchnie elektrod można przedstawić w postaci odcinka toru 1-2-4-3. Na rysunku 9 ten tor przedstawiony jest tak, jakby został on wygięty łukiem, w tym

celu, aby uwypuklić, że jest on „elektrycznie” długi w porównaniu z pozostałymi elementami układu.

Jak już wspomniano wcześniej, prąd I rozplywa się w punkcie 5 na dwie składowe, wpływające do toru 1-2-3-4 w punktach 2 i 4. Następnie przez pojemności cząstkowe przepływa do drugiego przewodu toru i wypływa zeń w punkcie 1. Jak widać, odcinek toru 1-2-3-4 jest zasilany w sposób bardzo specyficzny, nie brany pod uwagę w normalnych kursach teorii torów elektrycznych (por. np. [20], [28]).

Sprawa jest o tyle skomplikowana, że prądy wpływające do toru w punktach 2 i 4 zależą między innymi od napięcia między nimi, a to z kolei — od wszystkich prądów wpływających i wypływających z rozważanego odcinka toru. Jak się wydaje, najprostszym sposobem określenia właściwości takiego układu jest zbudowanie czwórnika równoważnego odcinkowi toru 1-2-3-4 i wykorzystanie metod znanych z teorii czwórników. Łatwo się jednak przekonać, że taki czwórnik równoważny nie może być „zwykłym” czwórnikiem o macierzy zawierającej trzy niezależne parametry. Jak wykazał m.in. Sigorski [56], musi to być czwórnik „ogólny” z trzema niezależnymi bramami ([56] str. 17). Taki czwórnik równoważny układowi o parametrach rozłożonych, który można by zastosować w omawianym przypadku — o ile autorowi wiadomo — nie był dotychczas w literaturze analizowany. Wyjątek stanowią pewne przesłanki podane przez Woronowa [47] przy omawianiu energetycznych napowietrznych linii długich, oraz przez Happa i Riddle’a [12] — przy omawianiu układów oporowo-pojemnościowych (RC) o stałych rozłożonych. Jednak i te wiadomości nie nadają się do bezpośredniego wykorzystania w rozpatrywanym przypadku. Powstała zatem konieczność przeprowadzenia takiej analizy. Sprawie tej jest poświęcona praca [38].

Jeżeli wziąć pod uwagę, że średnica zewnętrzna i wewnętrzna elektrody kondensatora są praktycznie równe sobie ($D_1 \approx D$), to dochodzimy do wniosku, że impedancje powierzchniowe są w przybliżeniu takie same dla obu powierzchni elektrody. Zakładamy przy tym, że wszystkie pozostałe przewodniki są oddalone i przenikalność magnetyczna ośrodka jest taka sama z obu stron elektrody.

Przy tych założeniach można wykazać, że

$$Z_5 = Z_{54} + Z_{52} = Z_0 \Gamma \mathcal{K} \quad (13)$$

W tych warunkach — jak wykazano w [39] — impedancja między punktami 1 i 5 wynosi dla toru współosiowego ($\sigma = 1$)

$$Z = Z_0 \cdot \left[\operatorname{cth} \Gamma + b \left(1 - \frac{b}{2} \right) \mathcal{K} \Gamma \right], \quad (14)$$

gdzie, podobnie jak we wzorze (9),

$$\Gamma = \gamma l; \quad b = \frac{Z_{52}}{Z_5} = \frac{l_1}{l},$$

a pozostałe oznaczenia — jak we wzorze (8).

(Zależność $b = \frac{l_1}{l}$ wynika bezpośrednio z faktu, że impedancja prostego odosobnionego przewodu jest w przybliżeniu wprost proporcjonalna do jego długości).

Jeżeli stratność kondensatora jest pomijalna, to po wprowadzeniu częstotliwości znormalizowanej (11) można otrzymać z (14)

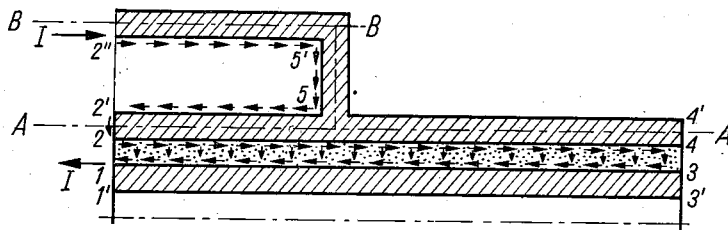
$$\frac{Z}{jZ_0} = -\operatorname{ctg}\left(x\frac{\pi}{2}\right) + b\left(1 - \frac{b}{2}\right)\frac{\pi}{2}x\kappa. \quad (15)$$

Powyższa zależność wskazuje, że impedancja kondensatora również i w tym przypadku jest sumą impedancji wejściowej rozwartego odcinka toru 1-2-3-4 i pewnej indukcyjności wynikającej z rozptywu prądu w zewnętrznej elektrodzie kondensatora. Odpowiedni układ równoważny, wynikający z zależności (15), przedstawiono na rysunku 13. Należy podkreślić, że przedstawiona na tym rysunku indukcyjność nie ma nic wspólnego z indukcyjnością samych końcówek, której tam nie uwidoczniiono.

2.3.3. Kondensator o grubych elektrodach z końcówkami o prostej konfiguracji geometrycznej

Weźmy teraz pod uwagę kondensator, w którym końcówki tworzą układ dwóch cylindrycznych przewodników, współosiowych z rurką kondensatora.

Odrzucając tę część tych przewodników, która znajduje się na lewo od przekroju 1-2 kondensatora i jest słabo sprzężona z elektrodą kondensa-



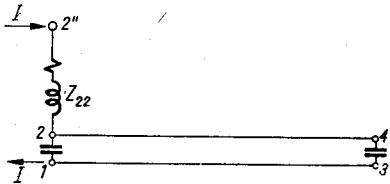
Rys. 9. Rozptyw prądów w kondensatorze o grubych elektrodach, z końcówką w kształcie cylindra (obraz uproszczony)

tora (i której wpływ został omówiony w rozdziale 2.2), otrzymamy układ przedstawiony schematycznie na rysunku 9. Przy odpowiednio wielkich częstotliwościach można uważać, że wszystkie uwidocznione tam przewodniki są „grube” i wewnątrz nich pole elektromagnetyczne nie istnieje. Wówczas, przecinając je powierzchniami, których ślady oznaczono na rys. 11 literami A-A i B-B, łatwo otrzymać układ równoważny, przedstawiony na rysunku 10. Na tym rysunku przez Z_{22} oznaczono impedancję powierzchni 2''-5'-5'-2'-2. Gdyby odległość między punktami 2'' i 2' była du-

zo mniejsza od odległości między punktami 2' i 5, to powierzchnia ta przedstawiałaby odcinek jednorodnego zwartego toru współosiowego. Ściślej rzecz biorąc założenia te nie są w praktyce spełnione. Mimo to w pierwszym przybliżeniu przyjmujemy dla prostoty, że impedancja Z_{22} jest wyrażona wzorem na impedancję wejściową zwartego odcinka toru elektrycznego

$$Z_{22} \approx Z_0 \tanh \Gamma' \approx Z_0 \Gamma', \quad (16)$$

gdzie Z_0 i Γ' oznaczają impedancję charakterystyczną i tamowność toru współosiowego, odpowiadającego powierzchniom 2''-5' i 2'-5. (Przybliżona



Rys. 10. Układ równoważny kondensatorowi o „grubych” elektrodach, z końcówką w kształcie cylindra (por. rys. 9)

równość wynika z faktu, że długość elektryczna” rozpatrywanego toru jest bardzo mała w całym rozważanym zakresie częstotliwości).

Wprowadźmy teraz parametry pomocnicze, określone jako stosunki

$$g = \frac{Z_0'}{Z_0} = \sqrt{\epsilon} \frac{\ln \frac{D'}{D_1}}{\ln \frac{D}{d}} \quad (17)$$

$$h = \frac{\Gamma'}{\Gamma} = \frac{l_1}{l} \frac{1}{\sqrt{\epsilon}} = b \cdot \frac{1}{\sqrt{\epsilon}}, \quad (18)$$

gdzie D' oznacza wewnętrzną średnicę zewnętrznego przewodnika cylindrycznego, tworzącego końcówkę kondensatora, a pozostałe oznaczenia — jak na rysunku 2 i we wzorach (1), (8), (9), (16).

Korzystając z wprowadzonych parametrów pomocniczych i z rysunku 10, możemy napisać wyrażenie na impedancję kondensatora

$$Z \approx Z_0 \cosh \Gamma + gh Z_0 \Gamma \quad (19)$$

(oznaczenia jak poprzednio).

W przypadku pomijalnych strat jest

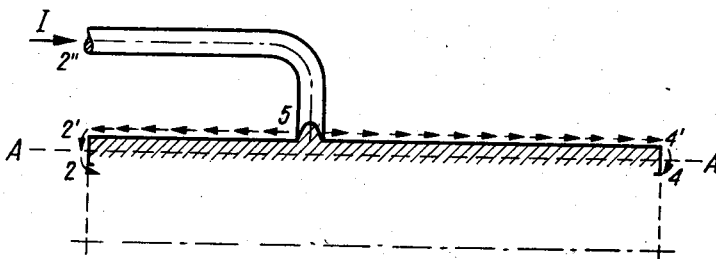
$$\frac{Z}{jZ_0} \approx -\operatorname{ctg} B + ghB \quad (20a)$$

albo, na podstawie związków (11), (12b),

$$\frac{Z}{jZ_0} \approx -\operatorname{ctg} \left(x \frac{\pi}{2} \right) + ghx \frac{\pi}{2}. \quad (20b)$$

Należy zauważyć, że w powyższej zależności jak również na rysunku 10, jest uwzględniony wpływ jedynie tej części końcówek, która jest bezpośrednio sprzężona z zewnętrzną elektrodą kondensatora.

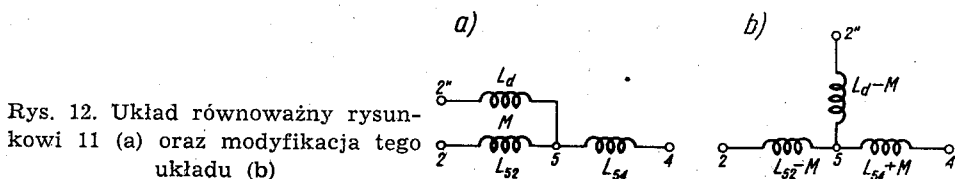
W zasadzie istnieje możliwość podobnej analizy wpływu sprzężenia końcówek z elektrodami kondensatora bez konieczności zakładania „grubych” elektrod. Można w tym przypadku wykorzystać np. teorię wieloprzewodowych torów elektrycznych [2], [26], [50], [51]. Pomijając już inne



Rys. 11. Rozpływ prądów w zewnętrznej elektrodzie kondensatora o „grubych” elektrodach, z końcówką drutową (por. rys. 9)

trudności rachunkowe (np. operowanie szesnastoelementowymi w naszym przypadku macierzami), praktyczne wykorzystanie tej teorii utrudnia fakt, że parametry charakterystyczne takiego układu zależą od częstotliwości i to w sposób raczej złożony. Przy takim postępowaniu uzyskane powyżej zależności byłyby wyrażeniami asymptotycznymi. Praktyczne zastosowanie teorii torów wieloprzewodowych (bez założeń uproszczających dotyczących grubości elektrod) do układu trzech przewodników współosiowych jest podane przez Raisbecka i innych [26].

Dotychczas rozpatrywaliśmy idealizowany przypadek graniczny, kiedy sprzężenie elektrody kondensatora z jego końcówką jest bardzo silne, oraz



Rys. 12. Układ równoważny rysunkowi 11 (a) oraz modyfikacja tego układu (b)

kiedy źródło prądu znajduje się wewnątrz szczelnego ekranu elektromagnetycznego, utworzonego przez jedną z końcówek kondensatora. Jeżeli idzie o inne kształty końcówek, to ograniczymy się jedynie do stwierdzenia, że przypadek spotykany najczęściej, przedstawiony schematycznie na rysunku 11 — jak się wydaje — można z pewnym przybliżeniem sprowadzić do omówionego wcześniej przypadku, z rysunków 5 i 9. W tym celu wystarcza rysunek 11 zastąpić równoważnym mu rysunkiem 12a. Układ z rys. 12a daje się łatwo sprowadzić do postaci przedstawionej na rysunku 12b. Ten ostatni w porównaniu z rys. 8 wskazuje, że wpływ

sprężenia końcówki z elektrodą kondensatora jest w przybliżeniu równoważny odpowiedniemu przesunięciu punktu 5, w którym końcówka łączy się z elektrodą.

2.4. Układ zastępczy kondensatora

2.4.1. Ogólny układ zastępczy

Z poprzedniego rozdziału wynika, że impedancja kondensatora — przynajmniej w rozpatrzonych przypadkach — może być wyrażona ogólną zależnością

$$Z \approx Z_0 \operatorname{cth} \Gamma + Z_0 \Gamma \xi, \quad (21)$$

gdzie:

Z_0 — impedancja charakterystyczna,

Γ — tamowność,

ξ — pewna rzeczywista funkcja nieujemna, zależna od umiejscowienia końcówek i kształtu kondensatora.

Istotne jest przy tym, że ξ jest niemalejącą funkcją parametru $b = \frac{l_1}{l}$,

$$\frac{\partial}{\partial b}[\xi] \geq 0.$$

W przypadku kondensatora o cienkich elektrodach jest

$$\xi = b \mathcal{K}. \quad (22a)$$

Dla kondensatora o grubych elektrodach mamy

$$\xi = b \left(1 - \frac{1}{2} b \right) \mathcal{K}, \quad (22b)$$

a dla tegoż kondensatora z końcówkami współosiowymi

$$\xi = gh, \quad g = \sqrt{\varepsilon} \frac{\ln \frac{D'}{D_1}}{\ln \frac{D}{d}}, \quad h = b \frac{1}{\sqrt{\varepsilon}}. \quad (22c)$$

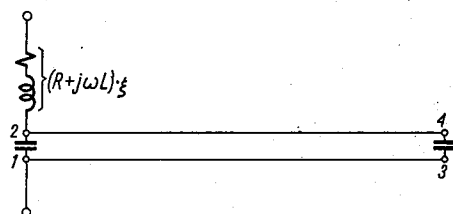
Iloczyn impedancji charakterystycznej i wypadkowej tamowności toru jest — jak wiadomo — równy wypadkowej impedancji wzdłużnej toru

$$Z_0 \Gamma = R + j\omega L.$$

Dlatego na rysunku 13, na którym przedstawiono ogólny układ zastępczy kondensatora (bez uwzględnienia końcówek), kondensator jest przedstawiony w postaci odcinka toru 1-2-3-4 połączonego w szereg z indukcyjnością i opornością, zależną między innymi od umiejscowienia końcówek.

W praktyce zarówno ta oporność, jak i pojemności rozproszone załączone na krańcach toru są pomijalne. Należy mieć na uwadze, że w rzeczywistości do uwidocznionej na rys. 13 indukcyjności dodają się jeszcze indukcyjności końcówek, jak wspomniano w rozdziale II — 2.

Jak widać, jeżeli tylko straty są małe, w układzie mogą zaistnieć rezonanse. W celu określenia częstotliwości rezonansowej założymy, że straty



Rys. 13. Ogólny układ zastępczy kondensatora. Końcówki pominięto

są pomijalne. Wówczas, po uwzględnieniu wprowadzonych wcześniej oznaczeń i po wprowadzeniu częstotliwości znormalizowanej z zależności (21) otrzymamy

$$\frac{Z}{jZ_0} = -\operatorname{ctg}\left(x\frac{\pi}{2}\right) + x\frac{\pi}{2}\xi. \quad (23)$$

Wynika stąd, że

$$\left|\frac{Z}{jZ_0}\right| \rightarrow \infty, \quad \text{gdy} \quad x \rightarrow x_{\infty}^{(k)} = 2k \quad (24)$$

($k = 0, 1, 2, \dots$)

$$\left|\frac{Z}{jZ_0}\right| \rightarrow x\frac{\pi}{2}\xi, \quad \text{gdy} \quad x \rightarrow x_{\xi}^{(k)} = (2k+1) \quad (25)$$

($k = 0, 1, 2, 3, \dots$)

$$\left|\frac{Z}{jZ_0}\right| \rightarrow 0, \quad \text{gdy zachodzi równość}$$

$$\frac{\pi}{2}\xi x_0^{(k)} = \operatorname{ctg}\left[\frac{\pi}{2}x_0^{(k)}\right]. \quad (26)$$

Można wykazać, że pierwiastki równania (26) spełniają nierówności

$$2k \leq x_0^{(k)} \leq (2k+1), \quad (27a)$$

przy tym, jeżeli $\xi x \rightarrow 0$,

to

$$x_0^{(k)} \rightarrow x_{\xi}^{(k)} = (2k+1), \quad (27b)$$

a jeżeli

$$\xi x \rightarrow \infty,$$

to

$$x_0^{(k)} \rightarrow x_{\infty}^{(k)} = 2k. \quad (27c)$$

Ze związku (26) wynika, że kondensator może wykazywać cały ciąg rezonansów i antyrezonansów. Dla uniknięcia wieloznaczności wyjaśnimy tutaj, że pod nazwą „częstotliwość rezonansowa” rozumiemy w niniejszym artykule częstotliwość pierwszego rezonansu (najniższą częstotliwość, przy której impedancja kondensatora osiąga minimum). Ta częstotliwość odpowiada najmniejszemu pierwiastkowi równania (26). Ponieważ równanie to jest przestępne i nie można określić jego rozwiązania dokładnego, pierwiastek ten wyznaczmy najpierw w sposób przybliżony dla dwóch szczególnych przypadków. Przy tym, dla uproszczenia zapisu, opuścimy dalej wskaźnik „1”, który — zgodnie z dotychczasowym sposobem oznaczenia — powinien figurować przy x .

a) Przypadek $\xi \ll 1$

W tym przypadku, jak łatwo się przekonać, $x_0 \approx 1$, i na mocy przybliżonej zależności

$$\operatorname{ctg}\left(x \frac{\pi}{2}\right) = \operatorname{tg}\left[(1-x) \frac{\pi}{2}\right] \approx (1-x) \frac{\pi}{2}$$

otrzymujemy ze wzoru (26)

$$x_0 \approx \frac{1}{1+\xi} \approx 1-\xi. \quad (27d)$$

b) Przypadek $\xi \gg 1$

Obecnie, jak można wykazać, $x_0 \ll 1$ i można wykorzystać inne przybliżenie

$$\operatorname{ctg}\left(x \frac{\pi}{2}\right) \approx \frac{2}{\pi x} - \frac{1}{3} \cdot \frac{\pi x}{2},$$

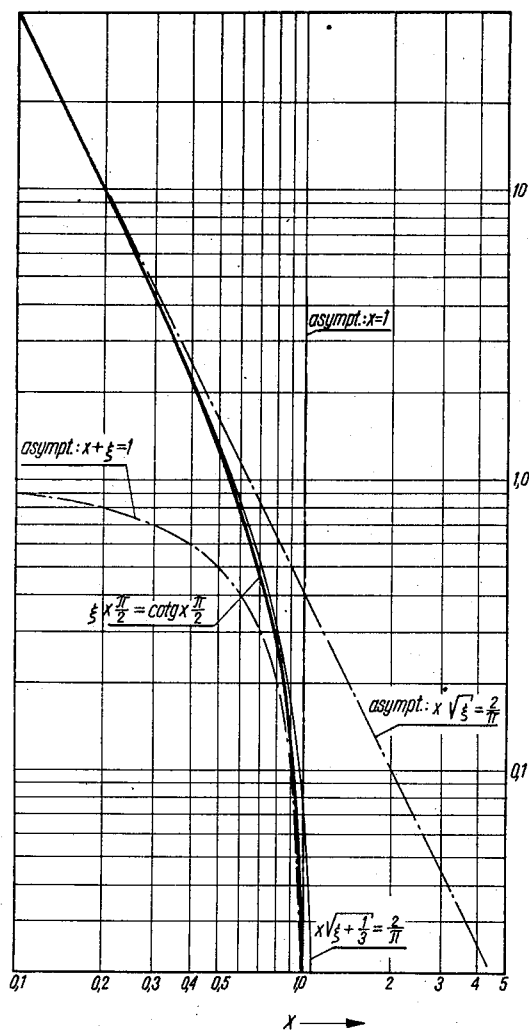
co, w połączeniu z zależnością (26), daje w wyniku

$$x_0 \approx \frac{2}{\pi} \frac{1}{\sqrt{\xi + \frac{1}{3}}}. \quad (27e)$$

Zależności (27d), (27e) przedstawiono wykreślenie na rysunku 14. Jednocześnie na tymże rysunku uwidoczniono „dokładną” wartość pierwszego pierwiastka równania przestępnego (26), wyznaczoną przy użyciu tablic wartości funkcji trygonometrycznych (np. [43]), ze związku

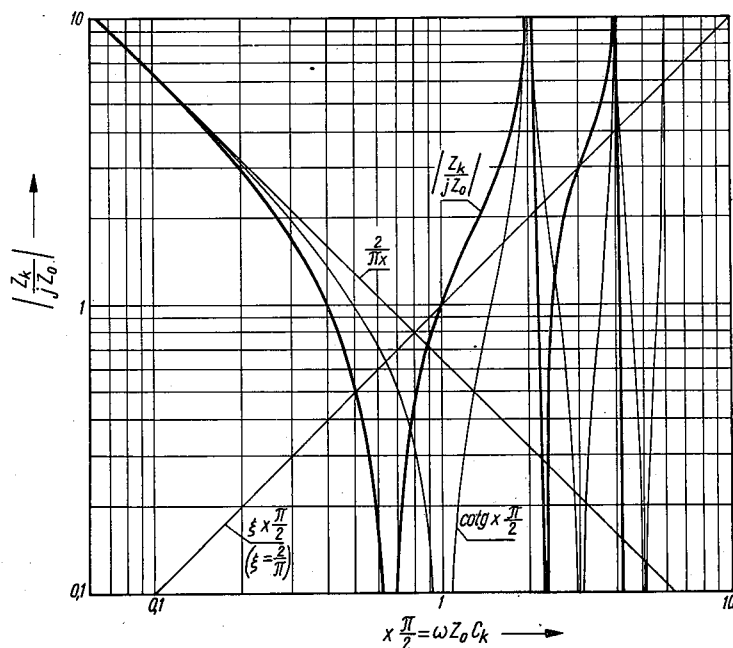
$$\xi = \frac{\operatorname{ctg}\left(x_0 \frac{\pi}{2}\right)}{x_0 \frac{\pi}{2}}. \quad (28)$$

Na zakończenie niniejszego punktu, dla bardziej poglądowego przedstawienia sprawy, na rysunku 15 podano wykres impedancji kondensa-



Rys. 14. Zależność częstotliwości (x) rezonansu wewnętrznego (pierwszego) od parametru ξ

tora bez strat w funkcji znormalizowanej częstotliwości (x), dla ustalonej wartości parametru $\xi \left(\xi = \frac{2}{\pi} \right)$ i bez uwzględnienia wpływu końcówek. Na rysunku tym zaznaczono również sposób konstrukcji geometrycznej



Rys. 15. Zależność impedancji kondensatora bez strat od częstotliwości (x). Przykład

$$\text{dla } \xi = \frac{2}{\pi}$$

wykresu. Mianowicie krzywa wykreślona linią cienką stanowi wykres funkcji $\text{ctg}\left(x \frac{\pi}{2}\right)$, jedna z prostych jest asymptotą tejże funkcji, a druga — odpowiada drugiemu członowi prawej strony zależności (23). Wykres modułu impedancji powstaje przez sumowanie (lub odejmowanie, zależnie od znaku funkcji ctg) rzędnych ostatniej ze wspomnianych prostych i krzywej $\text{ctg}\left(x \frac{\pi}{2}\right)$.

2.4.2. Przybliżony układ zastępczy dla małych częstotliwości

Jeżeli rozpatrywane są wyłącznie częstotliwości tak małe, że zachodzi nierówność

$$|\Gamma| \ll 1, \quad (29)$$

to wykorzystując przybliżoną zależność

$$\operatorname{ctg} \Gamma \approx \frac{1}{\Gamma} + \frac{1}{3} \Gamma$$

można wzór (21) sprowadzić do uproszczonej, przybliżonej postaci

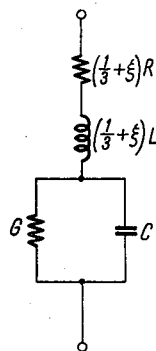
$$Z \approx \frac{Z_0}{\Gamma} + Z_0 \Gamma \left(\frac{1}{3} + \xi \right) \quad (30a)$$

albo

$$Z \approx \frac{1}{G + j\omega C} + (R + j\omega L) \left(\frac{1}{3} + \xi \right), \quad (30b)$$

gdzie Z_0 i Γ oznaczają — jak poprzednio — impedancję charakterystyczną i wypadkową tamowności toru, utworzonego przez powierzchnie elektrod, przylegające do ceramiki, a R , L , G i C oznaczają odpowiednio wypadkowe: oporność wzdłużną i indukcyjność oraz przewodność poprzeczną i pojemność tego toru.

Opierając się na zależności (30b) można narysować przybliżony układ zastępczy kondensatora, jak na rysunku 16. Należy podkreślić, że na tym



Rys. 16. Przybliżony układ zastępczy kondensatora dla małych częstotliwości. Końcówki pominięto

rysunku — podobnie jak i na poprzednich — nie uwidoczniło indukcyjności końcówek. W tym wszystkim istotny jest fakt, że nawet dla przybliżonego układu zastępczego indukcyjność własna kondensatora (a częściowo i jego stratność) zależy od umiejscowienia jego końcówek.

Jeżeli stratność kondensatora pominąć, to wykaże on — według zależności (30b) — rezonans przy częstotliwości f_r

$$f_r \approx \frac{1}{2\pi\sqrt{LC} \sqrt{\frac{1}{3} + \xi}}. \quad (31a)$$

Ponieważ

$$\sqrt{LC} = \sqrt{\frac{L}{C}} C = Z_0 C,$$

to po prostych przekształceniach i uwzględnieniu oznaczenia (1b) otrzymujemy

$$f_r \approx \frac{1,102}{\sqrt{1+3\xi}} f_0. \quad (31b)$$

Stąd możemy napisać

$$x_0 \approx \frac{1,102}{\sqrt{1+3\xi}}. \quad (32)$$

Łatwo sprawdzić, że zależnościami (32) i (27e) są w istocie identyczne. Już wcześniej wykazano, że częstotliwość rezonansowa nie przekracza nigdy częstotliwości f_0 . Wzór przybliżony (31b), (32) daje tymczasem w granicznym przypadku częstotliwość o około 10% większą. Jak już wcześniej wspomniano, w takich granicznych przypadkach, kiedy ξ jest małe, lepsze przybliżenie można uzyskać, posługując się zależnością (27d) lub korzystając ze wzoru dokładnego (26).

2.5. Indukcyjność i częstotliwość rezonansowa kondensatora

2.5.1. Optymalne umiejscowienie końcówek

W poprzednich rozdziałach wykazano, że umiejscowienie końcówek ma istotny wpływ na rozpyły prądów w elektrodach kondensatora, a tym samym i na jego częstotliwość rezonansową. W związku z tym nasuwa się natychmiast pytanie, jaki jest optymalny sposób umiejscowienia końcówek, zapewniający najszerszy, możliwy do uzyskania zakres częstotliwości użytkowych kondensatora?

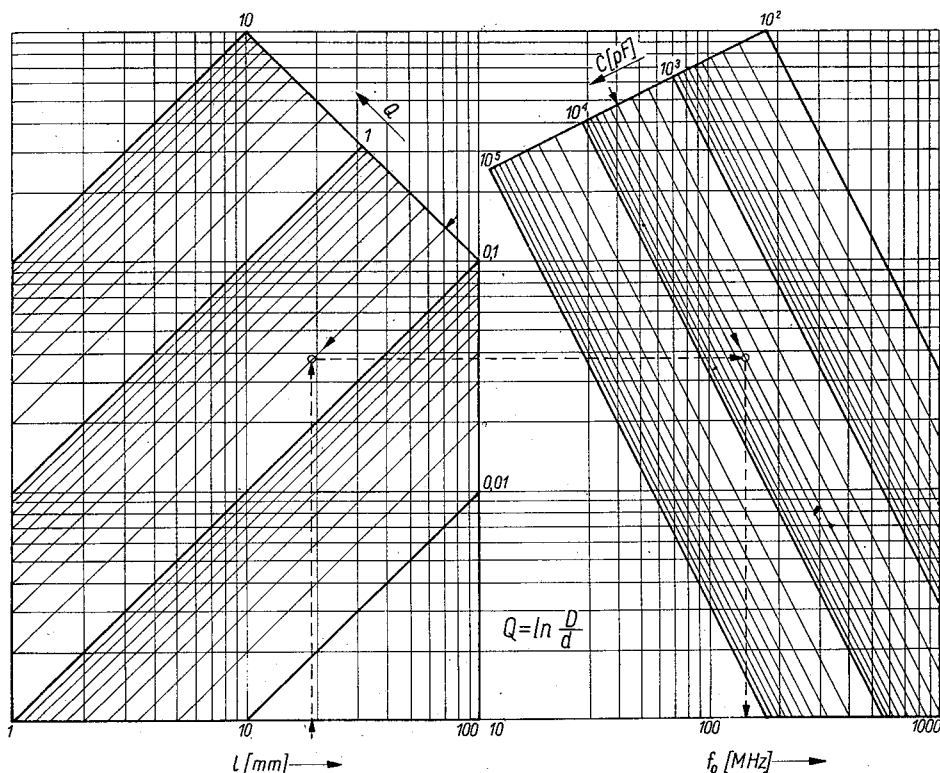
W oparciu o przeprowadzone wcześniej rozważania odpowiedzieć na to pytanie jest teraz sprawą zupełnie łatwą. Ponieważ iloczyn parametrów ξx jest nieujemny, przeto z zależności (23) wynika natychmiast, że maksymalna częstotliwość rezonansowa (minimalna indukcyjność własna) kondensatora odpowiada najmniejszej wartości funkcji ξ . W rozpatrywanych przypadkach, zgodnie ze wzorami (22a) — (22c) minimum ξ , równe zero zachodzi dla $b = 0$. Oznacza to, że częstotliwość rezonansowa kondensatora jest maksymalna, kiedy jego końcówki są dołączone do punktów „optymalnych”, leżących w jednej płaszczyźnie, prostopadłej do osi rurki ceramicznej. W tym i tylko w tym przypadku częstotliwość rezonansowa jest dokładnie równa częstotliwości f_0 , określonej wzorami (1a), (1c).

Przy optymalnym umiejscowieniu końcówek kierunki prądów płynących w elektrodach kondensatora są na całej jego długości przeciwne. Dzięki temu zewnętrzne pole elektromagnetyczne, skojarzone z nimi, a zatem i indukcyjność wewnętrzna kondensatora, osiąga najmniejszą możliwą wartość.

Z zależności (30b) wynika, że ta minimalna wartość indukcyjności, odpowiadająca „optymalnemu” umiejscowieniu końcówek, wynosi:

$$L_{min} = \frac{1}{3}L = \frac{1}{3} \frac{L}{C} C = \frac{1}{3} Z_0^2 C, \quad (33)$$

przy czym L_{min} oznacza minimalną wartość indukcyjności własnej kondensatora (bez uwzględnienia indukcyjności końcówek), L — wypadkową indukcyjność wzdłużną toru utworzonego przez powierzchnie elektrod przylegające do ceramiki (tj. iloczyn indukcyjności jednostkowej toru i je-



Rys. 17. Nomogram do wyznaczania maksymalnej częstotliwości rezonansowej kondensatorów rurkowych

$$f_0[\text{Hz}] = \frac{560}{\sqrt{Q l_{[m]} C_{[F]}}}; \quad Q = \ln \frac{D_2}{d_1}$$

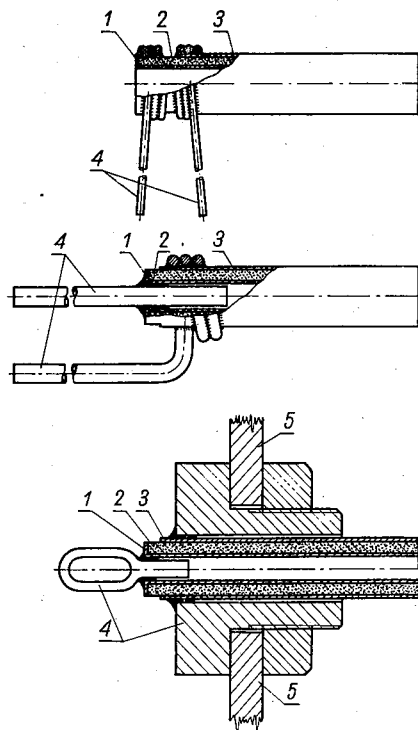
go długości czynnej), Z_0 — impedancję charakterystyczną tego toru, a C — pojemność wypadkową kondensatora.

Ze wzoru (33) wynika, że minimalna indukcyjność kondensatora przy ustalonej jego pojemności jest tym mniejsza, im mniejsza jest impedancja charakterystyczna toru Z_0 , tj. im większa jest przenikalność dielektryczna materiału ceramicznego i im mniejszy jest stosunek średnic rurki ceramicznej. Oprócz tego zależność (33) wskazuje, że wszystkie kondensatory o optymalnym umiejscowieniu końcówek, dla których iloczyn pojemności

i impedancji charakterystycznej w potęgze drugiej jest taki sam, mają jednakową indukcyjność własną.

Dla przykładu można podać, zgodnie z zależnością (33), że kondensator o impedancji charakterystycznej równej 1 om i pojemności 600 pF ma indukcyjność własną — przy optymalnym umiejscowieniu końcówek — równą ok. $2 \cdot 10^{-10}$ henra.

Jak już wspomniano, częstotliwość rezonansowa takiego kondensatora określona jest między innymi zależnością (1c). W celu ułatwienia posługiwania się tym wzorem przedstawiono go w postaci nomogramu (rys. 17).



Rys. 18. Przykłady konstrukcji kondensatorów o zmniejszonej indukcyjności wewnętrznej [34]

1 — elektroda wewnętrzna, 2 — dielektryk, 3 — elektroda zewnętrzna, 4 — końcówki dołączone do elektrod w punktach maksymalnie zbliżonych, 5 — chassis

Przykłady rozwiązań konstrukcyjnych kondensatora o optymalnym umiejscowieniu końcówek, zapewniającym minimalną indukcyjność własną, uwidoczniono na rysunku 18 ([34], [35]). W praktyce wymaganie dołączenia końcówek do punktów leżących w jednej płaszczyźnie, prostopadłej do osi rurki ceramicznej, najczęściej nie może być ściśle spełnione. Należy wówczas wziąć pod uwagę to, że maksymalne zbliżenie tych punktów jest zawsze korzystne. Znaczne zmniejszenie indukcyjności własnej kondensatora obserwuje się już dla wartości parametru b rzędu

$\frac{1}{3}$ [34]. W niektórych jednak przypadkach wymagane jest dołączenie końcówek do punktów umieszczonych jeszcze bliżej. Sprawa ta omawiana jest poniżej.

2.5.2. Dowlolne umiejscowienie końcówek

Kiedy końcówki kondensatora są dołączone do dowolnych punktów, to jego częstotliwość rezonansowa, na podstawie zależności (11), wynosi

$$f_r = x_0 f_0, \quad (34)$$

gdzie x_0 dane jest wzorami (26), (27), (32) (porównaj także rys. 14), a f_0 — określone wzorami (1) (porównaj rys. 17).

W tym przypadku indukcyjność kondensatora jest większa od minimalnej, określonej wzorem (33). Oznaczmy ją przez L_k i wyznaczmy stosunek

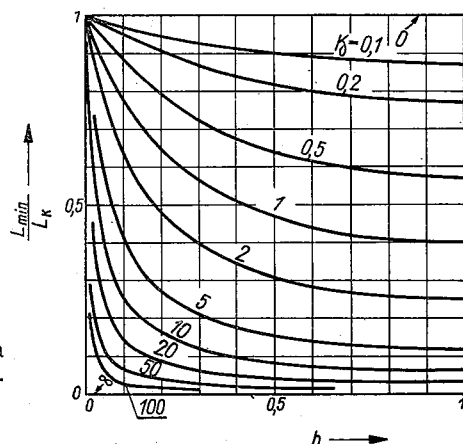
$$\frac{L_k}{L_{min}} = K. \quad (35)$$

Na mocy związku (30b) i (33) możemy napisać

$$K = 1 + 3\xi. \quad (36)$$

Zależność powyższa przedstawiona jest graficznie na rysunkach 19 i 20.

Na rysunku 19 w przypadku zastosowania go do modelu kondensatora

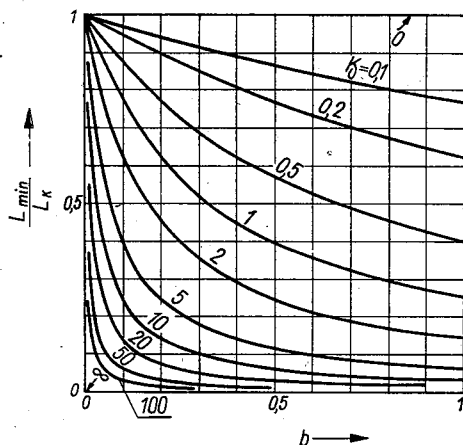


Rys. 19. Zależność indukcyjności wewnętrznej kondensatora od umiejscowienia końcówek. Model kondensatora o „cienkich” elektrodach

o grubych elektrodach z końcówkami w postaci cylindrycznych współosiowych przewodników, z których jeden otacza częściowo kondensator, należy podstawić — zgodnie z zależnością (22c) — zamiast $\mathcal{K}$ — stosunek

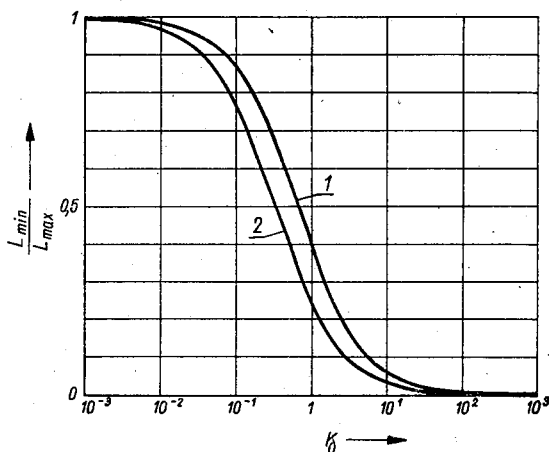
$$\frac{\ln \frac{D'}{D_1}}{\ln \frac{D}{d}}.$$

Jak wspomniano, kondensatory obecnie produkowane („konwencjonalne”) mają z reguły końcówki dołączone do przeciwnych, skrajnych punktów rurki (zob. rys. 1a, 1c, 1e, 1f). Taki sposób dołączenia końcówek powoduje, że kondensator ma indukcyjność własną największą, jaką tylko



Rys. 20. Zależność indukcyjności wewnętrznej kondensatora od umiejscowienia końcówek. Model kondensatora o „grubych” elektrodach

może mieć. Przez proste przesunięcie punktów dołączenia końcówek i maksymalne ich zbliżenie można znacznie zwiększyć częstotliwość rezonansową kondensatora (tj. zmniejszyć jego indukcyjność własną). Oznaczmy stosunek indukcyjności maksymalnej (kiedy końcówki są dołączone do punktów najbardziej odległych) do minimalnej (kiedy końcówki są do-



Rys. 21. Stosunek indukcyjności wewnętrznej kondensatora o optymalnej konstrukcji do indukcyjności wewnętrznej kondensatora o konwencjonalnej konstrukcji (przy takiej samej pojemności i wymiarach) w zależności od parametru k . Krzywa 1 — model teoretyczny kondensatora o grubych elektrodach, krzywa 2 — model teoretyczny kondensatora o cienkich elektrodach

łączone do punktów najbardziej zbliżonych) przez K_{max} . Otrzymamy wówczas z (36) i (22)

— dla cienkich elektrod:

$$K_{max} = 1 + 3\mathcal{K}, \quad (37a)$$

— dla grubych elektrod:

$$K_{max} = 1 + \frac{3}{2}\mathcal{K}, \quad (37b)$$

— dla grubych elektrod i końcówek cylindrycznych:

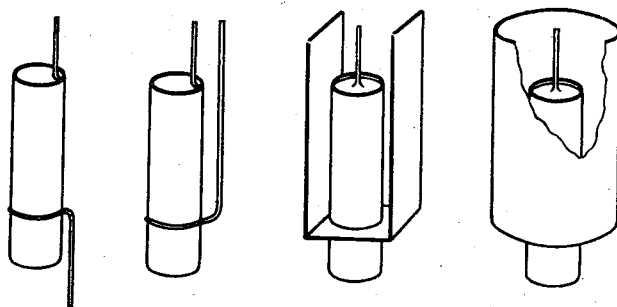
$$K_{max} = 1 + \frac{3 \ln \frac{D'}{D_1}}{\ln \frac{D}{d}}. \quad (37c)$$

Zależności (37a) i (37b) przedstawione są graficznie na rysunku 21. Przez prostą zamianę oznaczenia osi poziomej rysunek ten może odpowiadać także zależności (37c).

3. CZĘŚĆ EKSPERYMENTALNA

3.1. Opis i wyniki pomiarów

W celu sprawdzenia teorii przedstawionej w poprzedniej części wykonano szereg pomiarów częstotliwości rezonansowej kondensatorów o różnej pojemności, różnych wymiarach i różnym umiejscowieniu punktów



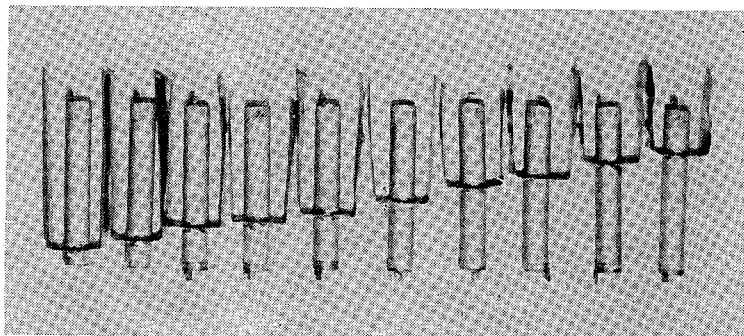
Rys. 22. Szkice badanych kondensatorów o różnych układach końcówek

dołączenia końcówek do elektrod (rys. 22). W sumie zbadano ponad 30 różnych kondensatorów. W mierzonych egzemplarzach częstotliwości rezonansowe mieściły się w zakresie od 30 MHz do 225 MHz.

Początkowo do pomiarów użyto kondensatorów w takim stanie, w jakim są one najczęściej stosowane w praktyce, tj. z dwoma prostymi końcówkami drutowymi o długości rzędu pojedynczych milimetrów. Jednakże

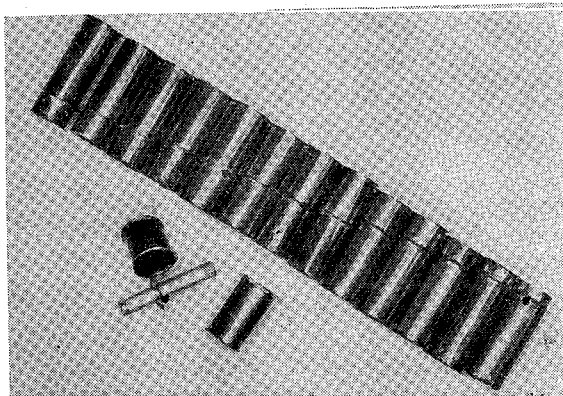
przy tym, na skutek niepożądanych i niekontrolowanych rozplywów prądów, wyniki pomiarów były niepewne, a ich interpretacja utrudniona zwłaszcza przy wyższych częstotliwościach. Ponadto indukcyjność końcówek maskowała w pewnym stopniu wpływ ich umiejscowienia.

W związku z tym zrezygnowano w dalszym ciągu z pomiarów kondensatorów w takim stanie i jedną z końcówek drutowych zastąpiono taśmą o szerokości około 25 mm, wykonaną z cienkiej folii miedzianej. Aby



Rys. 23. Widok przygotowanych do pomiarów 10 kondensatorów z końcówkami z pasków folii

zmniejszyć wpływ pozostałej końcówki, zastąpiono ją dwoma krótkimi, symetrycznie umieszczonymi odcinkami drutu. W celu zmniejszenia wrażliwości układu pomiarowego na niepożądane wpływy zewnętrzne i w celu skrócenia przewodów łączących, a nadto w celu zwiększenia powtarzalności pomiarów przy największych częstotliwościach, paski folii były za-

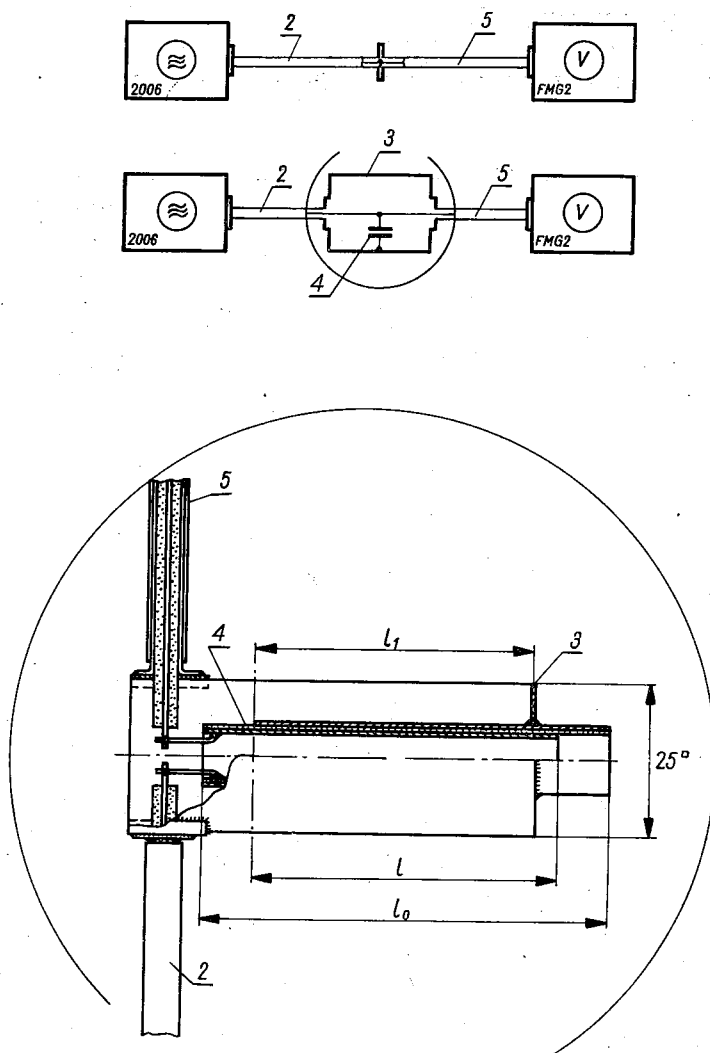


Rys. 24. Widok przygotowanych do pomiarów 14 kondensatorów z końcówkami tworzącymi szczelne ekrany cylindryczne (jeden z kondensatorów ze zdjętym ekranem)

gięte równolegle do osi kondensatora, tworząc swego rodzaju ekran elektromagnetyczny (co prawda nieuszczelny). Szczegóły te są widoczne na rysunku 23, na którym pokazano przygotowaną do pomiarów próbkę

składającą się z 10 takich kondensatorów. Taki sposób postępowania zapewnił dostateczną powtarzalność pomiarów aż do około 150—190 MHz.

Ponieważ przy wyższych częstotliwościach nadal obserwowano niepożądane i niekontrolowane rozprawy prądów na zewnętrznych płaszczech pomiarowych kabli koncentrycznych, powstała koncepcja całkowitego zae-



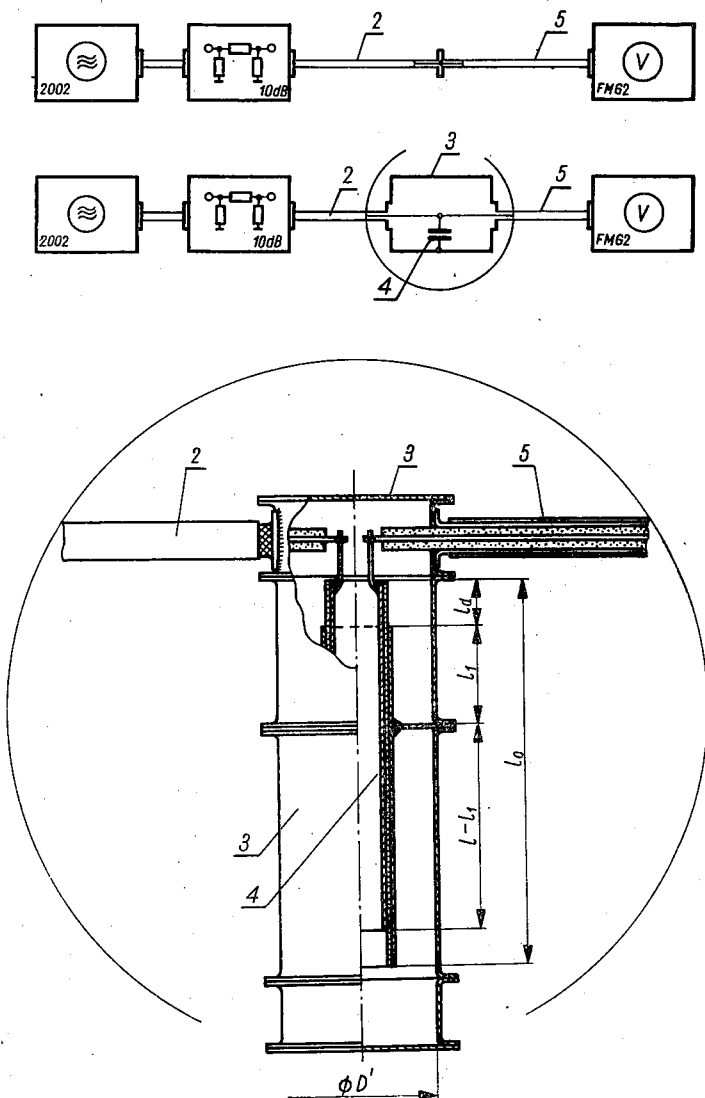
Rys. 25. Schemat blokowy układu pomiarowego I (stosowanego do badania kondensatorów pokazanych na rys. 23)

kranowania kondensatora przez zastosowanie jednej z końcówek kondensatora w kształcie walca otaczającego kondensator i tworzącej szczelny ekran elektromagnetyczny. Na rysunku 24 przedstawiono przygotowaną do pomiarów próbkę składającą się z 14 takich kondensatorów. Taki spo-

sób postępowania okazał się najbardziej dogodny: zapewnił bardzo dobrą powtarzalność pomiarów przy wszystkich częstotliwościach pomiarowych oraz — dzięki szczególnie prostej formie geometrycznej końcówek — ułatwił porównanie teorii z wynikami pomiarów.

Częstotliwość rezonansową kondensatora określono z pomiarów tłumienności (odbić) wprowadzanej przez badany kondensator do ustalonego układu pomiarowego [33].

Schematy zastosowanych układów pomiarowych przedstawiono na rysunkach 25 i 26. Przy każdej częstotliwości pomiarowej najpierw ustalano



Rys. 26. Schemat blokowy układu pomiarowego II (stosowanego do badania kondensatorów pokazanych na rys. 24)

odpowiedni poziom sygnału generatora i odpowiednią czułość odbiornika pomiarowego, łącząc je ze sobą bezpośrednio za pomocą dopasowanego kabla pomiarowego. Następnie badany kondensator załączano między przewodem zewnętrznym i wewnętrznym tego kabla. Na skutek tego w układzie warunki dopasowania zostały naruszone i napięcie wejściowe odbiornika było mniejsze niż poprzednio (w warunkach dopasowania). Gdyby impedancja kondensatora była dokładnie równa zeru, to kable pomiarowe byłyby zwarte i napięcie wejściowe odbiornika byłoby także równe zeru. Gdyby impedancja kondensatora była nieskończenie duża, to warunki dopasowania nie byłyby naruszone i napięcie na wejściu odbiornika nie ulegałoby zmianie.

Jeżeli kabel pomiarowy jest dopasowany odpowiednio do generatora i odbiornika (o impedancjach równych Z_c), to z pomiaru tłumienności (odbicia) wprowadzanej przez badany kondensator do układu pomiarowego, można wyznaczyć jego impedancję, wykorzystując np. zależność [33]

$$Z_k \approx \frac{1}{2} Z_c 10^{-\frac{A}{20}}. \quad (38)$$

Tutaj Z_c oznacza impedancję charakterystyczną układu pomiarowego (impedancję generatora lub odbiornika), zaś A — tłumienność wnoszoną przez

Tablica 1

Zestawienie pojemności, wymiarów i zmierzonych częstotliwości rezonansowych 10 kondensatorów z rys. 23 (układ pomiarowy z rys. 25)

Lp.	Nr	C	D	d	l_0	l	b	f
	*)	pF	mm	mm	mm	mm		MHz
1	1	500 ± 1	$12,4 \pm 0,4$	$11,1 \pm 0,4$	$80,3 \pm 0,7$	$64,25 \pm 0,55$	0,99	57
2	2						0,89	62
3	4						0,74	69
4	6						0,59	72
5	7						0,47	77
6	8						0,39	80
7	9						0,30	88
8	11						0,20	95
9	12						0,10	113
10	13						0,01	190

*) Numer kolejny kondensatora

Oznaczenia podano na rysunkach 1 i 25; $b = \frac{l_1}{l}$

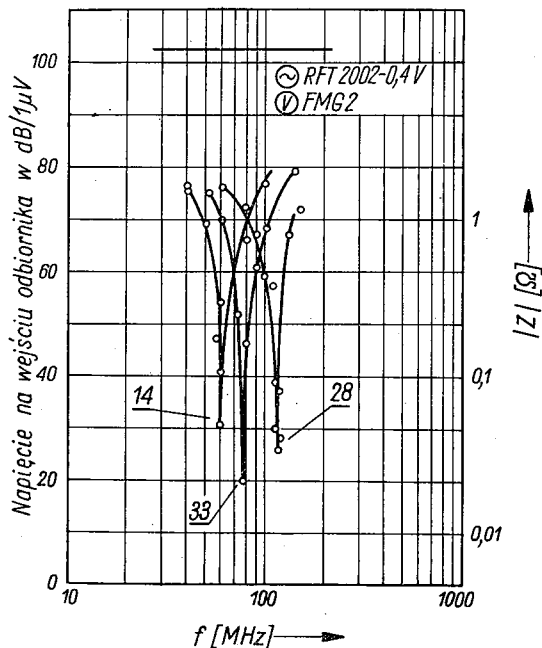
kondensator do tego układu, wyrażoną w decybelach. Zależność (38) daje wystarczającą dokładność dla tłumienności większych od 20 dB.

Należy tutaj podkreślić, że tylko w przypadku dopasowania kabla pomiarowego do generatora i odbiornika wyniki pomiarów nie zależą od

jego długości ani od miejsca włączenia kondensatora. (Sprawa ta jest omawiana m.in. w artykule [33]).

Z zależności (38) wynika, że w opisywanej procedurze pomiarowej najmniejsze napięcie wejściowe odbiornika, tj. największa tłumienność, odpowiada minimum impedancji, czyli rezonansowi kondensatora.

Pomiary wykonywano przy zastosowaniu generatora sygnałów firmy VEB Erfurt UKW Messgenerator typ 2006 (Nr 1401) i odbiornika pomiarowego firmy VEB Dresden, Storfeldstarkemessgerät typ FMG-2 (Nr 60151) (jest on szczegółowo opisany w artykule [22]). W części pomiarów wspomniany generator był zastąpiony generatorem mocy firmy



Rys. 27. Wybrane wyniki pomiarów kondensatorów opisanych w tablicy 3 i uwidocznionych na rys. 24

VEB Erfurt UKW Leistungsgenerator typ 2002 (Nr 1465). Przy stosowaniu tego generatora wykorzystywano oddzielający tłumik 10 dB, zapewniający dopasowanie. Dla dokładnego określenia częstotliwości stosowano częstotściomierz interferencyjny firmy VEB Erfurt, Präz. Frequenzmesser 20—300 MHz, typ 183 (Nr 1468).

Wyniki pomiarów kondensatorów uwidocznionych na rysunku 23 zestawiono w tablicy 1.

Wybrane wyniki pomiarów kondensatorów z rysunku 24 przedstawiono na rysunku 27. Na rysunku tym prosta pozioma 102 dB odpowiada napięciu na wejściu odbiornika (w decybelach względem 1 mikrowolta) przy bezpośrednim jego połączeniu z tłumikiem 10 dB i generatorem (na-

Tablica 2

Zestawienie pojemności, wymiarów oraz zmierzonych i obliczonych częstotliwości rezonansowych 14 kondensatorów z rysunku 24 (układ pomiarowy z rys. 26)

Lp.	Nr	C	D'	D	d	l	l _d	b	f _r ¹⁾	f _r ²⁾	f _r ³⁾	γ _t
	*)	pF	mm	mm	mm	mm	mm		MHz	MHz	MHz	
1	28	526	30,3	12,4	11,4	62,0	14,0	0,015	118,0	118,2	120,8	1,020
2	29	520	30,3	12,4	11,3	62,4	13,8	0,082	101,0	101,4	112,1	1,105
3	30	539	30,5	12,5	11,6	62,8	12,0	0,140	96,0	96,1	99,8	1,038
4	31	545	30,3	12,4	11,3	62,7	13,0	0,220	96,5	96,6	97,4	1,012
5	32	520	30,5	12,3	11,3	62,2	13,3	0,287	86,5	86,1	87,4	1,013
6	33	536	30,3	12,5	11,5	62,3	15,2	0,430	77,0	77,3	76,0	0,984
7	34	503	30,3	12,4	11,3	64,0	8,7	0,484	80,0	80,1	80,4	1,002
8	36	534	30,1	12,4	11,4	61,4	14,1	0,511	74,5	74,8	74,0	0,989
9	36	539	30,2	12,5	11,4	64,0	8,9	0,625	72,0	72,2	73,1	1,012
10	37	541	30,1	12,5	11,5	64,1	9,3	0,734	65,5	65,5	65,8	1,005
11	38	482	30,3	12,3	11,2	63,8	8,6	0,835	66,0	66,3	65,5	0,988
12	39	519	30,6	12,5	11,6	61,8	13,5	0,874	64,0	64,0	63,4	0,990
13	40	535	30,3	12,6	11,5	64,4	7,6	0,938	60,5	60,7	62,6	1,030
14	41	544	30,3	12,4	11,4	63,9	8,8	0,984	58,5	58,7	58,6	1,000

*) Numer kolejny kondensatora

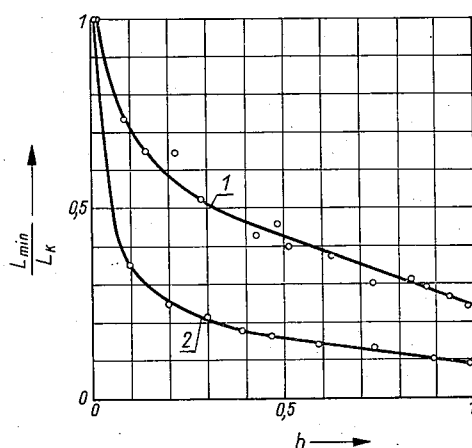
1) Częstotliwość odczytana ze skali odbiornika pomiarowego

2) Częstotliwość odczytana ze skali częstotlicznika

3) Częstotliwość obliczona

Oznaczenia podano na rysunkach 1 i 26; $b = \frac{l_1}{l}$

pięcie generatora utrzymywano przy wszystkich pomiarach na tym samym poziomie). Pozostałe krzywe odpowiadają napięciu na wejściu odbiornika przy włączeniu kondensatora w układ pomiarowy. Kółkami



Rys. 28. Empiryczna zależność indukcyjności wewnętrznej kondensatorów na umiejscowieniu końcówek

Krzywa 1 — kondensatory z rys. 23 i tablicy 1, krzywa 2 — kondensatory z rys. 24 i tablicy 2 (por. rys. 19 i rys. 20)

oznaczono wybrane wyniki pomiarów; cyfry obok krzywych odpowiadają numerom kondensatorów, których wymiary zestawiono w tablicy 2. Jednocześnie w tejże tablicy podano pojemności oraz częstotliwości rezonansowe tych kondensatorów. Indukcyjności kondensatorów obliczone z częstotliwości rezonansowych przedstawiono na rys. 28.

Dokładność pomiarów pojemności (mostek małej częstotliwości firmy Elektromatyka, typ 034/55 nr 1411) można określić na ± 1 pF. Długość kondensatorów i ich elektrod mierzono z dokładnością do $\pm 0,05$ mm. Należy jednak wziąć tu pod uwagę, że przy pomiarach długości wewnętrznych elektrod kondensatorów niekiedy trudno było uzyskać dokładność lepszą niż $\pm 0,5$ mm. Przy pomiarach średnic kondensatorów w wielu przypadkach stwierdzono, że przekrój poprzeczny rurki ceramicznej nie jest kołowy. W związku z tym w tablicach 1 i 2 podano wartości średnic uzyskane z kilku pomiarów. Dokładność pomiaru średnic wynosiła około $\pm 0,05$ mm. Dokładność pomiaru częstotliwości rezonansowej można oszacować na podstawie procedury pomiarowej, według której najbliższe punkty pomiarowe były odległe o 0,5 MHz, oraz na podstawie dokładności określenia częstotliwości. Według danych katalogowych skala odbiornika jest cechowana z dokładnością $\pm 1\%$, natomiast częstotlicznik pozwalał określać częstotliwość z dokładnością $\pm 2 \cdot 10^{-4}$.

3.2. Porównanie wniosków teoretycznych z wynikami pomiarów

W pracach opublikowanych dotychczas, związanych z rozważanym zagadnieniem, omawiane są rozmaite czynniki wpływające na częstotliwość rezonansową kondensatora, z wyjątkiem wpływu umiejscowienia końcówek. W związku z tym w niniejszej pracy główny nacisk położono na zbadanie i opis tego właśnie wpływu. Dlatego też w poprzednim punkcie zamieszczono tylko wybrane wyniki pomiarów, ilustrujące głównie wpływ umiejscowienia końcówek.

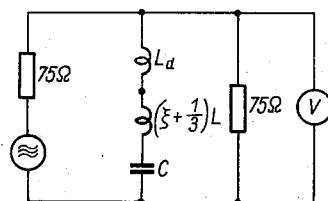
Dosyć złożony układ geometryczny końcówek i elektrod kondensatorów z rys. 23 powoduje pewne trudności rachunkowe przy wyznaczaniu indukcyjności wzajemnej między nimi. Dlatego zrezygnowano z bardziej dokładnego ilościowego porównania wyników pomiarów tych kondensatorów z teorią. Z tablicy 1 wynika, że jakościowo wyniki pomiarów i wnioski teoretyczne są w zasadzie zgodne.

Dla porównania ilościowego wykorzystano wyniki pomiarów kondensatorów uwidocznionych na rys. 24. W tych kondensatorach szczególnie prosty układ geometryczny końcówek i elektrod został wybrany specjalnie, nie tylko aby ułatwić pomiary, ale także aby ominąć wspomniane wyżej trudności rachunkowe.

Jak wynika z rozważań podanych w rozdziałach poprzednich, przy pominięciu strat układ pomiarowy wraz z badanym kondensatorem można

zastąpić układem równoważnym, uwidocznionym na rys. 29. Na tym rysunku generator pomiarowy został zastąpiony SEM-ną połączoną w szeregu z opornikiem o oporności $75\ \Omega$, a odbiornik pomiarowy został przedstawiony jako równoległe połączenie „idealnego” woltomierza i opornika $75\text{-}\Omega$ owego. C i L_k oznaczają, jak poprzednio, pojemność (przy m.cz.) i indukcyjność własną kondensatora, wywołaną rozplywem prądu w elektrodzie zewnętrznej. L_d oznacza indukcyjność końcówek, w tym przypadku równą indukcyjności części zewnętrznej elektrody kondensa-

Rys. 29. Przybliżony układ zastępczy kondensatora w układzie pomiarowym $\left(L_K = \left(\xi + \frac{1}{3}\right)L\right)$



tora o długości l_0 (zob. rys. 26). Przyjmując, że elektrody badanych kondensatorów są grube w rozpatrywanym zakresie częstotliwości, możemy wykorzystać wyprowadzone wcześniej zależności.

Na podstawie rysunku 29 możemy napisać natychmiast wzór (przybliżony) na częstotliwość rezonansową kondensatora wraz z końcówkami, która powinna być zaobserwowana w układzie pomiarowym.

$$f_{r0} \approx \frac{1}{2\pi \sqrt{C(L_k + L_d)}}. \quad (39)$$

Przez wprowadzenie parametrów pomocniczych

$$x'' = \frac{2}{\pi} \frac{1}{\sqrt{\xi'' + \frac{1}{3}}}, \quad (40)$$

$$\xi'' = \xi + \xi', \quad (41)$$

$$\xi' = \frac{L_d}{3L_{min}}, \quad (42)$$

gdzie ξ wynika z zależności (21) i (22), a L_{min} — ze wzoru (33), związek (39) można sprowadzić do postaci

$$f_{r0} \approx x'' f_0 \quad (43)$$

analogicznej do postaci (34). Teraz, do określenia częstotliwości rezonansowej kondensatora z końcówkami można wykorzystywać rys. 14 i 17, podstawiając tylko zamiast f_0 — f_{r0} , a zamiast x_0 — x'' .

W badanych kondensatorach elektrody były tak delikatne, że kilkukrotne dolutowywanie do nich końcówek mogłoby uszkodzić je. W związku z tym w każdym egzemplarzu kondensatora indywidualnie ustalono jedno tylko umiejscowienie końcówek. Określenie wpływu umiejscowie-

nia końcówek powinno w zasadzie odbywać się przez porównanie częstotliwości rezonansowych kondensatorów identycznych pod każdym względem, z wyjątkiem różnic w umiejscowieniu końcówek. Jednak dobranie odpowiedniej ilości identycznych kondensatorów okazało się w praktyce bardzo trudne.

Metoda porównania teorii z wynikami pomiarów, zastosowana w niniejszej pracy była następująca. Wybrano próbkę zawierającą $n = 14$ kondensatorów wykonanych z tego samego materiału ceramicznego.

Dla każdego z nich zmierzono pojemność, częstotliwość rezonansu i istotne wymiary geometryczne (wyniki zestawiono w Dodatku). Zakładając, że elektrody kondensatorów są „grube”, obliczono dla każdego z nich teoretyczną wartość częstotliwości rezonansowej. Dla zwiększenia dokładności wykorzystano przy tym fakt, że przenikalność dielektryczna materiału ceramicznego we wszystkich egzemplarzach powinna być taka sama. Następnie dla każdego kondensatora utworzono stosunek częstotliwości obliczonej do zmierzonej

$$y_i = \left(\frac{f_{r0}}{f_r} \right)_i \quad i = 1, 2, \dots, 14 \quad (44)$$

Dalej założono, że rozkład wartości y_i , traktowanej jako zmienna losowa, jest normalny [53], i wyznaczono jego parametry: wartość średnią

$$\bar{y} = \frac{1}{n} \sum_{i=1}^n y_i \quad (45)$$

i błąd

$$\frac{s}{\sqrt{n}} = \frac{1}{\sqrt{n}} \sqrt{\frac{1}{n-1} \sum_{i=1}^n (\bar{y} - y_i)^2}. \quad (46)$$

Wartość średnią $\bar{y}$ przyjęto za kryterium i miarę zgodności teorii z wynikami pomiarów.

Opisany sposób postępowania ma szereg zalet:

— pozwala wykorzystać wszystkie wyniki pomiarów w jednakowym stopniu, niezależnie od indywidualnych różnic między poszczególnymi egzemplarzami w próbce,

— przez łączne wykorzystanie jednocześnie wszystkich wyników pozwala zmniejszyć wpływ przypadkowych błędów i uzyskać dokładność znacznie większą niż dokładność indywidualnego pomiaru i obliczenia,

— stanowi obiektywną ocenę zarówno pomiarów (w sensie zespołu metody, aparatury i osoby wykonującej pomiar), jak i obliczeń,

— przy założeniu braku błędów systematycznych stanowi obiektywną ocenę i kryterium zgodności teorii i eksperymentu, wyrażone za pomocą jednej tylko (lub dwóch) liczb,

— pozwala łącznie sprawdzić wpływ wszystkich czynników, którymi różnią się badane egzemplarze, bez konieczności doboru specjalnych egzemplarzy różniących się tylko jedną określoną cechą, a poza tym identycznych.

Gdyby wszystkie pomiary i obliczenia były wykonane w sposób bezbłędny, a teoria adekwatna, to oczywiście wszystkie częstotliwości obliczone byłyby równe zmierzonym. Mielibyśmy wówczas

$$\bar{y} = y_i = 1$$

oraz

$$s = 0.$$

W warunkach realnych wartość s jest miarą rozrzutu wyników pomiarów i obliczeń, a wartość $(\bar{y} - 1)$ — miarą poprawności teorii. Należy zauważyć, że systematyczne błędy pomiarów czy obliczeń przejawiają się tutaj w taki sam sposób jak niezgodność teorii z eksperymentami.

Odpowiednie obliczenia numeryczne zostały podane w pracy [40]. Ponieważ nie wnoszą one w zasadzie nic nowego, zostały tutaj pominięte.

W celu zwiększenia dokładności za Mitropolskim [53] odrzucono rezultaty znacznie odbiegające od pozostałych (zaobserwowano jeden taki przypadek — kondensator nr 29). W tym celu przyjmowano liczbę

$$\chi = \text{Max}|\bar{y} - y_i| \quad (47)$$

i szacowano oczekiwaną (najbardziej prawdopodobną) wartość liczby (n_1) egzemplarzy w próbce, dla których wartość bezwzględna odchylenia jest większa lub równa χ

$$n_1 = n \cdot P(|\bar{y} - y_i| \geq \chi), \quad (48)$$

gdzie $P(|\bar{y} - y_i| \geq \chi)$ jest prawdopodobieństwem, że $|\bar{y} - y_i| \geq \chi$. Prawdopodobieństwo to szacowano przyjmując, że odchylenie średnie kwadratowe σ jest w przybliżeniu równe odchyleniu średniemu kwadratowemu s w próbce o liczności n . Wykorzystując fakt, że rozkład y_i jest normalny, prawdopodobieństwo P określano przy użyciu tablic funkcji Φ [53].

$$P(|\bar{y} - y_i| \geq \chi) = 1 - \Phi(x'), \quad (49)$$

gdzie

$$x' = \frac{\chi}{\sigma} \approx \frac{\chi}{s}. \quad (50)$$

W przypadku gdy wartość oczekiwana n_1' okazywała się znacznie mniejsza od zaobserwowanej w próbce ilości egzemplarzy n_1 , to odrzucano z próbki te egzemplarze, dla których $|\bar{y} - y_i| \geq \chi$, jako mało prawdopodobne i obciążone niedopuszczalnym błędem przypadkowym.

W wyniku obliczeń uzyskano ostatecznie

$$\bar{y} = 1,006, \quad (51)$$

$$\frac{s}{\sqrt{n}} = 0,005.$$

Zatem odchylenie od przypadku idealnej zgodności teorii i eksperymentu oraz idealnej dokładności jest rzędu 1‰. Taki rezultat można uznać za bardzo dobry.

Istniejąca rozbieżność może być spowodowana wieloma przyczynami:

- przybliżeniem układu zastępczego kondensatora i przybliżeniem wzorów teoretycznych,
- przybliżeniem obliczeń numerycznych,
- błędami systematycznymi pomiarów (mogły one powstać zwłaszcza przy pomiarach długości wewnętrznych, trudno dostępnych elektrod kondensatora).
- ograniczoną i niezbyt wielką ilością badanych egzemplarzy kondensatorów.

4. ZAKOŃCZENIE

Ze względu na grubość elektrod w porównaniu z głębokością wnika-
nia pola elektromagnetycznego wszystkie kondensatory rurkowe można
podzielić na trzy grupy: pierwszą, do której należą wszystkie kondensa-
tory o elektrodach bardzo cienkich, drugą — odpowiadającą grubości
elektrod porównywalnej z głębokością wnika-
nia i w końcu trzecią —
obejmującą kondensatory o elektrodach bardzo grubych. Kondensatory
przeznaczone do pracy w bardzo szerokim zakresie częstotliwości oczywi-
ście w miarę wzrostu częstotliwości należą do coraz to innej grupy, poczy-
nając od pierwszej przy bardzo małych częstotliwościach i kończąc na
ostatniej przy bardzo dużych.

Z kolei, zależnie od wzajemnego usytuowania geometrycznego końców-
wek i elektrod kondensatora, można rozróżnić dwa przypadki: pierwszy,
kiedy sprzężenie końcówek i elektrod jest pomijalne i drugi, w którym
to sprzężenie elektromagnetyczne jest znaczne i nie może być pominięte.

W ten sposób wszystkie możliwe przypadki można z grubsza podzielić
na sześć różnych grup. W niniejszej pracy z tych sześciu omówiono tylko
trzy grupy, najbardziej różniące się między sobą. Pomimo całkowicie od-
miennych modeli teoretycznych przyjętych do analizy tych skrajnych
przypadków granicznych uzyskano dla nich zaskakująco zbieżne rezultaty
[por. zał. (27) i (22)]. Dla końcówek dołączonych do elektrod w punktach
„optymalnych” wszystkie rozpatrzone modele teoretyczne prowadzą do
identycznych rezultatów numerycznych. Położenie tych punktów „opty-
malnych”, określone w oparciu o trzy zupełnie różne modele teoretycz-

ne, jest jedno i to samo. Największe rozbieżności odpowiadają końcówkom dołączonym do punktów najbardziej odległych. Jednak nawet wówczas częstotliwości rezonansowe określone w oparciu o różne modele teoretyczne, jak można oszacować, różnią się w pierwszym przybliżeniu nie więcej niż o około trzydziestu procent. Zatem można spodziewać się, że w przypadkach pośrednich, nie omówionych w niniejszej pracy, sytuacja przedstawia się podobnie.

Zarówno rozważania teoretyczne, jak i wyniki pomiarów potwierdzają, że częstotliwość rezonansowa kondensatora rurkowego zależy

- od jego pojemności określonej przy małej częstotliwości,
- od indukcyjności końcówek i indukcyjności wzajemnej między nimi,
- od impedancji charakterystycznej kondensatora,
- od umiejscowienia punktów dołączenia końcówek do jego elektrod,
- do indukcyjności wzajemnej między elektrodami i końcówkami,
- od zewnętrznej średnicy kondensatora i jego długości.

W szczególności wykazano, że wpływ umiejscowienia końcówek na najwyższą częstotliwość rezonansową kondensatora jest zasadniczy. W związku z tym wydaje się bardzo dziwne, że w dostępnej literaturze zależność ta jest pomijana milczeniem; można wymienić tutaj nie tylko poradniki encyklopedyczne [1], [14], [21], [23], [41] czy artykuły [5], [13], [15], [24], [42], [44], ale również poważne monografie [3], [6], [7], [45], [49], [55].

Być może przyczyna takiego stanu rzeczy tkwi w konkurencji, jak tłumaczy to Dummer i Nordenberg we wspomnianej monografii [7].

Należy jednak podkreślić, że w szczególnym przypadku, kiedy końcówki kondensatora są dołączone do przeciwnych krawędzi elektrod, rezultaty niniejszej pracy praktycznie rzecz biorąc pokrywają się z niektórymi opublikowanymi wcześniej wynikami badań. Na przykład układ zastępczy pokazany na rysunku 15 jest wówczas podobny do układu zastępczego proponowanego przez Weisa [44], a rysunek 17 — sprowadza się do układu zastępczego podanego przez Hartmanna [13]. Również wyniki pomiarów w tym szczególnym przypadku w pierwszym przybliżeniu pokrywają się z wynikami pomiarów opublikowanymi w literaturze [13], [14], [15], [24], [42], [46], [55].

Wykazano, że istnieje możliwość znacznego rozszerzenia zakresu częstotliwości użytkowych omawianych kondensatorów i zmniejszenia ich indukcyjności własnej bez zmiany ich pojemności, kształtu, wymiarów i materiału, a jedynie przez dołączenie końcówek do odpowiednich punktów. Oznacza to, że bez dodatkowych kosztów i istotnych zmian technologii przydatność i wartość użytkowa kondensatorów tego typu może być znacznie zwiększona.

Sprawa ta jest ważna o tyle, że takie kondensatory są stosowane w skali masowej we wszelkiego rodzaju aparaturze elektronicznej wiel-

kiej częstotliwości. Ponadto kondensatory o podobnej konstrukcji mogą być wykorzystywane w miernictwie wielkiej częstotliwości, na przykład jako wzorce pojemności lub jako elementy wzorcowych dzielników pojemnościowych: i w tym przypadku zmniejszenie ich indukcyjności wewnętrznej jest bardzo pożądane.

Na zakończenie pragnę złożyć wyrazy gorącej wdzięczności p. prof. Cz. Rajskiemu i p. prof. W. Rotkiewiczowi za owocne dyskusje, które pomogły mi w rozwiązaniu szeregu trudności napotkanych przy opracowywaniu niniejszego zagadnienia. Specjalne podziękowanie pragnę złożyć także p. prof. S. Ryżce za uprzejmą i życzliwą sugestię układu zastępczego kondensatora o „grubych” elektrodach jak również za unikliwe i krytyczne uwagi, które pozwoliły usunąć szereg usterek pierwszej wersji pracy [40].

Część eksperymentalna przedstawionej pracy została wykonana w Instytucie Łączności, w Oddziale we Wrocławiu. Za daleko idącą pomoc i życzliwość pragnę podziękować kierownictwu Instytutu w osobach p. dyr. Z. Szpiglera i p. dyr. W. Cetnera oraz Kolegom z Pracowni Z-10/2, a szczególnie p. mgrowi inż. H. Smorągowi.

Instytut Łączności
Oddział we Wrocławiu

WYKAZ LITERATURY

1. Antoniewicz J.: Poradnik radio- i teleelektryka. Część B. Elementy i podzespoły. Warszawa 1959.
2. Bayer von H. J., Knechtli R.: Über die Behandlung von Mehrleistersystemen mit transversal elektromagnetischen Wellen bei hohen Frequenzen. — Zeitschrift für Angewandte Mathematik und Physik, 1952, vol. III, nr 4, str. 271—286.
3. Brotherton M.: Capacitors — their Use in Electronic Circuits. London 1946.
4. Cirkler W., Kuny W.: Das Frequenzverhalten Keramischer Dielektrika im Bereich von 10^2 bis $5 \cdot 10^9$ Hz. — Frequenz, 1961, nr 9, str. 277—285.
5. Davidson R.: RF Characteristics of Capacitors. — Wireless World, 1952, vol. 68, August, str. 301—303 (zob. również korespondencję w tym samym roczniku: October, str. 419 i November, str. 457).
6. Dummer G.W.A.: Fixed Capacitors. London 1956.
7. Dummer G.W.A., Nordenberg H. M.: Fixed and Variable Capacitors. New York 1960.
8. Electronic Buyers Guide and Reference Issue 1960. A Mc Graw Hill Book Co Publ. (str. R28).
9. Electronic Buyers Guide and Reference Issue 1961. A Mc Graw Hill Book Co Publ. (str. R5).
10. Geiser D.T.: An Investigation of Lowest Resonant Frequency in Commercially Available by-pass Capacitors. — Conv. Rec. of the IRE, 1954, P3, str. 43.
11. Grobelny M.: Dobór kondensatorów blokujących opór przy bardzo wielkich częstotliwościach. — Przegląd Telekom., 1960, nr 4, str. 108—112.
12. Happ W. W., Riddle G. C.: Limitations of Film-Type Circuits Consisting Of Resistive and Capacitive Layers. — IRE Internat. Conv. Rec., P6. 1961, str. 141—165.
13. Hartmann W.: Über das Verhalten von Kondensatoren bei Hochfrequenz. — Bull. Schweiz. El., 1953, vol. 44, str. 258—262.

14. Henney K., Walsh C.: *Electronic Components Handbook. Resistors, Capacitors, Relays, Switches.* New York 1957.
15. Herfurth J.: *Eigenresonanzen von Kondensatoren.* — *Radio u. Fernsehen*, 1960, nr 18, str. 572—573.
16. Hippel von A. R.: *Dielectric Materials and Applications.* New York 1957.
17. International Electrotechnical Commission: doc. 40 (Central Office) 108: *Draft Specification for Polyester Film Dielectric Capacitors for Direct Current.* 1962, May.
18. International Electrotechnical Commission: doc. 40 (Central Office) 109: *Draft Specification for Fixed Metallized Paper Dielectric Capacitors.* 1962, May.
19. Commission Electrotechnique Internationale: doc. 40 (Bureau Central) 110: *Projet specification pour condensateurs à diélectrique en céramique — type 2:* 1962, Mai.
20. King R. W. P.: *Transmission-Line Theory.* New York 1955.
21. Landee R. W., Davis D. C., Albrecht A. P.: *Electronic Designers Handbook.* New York 1957.
22. Loreznz H.: *Ein Störfeldstärkemessgerät für den Frequenzbereich von 30 bis 225 MHz.* — *Nachrichtentechnik*, 1960, nr 1, str. 23—28.
23. Meinke H., Gundlach F. W.: *Taschenbuch der Hochfrequenztechnik.* Berlin, 1956, przekład rosyjski, t. 1, Moskwa 1960.
24. Mansfeld W.: *Schichtwiderstände, Kondensatoren, Drosselspulen, ihr Verhalten im Frequenzgebiet von 10 ... 200 MHz.* — *Nachrichtentechnik*, 1952, nr 4, str. 105—108, nr 6 str. 185—189, nr 7, str. 201—204, nr 8, str. 241—245, nr 10, str. 310—313.
25. Muckenhaupt J.: *Resonances in Capacitors.* — *Proc. IRE*, 1949, nr 5, str. 532—533.
26. Raisbeck G., Manley J. M.: *Transmission Characteristics of a Three Conductor Coaxial Transmission Line With Transpositions.* *The Bell Syst. Tech. Journ.*, 1958, nr 4, str. 835—877.
27. Rajska Cz.: *Teoria skośnika.* — *Archiwum Elektrotechniki*, t. II, 1953, nr 1—2, str. 137—165.
28. Ramo S., Whinnery J. R.: *Fields and Waves in Modern Radio.* New York 1944.
29. Rotkiewicz W.: *O indukcyjności kondensatorów przy bardzo wielkiej częstotliwości.* — *Wiadomości i Prace Instytutu Radiotechnicznego w Warszawie*, 1930, nr 5.
30. Ryżko S.: komunikat prywatny.
31. Slate M. W.: *Resonance Effects in Tubular Feed-Through Capacitors.* — *Tele-Tech. and Electronic Industries*, 1954, nr 6, str. 98—101, 180, 182, 374—378.
32. Strużak R. G.: *Zachowanie się kondensatorów zwijanych pracujących w szerokim pasmie częstotliwości.* — *Prace Instytutu Łączności*, 1959, nr 1(14), str. 75—94.
33. Strużak R. G.: *Charakterystyki kondensatorów przy wielkiej częstotliwości i metody ich pomiaru.* — *Prace Instytutu Łączności*, 1962, nr 3(28), str. 41—67.
34. Strużak R. G.: *Kondensator rurkowy stały o zmniejszonej indukcyjności.* Opis patentu PRL Nr 46168, 1.VII.1961.
35. Strużak R. G.: *Ceramiczny kondensator rurkowy o zmniejszonej indukcyjności.* — *Przegląd Telekomunikacyjny*, 1963, nr 3, str. 88—90.
36. Strużak R. G.: *Ogólny przypadek prądów i napięć w jednorodnym torze elektrycznym.* — *Archiwum Elektrotechniki*, 1963, nr 2, tom XII, str. 229—236.
37. Strużak R. G.: *Przyczynek do teorii dwójnika utworzonego z odcinka jednorodnego toru elektrycznego.* — *Archiwum Elektrotechniki*, tom XII, 1963, nr 2, str. 237—263.

38. Strużak R. G.: Ogólna teoria czwórnika w zastosowaniu do układów o parametrach rozłożonych jednorodnie. — Arch. Elektrotechn., 1965 (w druku).
39. Strużak R. G.: Impedancja kondensatora cylindrycznego przy wielkiej częstotliwości. — Arch. Elektrotechn., 1965 (w druku).
40. Strużak R. G.: Efekty rezonansowe w rurkowych kondensatorach ceramicznych (dysertacja), Politechnika Warszawska, 1962.
41. Terman F. E.: Radio Engineers Handbook. New York 1943.
42. Trankle K.: Die Eigenresonanz von Kondensatoren und ihre Messung. — Funkschau, 1959, nr 12, str. 290.
43. Warmus M.: Tablice funkcji elementarnych. Warszawa 1960.
44. Weis A.: Über den Scheinwiderstand von Drosseln, Widerständen und Kondensatoren bei Hochfrequenz. — Frequenz, 1955, nr 7, str. 221—227.
45. Аршанский М. Е.: Керамические конденсаторы малой реактивной мощности. Москва 1953.
46. Волгов В. А.: Детали контуров радиоаппаратуры. Москва 1954.
47. Воронов Р. А.: Общая теория четырёхполюсников и многополюсников. Москва 1951.
48. Говорков В. А.: Электрические и магнитные поля. Москва 1951.
49. Казарновский А. М.: Сегнетокерамические конденсаторы. Москва 1956.
50. Коваленков В. И.: Решение обобщенных телеграфных уравнений методом расщепления уравнений. — Сборник Научн. Работ по проводной связи, Академия Наук СССР, 1949, стр. 8—46.
51. Кузнецов В. Н.: О приведении к нормальному виду системы четырёх уравнений с частными производными по двум независимым переменным. — Сборник Научн. Работ по проводной связи, Академия Наук СССР, 1949, стр. 46—72.
52. Маш Д. Я.: Потери и диэлектрическая проницаемость титаната бария в поле ультравысокой частоты. — Журнал экспер. и теорет. физики, 1947, т. XVII, стр. 537—539.
53. Митропольский А. К.: Техника статистических вычислений. Москва 1961.
54. Новосильцев Н. С., Ходаков А. А.: Диэлектрические свойства титаната бария при высоких частотах. — Журнал технической физики, 1947, т. XVII, в. 7, стр. 651—656.
55. Ренне В. Т.: Электрические конденсаторы. Москва 1959.
56. Сигорский В. П.: Общая теория четырёхполюсника. Киев 1955.

R. G. STRUZAK

RESONANCES IN CERAMIC TUBE CAPACITORS

Summary

The author gives a theoretical and experimental analysis of the influence of various construction factors on the resonance frequency of a ceramic tube two-terminal capacitor. He takes into consideration:

- 1) the influence of the lead inductances,
- 2) the influence of the reflections of TEM wave (i.e. the influence of geometric dimensions and electrical properties of ceramic dielectric).
- 3) the influence of the location of terminals.

In the analysed capacitor the propagation of higher order waves (the effect of a waveguide), molecular resonances, piezo-electricity and other phenomena are of secondary importance [31].

Various theoretical models of the capacitor have been proposed for various frequency ranges:

1) for frequencies at which the thickness of capacitor electrodes is very small in comparison with penetration depth [36], [37], [formulae (8) and (10) and figs. 4 and 7] and

2) for frequencies at which the thickness of capacitor electrodes is big in comparison with penetration depth (skin-effect) [38], [39] [formulae (14), (15), (19), (20) and figs. 5 and from 8 to 12].

The general equivalent circuit for the analysed capacitor, ensuing from the above theoretical models (without considering lead inductances) is presented in the figs. 16 and 13; its impedance is given by the relations (21) — (32).

It follows from these relations that the impedance of the capacitor (without leads) depends to a considerable extent on the location of the terminals. In the sources known to the author this dependence has not been taken into consideration.

The author shows that there exists an optimum in placing the terminals on the capacitor electrodes which ensure maximum resonance frequency of the capacitor [formula (34)]. With the optimum location of the terminals, interior inductance of the capacitor is given by the relation (33); the resonance frequency — by the nomogram presented in the fig. 17. The examples of the construction of capacitors with the terminals joined in optimum places are shown in fig. 18.

With the arbitrary location of terminals capacitor inductance and resonance frequency (without considering lead inductances) are given by relations (34) — (37) or by figs. 14, 17, and 19—21. The influence of lead inductances is discussed in chapter 2.

Further the author gives the description and the results of his experiments. The measured capacitors are presented in the figs. 22—24; the measuring system — in figs. 25, 26 and 29; the selected results of the measurements in the figs. 27, 28 and in tables 1 and 2. The experiments have proved the correctness of the theory, at least in the cases which were tested.

Without the change of the material, shape and the dimensions of the capacitor and only by the change of the point joining the terminal with the electrodes, the resonance frequency may be changed in wide ranges e.g. from 57 MHz to 190 MHz when the influence of the lead inductances is suitably small. There exists, therefore, a new possibility for a considerable augmentation of the utility frequency range of the analysed capacitors without additional costs or serious alterations in technology.

R. G. STRUZAK

RÉSONANCES DANS CONDENSATEURS TUBULAIRES CÉRAMIQUES

Résumé

L'auteur présente l'analyse théorique et expérimentale de l'influence de différents facteurs constructifs sur la fréquence de résonance de condensateur tubulaire céramique à deux bornes. Il prend en considération:

1) influence de l'inductivité des extrémités,

2) influence des réflexions de l'onde électromagnétique transversale TEM, c. à d. l'influence des dimensions géométriques de condensateur et des propriétés de diélectrique.

La propagation des ondes de plus hauts modes (effet de guide d'onde) ainsi que les résonances moléculaires, l'effet piézo-électrique et les autres en cas de condensateur analysé jouent un rôle secondaire [31].

On a proposé de différents modèles théoriques de condensateur, en vigueur, dans de différents bandes de fréquences:

1) fréquences, pour lesquelles l'épaisseur des armatures de condensateur est très petite en comparaison avec la profondeur de pénétration des ondes [36], [37] [Formules concernant ce cas: (8) et (10) et figures 4 et 7].

2) fréquences, pour lesquelles l'épaisseur des armatures de condensateur est importante en comparaison avec la pénétration des ondes [38] [Formules concernant ce cas: (14), (15), (19), (20) et figures 5 et 8 jusqu'à 12].

Le schéma général équivalent de condensateur résultant des modèles théorique mentionnés ci-dessus (sans prendre en considération des extrémités) est présenté sur figures 13 et 16; son impédance est déterminée par des dépendances (21)–(32). Il résulte de ces dépendances, que l'impédance de condensateur (sans extrémités) dépend essentiellement entre autres de la localisation des bornes. Une telle dépendance n'est pas considérée par la littérature technique connue à l'auteur. L'auteur démontre qu'il y a des places optimales à raccorder des bornes aux armatures de condensateur, assurant la fréquence maximum de résonance de condensateur [34].

Quand la localisation des bornes et optimale, l'inductivité intérieure de condensateur est donnée par la dépendance (33), et la fréquence de résonance peut être déterminée de nomogramme présente sur la fig. 17. Les exemples de construction des condensateurs à bornes raccordées dans les places optimales sont présenté sur fig. 18.

Quand la localisation des bornes est librement choisie l'inductivité de condensateur et sa fréquence de résonance (sans prendre en considération l'inductivité des bornes) peut être déterminée des dépendances (34)–(37) ou des figures 14, 17 et 19–21.

Ensuite l'auteur présente la description et les résultats des expériences. Condensateurs essayés sont présentés sur figures 22–24; schéma de mesure — sur figures 25, 26 et 29; résultats choisis des mesures — sur figures 27 et 28 ainsi que dans tableaux 1 et 2. Ces expériences ont prouvé la régularité de la théorie, au moins dans les cas essayés.

Sans changer matériel, forme et dimensions de condensateur, changeant seulement les points de raccord des bornes aux armatures, la fréquence de résonance de condensateur peut être changée dans de larges limites, par ex. de 57 MHz à 190 MHz quand l'inductivité des extrémités est convenable petit. Ça veut dire qu'il existe une nouvelle possibilité d'importante amélioration des valeurs d'usage des condensateurs analysés en principe sans des frais additifs et changement de la technologie.

Р. Г. СТРУЖАК

РЕЗОНАНСЫ В КЕРАМИЧЕСКИХ ТРУБЧАТЫХ КОНДЕНСАТОРАХ

Резюме

Автор приводит теоретический и экспериментальный анализ влияния разных конструктивных факторов на резонансную частоту трубчатого керамического конденсатора с двумя выводами. При этом принимает во внимание:

- 1) влияние индуктивности выводов,
- 2) влияние отражений поперечной электромагнитной волны ТЕМ, т. е. влияние геометрических размеров конденсатора и свойств диэлектрика,
- 3) влияние расположения выводов.

Распространение волн высшего порядка (волноводный эффект), а также молекулярные резонансы, пьезоэлектрические явления и другие имеют в рассматриваемом конденсаторе второстепенное значение [31].

Созданы разные теоретические модели конденсатора, обязывающие в разных диапазонах частот:

1) при частотах, при которых толщина обкладок конденсатора очень мала по сравнению с глубиной проникновения [36], [37]. (Здесь обязывают формулы (8) и (10) а также рисунки 4 и 7);

2) при частотах, при которых толщина обкладок конденсатора большая по сравнению с глубиной проникновения [38]. (К этому случаю относятся формулы (14), (15), (19), (20) и рисунки 5, а также 8 и 12).

Общая эквивалентная схема конденсатора, вытекающая из вышеуказанных теоретических моделей (без учёта выводов) представлена на рис. 13 и 16; его импеданс определен зависимостями (21) — (23). Из этих зависимостей вытекает, что импеданс конденсатора (без выводов) существенным образом зависит между пр. от расположения выводов. В известной автору литературе такая зависимость не принимается во внимание. Автор доказывает что существуют оптимальные места приключения выводов к обкладкам конденсатора, обеспечивающие максимальную резонансную частоту конденсатора [34].

При оптимальном расположении выводов внутренняя индуктивность конденсатора дана зависимостью (33), а резонансная частота может быть определена из номограммы, представленной на рис. 17. Примеры конструкции конденсаторов с выводами, приключенными в оптимальных местах показаны на рис. 18.

При произвольном расположении выводов собственная индуктивность конденсатора и его резонансная частота (без учёта индуктивности выводов) может быть определена из зависимости (34) до (37) или из рисунков 14 и 17, а также 19—21.

Далее автором приводится описание и результаты экспериментальных исследований. Испытуемые конденсаторы представлено на рис. 22—24; измерительная схема — на рис. 25, 26 и 29; избранные результаты измерений — на рис. 27 и 28, а также в таблицах 1 и 2. Эксперименты подтвердили правильность теории, по крайней мере в исследуемых случаях.

Без изменения материала, формы и размеров конденсатора только лишь путем изменения точек приключения выводов к обкладкам, резонансная частота конденсатора может изменяться в широких пределах, нпр. от 57 мГц до 190 мГц. если индуктивность выводов соответственно мала. Существует затем новая возможность значительного повышения потребительского качества рассматриваемых конденсаторов в принципе без добавочных расходов и существенных изменений технологии производства.

Dodatek

Zestawienie częstotliwości rezonansowych ceramicznych kondensatorów rurkowych o dwóch końcówkach (wg literatury)

Lp.	Pojemność	Częstotliwość	Długość końcówek rurki		Srednica rurki	Material ceramiczny	Zródło
	pF	MHz	mm	mm	mm		
1	10	600	2,5				14
2	31	125					15
3	50	250	2,5				14
4	"	"	2×5	20—60	3—6		13
5	"	210	12				42
6	61	87					15
7	82	138	12				42
8	100	180	2,5				14
9	"	160	2×5	20—60	3—6		13
10	"	118	12				42
11	"	70					15
12	150	110	2×5	20—60	3—6		13
13	200	100	"	"	"		"
14	220	61	12				42
15	250	92	2×5	20—60	3—6		13
16	284	40					15
17	302	96	2×2	20	4	Condensa	24
18	350	71	2×5	20—60	3—6		13
19	500	120	2,5			„high ε”	14
20	"	75	"			„normal ε”	"
21	"	66	2×5	20—60	3—6		13
22	"	53	"	"	"		"
23	750	55	2,5			„high ε”	14
24	"	50	12				42
25	"	40	2,5			„normal ε”	14
26	"	46	2×5	20—60	3—6		13
27	1008	41,8	2×2	40	8	Condensa	24
28	1030	19,2					15
29	1450	52,4	2×2	7,5	4	Ultracond	24
30	1620	46,2	"	8,5	"	"	"
31	4910	22,6	"	20	"	Epsilan	"
32	4950	25,8	"	12	"	Ultracond	"
33	9200	18,5	"	17	"	"	"
34	10000	7,2	"				15
35	22800	9,7	2×2	30	4	Epsilan	24
36	małe, typu KTK1, KTK2	od 150 do 200					46 55
37	duże, typu KTK1, KTK2	od 50 do 70					46 55

534.4:621-39

JANUSZ KACPROWSKI

Zastosowanie analizy i syntezy mowy w telekomunikacji i automatyce

Rękopis dostarczono 30.9.1964

Praca zawiera krótki przegląd perspektywiczny nowych, lecz częściowo już obecnie realizowanych w skali laboratoryjnej, zastosowań sygnału mowy w telekomunikacji i automatyce, przy przekazywaniu informacji lingwistycznych w układach: człowiek — człowiek oraz człowiek — maszyna i maszyna — człowiek.

Omówiono trzy podstawowe etapy złożonego procesu automatyzacji mowy, a mianowicie: (a) analizę, (b) kodowanie i (c) syntezę dźwięków mowy, zwracając przede wszystkim uwagę na efektywność stosowanych metod z punktu widzenia sprawności przekazywania informacji, a następnie krótko scharakteryzowano kierunki docelowe prac naukowo-badawczych w tej dziedzinie, wraz z ich przyszłościowymi zastosowaniami technicznymi.

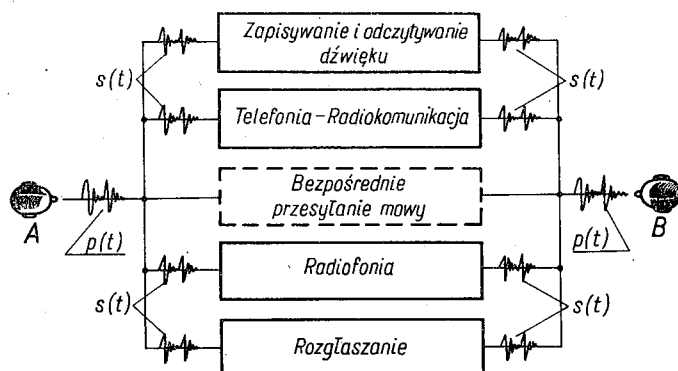
W zakończeniu przedstawiono obecny stan prac badawczych w zakresie analizy i syntezy mowy w Zakładzie Badania Drgan IPPT PAN.

1. WPROWADZENIE

Najprostszym, a jednocześnie najbardziej rozpowszechnionym sposobem porozumiewania się ludzi między sobą jest przekazywanie wiadomości lingwistycznych za pomocą dźwięków mowy. Nośnikiem informacji jest w tym przypadku sygnał akustyczny w postaci fali ciśnienia akustycznego $p(t)$. Przy przekazywaniu wiadomości konwencjonalnym systemem telekomunikacyjnym sygnał akustyczny $p(t)$ zostaje zakodowany w postaci sygnału elektrycznego $s(t)$, który jest jego mniej lub bardziej dokładnym odpowiednikiem. Na rys. 1 pokazano schemat przekazywania wiadomości lingwistycznych systemem konwencjonalnym. Jak wiadomo, system ten jest wysoce nieekonomiczny z punktu widzenia sprawności przekazywania informacji, gdyż wymaga stosowania kanałów telekomunikacyjnych o przepustowości informacyjnej rzędu 50 000 bitów na sekundę (bit/s), podczas gdy prędkość nadawania tych informacji przez źródło informacji (organ mowy) jest rzędu 50 bit/s. Wynika stąd, iż nadmiar informacji w naturalnym sygnale mowy jest rzędu 99,9%, co zresztą bywa niekiedy zjawiskiem pożądanym, gdyż stwarza szerokie możliwości techniczne częstotliwościowej, amplitudowej i czasowej kompresji lub ograniczania sygnału mowy przy zachowaniu dostatecznej zrozumiałości

przekazywanych wiadomości, a jednocześnie czyni sygnał mowy odpornym na wpływ zakłóceń.

W ciągu ostatniego dziesięciolecia daje się zaobserwować w skali światowej ogromnie intensywny i szybki rozwój prac badawczych nad problemem *automatyzacji mowy*. Problem ten, ogólnie biorąc, obejmuje



Rys. 1. Współczesne sposoby przekazywania informacji lingwistycznych za pomocą sygnału mowy

A — nadawca wiadomości, B — odbiorca wiadomości, $p(t)$ — akustyczny sygnał mowy, $s(t)$ — elektryczny sygnał mowy

trzy procesy: (a) analizę dźwięków mowy, (b) przetwarzanie ich w odpowiednio ekonomiczny sygnał kodowany, oraz (c) syntezę dźwięków mowy z elementów zastosowanego kodu, co stwarza ogromne możliwości techniczne usprawnienia przekazywania wiadomości lingwistycznych w układzie informacyjnym nadawczo-odbiorczym człowiek — człowiek. Jednocześnie otwierają się nowe możliwości techniczne zastosowania sygnału mowy do przekazywania informacji lingwistycznych w układzie: człowiek — maszyna lub maszyna — człowiek.

2. ANALIZA MOWY

Punktem wyjścia złożonego procesu automatyzacji mowy jest proces analizy dźwięków mowy, który w przyjętym tu wąskim znaczeniu tego pojęcia polega na automatycznym wydzieleniu z sygnału $p(t) \equiv s(t)$ pewnych parametrów charakterystycznych, zawierających w sobie cechy dystynktywne, na podstawie których poszczególne elementy fonetyczne, tworzące ciągły sygnał mowy, mogą być w sposób jednoznaczny zidentyfikowane i rozpoznane. Pod pojęciem elementów fonetycznych należy tu rozumieć zarówno elementy najniższego rzędu (fonemy, głoski), jak grupy fonemów pozbawione określonego znaczenia semantycznego (logatomy, sylaby), czy wreszcie elementy fonetyczne o określonym znaczeniu semantycznym (wyrazy) i treściowo-logicznym (frazy, zdania).

Techniczny proces syntezy mowy może być procesem ciągłym lub dyskretnym. *Analiza ciągła* polega na wydzieleniu z pierwotnego sygnału mowy $s(t)$ pewnego zespołu parametrów charakterystycznych (dystynktywnych), z których każdy jest wolnozmienną funkcją czasu, o prędkości zmian odpowiadającej prędkości zmian kształtu i rozmiarów fizjologicznego układu naturalnego organu mowy. Charakter fizyczny tych parametrów może być różny; mogą one opisywać mechaniczno-akustyczną strukturę wnętrza rezonansowych organu mowy przy wymawianiu poszczególnych głosek (fonemów), mogą dotyczyć widmowo-częstotliwościowych cech sygnału mowy, opisując rozkład gęstości energii akustycznej w widmie sygnału, jak to ma np. miejsce w przypadku klasycznego wokodera kanałowego Dudleya [1], mogą one opisywać kształt obwiedni widma za pomocą zespołu współczynników szeregu Fouriera, na który rozłożona została obwiednia (wokodery harmoniczne), czy wreszcie — wykorzystywać związki korelacyjne w sygnale mowy (wokodery korelacyjne). Szerokie możliwości techniczne stworzyła ostatnio analiza formantowa dźwięków mowy, oparta na ekstrakcji parametrów, określających formantowo-antyformantową strukturę sygnału mowy i opisanych matematycznie za pomocą rozkładu biegunów i zer funkcji przenoszenia (transmitancji) naturalnego organu mowy na płaszczyźnie częstotliwości zespolonej $\sigma \pm j\omega$, [2]. Podane wyżej przykłady analizy mowy na płaszczyźnie akustycznej można uzupełnić przykładami analizy na płaszczyźnie fonetycznej, do której zaliczyć należy proponowaną metodę określania cech dystynktywnych fonemów za pomocą odpowiednio dobranej grupy ich przeciwstawnych binarnych cech fonetycznych (Jakobson, Fant, Halle) [3].

Wspólną cechą wszystkich metod analizy ciągłej, bez względu na ich charakter i założenia teoretyczne, jest to, iż w wyniku otrzymuje się zespół n sygnałów parametrycznych, z których każdy jest wolnozmienną funkcją czasu $s_i(t)$ $_{i=1}^{i=n}$. Sygnały te mogą być przesłane dowolnym kanałem telekomunikacyjnym o małej przepustowości informacyjnej do syntetyzatora albo bezpośrednio, albo — po poddaniu ich procesom kwantyzacji, kodowania, szyfrowania itp. O całkowitej objętości informacyjnej zespołu sygnałów parametrycznych decyduje zarówno ich charakter fonetyczno-akustyczny, jak i ich ilość. Objętość informacyjna jest np. rzędu $2000 \div 4000$ bit/sek w przypadku wokoderów kanałowych, $500 \div 1000$ bit/sek przy wokoderach formantowych i rzędu $150 \div 300$ bit/sek przy systemach analizy fonetycznej, co pozwala na 2 do 10-krotne zawężenie pasma częstotliwości przesyłanych.

Jeszcze bardziej efektywne, z punktu widzenia kompresji sygnału mowy, są metody *analizy dyskretniej*, która może odbywać się na płaszczyźnie fonemów, sylab, wyrazów lub nawet zdań. Ogólnie biorąc im wyższy jest rząd elementów fonetycznych podlegających analizie dyskretniej, tym większa staje się liczebność zbioru pamięciowego układu rozpoznającego

i identyfikującego te elementy. Przy analizie fonemowej liczba elementów jest rzędu 40, przy sylabowej wzrasta co najmniej do około 100, a przy wyrazowej i zdaniowej zależy oczywiście od liczebności słownika, jaki bierze się pod uwagę przy eksploatacji systemu. Analiza fonemowa teoretycznie pozwala na obniżenie przepustowości informacyjnej kanału telekomunikacyjnego do wartości stosowanych w łączności telegraficznej (rzędu 60 bit/s) lub nawet mniej, jeżeli uwzględni się związki statystyczne języka, określające prawdopodobieństwo warunkowe następowania po sobie różnych fonemów. Przy analizie wyrazowej i liczebności słownika 32 wyrazów, co jest zresztą absolutnie niewystarczające do celów konwencjonalnej łączności telefonicznej, np. abonenckiej, przepustowość kanału może być rzędu 10 bit/s, ale wzrasta gwałtownie przy zwiększaniu słownika. Wydaje się, że najlepszym kompromisem między pojemnością układu pamięciowego a niezbędną przepustowością kanału jest system dyskretnej analizy sylabowej, (około 25 bit/s przy normalnej prędkości mówienia i dość ograniczonej, lecz praktycznie wystarczającej liczebności zestawu sylab).

3. KODOWANIE SYGNAŁU MOWY

Następnym podstawowym procesem automatyzacji mowy jest *kodowanie*, tj. przedstawienie zespołu parametrów dystynktywnych, otrzymanych w wyniku procesu analizy mowy, w postaci zespołu elementów odpowiedniego kodu, którego rodzaj i charakter zależą w dużym stopniu od najmniejszego elementu fonetycznego, poddanego procesowi analizy. Z technicznego punktu widzenia najprostsze byłoby kodowanie fonemów ze względu na ich ograniczoną ilość w danym języku. Okazuje się jednak, że charakter fonetyczny fonemu w dużej mierze zależy od wpływu fonemów sąsiednich, poprzedzających go i następujących po nim w kontekście sylab i wyrazów, co bardzo utrudnia i komplikuje proces kodowania fonemów. Znacznie wygodniejsze wydaje się być kodowanie zespołów fonemów, np. sylab, kiedy wpływ wzajemny elementów sąsiednich jest znacznie mniejszy, a operowanie zestawem 1000 ÷ 2000 elementów daje wyniki wystarczające do celów telefonii porozumiewawczej (handlowej), przy czym — do celów łączności specjalnej — może być ograniczone nawet do 100 ÷ 200 elementów. O wyborze rodzaju kodu decyduje ponadto charakter fonetyczny i akustyczny oraz ilość parametrów dystynktywnych, określających fonemy i sylaby. Stosując na przykład analizę pasmową rozkładu energii w widmie mowy, można ograniczyć się do stosunkowo niewielkiej ilości kanałów analizujących, rzędu 8, co w rezultacie prowadzi do 8-znakowego kodu binarnego, za pomocą którego można dostatecznie dokładnie opisać chwilowy kształt obwiedni widma sygnału oraz jego zmiany w czasie wymawiania sylaby. Stosując pięciostopniową kwantyzację czasu trwania jednej sylaby można

za pomocą macierzy o 8 rzędach i 5 kolumnach, czyli o 40 elementach, teoretycznie bez większych trudności technicznych, stworzyć układ pamięciowy o pojemności około 100 elementów fonetycznych (syllab), co w wielu przypadkach wystarcza do celów praktycznych.

Podany przykład kodu jest oczywiście przykładem najprostszym; w praktyce okazuje się często niezbędne rozbudowanie kodu wskutek uwzględnienia dodatkowych elementów informacyjnych sygnału mowy, otrzymywanych na przykład w wyniku akustycznej analizy mikrostruktury widma.

Ogólnie biorąc, proces kodowania dystynktywnych parametrów dźwięków mowy, mimo swojej specyfiki i konieczności uwzględnienia takich czynników, jak np. efektywność kodu (czyli: poszukiwanie kodu optymalnego dla danego języka), możliwość łatwego i skutecznego szyfrowania, niewrażliwość na wpływ zakłóceń i szumów itd., nie różni się wiele od problemu kodowania rozpatrywanego na płaszczyźnie teletransmisji i ogólnej teorii telekomunikacji.

4. SYNTEZA MOWY

Ostatnim wreszcie procesem automatyzacji mowy jest jej *synteza* z elementów kodu, tj. wytwarzanie w układzie technicznym przebiegów dźwiękowych, którym człowiek o normalnym słuchu potrafi przypisać znaczenie określonych elementów lub grup elementów fonetycznych.

Istnieje już obecnie wiele metod i systemów syntezy mowy, zarówno ciągłej, jak i dyskretniej. Metody *syntezy ciągłej* polegają na sterowaniu ciągłymi sygnałami parametrycznymi parametrów mechanicznych, akustycznych lub elektrycznych odpowiedniego układu, który wytwarza określone i kontrolowane przebiegi dźwiękowe. Układem takim może być np. mechaniczno-akustyczny model naturalnego organu mowy, elektryczny układ zastępczy organu o stałych $R(t)$, $L(t)$, $G(t)$, $C(t)$ rozłożonych w sposób ciągły i przestrajanych sygnałami parametrycznymi (tzw. układy analogowe, odpowiednio — mechaniczne, akustyczne lub elektryczne) [4], elektryczny układ zastępczy o stałych skupionych, który odtwarza rezonansowe właściwości wnęk naturalnego organu mowy (tzw. syntetyzator formantowy stosowany w wokoderach formantowych) [5], układ widmowo-pasmowy stosowany w wokoderach kanałowych [1] i wiele innych. Cechą wspólną wszystkich tych układów, bez względu na ich charakter, strukturę i zasadę działania, jest jednak to, że wytwarzanie założonych przebiegów dźwiękowych odbywa się w lokalnych, umieszczonych w syntezatorze źródłach dźwięku, które są jedynie sterowane bądź sygnałami parametrycznymi, wytworzonymi w analizatorze przez ich ekstrakcję z naturalnego sygnału mowy, bądź sygnałami parametrycznymi, odpowiednio z góry zaprogramowanymi.

Wspólną cechą wszystkich metod *dyskretnej syntezy* mowy, różniącą je od metod syntezy ciągłej, jest to, iż wymagają one stosowania w syntetyzatorze układu pamięciowego, zawierającego odpowiednio zarejestrowany, ściśle określony i skończony zbiór elementów fonetycznych odpowiedniego rzędu (fonemów, sylab, wyrazów zdań lub tp.). Sygnał kodowany, przychodzący z analizatora oraz układu kodującego, kolejno wyzwala zarejestrowane w układzie pamięciowym elementy fonetyczne, których sekwens składa się na przekazywaną wiadomość, odtwarzaną następnie przez przetwornik elektroakustyczny (źródło dźwięku) w postaci sygnału akustycznego $p(t)$ fali dźwiękowej.

Pod względem sprawności przekazywania informacji metody syntezy dyskretnej przewyższają oczywiście metody parametryczne (ciągłe), przy czym stosowanie dyskretnej syntezy sylabowej wydaje się być, podobnie jak w przypadkach analizy, rozsądnym kompromisem między wymaganą przepustowością informacyjną systemu transmisyjnego a niezbędną pojemnością układu pamięciowego i osiąganą zrozumiałością mowy syntetycznej.

Należy oczywiście zdawać sobie sprawę z tego, że w przypadku metod syntezy ciągłej z sygnałów parametrycznych, przekazywanie *informacji personalnych*, tj. dotyczących cech indywidualnych głosu nadawcy wiadomości i umożliwiających jego identyfikację, jest teoretycznie zupełnie możliwe, jakkolwiek technicznie dość trudne. Natomiast metody syntezy dyskretnej, operujące tworzeniem wiadomości z elementów fonetycznych uprzednio zarejestrowanych w układzie pamięciowym i sprowadzające się jedynie do ich właściwego wyboru przez elementy kodu, z natury rzeczy możliwość tę wykluczają, i ograniczają się do przekazywania wyłącznie *informacji lingwistycznych* dotyczących treści przekazywanej wiadomości. W związku z tym jako kryterium jakości transmisji mowy syntetycznej przy metodach dyskretnych może być przyjęta jedynie zrozumiałość mowy, umożliwiająca pojmowanie treści przesyłanej wiadomości.

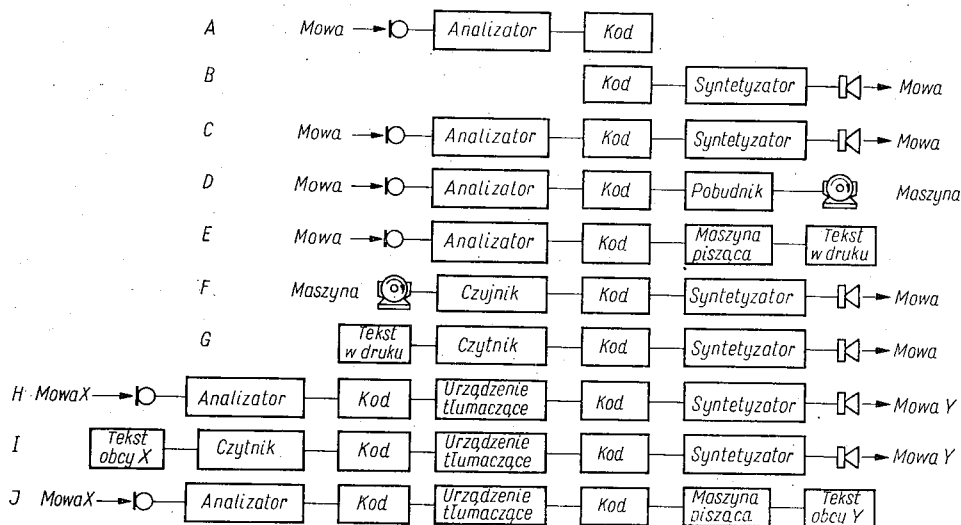
5. BIEŻĄCE I PERSPEKTYWICZNE ZASTOSOWANIA PROCESU AUTOMATYZACJI MOWY

Omówione poprzednio podstawowe procesy automatyzacji mowy, obejmujące jej: (a) analizę, (b) kodowanie i (c) syntezę, są przedstawione schematycznie w układach blokowych na rys. 2-A i 2-B.

Następne schematy blokowe, od 2-C do 2-J, przedstawiają bądź już realizowane, bądź teoretycznie możliwe i stanowiące przedmiot docelowy perspektywicznych prac naukowo-badawczych w tej dziedzinie możliwości zastosowań procesu automatyzacji mowy w telekomunikacji i automatyce, wytyczające kierunki ich rozwoju przyszłościowego [13].

Schemat na rys. 2-C stanowi kaskadowe połączenie schematów 2-A i 2-B i przedstawia stosunkowo najprostsze, technicznie najbardziej real-

ne, a obecnie w pewnym stopniu już eksploatowane zastosowanie procesu automatyzacji mowy w telekomunikacji fonicznej. Schemat ten symbolizuje kompletny, nadawczo-odbiorczy telefoniczny układ transmisyjny o zmniejszonej przepustowości informacyjnej, umożliwiającą posługiwanie się kanałem łączności o wąskim pasmie częstotliwości i ograniczonym zakresie dynamiki sygnału. Wydaje się, że osiągnęte obecnie w warunkach



Rys. 2. Perspektywiczne możliwości zastosowań sygnału mowy w telekomunikacji i automatyce do przekazywania informacji lingwistycznych w układach: człowiek — człowiek, człowiek — maszyna i maszyna — człowiek

produkcyjnych i eksploatacyjnych przepustowości rzędu $2000 \div 1000$ bitów na sekundę (np. wokodery kanałowe i formantowe) będą mogły być w przyszłości obniżone nawet do rzędu kilkudziesięciu bit/s w systemach dyskretnej analizy-syntezy, i to przy normalnej prędkości mówienia.

Schematy na rys. 2-D i 2-E przedstawiają podstawowe możliwości zastosowań procesu automatyzacji mowy przy przekazywaniu informacji lingwistycznych w układzie: człowiek — maszyna. Schemat 2-D dotyczy sterowania układów mechanicznych głosem, kiedy ograniczenie posługiwania się przez obsługę rękami i nogami jest niezbędne lub przynajmniej pożądane. Schemat 2-E jest szczególnym przypadkiem schematu 2-D, kiedy układem mechanicznym sterowanym głosem jest maszyna do pisania pod dyktando.

Schematy 2-F i 2-G są odwróconymi w swoim działaniu układami 2-D i 2-E, i służą do przekazywania informacji lingwistycznych w układzie: maszyna — człowiek. Układ 2-F polega na przekazywaniu za pomocą dźwięków mowy informacji o aktualnym stanie funkcjonalnym układu mechanicznego, kiedy ograniczenie posługiwania się przez operatora zmysłem wzroku lub dotyku jest niezbędne lub co najmniej pożądane. Układ

2-G jest układem, który przetwarza tekst drukowany w tekst wymawiany głosem.

Należy dodać, że układy funkcjonalne przedstawione schematycznie na rys. 2-D, 2-E, 2-F i 2-G, są przedmiotem aktualnie prowadzonych, intensywnych badań w skali laboratoryjnej. Systemy przekazywania informacji lingwistycznych w układzie: człowiek — maszyna są najbardziej realne technicznie; próby pracy maszyn piszących pod dyktando (Olson w U.S.A [6], Dreyfus-Graf w Szwajcarii [7]) dają wyniki zadowalające. Wreszcie schematy 2-H, 2-I i 2-J przedstawiają układy, których opracowanie i rozwiązanie jest najtrudniejsze zarówno z technicznego, jak i teoretycznego punktu widzenia, gdyż wymaga uprzedniego rozwiązania szeregu problemów z dziedziny logiki matematycznej oraz analizy statystycznej, porównawczej fonetyki i porównawczej gramatyki i składni różnych języków.

Schemat 2-H przedstawia kompletny układ do bezpośredniego tłumaczenia tekstu słownego z języka X na język Y, a schematy 2-I oraz 2-J są jego modyfikacjami w postaci układów do natychmiastowego tłumaczenia tekstu drukowanego w języku X na tekst wymawiany w języku Y (rys. 2-I, porównaj rys. 2-G) lub odwrotnie — do tłumaczenia tekstu wymawianego w języku X na tekst drukowany w języku Y (rys. 2-J, porównaj rys. 2-E).

6. STAN BIEŻĄCY ORAZ PERSPEKTYWY ROZWOJOWE PRAC BADAWCZYCH NAD AUTOMATYZACJĄ MOWY

Podany w wielkim skrócie przegląd zastosowań procesów automatyzacji mowy w telekomunikacji i automatyce należy traktować wyłącznie jako docelowy i perspektywiczny plan badań naukowych w tej dziedzinie, który obecnie znajduje się w stanie zaawansowania jeszcze dalekim od realizacji technicznej nawet w skali laboratoryjnej, i to w przodujących ośrodkach badawczych w świecie. Obecnie prowadzone są głównie prace podstawowe w zakresie poznawania praw i zjawisk związanych z procesami analizy i syntezy mowy (schematy na rys. 2-A, 2-B i 2-C), od których rozwiązania uzależnia się bezpośrednio prace układowo-techniczne w zakresie zastosowań w telekomunikacji i automatyce (schematy rys. 2-D ÷ ÷ 2-J). To stosunkowo słabe zaawansowanie (w skali całego problemu automatyzacji mowy) prac techniczno-badawczych wynika zarówno ze skomplikowanego charakteru procesów fizjologiczno-fizycznych związanych z wytwarzaniem i percepcją dźwięków mowy, jak i z kompleksowości zagadnień, które można rozwiązać jedynie przy ścisłej współpracy wielu specjalistów z różnych dziedzin: akustyki mowy i słuchu, fonetyki stosowanej, elektroakustyki, teorii układów elektronicznych, teletransmisji i ogólnej teorii informacji. W taki właśnie kompleksowy sposób prowadzone są badania w zakresie automatyzacji mowy w krajach przodu-

jących w tej dziedzinie: Stanach Zjednoczonych A.P., Szwecji, ZSRR, Wielkiej Brytanii.

W tym stanie rzeczy prace badawcze nad automatyzacją mowy prowadzone w Polsce od mniej więcej 3 lat, głównie w Zakładzie Badania Drgań I.P.P.T. — PAN, nie mogłyby wyjść poza zakres prac podstawowych w dziedzinie analizy i syntezy mowy. Prace prowadzone są równolegle na dwóch płaszczyznach: fonetyczno-akustycznej (Pracownia Fonetyki Akustycznej) i elektroakustyczno-układowej (Pracownia Elektroakustyki).

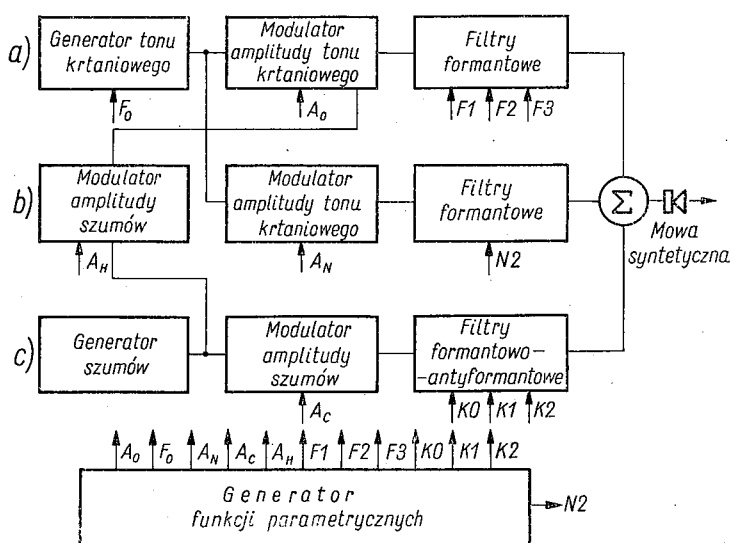
Prace fonetyczno-akustyczne dotyczą statystycznej i strukturalnej analizy dźwięków mowy polskiej w stanie ustalonym i nieustalonym. Bezpośrednim celem tych prac jest określenie akustycznych cech dystynktywnych dźwięków mowy polskiej, które w sposób optymalny mogłyby być wykorzystane przy automatycznym rozpoznawaniu i identyfikacji elementów fonetycznych (fonemów, sylab) języka polskiego. Swymi bezpośrednimi zastosowaniami prace te zmierzają do technicznej realizacji układów do automatycznej ekstrakcji z naturalnego sygnału mowy parametrów akustycznych dystynktywnych, które, odpowiednio zakodowane, mogłyby być przesłane kanałem telekomunikacyjnym o minimalnej pojemności informacyjnej. W obecnym stanie zaawansowania prac przedmiotem badań są objęte przede wszystkim parametry określające formantowo-antyformantową strukturę sygnału mowy.

Jednocześnie prace te stwarzają podstawy fonetyczne prac elektroakustyczno-układowych nad syntezą dźwięków mowy w rezonansowych układach formantowych o stałych skupionych. W oparciu o opracowaną teorię syntezy samogłosek [8] zaprojektowano; zbudowano i zbadano doświadczalny syntetyzator formantowy samogłosek i dwugłosek (dyftongów) polskich, sterowany pięcioma sygnałami parametrycznymi o kształcie trapezowym, programowanymi i wytwarzanymi w prowizorycznych generatorach funkcji parametrycznych [9], [10]. Zbadano analitycznie formantowo-antyformantową strukturę polskich spółgłosek nosowych [11] i na tej podstawie zaprojektowano i zbudowano kanał nosowy syntetyzatora formantowego [12], który współpracuje z opracowanym poprzednio kanałem samogłoskowym i pozwala na syntezę spółgłosek nosowych oraz sylab w zestawie V-NC-V oraz NC-V gdzie V oznacza samogłoskę, a NC — spółgłoskę nosową. Prowadzone obecnie prace nad syntezą spółgłosek zmierzają do opracowania modelu uniwersalnego syntetyzatora formantowego SYNFOR I, który sterowany 12 sygnałami parametrycznymi, programowanymi przez ulepszony generator ciągłych funkcji parametrycznych pozwalałby na syntezę słów i tekstów ciągłych (zdań).

W tej postaci syntetyzator ten będzie wykorzystany do badań fonetycznych, tzn. metodą „analizy przez syntezę”, a jednocześnie stanowić będzie prototyp laboratoryjny części odbiorczej (tj. układu syntezy) wokodera formantowego o małej pojemności informacyjnej kanału łączności.

Uproszczony funkcjonalny schemat blokowy syntetyzatora SYNFOR I podano na rys. 3.

W miarę postępu prac nad analizą mowy, a w szczególności nad ekstrakcją parametrycznych funkcji formantowych z naturalnego sygnału mowy, przewidziana jest współpraca syntetyzatora SYNFOR I z analizatorem formantowym, co prowadzi do stworzenia w skali laboratoryjnej kompletnego nadawczo-odbiorczego układu transmisyjnego o małej prze-



Rys. 3. Uproszczony schemat blokowy parametrycznego syntetyzatora formantowego mowy SYNFOR I

A — tor syntezy samogłosek i spółgłosek płynnych, B — tor syntezy spółgłosek nosowych, C — tor syntezy pozostałych spółgłosek

pułtowości informacyjnej, pracującego na zasadzie ciągłej analizy-syntezy formantowej.

Jednocześnie są podejmowane prace wstępne nad automatycznym rozpoznawaniem elementów fonetycznych (fonemów, w dalszym etapie — sylab i wyrazów) języka polskiego. Badania zmierzają będą perspektywnie do opracowywania prostych systemów przekazywania informacji w układzie: człowiek — maszyna.

*Instytut Podstawowych Problemów Techniki PAN
Zakład Badania Drgan*

WYKAZ LITERATURY

1. Dudley H. W.: The Vocoder. — Bell Labs. Rec. t. 18 (1936), s. 122—126.
2. S a p o ż k o w M. A.: Rieczewoj signal w kibiernetikke i swjazi. — Swiezizdat, Moskwa 1963.

3. Jakobson R., Fant C. G., Halle M.: Preliminaries to Speech Analysis. The Distinctive Features and Their Correlates. — Acoustics Laboratory, Mass. Inst. of Technol., Techn. Rep. Nr 13, 1952.
4. Rosen G.: Dynamic Analog Speech Synthesizer. — JASA, t. 30 (1958), s. 201—203.
5. Weibel E. S.: Vowel Synthesis by Means of Resonant Circuits. — JASA, t. 27 (1955), s. 858—865.
6. Olson H. F., Belar H.: Phonetic Typewriter III. — JASA, t. 33 (1961), s. 1610—1615.
7. Dreyfus, Graf J.: Phonétographs: Présent et Future. — Bull. Techn. PTT, Nr 5, 1961, s. 160—170.
8. Kacprowski J.: Teoretyczne podstawy syntezy samogłosek polskich w rezonansowych układach formantowych. — Rozprawy Elektrotechniczne, t. VIII (1962), Nr 1, s. 127—203.
9. Kacprowski J., Mikiel W.: Preliminary Synthesis of Polish Vowels by Means of Recurrently Impulsed Formant Filters. — Proc. Vibr. Problems, 1, 4, 1963, s. 27—41.
10. Kacprowski J.: Preliminary Synthesis of Polish Vowels by Means of „Terminal — Analog” Speech Synthesizer. — Materiały IV. Międzynarodowego Kongresu Akustyki, Kopenhaga 1962, referat G 55.
11. Kacprowski J.: An Approach to the Synthesis of Polish Nasal Consonants by Means of „Terminal — Analog” Speech Synthesizer. — Proc. Vibr. Problems, 3, 4, 1963, s. 235—254.
12. Kacprowski J.: Synteza polskich spółgłosek nosowych w rezonansowych układach formantowych. — Rozprawy Elektrotechniczne, t. IX (1963), Nr 3, s. 439—465.
13. Olson H. F.: Processing of Sound. — Proc. JRE, t. 50 (1962), Nr 5, s. 599—600.

J. KACPROWSKI

THE APPLICATION OF SPEECH ANALYSIS AND SYNTHESIS IN TELECOMMUNICATION AND AUTOMATION

Summary

This paper contains a short perspective survey of investigated at present new applications of speech analysis and synthesis methods in telecommunication and automation. These applications concern the conveying of linguistic information in the following transmission systems: man-to-man, man-to-machine and machine-to-man.

The three fundamental stages of the complex process of speech automation are discussed, namely: the speech analysis, the conversion of speech signal to a code and the speech synthesis form a code. Special attention is paid to the problem of speech signal compression and transmission of speech at low information rate. The new trends of research work in the field of speech automation are described, together with their future applications.

Finally, the present research work in the field of speech analysis and synthesis, being carried on at the Department for Vibration Research of the Institute of

Fundamental Technical Problems of the Polish Academy of Sciences, is shortly reported.

J. KACPROWSKI

LES APPLICATIONS DE L'ANALYSE ET DE LA SYNTHÈSE DE LA PAROLE DANS LE DOMAINE DES TELECOMMUNICATIONS ET DE L'AUTOMATISATION

Résumé

L'oeuvre contient une courte revue perspective de nouvelles et pourtant déjà réelles applications techniques du signal de la parole humaine dans le domaine des télécommunications et de l'automatisation. Ces applications concernent surtout la transmission des informations linguistiques dans les systèmes suivants: être humain — être humain, être humain — machine et machine — être humain.

On a discuté les trois étapes fondamentales d'un procès complexe d'automatisation de la parole, à savoir: l'analyse et la synthèse des sons de la parole humaine ainsi que la transformations de leurs paramètres phonétiques et acoustiques dans des signaux d'un code effectif. On a accentué le problème d'efficacité des méthodes applicables au point de vue de l'économie de transmission des informations linguistiques. Puis on a caractérisé en bref le développement des recherches scientifiques dans ce domaine ainsi que leurs applications techniques à l'avenir.

En terminant on a présenté les résultats des recherches scientifiques dans le domaine de l'analyse et de la synthèse de la parole, réalisées actuellement dans le Département des Recherches des Vibrations de l'Institut des Problèmes Techniques Fondamentaux de l'Académie Polonaise des Sciences.

J. KACPROWSKI

DIE ANWENDUNGEN DER ANALYSE UND SYNTHESE DER SPRACHE IN FERNMELDETECHNIK UND AUTOMATIK

Zusammenfassung

Die Arbeit enthält eine kurze perspektivische Übersicht der neuen, doch teilweise schon im laboratorischen Massstabe verwirklichten technischen Anwendungen der Analyse und Synthese des Sprachsignals in Fernmeldetechnik und Automatik. Diese Anwendungen betreffen im Allgemeinen die Übertragung der linguistischen Informationen in Übertragungssystemen: Mensch — Mensch, Mensch — Maschine und Maschine — Mensch.

Die drei Hauptetappen des komplexen Prozesses der Automatisierung der Sprache, nämlich die Analyse, die Kodierung und die Synthese werden kurz besprochen. Dabei zieht man in Betracht die Wirksamkeit der angewandten Methoden bezüglich der Kompression des Informationsinhaltes des Sprachsignals. Ferner werden auch die Richtungen der wissenschaftlichen Forschungsarbeiten auf diesem Gebiete kurz charakterisiert.

Zum Schluss bespricht man den aktuellen Stand der Forschungsarbeiten auf dem Gebiete der Analyse und Synthese der Sprache in der Abteilung für Schwingungsvorschung des Institutes für Grundprobleme der Technik der Polnischen Akademie der Wissenschaften.

Я. КАЦПРОВСКИ

ПРИМЕНЕНИЕ АНАЛИЗА И СИНТЕЗА РЕЧИ В ЭЛЕКТРОСВЯЗИ
И В АВТОМАТИКЕ

Резюме

Работа содержит краткий перспективный обзор новых, но уже частично теперь осуществляемых в лабораторном масштабе, применений сигнала речи в области электросвязи и автоматики при передаче лингвистических информации по схемам: человек — человек, а также человек — машина и машина — человек.

Описано три основные этапы сложного процесса автоматизации речи, а именно: (а) анализ, (б) кодирование и (в) синтез звуков речи, обращая прежде всего внимание на эффективность применяемых методов с точки зрения исправности передачи информации, а далее кратко представлено направления ведения научно-исследовательских работ в этой области совместно с их будущими техническими применениями.

В окончании представлено современное состояние исследовательских работ по анализу и синтезу речи в Отделе исследования колебаний Института основных проблем техники Польской Академии наук.

R. G. Strużak: Resonances in Ceramic Tube Capacitors	474
Résonances dans condensateurs tubulaires céramiques	475
J. Kacprowski: The Application of Speech Analysis and Synthesis in Telecommunication and Automation	489
Les applications de l'analyse et de la synthèse de la parole dans le do- maine des télécommunications et de l'automatisation	490
Die Anwendungen der Analyse und Synthese der Sprache in Fernmelde- technik und Automatik	490

СОДЕРЖАНИЕ

К. Биштыга: Динамическое торможение ас. двигателя как крайний случай частотного управления	300
М. Домбровски: Электромагнитное поле в анизотропном прямоуголь- ном стержне	318
З. Пэрковски: Измерение первичных параметров цепей проводной связи непосредственным методом	350
В. Жоховски: Продольные колебания пластинки из пьезоэлектриче- ского материала	379
В. Т. Беньковска: Магнитный переключатель каналов для телефон- ного коммутатора с разделением во времени	404
К. Грабовски: Характеристики накачиваемого $p - n$ перехода	425
Р. Г. Стружак: Резонансы в керамических трубчатых конденсаторах	476
Я. Кацпровски: Применение анализа и синтеза речи в электросвязи и в автоматике	491

Rozprawy Elektrotechniczne

Kwartalnik

WARUNKI PRENUMERATY CZASOPISMA

Cena prenumeraty rocznej	zł 100,—
„ „ półrocznej	zł 50,—

Zamówienia i wpłaty przyjmują:

1. Centrala Kolportażu Prasy i Wydawnictw „Ruch”,
Warszawa, ul. Wronia 23, konto PKO nr 1-6-100.020.
2. Oddziały i Delegatury „Ruchu”.
3. Urzędy pocztowe i listonosze.
4. Księgarnie „Domu Książki”.

Zamówienia przyjmowane są do dnia 15 miesiąca poprzedzającego okres prenumeraty.

Zamówienia dla zagranicy przyjmuje Biuro Kolportażu Wydawnictw Zagranicznych „Ruch”, Warszawa, Wronia 23 (tel. 20-46-88), konto PKO nr 1-6-100.024. Koszt prenumeraty ze zleceniem wysyłki za granicę jest o 40% wyższy.

Bieżące oraz archiwalne numery można nabywać lub zamawiać w księgarniach „Domu Książki” oraz we Wzorcowni Wydawnictw Naukowych PAN — Ossolineum — PWN, Warszawa, Pałac Kultury i Nauki (wysoki parter).

Archiwalne egzemplarze można nabywać także w Punkcie Wysyłkowym Prasy Archiwalnej „Ruch” — Warszawa, ul. Srebrna 12, konto PKO nr 114-6-700041 VII O/M.



Tylko prenumerata zapewnia regularne otrzymywanie czasopisma.

INSTYTUT PODSTAWOWYCH PROBLEMÓW TECHNIKI
POLSKIEJ AKADEMII NAUK

ROZPRAWY ELEKTROTECHNICZNE

TOM XI · ZESZYT 3

KWARTALNIK

PAŃSTWOWE WYDAWNICTWO NAUKOWE
WARSZAWA 1965

TREŚĆ

S. Pogorzelski: Współczynnik zbieżności fali w technice antenowej . . .	495
M. Nałęcz: Dokładność i czułość przetwornika halotronowego przy pomiarach statycznych przemieszczeń mechanicznych . . .	521
M. Dąbrowski: Wpływ zniekształceń opóźnieniowych kanałów teletransmisyjnych na transmisję danych . . .	545
W. Kompało: Technologia oraz podstawowe parametry germanowych tranzystorów stopowych $n-p-n$. . .	571
I. Wójcik: Germanowy tranzystor mocy o wielkiej częstotliwości granicznej wykonany techniką stopowo-dyfuzyjną . . .	599
S. Piątek: Możliwości zastosowania aluminium na uzwojenie wirników turbogeneratorów przy różnych systemach chłodzenia stojana i wirnika . . .	643
Z. Kowalski: Zasada wzajemności oporności w układzie składowych symetrycznych . . .	653
S. Góra: Metoda wyboru struktury mocy elektrowni w systemie elektroenergetycznym z zastosowaniem matematycznych maszyn cyfrowych . . .	677

CONTENTS — TABLE DES MATIÈRES — INHALT

S. Pogorzelski: Wave Convergence Coefficient in Antenna Problems . . .	517
Coefficient de convergence d'onde dans la technique des antennes . . .	518
Wellenkonvergenzkoeffizient in Problemen der Antennentechnik . . .	519
M. Nałęcz: On Accuracy and Sensitivity of the Hall Generator Transducer in Static Mechanical Measurement . . .	541
Précision et sensibilité des convertisseurs de Hall en cas des mesures mécaniques statiques . . .	542
M. Dąbrowski: The Influence of Delay-Distortions in Telecommunication Channels on Data-Transmission Systems . . .	569
Influence des distortions de retard des canaux des télécommunications sur la transmission des données . . .	569
W. Kompało: Technology and Basic Parameters of N-P-N Germanium Alloy Transistors . . .	596
La technologie et les paramètres fondamentaux des transistors $n-p-n$ au germanium par alliage . . .	597
Technologie und Grundkenngrößen der $n-p-n$ Germanium-Legierungstransistoren . . .	597
I. Wójcik: High Frequency Alloy-Diffused Germanium Power Transistor Transistors de puissance au germanium pour fréquences élevées à jonctions par alliage-diffusion . . .	641
S. Piątek: Possibilities of Application of Aluminium for Rotor Windings of Turbogenerators Equipped with Various Cooling Systems for Stators and Rotors . . .	651
Possibilités d'application d'aluminium aux enroulements rotoriques des turboalternateurs pour différents systèmes de refroidissement du stator et du rotor . . .	651
Die Anwendungsmöglichkeit der Aluminiumwicklung bei Turbogeneratorenankern mit verschiedenen Kühlsystemen des Stators und Ankers . . .	652
Z. Kowalski: Mutuality Principle of Resistances in Symmetrical Components System . . .	674
Le principe de la réciprocité des résistances et réactances dans le système de composants symétriques . . .	675
Gegenseitigkeitsprinzip der Widerstände in symmetrischem Komponentensystem . . .	675
S. Góra: Method of Pattern Selection of Generating Plants in Power System with Use of Digital Computers . . .	733
La méthode du choix de la puissance des usines génératrices dans un réseau interconnecté à l'aide des calculateurs électroniques . . .	734
Methode zur Auswahl der Leistungskomposition von Kraftwerken im Verbundnetz mit Hilfe des Digitalrechners . . .	735

INSTYTUT PODSTAWOWYCH PROBLEMÓW TECHNIKI
POLSKIEJ AKADEMII NAUK

ROZPRAWY ELEKTROTECHNICZNE

TOM XI • ZESZYT 3

KWARTALNIK

PAŃSTWOWE WYDAWNICTWO NAUKOWE
WARSZAWA 1965

RADA REDAKCYJNA

PROF. JANUSZ GROSZKOWSKI, PROF. JANUSZ LECH JAKUBOWSKI,
PROF. BOLESŁAW KONORSKI, PROF. IGNACY MALECKI
PROF. PAWEŁ NOWACKI, PROF. PAWEŁ SZULKIN

KOMITET REDAKCYJNY

Redaktor Naczelny

PROF. WITOLD NOWICKI

Sekretarz

JULIUSZ MIERZEJEWSKI

ADRES REDAKCJI

*Warszawa, Politechnika, Plac Jedności Robotniczej 1,
Katedra Teletransmisji Przewodowej, Gmach Mechaniki*

PANSTWOWE WYDAWNICTWO NAUKOWE
Warszawa, Miodowa 10

Nakład 720 (590 + 130)	Oddano do składania 8.VI.1965 r.
Ark. wyd. 17,5, ark. druk. 15,5	Podpisano do druku we wrześniu 1965 r.
Papier druk sat. kl. III, 80 g, 70×100	Druk ukończono we wrześniu 1965 r.
Zamówienie nr 864/65	Cena zł 330.— E-48

Drukarnia im. Rewolucji Październikowej, Warszawa

621.396.677

SEWERYN POGORZELSKI

Współczynnik zbieżności fali w technice antenowej

Rękopis dostarczono 23.10.1964

Wyprowadzono zależności ogólne pozwalające na obliczanie współczynnika zbieżności fali odbitej od reflektora dla znanych typów anten reflektorowych. Podano przykłady obliczeń dla reflektora sferycznego oraz dla rozogniskowanej anteny z reflektorem parabolicznym.

1. WSTĘP

Współczynnik zbieżności fali elektromagnetycznej jest dobrze znanym pojęciem w dziedzinie optyki geometrycznej. Jest to wielkość opisująca zmianę amplitudy fali wzdłuż promienia. Niech s oznacza długość łuku na promieniu, $A(s)$ — amplitudę fali. Mamy wtedy zależność

$$A(s) = A(s_0)D(s), \quad (1)$$

gdzie $D(s)$ jest współczynnikiem zbieżności, zaś $A(s_0)$ oznacza początkową wartość amplitudy (dla $s = s_0$). Dowodzi się w optyce geometrycznej, że w izotropowym jednorodnym ośrodku

$$D(s) = \sqrt{\frac{K(s)}{K(s_0)}}. \quad (2)$$

W wyrażeniu tym $K(s)$ jest krzywizną Gaussa powierzchni równego eikonalu.

Współczynnik zbieżności posiada znaczenie w technice anten reflektorowych jak również w tych zagadnieniach propagacji, w których rozpatruje się odbicia fali od powierzchni zakrzywionych.

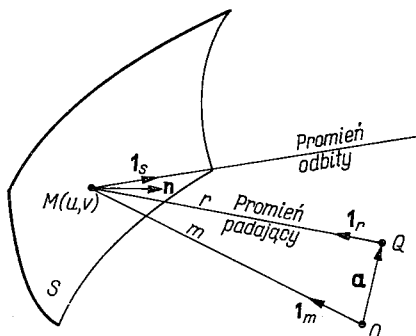
Praca niniejsza jest poświęcona zagadnieniu obliczania współczynnika zbieżności w technice anten reflektorowych. Wyprowadza się zależności ogólne pozwalające na obliczenie współczynnika D dla znanych typów anten reflektorowych. Otrzymane wzory można stosować również w przypadku anten dowolnie rozogniskowanych. Ułatwia to w znacznym stopniu konstrukcję ściśłych — w sensie optyki geometrycznej — wyrażeń dla pola w aperturze anteny. Rozważania teoretyczne ilustrowane są przykładami obliczeń dla reflektora sferycznego oraz dla rozogniskowanej anteny z reflektorem parabolicznym.

2. ZALEŻNOŚCI PODSTAWOWE DLA ANTEN REFLEKTOROWYCH

Wielkości geometryczne, z jakimi mamy do czynienia w analizie anteny reflektorowej, są pokazane na rys. 1.

Punkt O reprezentuje początek układu współrzędnych, punkt Q — źródło fali kulistej oświetlającej reflektor S^1 . Powierzchnia reflektora jest opisana równaniem

$$\mathbf{m} = \mathbf{m}(u, v) = m(u, v) \mathbf{l}_m. \quad (3)$$



Rys. 1. Szkic do analizy anteny reflektorowej

Można też w razie potrzeby stosować wektor pozycyjny $\mathbf{r}(u, v)$ traktując punkt Q jako początek układu współrzędnych. Wektory $\mathbf{m}$ oraz $\mathbf{r}$ spełniają zależność

$$\mathbf{m} = \mathbf{a} + \mathbf{r}, \quad (4)$$

gdzie $\mathbf{a} = \text{const}$ określa położenie źródła Q względem początku układu współrzędnych O .

Pokazane na rys. 1 wektory jednostkowe $\mathbf{l}_r$ oraz $\mathbf{l}_s$ są zwierciadlane względem płaszczyzny stycznej do reflektora w punkcie odbicia $M(u, v)$, tj.

$$\begin{aligned} \mathbf{l}_r &= {}^*\mathbf{l}_s, \\ \mathbf{l}_s &= {}^*\mathbf{l}_r. \end{aligned} \quad (5)$$

Własności wektorów zwierciadlanych podane są w dodatku I.

Niech s oznacza odległość od punktu odbicia na reflektorze mierzoną wzdłuż promienia odbitego. Powierzchnie równego eikonu fali odbitej opisane są równaniem

$$\mathbf{e}(u, v) = \mathbf{m} + s \mathbf{l}_s, \quad (6)$$

gdzie $s = \text{const} - r(u, v)$.

¹⁾ Umieszczenie źródła oświetlającego nie w początku układu współrzędnych nadaje rachunkom większą ogólność, co jest korzystne w zastosowaniach.

Przejdziemy teraz do wyrażenia dla pola padającego i odbitego. Dla ustalenia uwagi będziemy zajmować się jedynie polem elektrycznym. Niech pole padające na reflektor ma postać

$$E^{pad} = \frac{e^{ikr}}{r} F_o(1_r), \quad (7)$$

gdzie funkcja $F_o(1_r)$ jest charakterystyką promieniowania źródła oświetlającego. Pole odbite wyrazi się następująco:

$$E^{odb} = - \frac{e^{ik(r+s)}}{r} {}^*F_o(1_r) D(s). \quad (8)$$

${}^*F_o(1_r)$ jest odbiciem zwierciadlanym $F_o(1_r)$ w płaszczyźnie stycznej do reflektora. $D(s)$ reprezentuje współczynnik zbieżności fali.

W niektórych zagadnieniach właściwe jest zapisanie pola odbitego w nieco innej postaci

$$E^{odb} = - \frac{e^{ik(r+s)}}{r+s} {}^*F_o(1_r) D_k(s). \quad (9)$$

$D_k(s)$ jest tu współczynnikiem zbieżności fali odmiennym od poprzedniego. Różnicę między obu współczynnikami można wyrazić poglądowo w sposób następujący: $D(s)$ jest miarą „odchylenia” fali odbitej od fali płaskiej, natomiast $D_k(s)$ — miarą „odchylenia” fali odbitej od fali kulistej promieniowanej przez źródło oświetlające. Oba współczynniki są związane prostą zależnością

$$D_k(s) = \frac{r+s}{r} D(s). \quad (10)$$

Należy zwrócić uwagę, że przyjęty w niniejszej pracy termin „współczynnik zbieżności” nie ma jednolitego znaczenia w literaturze. Niektórzy autorowie posługują się wielkościami odpowiadającymi współczynnikom D lub D_k używając nazw „współczynnik zbieżności” oraz „współczynnik rozbieżności” zależnie od tego, jak w danym przypadku szczególnie zmienia się zbieżność wiązki promieni po odbiciu. Nazwy takie spotyka się również dla określenia pewnych współczynników o charakterze ogólniejszym niż rozważane w tej pracy. Jako przykład można podać rozpatrywany w [2] współczynnik zbieżności dla fali wielokrotnie odbitej od jonosfery.

3. WSPÓŁCZYNNIK ZBIEŻNOŚCI FALI ODBITEJ

Zgodnie ze wzorem (2) współczynnik zbieżności $D(s)$ fali odbitej od reflektora anteny wyraża się następująco:

$$D(s) = \sqrt{\frac{K(s)}{K(0)}}, \quad (11)$$

gdzie $K(s)$ jest krzywizną Gaussa powierzchni równego eikonu.

Niech R_1 oraz R_2 oznaczają główne promienie krzywizny powierzchni fali odbitej w punkcie odbicia¹⁾. Mamy wtedy

$$K(0) = \frac{1}{R_1 R_2} \quad (12)$$

oraz

$$K(s) = \frac{1}{(R_1 - s)(R_2 - s)}. \quad (13)$$

Korzystając z tych zależności można zapisać wzór (11) w następującej postaci:

$$D(s) = \frac{1}{\sqrt{K_0 s^2 - 2H_0 s + 1}}; \quad (14)$$

K_0 jest tu krzywizną Gaussa powierzchni fali odbitej w punkcie odbicia, zaś H_0 — krzywizną średnią tej powierzchni. K_0 jest określone wzorem (12), natomiast dla H_0 mamy zależność

$$H_0 = \frac{1}{2} \left(\frac{1}{R_1} + \frac{1}{R_2} \right). \quad (15)$$

Obliczanie parametrów K_0 i H_0 będzie przedmiotem dalszych rozważań.

Zwróćmy uwagę na interesujące przypadki szczególne:

1) $K_0 = 0$ i $H_0 = 0$. Wtedy $D = 1$. Fala odbita jest falą płaską, amplituda fali nie zmienia się wzdłuż promienia.

2) $K_0 = 0$, $H_0 \neq 0$. W tym przypadku mamy

$$D = \frac{1}{\sqrt{-2H_0 s + 1}}. \quad (16)$$

Taka sytuacja zachodzi np. wtedy, gdy fala odbita jest falą cylindryczną lub stożkową.

3) $K_0 = H_0^2$. Mamy wtedy

$$D = \frac{1}{H_0 s - 1}. \quad (17)$$

Fala odbita jest falą kulistą.

Punkty, w których

$$K_0 s^2 - 2H_0 s + 1 = 0, \quad (18)$$

leżą na obwiedni promieni fali odbitej, tj. na kaustyce. W otoczeniu ta-

¹⁾ Promienie krzywizny R_1 , R_2 mogą przyjmować wartości dodatnie lub ujemne. Przyjmujemy, że $R > 0$, jeśli wektor $\mathbf{l}_s$ jest skierowany w stronę środka krzywizny odpowiedniego przekroju normalnego powierzchni równego eikonu.

kich punktów przybliżenie optyki geometrycznej zawodzi i równania (8) nie można stosować. Z zależności (18) wynika, że w punktach kaustyki

$$s_{1,2} = R_{1,2} = \frac{1}{K_o} (H_o \pm \sqrt{H_o^2 - K_o}). \quad (19)$$

Ponieważ krzywizna średnia H_o spełnia nierówność $H_o^2 \geq K_o$, równanie (19) daje zawsze rzeczywiste wartości s , dodatnie lub ujemne, zależnie od położenia kaustyki względem reflektora.

Z innego rodzaju kaustyką mamy do czynienia, gdy współczynniki K_o i H_o nieograniczenie rosną tak, że $D \rightarrow 0$. Jak zobaczymy, przypadek taki zachodzi, gdy fala pada stycznie do powierzchni reflektora.

4. KRZYWIZNA POWIERZCHNI FALI ODBITEJ

4.1. Równanie dla krzywizny Gaussa

Znajdziemy krzywiznę Gaussa $K(s)$ powierzchni równego eikonu fali odbitej. Obierzemy jako punkt wyjściowy równanie (6), które przepisujemy w postaci

$$\mathbf{e} = \mathbf{r} + a + s\mathbf{1}_s. \quad (20)$$

Jednostkowy wektor $\mathbf{1}_s$ jest normalny do powierzchni $e(u, v)$. Wobec tego dogodnie jest zastosować tu znany z geometrii różniczkowej związek

$$K \left(\frac{\partial \mathbf{e}}{\partial u} \times \frac{\partial \mathbf{e}}{\partial v} \right) = \frac{\partial \mathbf{1}_s}{\partial u} \times \frac{\partial \mathbf{1}_s}{\partial v}. \quad (21)$$

Mnożąc obie strony tego równania skalarnie przez $\mathbf{1}_s$ otrzymamy

$$K = \frac{\left(\mathbf{1}_s, \frac{\partial \mathbf{1}_s}{\partial u}, \frac{\partial \mathbf{1}_s}{\partial v} \right)}{\left(\mathbf{1}_s, \frac{\partial \mathbf{e}}{\partial u}, \frac{\partial \mathbf{e}}{\partial v} \right)}. \quad (22)$$

Przekształcimy mianownik tego wyrażenia. Z równania (20) dostaniemy kolejno

$$\begin{aligned} \frac{\partial \mathbf{e}}{\partial u} &= \frac{\partial \mathbf{r}}{\partial u} + s \frac{\partial \mathbf{1}_s}{\partial u} + \dots, \\ \frac{\partial \mathbf{e}}{\partial v} &= \frac{\partial \mathbf{r}}{\partial v} + s \frac{\partial \mathbf{1}_s}{\partial v} + \dots \end{aligned}$$

Wielokropkami oznaczono składniki nie wpływające na wynik końcowy. W dalszym ciągu mamy

$$\frac{\partial \mathbf{e}}{\partial u} \times \frac{\partial \mathbf{e}}{\partial v} = s^2 \left(\frac{\partial \mathbf{1}_s}{\partial u} \times \frac{\partial \mathbf{1}_s}{\partial v} \right) + s \left(\frac{\partial \mathbf{1}_s}{\partial u} \times \frac{\partial \mathbf{r}}{\partial v} + \frac{\partial \mathbf{r}}{\partial u} \times \frac{\partial \mathbf{1}_s}{\partial v} \right) + \frac{\partial \mathbf{r}}{\partial u} \times \frac{\partial \mathbf{r}}{\partial v} + \dots$$

Podstawiając to wyrażenie do (22) otrzymamy

$$K = \frac{\left(\mathbf{1}_s, \frac{\partial \mathbf{1}_s}{\partial u}, \frac{\partial \mathbf{1}_s}{\partial v} \right)}{s^2 \left(\mathbf{1}_s, \frac{\partial \mathbf{1}_s}{\partial u}, \frac{\partial \mathbf{1}_s}{\partial v} \right) + s \left[\left(\mathbf{1}_s, \frac{\partial \mathbf{1}_s}{\partial u}, \frac{\partial \mathbf{r}}{\partial v} \right) + \left(\mathbf{1}_s, \frac{\partial \mathbf{r}}{\partial u}, \frac{\partial \mathbf{1}_s}{\partial v} \right) \right] + \left(\mathbf{1}_s, \frac{\partial \mathbf{r}}{\partial u}, \frac{\partial \mathbf{r}}{\partial v} \right)}. \quad (23)$$

Kładąc w tym równaniu $s = 0$ otrzymujemy krzywiznę Gaussa w punkcie odbicia

$$K_o = \frac{\left(\mathbf{1}_s, \frac{\partial \mathbf{1}_s}{\partial u}, \frac{\partial \mathbf{1}_s}{\partial v} \right)}{\left(\mathbf{1}_s, \frac{\partial \mathbf{r}}{\partial u}, \frac{\partial \mathbf{r}}{\partial v} \right)}. \quad (24)$$

Stąd

$$\frac{K}{K_o} = \frac{1}{s^2 \frac{\left(\mathbf{1}_s, \frac{\partial \mathbf{1}_s}{\partial u}, \frac{\partial \mathbf{1}_s}{\partial v} \right)}{\left(\mathbf{1}_s, \frac{\partial \mathbf{r}}{\partial u}, \frac{\partial \mathbf{r}}{\partial v} \right)} + s \frac{\left[\left(\mathbf{1}_s, \frac{\partial \mathbf{1}_s}{\partial u}, \frac{\partial \mathbf{r}}{\partial v} \right) + \left(\mathbf{1}_s, \frac{\partial \mathbf{r}}{\partial u}, \frac{\partial \mathbf{1}_s}{\partial v} \right) \right]}{\left(\mathbf{1}_s, \frac{\partial \mathbf{r}}{\partial u}, \frac{\partial \mathbf{r}}{\partial v} \right)} + 1}. \quad (25)$$

Z porównania wzorów (11), (14) i (25) wynika, że

$$2H_o = - \frac{\left(\mathbf{1}_s, \frac{\partial \mathbf{1}_s}{\partial u}, \frac{\partial \mathbf{r}}{\partial v} \right) + \left(\mathbf{1}_s, \frac{\partial \mathbf{r}}{\partial u}, \frac{\partial \mathbf{1}_s}{\partial v} \right)}{\left(\mathbf{1}_s, \frac{\partial \mathbf{r}}{\partial u}, \frac{\partial \mathbf{r}}{\partial v} \right)}. \quad (26)$$

W przypadku gdy fala pada stycznie do powierzchni reflektora, mamy

$$\left(\mathbf{1}_s, \frac{\partial \mathbf{r}}{\partial u}, \frac{\partial \mathbf{r}}{\partial v} \right) = 0.$$

Wzory (24) i (26) dają wówczas nieskończenie wielkie wartości K_o i H_o .

4.2. Krzywizna Gaussa w punkcie odbicia

Przekształcimy równanie (24) wprowadzając do niego parametry krzywizny reflektora w punkcie odbicia. Korzystając z zależności

$$\mathbf{1}_s = {}^*\mathbf{1}_r = \frac{{}^*\mathbf{r}}{r} \quad (27)$$

można równanie (24) przepisać w postaci

$$K_o = \frac{({}^*\mathbf{r}, {}^*\mathbf{r}_u, {}^*\mathbf{r}_v)}{r^2({}^*\mathbf{r}, \mathbf{r}_u, \mathbf{r}_v)},$$

gdzie dla skrócenia zapisu zastąpiono symbole pochodnych cząstkowych $\frac{\partial}{\partial u}$ i $\frac{\partial}{\partial v}$ indeksami u oraz v .

Biorąc pod uwagę łatwą do sprawdzenia zależność

$$(*\mathbf{r}, \mathbf{r}_u, \mathbf{r}_v) = -\sqrt{g}(\mathbf{n} \cdot \mathbf{r}), \quad (28)$$

gdzie

$$\sqrt{g} = |\mathbf{r}_u \times \mathbf{r}_v|,$$

otrzymamy następnie

$$K_o = -\frac{(*\mathbf{r}, *\mathbf{r}_u, *\mathbf{r}_v)}{\sqrt{g}r^2(\mathbf{n} \cdot \mathbf{r})}. \quad (29)$$

Bezpośrednie przekształcanie licznika w tym wyrażeniu jest rachunkowo uciążliwe. Dlatego zastosujemy tu metodę pośrednią biorąc jako punkt wyjściowy otrzymane w dodatku II wyrażenie dla krzywizny przekroju normalnego powierzchni reflektora w płaszczyźnie padania (wzór DII-7).

$$H^{\parallel} = -\frac{1}{4\sqrt{g}(\mathbf{n} \times \mathbf{r})^2} [(\mathbf{r}, *\mathbf{r}_u, \mathbf{r}_v) + (\mathbf{r}, \mathbf{r}_u, *\mathbf{r}_v) + (*\mathbf{r}, *\mathbf{r}_u, \mathbf{r}_v) + (*\mathbf{r}, \mathbf{r}_u, *\mathbf{r}_v)].$$

Będziemy teraz przekształcać iloczyny mieszane zawarte w nawiasie kwadratowym wykorzystując odpowiednie zależności dla wektorów zwierciedlanych. Mamy kolejno

$$\begin{aligned} (\mathbf{r}, *\mathbf{r}_u, \mathbf{r}_v) &= (\mathbf{r}, \mathbf{r}_u, \mathbf{r}_v) - 2(\mathbf{n}_u \cdot \mathbf{r})(\mathbf{r}, \mathbf{n}, \mathbf{r}_v) - 2(\mathbf{n} \cdot \mathbf{r})(\mathbf{r}, \mathbf{n}_u, \mathbf{r}_v), \\ (\mathbf{r}, \mathbf{r}_u, *\mathbf{r}_v) &= (\mathbf{r}, \mathbf{r}_u, \mathbf{r}_v) - 2(\mathbf{n}_v \cdot \mathbf{r})(\mathbf{r}, \mathbf{r}_u, \mathbf{n}) - 2(\mathbf{n} \cdot \mathbf{r})(\mathbf{r}, \mathbf{r}_u, \mathbf{n}_v), \\ (*\mathbf{r}, *\mathbf{r}_u, \mathbf{r}_v) &= (*\mathbf{r}, *\mathbf{r}_u, *\mathbf{r}_v) + 2(\mathbf{n}_v \cdot \mathbf{r})(*\mathbf{r}, *\mathbf{r}_u, \mathbf{n}) + 2(\mathbf{n} \cdot \mathbf{r})(*\mathbf{r}, *\mathbf{r}_u, \mathbf{n}_v), \\ (*\mathbf{r}, \mathbf{r}_u, *\mathbf{r}_v) &= (*\mathbf{r}, *\mathbf{r}_u, *\mathbf{r}_v) + 2(\mathbf{n}_u \cdot \mathbf{r})(*\mathbf{r}, \mathbf{n}, *\mathbf{r}_v) + 2(\mathbf{n} \cdot \mathbf{r})(*\mathbf{r}, \mathbf{n}_u, *\mathbf{r}_v). \end{aligned} \quad (30)$$

Biorąc pod uwagę powyższe związki przepisujemy równanie dla $H^{\parallel}$, jak następuje

$$H^{\parallel} = -\frac{1}{4\sqrt{g}(\mathbf{n} \times \mathbf{r})^2} [2\mathbf{r}, \mathbf{r}_u, \mathbf{r}_v) + 2(*\mathbf{r}, *\mathbf{r}_u, *\mathbf{r}_v) + T], \quad (31)$$

gdzie oznaczono dla skrócenia zapisu

$$\begin{aligned} T &= 2(\mathbf{n}_v \cdot \mathbf{r})[(*\mathbf{r}, *\mathbf{r}_u, \mathbf{n}) - (\mathbf{r}, \mathbf{r}_u, \mathbf{n})] + \\ &+ 2(\mathbf{n}_u \cdot \mathbf{r})[(*\mathbf{r}, \mathbf{n}, *\mathbf{r}_v) - (\mathbf{r}, \mathbf{n}, \mathbf{r}_v)] + \\ &+ 2(\mathbf{n} \cdot \mathbf{r})[(*\mathbf{r}, *\mathbf{r}_u, \mathbf{n}_v) - (\mathbf{r}, \mathbf{r}_u, \mathbf{n}_v)] + \\ &+ 2(\mathbf{n} \cdot \mathbf{r})[(*\mathbf{r}, \mathbf{n}_u, *\mathbf{r}_v) - (\mathbf{r}, \mathbf{n}_u, \mathbf{r}_v)]. \end{aligned} \quad (32)$$

Z kolei mamy

$$(\mathbf{r}, \mathbf{r}_u, \mathbf{r}_v) = \sqrt{g}\mathbf{n} \cdot \mathbf{r}, \quad (33)$$

$$(*\mathbf{r}, *\mathbf{r}_u, *\mathbf{r}_v) = -\sqrt{g}r^2(\mathbf{n} \cdot \mathbf{r})K_o. \quad (34)$$

Wobec tego otrzymujemy z (31)

$$H^{\parallel} = -\frac{\mathbf{n} \cdot \mathbf{r}}{2(\mathbf{n} \times \mathbf{r})^2} + \frac{r^2(\mathbf{n} \cdot \mathbf{r})}{2(\mathbf{n} \times \mathbf{r})^2} K_o - \frac{T}{4\sqrt{g}(\mathbf{n} \times \mathbf{r})^2}, \quad (35)$$

a stąd

$$K_o = \frac{1}{r^2} + \frac{T}{2\sqrt{g}} + \frac{2(\mathbf{n} \times \mathbf{r})^2}{r^2(\mathbf{n} \cdot \mathbf{r})} H^{\parallel}. \quad (36)$$

Podstawiając teraz w wyrażeniu (32) odpowiednie zależności dla wektorów zwierciadlanych

$$\begin{aligned} {}^*\mathbf{r} &= \mathbf{r} - 2(\mathbf{n} \cdot \mathbf{r})\mathbf{n}, \\ {}^*\mathbf{r}_u &= \mathbf{r}_u - 2(\mathbf{n}_u \cdot \mathbf{r})\mathbf{n} - 2(\mathbf{n} \cdot \mathbf{r})\mathbf{n}_u, \\ {}^*\mathbf{r}_v &= \mathbf{r}_v - 2(\mathbf{n}_v \cdot \mathbf{r})\mathbf{n} - 2(\mathbf{n} \cdot \mathbf{r})\mathbf{n}_v, \end{aligned} \quad (37)$$

łatwo sprawdzić, że

$$\frac{T}{2\sqrt{g}} = 4K_s + \frac{4(\mathbf{n} \cdot \mathbf{r})}{r^2} H_s. \quad (38)$$

Oznaczono tu przez K_s i H_s odpowiednio krzywiznę Gaussa i krzywiznę średnią powierzchni reflektora w punkcie odbicia. Wyrażenia analityczne dla tych krzywizn są następujące:

$$K_s = \frac{(\mathbf{n}, \mathbf{n}_u, \mathbf{n}_v)}{\sqrt{g}}, \quad (39)$$

$$H_s = \frac{(\mathbf{n}, \mathbf{r}_v, \mathbf{n}_u) - (\mathbf{n}, \mathbf{r}_u, \mathbf{n}_v)}{2\sqrt{g}}. \quad (40)$$

Wyrażenie (38) wprowadzamy do (36) i otrzymujemy

$$K_o = \frac{1}{r^2} + 4K_s + \frac{4(\mathbf{n} \cdot \mathbf{r})}{r^2} H_s + \frac{2(\mathbf{n} \times \mathbf{r})^2}{r^2(\mathbf{n} \cdot \mathbf{r})} H^{\parallel}. \quad (41)$$

Wzór ten można doprowadzić do bardziej przejrzystej postaci kładąc

$$2H_s = H^{\perp} + H^{\parallel}; \quad (42)$$

$H^{\perp}$ oznacza tu krzywiznę przekroju normalnego powierzchni reflektora w płaszczyźnie prostopadłej do płaszczyzny padania. Metoda obliczenia tej krzywizny jest podana w dodatku III.

Otrzymujemy ostatecznie

$$K_o = \frac{1}{r^2} + 4K_s + \frac{2}{r} \left(\frac{r}{\mathbf{n} \cdot \mathbf{r}} H^{\parallel} + \frac{\mathbf{n} \cdot \mathbf{r}}{r} H^{\perp} \right). \quad (43)$$

Wzór ten — w odmiennej, ale równoważnej postaci — jest podany w [3]. Pierwszy składnik prawej strony reprezentuje krzywiznę Gaussa fali padającej, pozostałe składniki związane są z krzywizną powierzchni reflektora w punkcie odbicia.

Warto zwrócić uwagę na przypadek szczególny, gdy $r \rightarrow \infty$. Wówczas mamy do czynienia z falą płaską padającą na reflektor. Wtedy krzywizna Gaussa fali odbitej spełnia prosty związek

$$K_o = 4K_s \quad (44)$$

niezależnie od kąta padania fali. W szczególności gdy fala płaska pada na powierzchnię, której krzywizna Gaussa jest równa zero, krzywizna Gaussa powierzchni fali odbitej jest też równa zero.

4.3. Krzywizna średnia w punkcie odbicia

Krzywizna średnia powierzchni fali odbitej jest określona w punkcie odbicia wzorem (26). Zależność tę przekształcimy w podobny sposób jak w przypadku krzywizny Gaussa. Na wstępie wykorzystamy stosowaną już poprzednio zależność (27)

$$\mathbf{l}_s = {}^*\mathbf{l}_r = \frac{{}^*\mathbf{r}}{r}$$

i przepiszemy wzór (26) w postaci

$$H_o = -\frac{1}{2r} \frac{({}^*\mathbf{r}, {}^*\mathbf{r}_u, \mathbf{r}_v) + ({}^*\mathbf{r}_o, \mathbf{r}_u, {}^*\mathbf{r}_v)}{({}^*\mathbf{r}, \mathbf{r}_u, \mathbf{r}_v)} \quad (45)$$

Tak jak poprzednio mamy [wzór (28)]

$$({}^*\mathbf{r}, \mathbf{r}_u, \mathbf{r}_v) = -\sqrt{g}(\mathbf{n} \cdot \mathbf{r})$$

i otrzymamy zamiast (45)

$$H_o = \frac{({}^*\mathbf{r}, {}^*\mathbf{r}_u, \mathbf{r}_v) + ({}^*\mathbf{r}, \mathbf{r}_u, {}^*\mathbf{r}_v)}{2\sqrt{g}r(\mathbf{n} \cdot \mathbf{r})} \quad (46)$$

Podobnie jak przy obliczaniu krzywizny Gaussa wykorzystamy wyprowadzone w dodatku II równanie (DII-7) dla krzywizny $H^{\parallel}$:

$$H^{\parallel} = -\frac{1}{4\sqrt{g}(\mathbf{n} \times \mathbf{r})^2} [(\mathbf{r}, {}^*\mathbf{r}_u, \mathbf{r}_v) + (\mathbf{r}, \mathbf{r}_u, {}^*\mathbf{r}_v) + ({}^*\mathbf{r}, {}^*\mathbf{r}_u, \mathbf{r}_v) + ({}^*\mathbf{r}, \mathbf{r}_u, {}^*\mathbf{r}_v)].$$

Pierwsze dwa składniki w nawiasie kwadratowym przekształcimy, jak następuje:

$$\begin{aligned} (\mathbf{r}, {}^*\mathbf{r}_u, \mathbf{r}_v) &= ({}^*\mathbf{r}, {}^*\mathbf{r}_u, \mathbf{r}_v) + 2(\mathbf{n} \cdot \mathbf{r})(\mathbf{n}, {}^*\mathbf{r}_u, \mathbf{r}_v), \\ (\mathbf{r}, \mathbf{r}_u, {}^*\mathbf{r}_v) &= ({}^*\mathbf{r}, \mathbf{r}_u, {}^*\mathbf{r}_v) + 2(\mathbf{n} \cdot \mathbf{r})(\mathbf{n}, \mathbf{r}_u, {}^*\mathbf{r}_v). \end{aligned} \quad (47)$$

Biorąc nadto pod uwagę, że

$$({}^*\mathbf{r}, {}^*\mathbf{r}_u, \mathbf{r}_v) + ({}^*\mathbf{r}, \mathbf{r}_u, {}^*\mathbf{r}_v) = 2\sqrt{g}r(\mathbf{n} \cdot \mathbf{r})H_o, \quad (48)$$

otrzymamy z (DII-7)

$$H^{\parallel} = -\frac{r(\mathbf{n} \cdot \mathbf{r})}{(\mathbf{n} \times \mathbf{r})^2} H_o - \frac{(\mathbf{n} \cdot \mathbf{r})}{2\sqrt{g}(\mathbf{n} \times \mathbf{r})^2} [(\mathbf{n}, {}^*\mathbf{r}_u, \mathbf{r}_v) + (\mathbf{n}, \mathbf{r}_u, {}^*\mathbf{r}_v)]. \quad (49)$$

Dla wyrażenia w nawiasie kwadratowym mamy

$$(\mathbf{n}, * \mathbf{r}_u, \mathbf{r}_v) + (\mathbf{n}, \mathbf{r}_u, * \mathbf{r}_v) = 2(\mathbf{n}, \mathbf{r}_u, \mathbf{r}_v) - 2(\mathbf{n} \cdot \mathbf{r}) [(\mathbf{n}, \mathbf{n}_u, \mathbf{r}_v) + (\mathbf{n}, \mathbf{r}_u, \mathbf{n}_v)] = \\ = 2\sqrt{g} + 4\sqrt{g}(\mathbf{n} \cdot \mathbf{r})H_s. \quad (50)$$

Wobec tego dostaniemy z (49)

$$H^{\parallel} = -\frac{r(\mathbf{n} \cdot \mathbf{r})}{(\mathbf{n} \times \mathbf{r})^2} H_o - \frac{\mathbf{n} \cdot \mathbf{r}}{(\mathbf{n} \times \mathbf{r})^2} \cdot \frac{2(\mathbf{n} \cdot \mathbf{r})^2}{(\mathbf{n} \times \mathbf{r})^2} H_s. \quad (51)$$

Stąd uzyskujemy wyrażenie dla H_o w postaci

$$H_o = -\frac{1}{r} - \frac{2(\mathbf{n} \cdot \mathbf{r})}{r} H_s - \frac{(\mathbf{n} \times \mathbf{r})^2}{r(\mathbf{n} \cdot \mathbf{r})} H^{\parallel}. \quad (52)$$

Podobnie jak w przypadku krzywizny Gaussa wykorzystamy związek (42)

$$2H_s = H^{\perp} + H^{\parallel}$$

i otrzymamy następującą postać wzoru dla krzywizny średniej H_o :

$$H_o = -\frac{1}{r} - \left(\frac{r}{\mathbf{n} \cdot \mathbf{r}} H^{\parallel} + \frac{\mathbf{n} \cdot \mathbf{r}}{r} H^{\perp} \right). \quad (53)$$

Pierwszy składnik prawej strony przedstawia krzywiznę średnią powierzchni fali padającej, pozostałe składniki są związane z krzywizną powierzchni reflektora.

Niech $\cos \Theta = -\mathbf{n} \cdot \mathbf{l}_r$ określa kąt padania fali. Wzór (53) daje wtedy

$$H_o = -\frac{1}{r} + \left(\frac{1}{\cos \Theta} H^{\parallel} + \cos \Theta H^{\perp} \right). \quad (54)$$

Dla $r \rightarrow \infty$ otrzymujemy wartość H_o dla płaskiej fali padającej

$$H_o = \frac{1}{\cos \Theta} H^{\parallel} + \cos \Theta H^{\perp}. \quad (55)$$

Analizując wzory (43) i (53) łatwo sprawdzić, że krzywizny K_o oraz H_o są związane następującą prostą zależnością:

$$K_o = 4K_s - \frac{2H_o}{r} - \frac{1}{r^2}. \quad (56)$$

Wzory (43) i (53) pozwalają nadto uzyskać nowe wyrażenia dla współczynnika zbieżności. Otrzymujemy po prostych przekształceniach odpowiednio z (14) i (10)

$$D(s) = \frac{1}{s \sqrt{\left(\frac{r+s}{rs} \right)^2 + 2 \left(\frac{r+s}{rs} \right) \left(\frac{r}{\mathbf{n} \cdot \mathbf{r}} H^{\parallel} + \frac{\mathbf{n} \cdot \mathbf{r}}{r} H^{\perp} \right) + 4K_s}} \quad (57)$$

oraz

$$D_k(s) = \frac{1}{\sqrt{4K_s \left(\frac{rs}{r+s} \right)^2 + \frac{2rs}{r+s} \left(\frac{r}{\mathbf{n} \cdot \mathbf{r}} H^{\parallel} + \frac{\mathbf{n} \cdot \mathbf{r}}{r} H^{\perp} \right) + 1}}. \quad (58)$$

Wyrażenia

$$\frac{e^{ik(r+s)}}{r} D(s) \quad \text{oraz} \quad \frac{e^{ik(r+s)}}{r+s} D_k(s)$$

są symetryczne względem zmiennych r , s . Jest to potwierdzeniem dobrze znanej w optyce geometrycznej zasady odwracalności kierunku promieni. Z otrzymanych wzorów wynika bezpośrednio, że gdy kąt padania fali zbliża się do kąta prostego, oba współczynniki $D(s)$ i $D_k(s)$ maleją do zera. W granicy, gdy fala pada stycznie do powierzchni reflektora, punkt odbicia staje się punktem kaustycznym i pole nie da się tu obliczyć przy pomocy równań optyki geometrycznej.

4.4. Transformacja wyrażenia zawierającego krzywizny $H^{\parallel}$ i $H^{\perp}$

Wyrażenie

$$Q = \frac{r}{\mathbf{n} \cdot \mathbf{r}} H^{\parallel} + \frac{\mathbf{n} \cdot \mathbf{r}}{r} H^{\perp} \quad (59)$$

wchodzi w skład wzorów (43) i (53) określających krzywiznę Gaussa oraz krzywiznę średnią powierzchni fali odbitej. Wyrażenie to zawierają również wzory (57) i (58), jakie otrzymaliśmy dla współczynników zbieżności fali odbitej. Obliczanie tego wyrażenia jest na ogół kłopotliwe. Tylko w niektórych przypadkach szczególnych, np. dla reflektora sferycznego, obliczenie krzywizn $H^{\parallel}$ i $H^{\perp}$ nie sprawia trudności. W ogólności chcąc obliczyć te krzywizny jesteśmy zmuszeni korzystać z uciążliwych zależności podanych w dodatkach II i III. Jest więc rzeczą bardzo pożyteczną doprowadzić wyrażenie (59) do takiej postaci, która pozwalałaby w zastosowaniach na możliwie szybki rachunek. W tym celu podstawimy do (59) podane w dodatku III równania

$$H^{\perp} = - \frac{\mathbf{n} \times \mathbf{r}}{\sqrt{g}(\mathbf{n} \times \mathbf{r})^2} \cdot [(\mathbf{r} \cdot \mathbf{r}_u) \mathbf{n}_v - (\mathbf{r} \cdot \mathbf{r}_v) \mathbf{n}_u], \quad (\text{DIII-7})$$

$$H^{\parallel} = 2H_s - H^{\perp}, \quad (\text{DIII-8})$$

gdzie wg (40)

$$H_s = \frac{(\mathbf{n}, \mathbf{r}_v, \mathbf{n}_u) - (\mathbf{n}, \mathbf{r}_u, \mathbf{n}_v)}{2\sqrt{g}}.$$

Otrzymamy po łatwych rachunkach

$$Q = \frac{1}{(\mathbf{n} \cdot \mathbf{l}_r) \sqrt{g}} [(\mathbf{n}, \mathbf{r}_v, \mathbf{n}_u) - (\mathbf{n}, \mathbf{r}_u, \mathbf{n}_v) + (\mathbf{n}, \mathbf{l}_r, \mathbf{n}_v) (\mathbf{l}_r \cdot \mathbf{r}_u) + \\ - (\mathbf{n}, \mathbf{l}_r, \mathbf{n}_u) (\mathbf{l}_r \cdot \mathbf{r}_v)]. \quad (60)$$

Wyrażenie w nawiasie kwadratowym zawiera pochodne wektora jednostkowego $\mathbf{n}$. Zastąpimy je przez pochodne wektora

$$\mathbf{N} = \mathbf{r}_u \times \mathbf{r}_v$$

korzystając z zależności

$$\sqrt{g} \mathbf{n} = \mathbf{N}.$$

Otrzymamy zamiast (60)

$$Q = \frac{1}{(\mathbf{n} \cdot \mathbf{l}_r) g} [(\mathbf{n}, \mathbf{r}_v, \mathbf{N}_u) - (\mathbf{n}, \mathbf{r}_u, \mathbf{N}_v) + (\mathbf{n}, \mathbf{l}_r, \mathbf{N}_v) (\mathbf{l}_r \cdot \mathbf{r}_u) + \\ - (\mathbf{n}, \mathbf{l}_r, \mathbf{N}_u) (\mathbf{l}_r \cdot \mathbf{r}_v)]. \quad (61)$$

Łatwo sprawdzić, że

$$\begin{aligned} (\mathbf{n}, \mathbf{r}_v, \mathbf{N}_u) &= \mathbf{r}_v^2 (\mathbf{n} \cdot \mathbf{r}_{uu}) - (\mathbf{r}_u \cdot \mathbf{r}_v) (\mathbf{n} \cdot \mathbf{r}_{uv}), \\ (\mathbf{n}, \mathbf{r}_u, \mathbf{N}_v) &= (\mathbf{r}_u \cdot \mathbf{r}_v) (\mathbf{n} \cdot \mathbf{r}_{uv}) - \mathbf{r}_u^2 (\mathbf{n} \cdot \mathbf{r}_{vv}), \\ (\mathbf{n}, \mathbf{l}_r, \mathbf{N}_v) &= (\mathbf{l}_r \cdot \mathbf{r}_v) (\mathbf{n} \cdot \mathbf{r}_{uv}) - (\mathbf{l}_r \cdot \mathbf{r}_u) (\mathbf{n} \cdot \mathbf{r}_{vv}), \\ (\mathbf{n}, \mathbf{l}_r, \mathbf{N}_u) &= (\mathbf{l}_r \cdot \mathbf{r}_v) (\mathbf{n} \cdot \mathbf{r}_{uu}) - (\mathbf{l}_r \cdot \mathbf{r}_u) (\mathbf{n} \cdot \mathbf{r}_{uv}). \end{aligned}$$

Podstawienie tych wzorów do (61) prowadzi bezpośrednio do następującego wyniku

$$Q = \frac{1}{(\mathbf{n} \cdot \mathbf{l}_r) g} [(\mathbf{n} \cdot \mathbf{r}_{uu}) (\mathbf{l}_r \times \mathbf{r}_v)^2 - 2 (\mathbf{n} \cdot \mathbf{r}_{uv}) (\mathbf{l}_r \times \mathbf{r}_u) (\mathbf{l}_r \times \mathbf{r}_v) + \\ + (\mathbf{n} \cdot \mathbf{r}_{vv}) (\mathbf{l}_r \times \mathbf{r}_u)^2]. \quad (62)$$

Wzór ten o budowie symetrycznej względem u, v ułatwia obliczenie wyrażenia (59) w tych przypadkach, w których bezpośrednie wyznaczenie krzywizn $H^{\parallel}$ i $H^{\perp}$ jest uciążliwe.

5. WSPÓŁCZYNNIK ZBIEŻNOŚCI I METODA STACJONARNEJ FAZY

Pole odbite od reflektora anteny może być obliczone przy pomocy następującego wzoru całkowego

$$\mathbf{E}^{odb} = -\frac{1}{2\pi i k} \nabla \times \nabla \times \int_S [\mathbf{n} \times (\mathbf{l}_\rho \times \mathbf{F}_\sigma)] \frac{e^{ik(\rho+\sigma)}}{\rho\sigma} dS. \quad (63)$$

W wyrażeniu podcałkowym wprowadzono tu zmienne ρ oraz σ dla oznaczenia odległości od źródła fali Q do punktu całkowania na reflektro-

rze oraz od punktu całkowania do punktu obserwacji P . Oznaczenia r oraz s są zarezerwowane dla punktu geometrycznego odbicia (por. rys. 2).

Całka (63) może być obliczona metodą stacjonarnej fazy, co przy odpowiednich założeniach daje przybliżenie optyki geometrycznej. Nie będziemy zajmować się tym rachunkiem, gdyż przekracza on ramy pracy, zwrócimy jedynie uwagę na wynik końcowy, który można zapisać w postaci

$$E^{odb} \sim - \frac{e^{ik(r+s)}}{r} {}^*F_o(\mathbf{1}_r) \frac{(\mathbf{n} \cdot \mathbf{1}_s)}{s} \sqrt{\frac{g(\mathbf{r})}{\mathcal{H}_o(\varrho + \sigma)}} + \dots, \quad (64)$$

gdzie:

$g(\mathbf{r})$ — wyróżnik I formy kwadratowej powierzchni reflektora w punkcie odbicia,

$\mathcal{H}_o$ — hesjan funkcji $\varrho + \sigma$ w punkcie odbicia.

Wielokropek zastępuje wyrazy niższego rzędu względem k . Wyraz przytoczony reprezentuje przybliżenie optyki geometrycznej.

Konfrontacja równań (8) i (64) prowadzi do następującego interesującego wyrażenia dla współczynnika zbieżności:

$$D(s) = \frac{\mathbf{n} \cdot \mathbf{1}_s}{s} \sqrt{\frac{g(\mathbf{r})}{\mathcal{H}_o(\varrho + \sigma)}}. \quad (65)$$

Wykażemy bezpośrednio, że wyrażenia (57) i (65) dla współczynnika zbieżności są równoważne. Rachunek sprowadza się właściwie do wyznaczenia wartości hesjanu funkcji $\varrho + \sigma$ w punkcie odbicia.

Zgodnie z definicją hesjanu mamy

$$\mathcal{H}(\varrho + \sigma) = \frac{\partial^2(\varrho + \sigma)}{\partial u^2} \frac{\partial^2(\varrho + \sigma)}{\partial v^2} - \left(\frac{\partial^2(\varrho + \sigma)}{\partial u \partial v} \right)^2. \quad (66)$$

Obliczymy wartości drugich pochodnych występujących w tym wyrażeniu. Niech

$$\mathbf{p} = \varrho \mathbf{1}_p.$$

Stąd

$$\frac{\partial \mathbf{p}}{\partial u} = \mathbf{1}_p \frac{\partial \varrho}{\partial u} + \varrho \frac{\partial \mathbf{1}_p}{\partial u}.$$

Mnożąc skalarnie to równanie przez $\mathbf{1}_p$ otrzymamy

$$\frac{\partial \varrho}{\partial u} = \mathbf{1}_p \cdot \frac{\partial \mathbf{p}}{\partial u}$$

i następnie

$$\frac{\partial^2 \varrho}{\partial u^2} = \mathbf{1}_p \cdot \frac{\partial^2 \mathbf{p}}{\partial u^2} + \frac{\partial \mathbf{1}_p}{\partial u} \cdot \frac{\partial \mathbf{p}}{\partial u}, \quad (67)$$

Analogicznie dostaniemy

$$\frac{\partial \sigma}{\partial u} = \mathbf{1}_\sigma \cdot \frac{\partial \sigma}{\partial u}.$$

Ponieważ $\rho + \sigma = \text{const}$, to

$$\frac{\partial \sigma}{\partial u} = -\frac{\partial \rho}{\partial u}.$$

Zgodnie z tym

$$\frac{\partial \sigma}{\partial u} = -\mathbf{1}_\sigma \cdot \frac{\partial \rho}{\partial u}$$

oraz

$$\frac{\partial^2 \sigma}{\partial u^2} = -\mathbf{1}_\sigma \cdot \frac{\partial^2 \rho}{\partial u^2} - \frac{\partial \mathbf{1}_\sigma}{\partial u} \cdot \frac{\partial \rho}{\partial u}. \quad (68)$$

Dodając stronami (67) i (68) dostaniemy

$$\frac{\partial^2 (\rho + \sigma)}{\partial u^2} = \mathbf{a} \cdot \frac{\partial^2 \rho}{\partial u^2} + \frac{\partial \mathbf{a}}{\partial u} \cdot \frac{\partial \rho}{\partial u}, \quad (69)$$

gdzie

$$\mathbf{a} = \mathbf{1}_\rho - \mathbf{1}_\sigma. \quad (70)$$

W podobny sposób otrzymamy

$$\frac{\partial^2 (\rho + \sigma)}{\partial v^2} = \mathbf{a} \cdot \frac{\partial^2 \rho}{\partial v^2} + \frac{\partial \mathbf{a}}{\partial v} \cdot \frac{\partial \rho}{\partial v}, \quad (71)$$

$$\frac{\partial^2 (\rho + \sigma)}{\partial u \partial v} = \mathbf{a} \cdot \frac{\partial^2 \rho}{\partial u \partial v} + \frac{\partial \mathbf{a}}{\partial v} \cdot \frac{\partial \rho}{\partial u}. \quad (72)$$

Z kolei obliczymy pochodne $\frac{\partial \mathbf{a}}{\partial u}$ oraz $\frac{\partial \mathbf{a}}{\partial v}$. Zgodnie z (70) mamy

$$\mathbf{a} = \frac{\rho}{\varrho} - \frac{\sigma}{\sigma},$$

stąd

$$\frac{\partial \mathbf{a}}{\partial u} = \frac{1}{\varrho} \frac{\partial \rho}{\partial u} - \frac{\rho}{\varrho^2} \frac{\partial \varrho}{\partial u} - \frac{1}{\sigma} \frac{\partial \sigma}{\partial u} + \frac{\sigma}{\sigma^2} \frac{\partial \sigma}{\partial u}.$$

Kładąc w tym wyrażeniu stosowane poprzednio zależności dla $\frac{\partial \rho}{\partial u}$, $\frac{\partial \sigma}{\partial u}$

i $\frac{\partial \sigma}{\partial u}$ dostaniemy

$$\frac{\partial \mathbf{a}}{\partial u} = \frac{1}{\varrho} \frac{\partial \rho}{\partial u} - \frac{\mathbf{1}_\rho}{\varrho} \left(\mathbf{1}_\rho \cdot \frac{\partial \rho}{\partial u} \right) + \frac{1}{\sigma} \frac{\partial \rho}{\partial u} - \frac{\mathbf{1}_\sigma}{\sigma} \left(\mathbf{1}_\sigma \cdot \frac{\partial \rho}{\partial u} \right) \quad (73)$$

i analogicznie

$$\frac{\partial \mathbf{a}}{\partial v} = \frac{1}{\varrho} \frac{\partial \mathbf{p}}{\partial v} - \frac{\mathbf{1}_\rho}{\varrho} \left(\mathbf{1}_\rho \cdot \frac{\partial \mathbf{p}}{\partial v} \right) + \frac{1}{\sigma} \frac{\partial \mathbf{p}}{\partial v} - \frac{\mathbf{1}_\sigma}{\sigma} \left(\mathbf{1}_\sigma \cdot \frac{\partial \mathbf{p}}{\partial v} \right). \quad (74)$$

Otrzymane zależności, słuszne dla dowolnego punktu na reflektorze, należy teraz odnieść do punktu odbicia. W tym punkcie $\varrho = r$, $\sigma = s$, $\mathbf{1}_\sigma = \mathbf{1}_s = * \mathbf{1}_r$, nadto

$$\mathbf{a} = \mathbf{1}_r - \mathbf{1}_s = 2(\mathbf{n} \cdot \mathbf{1}_r) \mathbf{n}, \quad (75)$$

$$\left(\frac{\partial \mathbf{a}}{\partial u} \right)_0 = \frac{r+s}{rs} \frac{\partial \mathbf{r}}{\partial u} - \frac{\mathbf{1}_r}{r} \left(\mathbf{1}_r \cdot \frac{\partial \mathbf{r}}{\partial u} \right) - \frac{* \mathbf{1}_r}{s} \left(* \mathbf{1}_r \cdot \frac{\partial \mathbf{r}}{\partial u} \right), \quad (76)$$

$$\left(\frac{\partial \mathbf{a}}{\partial v} \right)_0 = \frac{r+s}{rs} \frac{\partial \mathbf{r}}{\partial v} - \frac{\mathbf{1}_r}{r} \left(\mathbf{1}_r \cdot \frac{\partial \mathbf{r}}{\partial v} \right) - \frac{* \mathbf{1}_r}{s} \left(* \mathbf{1}_r \cdot \frac{\partial \mathbf{r}}{\partial v} \right). \quad (77)$$

Indeks zero oznacza wartości pochodnych w punkcie odbicia. Podstawiając (75), (76) i (77) do (69), (71) i (72) dostaniemy po wykonaniu elementarnych operacji

$$\left(\frac{\partial^2 (\varrho + \sigma)}{\partial u^2} \right)_0 = 2(\mathbf{n} \cdot \mathbf{1}_r) (\mathbf{n} \cdot \mathbf{r}_{uu}) + \frac{r+s}{rs} (\mathbf{1}_r \times \mathbf{r}_u)^2, \quad (78)$$

$$\left(\frac{\partial^2 (\varrho + \sigma)}{\partial v^2} \right)_0 = 2(\mathbf{n} \cdot \mathbf{1}_r) (\mathbf{n} \cdot \mathbf{r}_{vv}) + \frac{r+s}{rs} (\mathbf{1}_r \times \mathbf{r}_v)^2, \quad (79)$$

$$\left(\frac{\partial^2 (\varrho + \sigma)}{\partial u \partial v} \right)_0 = 2(\mathbf{n} \cdot \mathbf{1}_r) (\mathbf{n} \cdot \mathbf{r}_{uv}) + \frac{r+s}{rs} (\mathbf{1}_r \times \mathbf{r}_u) \cdot (\mathbf{1}_r \times \mathbf{r}_v). \quad (80)$$

Wprowadzając te wartości do równania (66) otrzymujemy następujące wyrażenie dla hesjanu w punkcie odbicia:

$$\begin{aligned} \mathcal{H}_o(\varrho + \sigma) &= \left(\frac{r+s}{rs} \right)^2 \left[(\mathbf{1}_r \times \mathbf{r}_u)^2 (\mathbf{1}_r \times \mathbf{r}_v)^2 - [(\mathbf{1}_r \times \mathbf{r}_u) \cdot (\mathbf{1}_r \times \mathbf{r}_v)]^2 \right] + \\ &+ 2(\mathbf{n} \cdot \mathbf{1}_r) \left(\frac{r+s}{rs} \right) [(\mathbf{n} \cdot \mathbf{r}_{uu}) (\mathbf{1}_r \times \mathbf{r}_v)^2 - 2(\mathbf{n} \cdot \mathbf{r}_{uv}) (\mathbf{1}_r \times \mathbf{r}_u) \cdot (\mathbf{1}_r \times \mathbf{r}_v) + (\mathbf{n} \cdot \mathbf{r}_{vv}) (\mathbf{1}_r \times \mathbf{r}_u)^2] + \\ &+ 4(\mathbf{n} \cdot \mathbf{1}_r)^2 [(\mathbf{n} \cdot \mathbf{r}_{uu}) (\mathbf{n} \cdot \mathbf{r}_{vv}) - (\mathbf{n} \cdot \mathbf{r}_{uv})^2] \end{aligned} \quad (81)$$

lub krócej

$$\mathcal{H}_o(\varrho + \sigma) = a_2 \left(\frac{r+s}{rs} \right)^2 + 2a_1(\mathbf{n} \cdot \mathbf{1}_r) \frac{r+s}{rs} + 4a_0(\mathbf{n} \cdot \mathbf{1}_r)^2,$$

gdzie współczynniki a_0 , a_1 , a_2 reprezentują odpowiednie wyrażenia w nawiasach kwadratowych.

Łatwo sprawdzić bezpośrednim rachunkiem, że

$$a_2 = g(\mathbf{n} \cdot \mathbf{1}_r)^2.$$

Następnie biorąc pod uwagę (62) widzimy, że

$$a_1 = Qg(\mathbf{n} \cdot \mathbf{l}_r) = \left(\frac{r}{\mathbf{n} \cdot \mathbf{r}} H^{\parallel} + \frac{\mathbf{n} \cdot \mathbf{r}}{r} H^{\perp} \right) g(\mathbf{n} \cdot \mathbf{l}_r).$$

Współczynnik a_0 jest wyróżnikiem II formy kwadratowej powierzchni reflektora; wobec tego

$$a_0 = gK_s.$$

Otrzymujemy więc zamiast (81)

$$\mathcal{H}_0(\varrho + \sigma) = g(\mathbf{n} \cdot \mathbf{l}_r)^2 \left[\left(\frac{r+s}{rs} \right)^2 + 2 \frac{r+s}{rs} \left(\frac{r}{\mathbf{n} \cdot \mathbf{r}} H^{\parallel} + \frac{\mathbf{n} \cdot \mathbf{r}}{r} H^{\perp} \right) + 4K_s \right]. \quad (82)$$

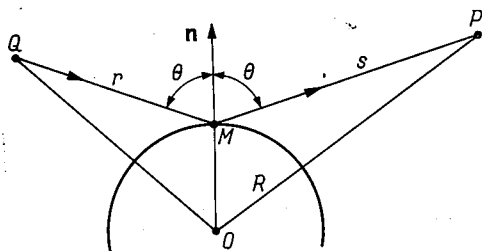
Podstawienie tej wartości hesjanu do (65) daje bezpośrednio wyrażenie (57) dla współczynnika zbieżności. Równoważność wzorów (65) i (57) została więc wykazana.

6. PRZYKŁADY

6.1. Odbicie od powierzchni sferycznej

Rozpatrzmy przypadek, gdy reflektor jest powierzchnią sferyczną o promieniu R (rys. 2).

Źródło fali kulistej umieszczone jest w punkcie Q . Punkt P jest punktem obserwacji.



Rys. 2. Odbicie od powierzchni sferycznej

Obliczymy współczynnik zbieżności D_k . Obieramy jako punkt wyjściowy równanie (58)

$$D_k = \frac{1}{\sqrt{\frac{4r^2s^2}{(r+s)^2} K_s + \frac{2rs}{r+s} \left(\frac{r}{\mathbf{n} \cdot \mathbf{r}} H^{\parallel} + \frac{\mathbf{n} \cdot \mathbf{r}}{r} H^{\perp} \right) + 1}}.$$

Dla reflektora sferycznego mamy

$$\begin{cases} K_s = \frac{1}{R^2}. \\ H^{\parallel} = H^{\perp} = -\frac{1}{R}. \end{cases} \quad (83)$$

Nadto $\mathbf{n} \cdot \mathbf{r} = -r \cos \theta$, gdzie θ jest kątem padania fali. Dostaniemy

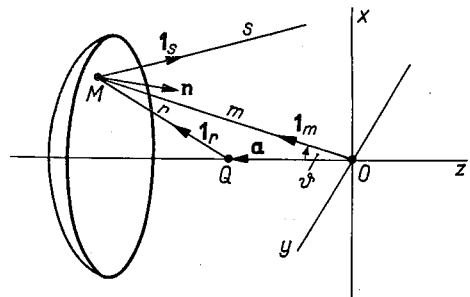
$$D_k = \frac{1}{\sqrt{\frac{4r^2s^2}{R^2(r+s)^2} + \frac{2rs}{R(r+s)} \left(\frac{1}{\cos \theta} + \cos \theta \right) + 1}}. \quad (84)$$

Powyższy wzór (w nieco odmiennej postaci) jest znany w literaturze [2]. Wzór ten znajduje zastosowanie w problemie propagacji w troposferze z uwzględnieniem odbić od powierzchni ziemi.

6.2. Współczynnik zbieżności dla przypadku rozogniskowanej anteny parabolicznej

Układ rozogniskowanej anteny parabolicznej pokazany jest na rys. 3. Dla uproszczenia rachunku zakładamy, że źródło oświetlające pozostaje na osi reflektora. Położenie źródła jest określone wektorem

$$\mathbf{a} = a(\mathbf{1}_m \cos \vartheta - \mathbf{1}_s \sin \vartheta). \quad (85)$$



Rys. 3. Antena paraboliczna w układzie współrzędnych

Kąt ϑ mierzymy względem osi reflektora jak pokazano na rys. 3. Powierzchnia reflektora spełnia równanie

$$\mathbf{m} = m \mathbf{1}_m = \frac{f}{\cos^2 \frac{\vartheta}{2}} \mathbf{1}_m, \quad (86)$$

gdzie f jest ogniskową. Jednostkowy wektor normalny $\mathbf{n}$ wyrazi się następująco:

$$\mathbf{n} = -\mathbf{1}_m \cos \frac{\vartheta}{2} + \mathbf{1}_s \sin \frac{\vartheta}{2}. \quad (87)$$

Wektor pozycyjny $\mathbf{r}$ określający położenie punktu na reflektorze względem źródła Q spełnia zależności

$$\begin{cases} \mathbf{r} = \mathbf{m} - \mathbf{a}, \\ r = \sqrt{m^2 + a^2 - 2macos \vartheta}. \end{cases} \quad (88)$$

Współczynnik zbieżności fali odbitej od reflektora wyrazi się wzorem (57)

$$D(s) = \frac{1}{s \sqrt{\left(\frac{r+s}{rs}\right)^2 + 2 \frac{r+s}{rs} \left(\frac{r}{\mathbf{n} \cdot \mathbf{r}} H^{\parallel} + \frac{\mathbf{n} \cdot \mathbf{r}}{r} H^{\perp}\right) + 4K_s}}.$$

Wykorzystując wzory (39), (40), (42) i (DIII-7) łatwo sprawdzić, że

$$K_s = \frac{1}{4m^2}, \quad (89)$$

$$H^{\parallel} = \frac{\cos \frac{\vartheta}{2}}{2m}, \quad (90)$$

$$H^{\perp} = \frac{1}{2m \cos \frac{\vartheta}{2}}. \quad (91)$$

Nadto mamy

$$\mathbf{n} \cdot \mathbf{r} = \mathbf{n} \cdot (\mathbf{m} - \mathbf{a}) = -(m-a) \cos \frac{\vartheta}{2}. \quad (92)$$

Podstawiając powyższe wyrażenia do (57) dostaniemy następujący ścisły wzór dla współczynnika zbieżności:

$$D(s) = \frac{1}{s \sqrt{\left(\frac{r+s}{rs}\right)^2 - \left(\frac{r+s}{rs}\right) \frac{r^2 + (m-a)^2}{rm(m-a)} + \frac{1}{m^2}}}, \quad (93)$$

gdzie

$$r = \sqrt{m^2 + a^2 - 2ma \cos \vartheta}.$$

Znajdziemy teraz parametry krzywizny fali odbitej K_o oraz H_o . W tym celu najprościej jest przekształcić prawą stronę równania (93) tak, aby uzyskać pod pierwiastkiem trójmian kwadratowy względem s . Otrzymamy

$$D(s) = \frac{1}{\sqrt{\frac{a^2[m(2\cos\vartheta-1)-a]}{m^2r^2(m-a)} s^2 + \frac{2a(m\cos\vartheta-a)}{mr(m-a)} s + 1}}. \quad (94)$$

Stąd bezpośrednio dostaniemy

$$K_o = \frac{a^2}{m^2r^2(m-a)} [m(2\cos\vartheta-1)-a], \quad (95)$$

$$H_o = -\frac{a(m\cos\vartheta-a)}{mr(m-a)}. \quad (96)$$

Na uwagę zasługuje przypadek małych wartości a . Z (94), (95) oraz (96) otrzymujemy wtedy przybliżone wzory

$$D(s) \approx 1 - \frac{a \cos \vartheta}{m^2}, \quad (97)$$

$$K_o \approx \frac{a^2}{m^4} (2 \cos \vartheta - 1), \quad (98)$$

$$H_o \approx -\frac{a \cos \vartheta}{m^2}. \quad (99)$$

W elementarnej optyce geometrycznej dowodzi się, że jeśli ograniczyć się do rozpatrywania promieni biegnących blisko osi zwierciadła parabolicznego, to odbita jest kulista, przy czym zachodzi związek

$$\frac{1}{f-a} = \frac{1}{R} + \frac{1}{f}, \quad (100)$$

gdzie R jest promieniem krzywizny powierzchni tej fali. Ten sam rezultat otrzymujemy z naszych rozważań kładąc we wzorach (95) i (96) $\cos \vartheta \approx 1$. Wtedy spełniony jest warunek $H_o^2 = K_o$, tj. fala odbita jest kulistą, zaś promień krzywizny R powierzchni tej fali spełnia związek (100).

Współczynnik zbieżności D jest funkcją odległości s . Aby nawiązać do potrzeb techniki antenowej, przyjmijmy, że płaszczyzna $z = 0$ jest aperturą anteny i obliczymy $D(s)$ w punktach tej apertury. Można wykazać, że dla dowolnego rozogniskowania a wartość s w aperturze $z = 0$ wynosi

$$s = \frac{m^2 r \cos \vartheta}{m^2 - \mathbf{m} \cdot \mathbf{a}}. \quad (101)$$

W naszym przypadku

$$\mathbf{m} \cdot \mathbf{a} = m a \cos \vartheta \quad (102)$$

i otrzymamy

$$s = \frac{m r \cos \vartheta}{m - a \cos \vartheta}. \quad (103)$$

Podstawiając tę wartość do wzoru (94) otrzymamy

$$D|_{z=0} = \frac{1}{\sqrt{\frac{a^2 \cos^2 \vartheta [m(2 \cos \vartheta - 1) - a]}{(m - a \cos \vartheta)^2 (m - a)} + \frac{2 a \cos \vartheta (m \cos \vartheta - a)}{(m - a \cos \vartheta)(m - a)} + 1}}, \quad (104)$$

lub przybliżenie dla małych wartości a

$$D|_{z=0} \approx 1 - \frac{a \cos^2 \vartheta}{m}. \quad (105)$$

Minimalna wartość współczynnika zbieżności wypada dla $\vartheta = 0$ i wynosi

$$D_{min} = 1 - \frac{a}{f}. \quad (106)$$

W miarę wzrostu kąta ϑ współczynnik D rośnie. Jest to zrozumiałe, gdyż wtedy maleje długość promienia s . Dla $\vartheta = \frac{\pi}{2}$ mamy $s = 0$ i $D = 1$.

7. ZAKOŃCZENIE

W pracy pokazano ogólną metodę obliczania współczynnika zbieżności. Wybrano proste przykłady obliczeniowe ilustrujące możliwości zastosowań. Przykłady bardziej skomplikowane — ze względu na obszerność rachunków — powinny być przedmiotem oddzielnych opracowań.

Przemysłowy Instytut Telekomunikacji

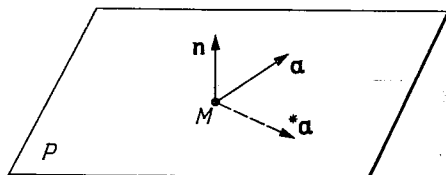
WYKAZ LITERATURY

1. Biernacki M.: Geometria różniczkowa t. I/II. PWN, Warszawa 1955.
2. Bremmer H.: Terrestrial Radiowaves. Elsevier Publish. Comp. 1949.
3. Silver S.: Microwave Antenna Theory and Design. Mc Graw-Hill Book Comp. Inc. 1949.

Dodatek I

WEKTORY ZWIERCIADLANE

Analiza pola odbitego od reflektora anteny nabiera znacznie większej przejrzystości i zwięzłości, jeśli wprowadzić pojęcie wektorów zwierciadlanych. Obierzmy pewną płaszczyznę P , $\mathbf{n}$ jest jednostkowym wektorem normalnym do tej płaszczyzny (rys. 4).



Rys. 4. Odbicie zwierciadlane wektora $\mathbf{a}$ w płaszczyźnie P

Niech $\mathbf{a}$ będzie dowolnym wektorem. Wektorowi $\mathbf{a}$ można przyporządkować wektor $^*\mathbf{a}$ będący odbiciem zwierciadlanym $\mathbf{a}$ w płaszczyźnie P . Wektor $^*\mathbf{a}$ nazwiemy wektorem zwierciadlanym. Spełnia on zależność

$$^*\mathbf{a} = \mathbf{a} - 2(\mathbf{n} \cdot \mathbf{a})\mathbf{n}. \quad (\text{DI-1})$$

Opierając się na tej definicji łatwo sprawdzić następujące znane własności wektorów zwierciadlanych:

- 1) $^*(^*\mathbf{a}) = \mathbf{a}$;

2) $|\mathbf{a}| = |\mathbf{a}|$; w szczególności wektor zwierciadlany względem wektora jednostkowego pozostaje wektorem jednostkowym;

3) $\mathbf{a} \cdot \mathbf{b} = \mathbf{a} \cdot \mathbf{b}$;

4) $\mathbf{a} \times \mathbf{b} = -(\mathbf{a} \times \mathbf{b})$;

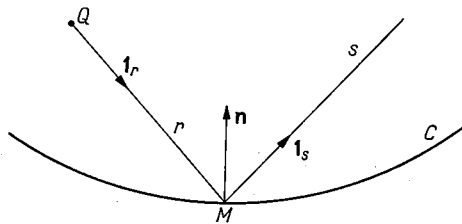
5) $\mathbf{a} \times (\mathbf{b} \times \mathbf{c}) = [\mathbf{a} \times (\mathbf{b} \times \mathbf{c})]$;

6) $(\mathbf{a}, \mathbf{b}, \mathbf{c}) = -(\mathbf{a}, \mathbf{b}, \mathbf{c})$.

Dodatek II

KRZYWIZNA PRZEKROJU NORMALNEGO POWIERZCHNI REFLEKTORA W PŁASZCZYZNIE PADANIA

Przekrój normalny powierzchni reflektora w płaszczyźnie padania pokazany jest na rys. 5.



Rys. 5. Przekrój normalny reflektora w płaszczyźnie padania

Oznaczmy przez σ długość łuku na linii przekroju C. Wykorzystamy znany z geometrii różniczkowej związek dla krzywizny linii

$$H_{||} = -\frac{d\mathbf{r}}{d\sigma} \cdot \frac{d\mathbf{n}}{d\sigma}. \quad (\text{DII-1})$$

Symbol $H_{||}$ oznacza tu krzywiznę przekroju normalnego w płaszczyźnie padania. Jednostkowy wektor $\mathbf{n}$ jest normalny do linii i skierowany w stronę środka krzywizny linii. W naszym przypadku jest to jednocześnie wektor normalny do powierzchni reflektora.

Obliczmy na wstępie $\frac{d\mathbf{r}}{d\sigma}$. Wektor ten spełnia następujący układ równań:

$$\begin{cases} \frac{d\mathbf{r}}{d\sigma} \cdot (\mathbf{n} \times \mathbf{r}) = 0, \\ \frac{d\mathbf{r}}{d\sigma} \cdot \mathbf{n} = 0, \\ \left(\frac{d\mathbf{r}}{d\sigma} \right)^2 = 1. \end{cases} \quad (\text{DII-2})$$

Stąd

$$\frac{d\mathbf{r}}{d\sigma} = \frac{(\mathbf{n} \times \mathbf{r}) \times \mathbf{n}}{|\mathbf{n} \times \mathbf{r}|} = \frac{\mathbf{r} + \mathbf{r}}{2|\mathbf{n} \times \mathbf{r}|}. \quad (\text{DII-3})$$

Następnie mamy

$$\mathbf{n} = \frac{1}{\sqrt{g}} \left(\frac{\partial \mathbf{r}}{\partial u} \times \frac{\partial \mathbf{r}}{\partial v} \right). \quad (\text{DII-4})$$

Stąd dostaniemy

$$\frac{d\mathbf{n}}{d\sigma} = \frac{1}{\sqrt{g}} \left[\frac{\partial}{\partial u} \left(\frac{d\mathbf{r}}{d\sigma} \right) \times \frac{\partial \mathbf{r}}{\partial v} + \frac{\partial \mathbf{r}}{\partial u} \times \frac{\partial}{\partial v} \left(\frac{d\mathbf{r}}{d\sigma} \right) \right] + \dots \quad (\text{DII-5})$$

Wielokropkiem oznaczono składniki nie wpływające na wynik końcowy.

Podstawiając w równaniu (5) wyrażenie (3) otrzymamy

$$\frac{d\mathbf{n}}{d\sigma} = \frac{1}{2\sqrt{g}|\mathbf{n} \times \mathbf{r}|} \left[\frac{\partial^* \mathbf{r}}{\partial u} \times \frac{\partial \mathbf{r}}{\partial v} + \frac{\partial \mathbf{r}}{\partial u} \times \frac{\partial^* \mathbf{r}}{\partial v} \right] + \dots \quad (\text{DII-6})$$

Podstawiając (3) i (6) do (1) dostaniemy wzór dla krzywizny $H^{\parallel}$ w postaci

$$H^{\parallel} = - \frac{1}{4\sqrt{g}(\mathbf{n} \times \mathbf{r})^2} [(\mathbf{r}, {}^* \mathbf{r}_u, \mathbf{r}_v) + (\mathbf{r}, \mathbf{r}_u, {}^* \mathbf{r}_v) + ({}^* \mathbf{r}, {}^* \mathbf{r}_u, \mathbf{r}_v) + ({}^* \mathbf{r}, \mathbf{r}_u, {}^* \mathbf{r}_v)]. \quad (\text{DII-7})$$

Dodatek III

KRZYWIZNA PRZEKROJU NORMALNEGO POWIERZCHNI REFLEKTORA W PŁASZCZYZNIE PROSTOPADŁEJ DO PŁASZCZYZNY PADANIA

Rachunek przebiega podobnie jak w przypadku krzywizny $H^{\parallel}$. Wykorzystujemy wzór

$$H^{\perp} = - \frac{d\mathbf{r}}{d\sigma} \cdot \frac{d\mathbf{n}}{d\sigma}, \quad (\text{DIII-1})$$

w którym symbol $H^{\perp}$ oznacza krzywiznę przekroju normalnego w płaszczyźnie prostopadłej do płaszczyzny padania.

Obliczamy $\frac{d\mathbf{r}}{d\sigma}$, Wektor ten spełnia następujący układ równań:

$$\begin{cases} \frac{d\mathbf{r}}{d\sigma} \cdot \mathbf{n} = 0, \\ \frac{d\mathbf{r}}{d\sigma} \cdot \mathbf{r} = 0, \\ \left(\frac{d\mathbf{r}}{d\sigma} \right)^2 = 1. \end{cases} \quad (\text{DIII-2})$$

Stąd otrzymujemy

$$\frac{d\mathbf{r}}{d\sigma} = \frac{\mathbf{n} \times \mathbf{r}}{|\mathbf{n} \times \mathbf{r}|}. \quad (\text{DIII-3})$$

Następnie mamy

$$\mathbf{n} = \frac{1}{\sqrt{g}} \left(\frac{\partial \mathbf{r}}{\partial u} \times \frac{\partial \mathbf{r}}{\partial v} \right), \quad (\text{DIII-4})$$

co pozwala napisać

$$\frac{dn}{d\sigma} = \frac{1}{\sqrt{g}} \left[\frac{\partial}{\partial u} \left(\frac{d\mathbf{r}}{d\sigma} \right) \times \frac{\partial \mathbf{r}}{\partial v} + \frac{\partial \mathbf{r}}{\partial u} \times \frac{\partial}{\partial v} \left(\frac{d\mathbf{r}}{d\sigma} \right) \right] + \dots \quad (\text{DIII-5})$$

Wielokropkiem oznaczono składniki nie wpływające na wynik końcowy.

Podstawiając w (5) wyrażenie (3) otrzymamy

$$\frac{dn}{d\sigma} = \frac{1}{\sqrt{g}|\mathbf{n} \times \mathbf{r}|} \left[\frac{\partial}{\partial u} (\mathbf{n} \times \mathbf{r}) \times \frac{\partial \mathbf{r}}{\partial v} + \frac{\partial \mathbf{r}}{\partial u} \times \frac{\partial}{\partial v} (\mathbf{n} \times \mathbf{r}) \right] + \dots \quad (\text{DIII-6})$$

Biorąc pod uwagę wzory (3) i (6), które podstawiamy do (1), łatwo sprawdzić, że krzywizna $H^\perp$ wyrazi się następująco:

$$H^\perp = - \frac{\mathbf{n} \times \mathbf{r}}{\sqrt{g}(\mathbf{n} \times \mathbf{r})^2} \cdot [(\mathbf{r} \cdot \mathbf{r}_u) \mathbf{n}_v - (\mathbf{r} \cdot \mathbf{r}_v) \mathbf{n}_u]. \quad (\text{DIII-7})$$

Wzór ten można wykorzystać również do obliczenia krzywizny $H^\parallel$ zgodnie z zależnością

$$H^\parallel = 2H_s - H^\perp. \quad (\text{DIII-8})$$

S. POGORZELSKI

WAVE CONVERGENCE COEFFICIENT IN ANTENNA PROBLEMS

Summary

This paper deals with calculation of the wave convergence coefficient arising in reflector antenna theory and in other similar problems. A spherical wave incident upon the reflector is assumed. Two calculation methods are presented. In the first case a well known relation in geometrical optics is taken as a starting point

$$D(s) = \sqrt{\frac{K(s)}{K(s_0)}}.$$

This equation relates the wave convergence coefficient D to the Gaussian curvature K of the equieiconal surface. The variable s denotes arc length on a ray. After a simple transformation a new formula is easily to obtain

$$D(s) = \frac{1}{\sqrt{K_0 s^2 - 2H_0 s + 1}}.$$

The parameters K_0 and H_0 respectively denote Gaussian curvature and mean curvature of the equieiconal surface of the reflected field. Both curvatures are to be taken at the reflection point. The values of K_0 and H_0 are determined in relation to local curvature parameters of the reflector surface. There is often difficult to determine the curvature parameters explicitly. To avoid the difficulties a transformed expression for $D(s)$ is established which permits relatively easy computations in any particular case.

To calculate the convergence coefficient another method can be also used. This method bases upon a known integral formula which determines the reflected field. Now the stationary phase method can be applied to evaluate the field integral. This procedure leads to a new interesting expression for $D(s)$. The most significant factor occurring in this expression is the Hessian determinant of the phase function taken at the reflection point. It is directly proved that all the obtained expressions for $D(s)$ are equivalent.

For illustration two particular examples are considered. The first example deals with a problem of spherical wave reflection from a spherical surface. A well known formula is obtained which is frequently used in tropospheric propagation. The other example belongs to spherical wave reflection from a parabolic antenna reflector in a case of defocalisation.

S. POGORZELSKI

COEFFICIENT DE CONVERGENCE D'ONDE DANS LA TECHNIQUE DES ANTENNES

Résumé

Le sujet du travail présenté c'est le calcul de coefficient de convergence d'onde, qui se trouve dans la théorie des antennes à réflecteurs ainsi que dans d'autres problèmes. On suppose une onde sphérique tombant sur un réflecteur. On utilise deux méthodes de calcul. En premier cas le point d'issue c'est une dépendance bien connue de l'optique géométrique:

$$D(s) = \sqrt{\frac{K(s)}{K(s_0)}}.$$

Cette équation donne la liaison du coefficient nommé ci-dessus D avec la courbure de Gauss K de la surface d'eiconal égal. La variable s c'est la longueur de l'arc sur le rayon. Après une simple transformation on obtient légèrement une formule nouvelle:

$$D(s) = \frac{1}{\sqrt{K_0 s^2 - 2H_0 s + 1}}.$$

Les caractéristiques K_0 et H_0 déterminent courbure de Gauss et courbure moyenne de surface d'eiconal égal du champ réfléchi. Il faut prendre les deux courbures ci-dessus dans le point de réflexion. On détermine les valeurs K_0 et H_0 en dépendance des caractéristiques locales des courbures de la surface du réflecteur. Souvent il est difficile de déterminer explicitement ces caractéristiques de courbure. Pour éviter les difficultés on a utilisé une formule transformée de $D(s)$, qui permet un calcul relativement plus facile en chaque cas particulier.

Le calcul de coefficient de convergence d'onde peut être effectué aussi à l'aide d'une autre méthode. Celle-ci se base sur une formule intégrale bien connue déterminant le champ réfléchi. Pour calculer l'intégral on peut employer la méthode de la phase stationnaire. Ce procédé donne une très intéressante nouvelle formule de $D(s)$. Le facteur le plus important se trouvant dans cette formule c'est le déterminant de Hess de la fonction de phase, pris dans le point de réflexion. On prouve directement que toutes les formules obtenues de $D(s)$ sont équivalentes.

Pour illustrer le problème on présente deux exemples particuliers. Le premier traite le problème de la réflexion de l'onde sphérique tombant sur la surface sphérique. On obtient une formule bien connue utilisée souvent dans les problèmes de propagation troposphérique. Le second concerne la réflexion de l'onde sphérique tombant sur un réflecteur parabolique d'antenne en cas de défocalisation.

S. POGORZELSKI

WELLENKONVERGENZKOEFFIZIENT IN PROBLEMEN DER ANTENNENTECHNIK

Zusammenfassung

Die Aufgabe dieser Arbeit besteht darin, den Wellenkonvergenzkoeffizienten, welcher in der Theorie der Reflektorantennen und in ähnlichen Problemen erscheint, zu bestimmen. Es wird eine sphärische Welle fallend auf die Reflektorfläche angenommen. Es sind zwei Rechnungsmethoden verfügbar. Die erste Methode fängt von einer in der geometrischen Optik gut bekannten Relation an

$$D(s) = \sqrt{\frac{K(s)}{K(s_0)}}.$$

Diese Gleichung knüpft den Wellenkonvergenzkoeffizienten mit der Gausschen Krümmung K der Equeikonalfäche zusammen. Die Veränderliche s bezeichnet darin die Bogenlänge des Strahles. Nach einfacher Transformation eine neue Gleichung für $D(s)$ ist leicht zu bekommen

$$D(s) = \frac{1}{\sqrt{K_0 s^2 - 2H_0 s + 1}}.$$

Die Parameter K_0 und H_0 bezeichnen darin die Gaussche und die mittlere Krümmung beziehungsweise. Beide Krümmungen betreffen die Equeikonalfäche der reflektierten Welle im Punkte der Reflexion. Die Werte der Krümmungen K_0 und H_0 werden in Abhängigkeit von den lokalen Krümmungsparametern der Reflektorfläche bestimmt. Diese Resultate führen zur Aufstellung eines neuen Ausdrucks für $D(s)$. Dieser ist von den lokalen Krümmungsparametern abhängig. Es ist in vielen Fällen schwierig die Krümmungsparameter explicite zu berechnen. Um diese Schwierigkeiten zu beseitigen wird der Ausdruck für $D(s)$ weiter transformiert, so dass relativ leichte Rechnungen bleiben in jedem besonderen Fall übrig.

Die zweite Methode der Bestimmung des Wellenkonvergenzkoeffizienten hat zum Ausgangspunkt eine bekannte Integralformel, welche das reflektierte Feld beschreibt. Man kann das Integral mittels der Methode der stationären Phase berechnen. Dieses Verfahren führt zu einem neuen interessanten Ausdruck für $D(s)$. Der wichtigste Faktor, welcher in diesem Ausdruck erscheint, ist die Hessesche Determinante der Phasenfunktion bezogen auf den Reflexionspunkt. Es wird unmittelbar bewiesen, dass alle erhaltenen Formeln für $D(s)$ gleichbedeutend sind.

Es werden zwei Beispiele als Illustration herangeführt. Das erste Beispiel beschäftigt sich mit der Reflexion einer Kugelwelle an einer sphärischen Fläche. Es wird eine bekannte Formel aufgestellt, die in Problemen der Wellenausbreitung in der Troposphäre oft benutzt ist. Das zweite Beispiel betrifft die Reflexion einer Kugelwelle an einem parabolischen Antennenreflektor im Falle einer Defokalisierung.

С. ПОГОЖЭЛЬСКИ

КОЭФФИЦИЕНТ СХОДИМОСТИ ВОЛНЫ В АНТЕННОЙ ТЕХНИКЕ

Резюме

Темой разработки является расчёт коэффициента сходимости волны, который выступает в теории рефлекторных антенн и в других похожих вопросах. Принимается сферическая волна, падающая на рефлектор. Представлено два метода вычисления. В первом случае исходной точкой является хорошо известная в геометрической оптике зависимость:

$$D(s) = \sqrt{\frac{K(s)}{K(s_0)}}.$$

Это уравнение связывает коэффициент сходимости волны D с кривизной Гаусса K поверхности равного эйконала. Переменная s обозначает длину дуги на радиусе. После простого преобразования мы легко получаем новую формулу:

$$D(s) = \frac{1}{\sqrt{K_0 s^2 - 2H_0 s + 1}}.$$

Параметры K_0 и H_0 обозначают соответственно кривизну Гаусса и кривизну средней поверхности равного эйконала отраженного поля. Обе кривизны следует принимать в точке отражения. Значения K_0 и H_0 определяются в зависимости от локальных параметров кривизны поверхности рефлектора. Часто трудно определить *explicite* эти параметры кривизны. С целью избежать трудностей, введено преобразованное выражение на $D(s)$, которое позволяет относительно легко произвести вычисления в каждом частном случае.

Для вычисления коэффициента сходимости можно применить еще другой метод. Этот метод основан на известной интегральной формуле, определяющей отраженное поле. С целью расчета интеграла можно применить метод стационарной фазы. Это приводит к новому интересному выражению на $D(s)$. Самым важным фактором, выступающим в этом выражении, является определитель Гесса функции фазы, принимаемый в точке отражения. Доказывается непосредственно, что все полученные выражения на $D(s)$ эквивалентны.

С целью иллюстрирования обсуждаются два частных примера. Первый из них занимается проблемой отражения сферической волны от сферической поверхности. Получается хорошо известная формула, часто применяемая в проблемах тропосферного распространения волн. Второй пример касается отражения сферической волны от антенного параболического рефлектора в случае дефокусировки.

MACIEJ NAŁĘCZ

Dokładność i czułość przetwornika halotronowego¹⁾ przy pomiarach statycznych przemieszczeń mechanicznych

Rękopis dostarczono 16.11.1964

Zjawisko Halla umożliwia dokonanie w prosty sposób pomiaru przesunięć i wibracji mechanicznych. Jeśli halotron jest przesuwany mechanicznie w niejednorodnym polu magnetycznym, równocześnie będzie się zmieniać napięcie Halla. Dla osiągnięcia najkorzystniejszych warunków pracy, gradient pola magnetycznego powinien być jak największy i stały.

Wykonano badania w celu określenia dokładności i czułości przetworników halotronowych przy statycznych pomiarach mechanicznych. Prześledzono szereg uchybów, jak (a) niesymetryczne położenie elektrod Halla, (b) napięcie indukowane w przewodach pętli napięciowej, (c) napięcie wywołane polem ziemskim lub zewnętrznym polem magnetycznym, (d) nieliniowość funkcji $U_H = f(B)$, (e) wpływ temperatury na liniowość funkcji $U_H = f(B)$, (f) nieliniowość funkcji $B = f(x)$, (g) wpływ temperatury na stałą Halla i oporność wewnętrzną, (h) wpływ elementów ferromagnetycznych, (i) uchyb wywołany przez urządzenie mechaniczne. Wyniki wskazują na to, że istnieje możliwość budowy przetworników halotronowych o dokładności lepszej niż 1,5% i czułości większej, niż wykazują inne przetworniki obecnie stosowane. Stosując specjalne obwody magnetyczne o gradiencie indukcji magnetycznej 10^4 Gs/mm i halotроны o czułości (0,1 do 1,0) V/ 10^4 Gs w specjalnych układach, można uzyskać czułość rzędu (0,1 do 10,0) V/mm.

1. WSTĘP

Artykuł niniejszy jest poświęcony rozważaniu podstawowych uchybów występujących przy pomiarach statycznych przemieszczeń mechanicznych za pomocą halotyonu. Wprowadzono pojęcie klasy przetwornika halotronowego, jego czułości oraz podano wyniki pomiarów, określające możliwości zastosowań tego typu przetworników.

Pomiar przemieszczeń odbywa się poprzez pomiar indukcji magnetycznej. Jeśli bowiem płytkę z materiału przewodzącego (najczęściej półprzewodnika) umieścimy w polu magnetycznym o indukcji B prostopadłej do

¹⁾ Nazwę halotron przyjęto [8] w Polsce używać w miejsce amerykańskiej czy niemieckiej nazwy generator, rosyjskiej — датчик, czeskiej — sonda lub angielskiej — unit. Nazwę przetwornik halotronowy wprowadza autor, aby podkreślić, że tu halotron jest zastosowany do przetwarzania wielkości nieelektrycznej w elektryczną.

jej powierzchni oraz przepuścimy przez nią prąd I_p , wówczas na krawędziach płytki powstanie napięcie zwane napięciem Halla U_H (rys. 1):

$$U_H = R_H \frac{I_p \cdot B}{d}, \quad (1)$$

gdzie:

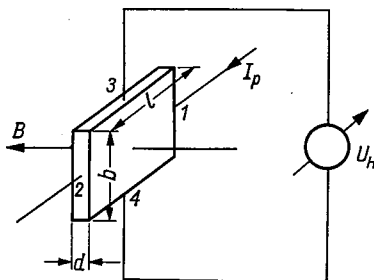
$R_H = \frac{3\pi}{8} \cdot \frac{1}{n \cdot e}$ — stała Halla (dodatnia dla dziur, ujemna dla elektronów),

n — liczba nośników prądu w jednostce objętości,

e — ładunek dziury lub elektronu,

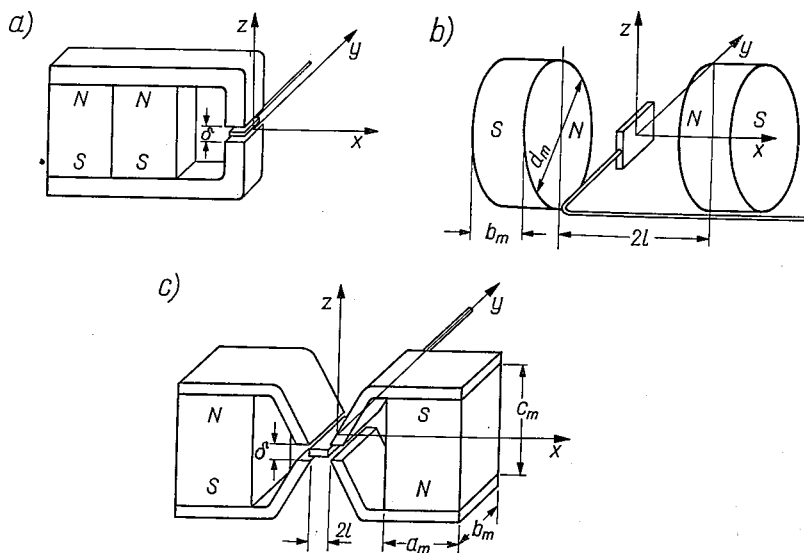
d — grubość płytki.

Gdy równocześnie występują w półprzewodniku dziury i elektrony, wyrażenie na R_H jest bardziej złożone.



Rys. 1. Zjawisko Halla

Pomimo, że niniejsza zasada została odkryta przez E. H. Halla już w 1879 roku [1], właściwy rozwój zastosowań efektu Halla rozpoczął się z chwilą odkrycia półprzewodników.



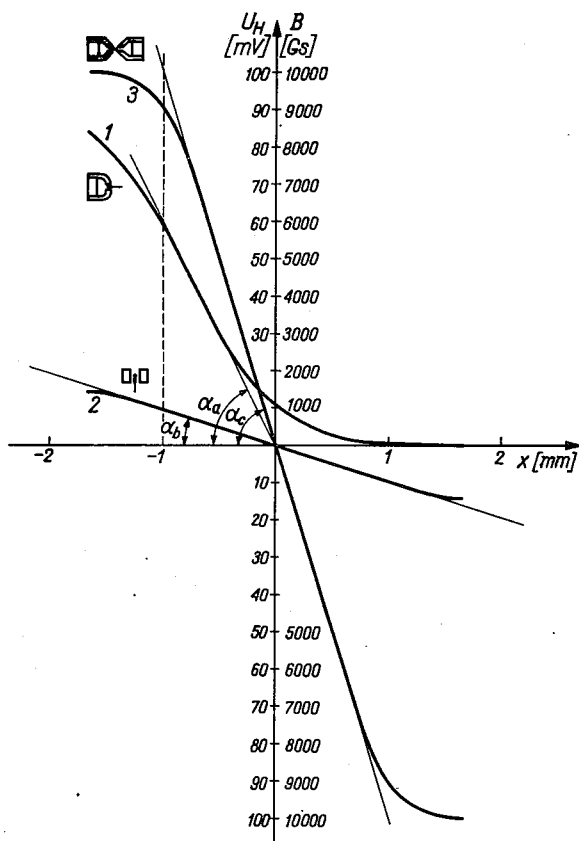
Rys. 2. Obwody magnetyczne przetwornika halotronowego

W 1948 roku Pearson [2] opublikował konstrukcję gausmiera z halotronem germanowym. Obecnie wiele zakładów produkuje na skalę przemysłową gausmiery umożliwiające pomiar coraz to mniejszych pól magnetycznych [3].

Jeśli halotron będzie przemieszczany w niejednorodnym polu magnetycznym, wówczas napięcie Hala U_H będzie funkcją tego przemieszczenia x .

$$U_H = f(x). \quad (2)$$

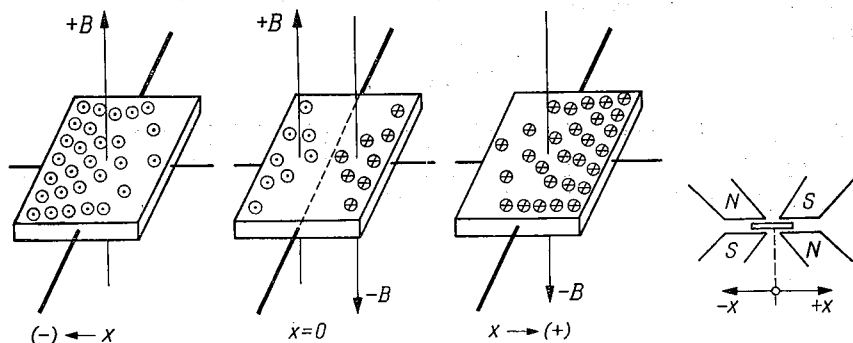
W poprzednich artykułach [4], [5] została bardziej szczegółowo rozpatrzona zasada pomiaru oraz niektóre układy do wytwarzania pól magnetycznych o stałym gradiencie w określonym przedziale przemieszczania



Rys. 3. Funkcje $B = f(x)$ i $U_H = f(x)$ dla obwodów magnetycznych przedstawionych na rys. 2 (przebieg 1 odpowiada rysunkowi 2a, przebieg 2 — rysunkowi 2b, przebieg 3 — rysunkowi 2c)

halotronu. Układy te, przedstawione na rys. 2a, b, c, pozwalają na używanie charakterystycznych przebiegów indukcji magnetycznej lub napięcia Hala (rys. 3), w funkcji przemieszczania płytki x .

Przy przesuwaniu płytki w pobliżu zera indukcji magnetycznej B (rys. 2c) przez płytkę przenikają przeciwnie skierowane strumienie magnetyczne (rys. 4). W ten sposób płytka pracuje w układzie różnicowym. W tym układzie istnieje możliwość pomiaru bardzo małych przemieszczeń lub drgań mechanicznych.



Rys. 4. Halotron w układzie pola magnetycznego z rys. 2c

2. UCHYBY PRZETWORNIKA HALOTRONOWEGO

Uchyby przetwornika halotronowego wynikające ze zjawisk fizycznych występujących w samej płytce półprzewodnikowej były rozpatrywane wielokrotnie tak w literaturze krajowej [6], [7], [8], [9], jak i zagranicznej [10], [11], [12], [13], [14], [15], [16], [17]. Dodatkowo należy rozpatrzeć uchyby wynikające z konstrukcji mechanicznej i układu magnetycznego przetwornika.

Przy rozważaniu uchybów halotronu oparto się głównie na artykule Wasilczenki, Sientiurina i Sockowa [15].

Aby skonstruować przetwornik o określonej klasie dokładności N_p , należy rozpatrzeć następujące uchyby warunkujące tę dokładność:

$$N_p = \frac{U_{Ha} + U_{Hb} + U_{Hc} + \Delta U_{H1} + \Delta U_{H2} + \Delta U_{H3} + \Delta U_{H4} + \Delta U_{H5} + \Delta U_{H6}}{U_{obc\ max}}, \quad (3)$$

gdzie:

U_{Ha} — napięcie wywołane asymetrią elektrod napięciowych halotronu,

U_{Hb} — napięcie indukowane w przewodach obwodu napięciowego,

U_{Hc} — napięcie wywołane natężeniem pola ziemskiego lub wpływem stałych pól magnetycznych zewnętrznych,

ΔU_{H1} — uchyb wywołany przez nieliniowość funkcji $U_{obc} = f(B)$, (dobór oporności w obwodzie napięciowym halotronu, efekt Gaussa),

ΔU_{H2} — uchyb wywołany wpływem temperatury na nieliniowość funkcji $U_{obc} = f(B)$,

- ΔU_{H3} — uchyb wywołany przez nieliniowość funkcji $B = f(x)$ (gdzie x — przemieszczenie halotronu), spowodowaną układem pola magnetycznego,
- ΔU_{H4} — uchyb wywołany wpływem temperatury na stałą Halla R_H i oporność wewnętrzną r_{wew} ,
- ΔU_{H5} — uchyb wywołany przez elementy ferromagnetyczne,
- ΔU_{H6} — uchyb wywołany przez układ mechaniczny,
- $U_{obc\ max}$ — maksymalne napięcie Halla na obciążeniu w danych warunkach pomiaru.

Poniżej pokrótce zostaną rozpatrzone poszczególne z wymienionych uchybów.

Napięcie asymetrii, U_{Ha} , jest związane przede wszystkim z własnościami elektrod napięciowych halotronu. Elektrody napięciowe powinny mieć jak najmniejszą oporność, brak własności prostujących (istotnych przy pracy przy zmiennych przebiegach) oraz powinny być przylutowane symetrycznie.

Asymetrię kontaktów napięciowych można wyrazić przy pomocy współczynnika r_a określonego zależnością

$$U_{Ha} = r_a I_p, \quad (4)$$

gdzie:

U_{Ha} — napięcie na elektrodach Halla bez pola magnetycznego,

I_p — prąd pierwotny.

W powyższym równaniu współczynnik asymetrii r_a wyraża oporność płytki wzdłuż kierunku prądu I_p między elektrodami Halla. Wynika stąd wniosek, że napięcie asymetrii (przy tej samej technologii elektrod) jest większe dla płytek wykonanych z materiału o większej oporności właściwej.

Istnieje wiele metod kompensacji napięcia asymetrii [7].

Produkowane są już obecnie halotроны o $r_a < 0,1$ mV/A. Najczęściej jednak współczynnik asymetrii $r_a \approx (0,1 \text{ do } 5)$ mV/A, co odpowiada $U_{Ha} \approx (0,1 \text{ do } 0,5) \cdot 10^{-3}$ V.

Napięcie indukowane w przewodach, U_{Hb} , występuje przy pracy halotronu przy prądzie zmiennym. Powstaje również przy szybkozmiennych drganiach mechanicznych samej płytki.

Przewody odprowadzające napięcie Halla tworzą pętlę, przez którą przenika zmienny strumień magnetyczny wywołując napięcie indukowane o wartości

$$U_{Hb} = -\frac{d\Phi}{dt} = -S \frac{dB}{dt}, \quad (5)$$

gdzie:

U_{Hb} — napięcie indukowane w pętli,

Φ — strumień indukcji magnetycznej,

S — powierzchnia skuteczna pętli.

Powierzchnia skuteczna pętli S jest istotnym współczynnikiem warunkującym szkodliwy wpływ napięcia U_{Hb} i dlatego powinna być jak najmniejsza.

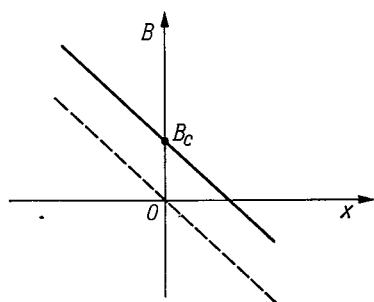
W produkowanych halotronach $S = (0,02 \text{ do } 0,04) \text{ cm}^2$, co przy indukcjach rzędu 3.000 Gs i częstotliwości $f = 50 \text{ Hz}$ powoduje powstanie napięć $U_{Hb} = (0,15 \text{ do } 0,3) \cdot 10^{-3} \text{ V}$.

Napięcie od pola magnetycznego ziemskiego, U_{Hc} , lub od wpływu innego stałego pola magnetycznego powoduje powstanie stałej różnicy potencjałów na elektrodach Halla o wartości:

$$U_{Hc} = K_c \cdot B_c. \quad (6)$$

Wartość tego napięcia zależy od wartości składowej pola magnetycznego prostopadłej do płytki B_c . Wpływ ten jest zazwyczaj pomijalny, w szczególności przy pracy przetwornika w silnych polach magnetycznych.

Na przykład w przetworniku przedstawionym na rys. 2c zewnętrzne stałe pole magnetyczne może przesunąć punkt przejścia przez zero funkcji $B = f(x)$. Zostało to przedstawione na rys. 5 przy założeniu, że jednorodność stałego pola magnetycznego nie została zniekształcona w pobliżu magnesów trwałych.



Rys. 5. Wpływ zewnętrznego jednorodnego pola magnetycznego na przesunięcie funkcji $B = f(x)$

Przy założeniu, że $B_c = 0,5 \text{ Gs}$ oraz że halotron jest wykonany z InAs, $U_{Hc} = 0,005 \text{ do } 0,01) \cdot 10^{-3} \text{ V}$.

Uchyb od nieliniowości funkcji $U_H = f(B)$, ΔU_{H1} .

Dla obwodu napięciowego halotyonu można napisać następujące zależności:

$$U_H = r_{wev B} \cdot I_H + U_{obc}, \quad (7)$$

$$U_H = U_{obc} \left(\frac{r_{wev B}}{r_{obc}} + 1 \right), \quad (8)$$

$$U_{obc} = \frac{1}{\frac{r_{wev B}}{r_{obc}} + 1} U_H, \quad (9)$$

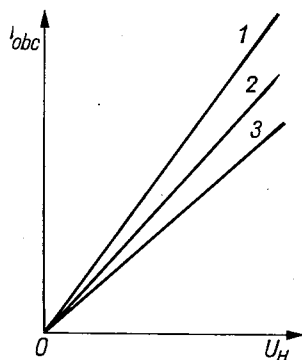
gdzie:

- I_H — prąd w obwodzie napięciowym halotronu,
- U_{obc} — napięcie na obciążeniu halotronu,
- r_{obc} — oporność obciążenia,
- $r_{wew B}$ — oporność wewnętrzna halotronu zależna od natężenia pola magnetycznego B .

Z uproszczonego równania (9) wynika,

$$U_{obc} = k_1 \cdot U_H. \quad (10)$$

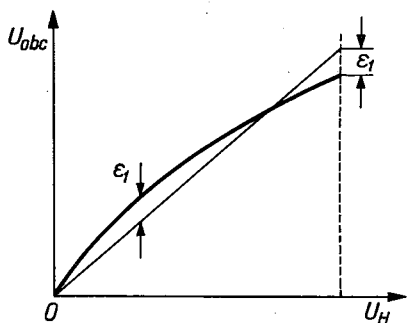
Znaczy to, że gdy oporność obciążenia r_{obc} rośnie, wówczas współczynnik k_1 również rośnie. Stąd linie proste $U_{obc} = k_1 U_H$ obracają się ku górze



Rys. 6. Funkcja $U_{obc} = f(U_H)$ dla różnych wartości r_{obc} ; 1 — r_{obc1} , 2 — $r_{obc2} < r_{obc1}$, 3 — $r_{obc3} < r_{obc2}$

(rys. 6). Lecz jednocześnie $r_{wew B}$ rośnie również ze względu na zjawisko Gaussa (magnetorezystancji). Dlatego funkcja $U_{obc} = f(U_H)$ obraca się ku dołowi.

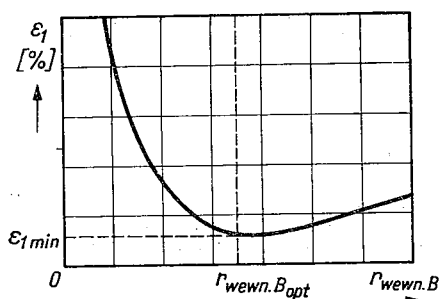
Ponieważ funkcja $r_{wew} = f(B)$ wynikająca ze zjawiska Gaussa jest nieliniowa, powoduje również nieliniowy uchyb ε_1 (rys. 7).



Rys. 7. Funkcja $U_{obc} = f(U_H)$ i jej odchylenie od liniowości

Uchyb ε_1 jest większy, gdy oporność obciążenia r_{obc} jest mniejsza, ponieważ wówczas funkcja $U_{obc} = f(U_H)$ jest bardziej czuła na zmiany oporności wewnętrznej $r_{wew B}$.

Można stwierdzić z wykresu $\varepsilon_1 = f(r_{wewB})$ (rys. 8), że gdy opór wewnętrzny rośnie, ε_1 maleje. Dla pewnej wartości $r_{wew\ opt}$ wartość ε_1 jest minimalna.



Rys. 8. Zależność współczynnika ε_1 od oporności wewnętrznej $r_{wew\ B}$ [15]

Gdy wartość r_{wewB} dalej wzrasta, funkcja $\varepsilon_1 = f(r_{wewB})$ nieco rośnie [15]. Odchylenie od liniowości będzie tym mniejsze, im mniejszy dla danej próbki będzie stosunek

$$\frac{\Delta \varrho}{\varrho} = f(B), \quad (11)$$

gdzie:

- ϱ — oporność właściwa próbki, z której wykonany jest halotron,
- $\Delta \varrho$ — zmiana oporności właściwej, gdy halotron zostanie umieszczony w polu B .

Uchyb pomiaru pola magnetycznego w określonych granicach, np. $B = 0$ i $B = 10^4$ Gs, wynikający z określonej nieliniowości funkcji $U_H = f(B)$ wynosi

$$\Delta U_{H1} = \varepsilon_1 U_{H\ max}. \quad (12)$$

ΔU_{H1} odpowiada pewnej określonej wartości r_{obc} powiązanej z ε_1 zależnością przedstawioną na rys. 8. Uchyb ΔU_{H1} zwykle jest rozpatrywany łącznie z ΔU_{H2} , gdyż na oporność wewnętrzną halotronu równocześnie wpływają zmiany temperatury.

Uchyb od wpływu temperatury na nieliniowość funkcji $U_{obc} = f(B)$, ΔU_{H2} , wynika ze zmiany oporności wewnętrznej halotronu r_{wew} w funkcji temperatury.

$$r_{wew} = r_{wew20} [1 + \alpha_t (\Theta - 20^\circ)], \quad (13)$$

gdzie:

- r_{wew20} — temperatura pracy,
- α_t — współczynnik cieplny oporności,
- Θ — wewnętrzna oporność halotronu przy 20°C ,

Zmiana r_{wev} prowadzi do zmiany współczynnika ε_1 (rys. 8) i podwyższa uchyb nieliniowy ΔU_{H1} .

Aby uwzględnić powyższy wpływ temperatury, należy określić wielkość przypuszczalnej zmiany r_{wev} i stąd znaleźć nowy współczynnik nieliniowości ε_2 .

$$\Delta U_{H12} = \varepsilon_2 U_{obc max}. \quad (14)$$

Dla stosowanych halotronów $\varepsilon_2 \approx$ od 0,5‰ do 1‰, tzn., gdy np. $U_{obc max} = 500$ mV, $\Delta U_{H12} \approx (2,5 \text{ do } 5) \cdot 10^{-3}$ V.

Uchyb od wpływu nieliniowości $B = f(x)$, ΔU_{H3} , wynika z faktu, że gradient indukcji magnetycznej B na drodze przemieszczania halotronu nie jest idealnie stały. To z kolei wynika z nieprecyzyjnego wykonania nabiegunków lub samych magnesów, niewłaściwego ich układu oraz z niejednorodności materiału magnetycznego tworzącego obwód magnetyczny przetwornika.

Uchyb ten można wyrazić jako

$$\Delta U_{H3} = \varepsilon_3 U_{obc max}, \quad (15)$$

gdzie ε_3 — maksymalne odchylenie w górę i w dół od linii prostej przebiegu $B = f(x)$ w zakresie liniowej zależności przy $x = \pm x_{max}$.

Wartość tego uchybu jest ściśle uzależniona od wielkości powierzchni samego halotronu, którym indukcja była mierzona. Przy małych powierzchniach halotronu w stosunku do obszaru, w którym gradient pola jest stały, uzyskuje się pomiary bardziej „punktowe”. W tym przypadku uchyb ΔU_{H3} jest większy niż wtedy, gdy powierzchnia halotronu jest relatywnie większa, gdyż wówczas drobne odchylenia od stałości gradientu pola są niwelowane, ponieważ napięcie Halla jest proporcjonalne do wartości średniej indukcji magnetycznej na całej powierzchni halotronu.

Ponieważ sam pomiar indukcji w szczelinie przetwornika jest wykonywany przy pomocy halotronu o uchybie nieliniowym ΔU_{H12} , stąd uchybów tych nie można rozdzielić i sumaryczny uchyb od nieliniowości

$$\Delta U_{H123} = \varepsilon \cdot U_{obc max}, \quad (16)$$

gdzie ε maksymalne odchylenie w górę i w dół od linii prostej przebiegu $U_H = f(x)$ w zakresie liniowej zależności przy $x = \pm x_{max}$ zdjęte doświadczalnie przy prądach stałych dla zadanej temperatury granicznej.

Uchyb od wpływu temperatury na stałą Halla R_H i oporność wewnętrzną halotronu r_{wev} , ΔU_{H4} , można określić rozpatrując równanie (13) oraz równanie na zmianę stałej Halla

$$R_H = R_{H20} [1 + \alpha'_t (\Theta - 20^\circ)], \quad (17)$$

gdzie:

R_{H20} — stała Halla dla materiału halotronu przy 20°C ,

α'_t — współczynnik zmiany stałej Halla z temperaturą.

W wyniku otrzymujemy równanie:

$$\Delta U_{H4} = U_{obc\ max}(\Theta - 20^\circ) \left(\alpha'_t - \alpha_t \frac{r_{wew}}{r_{wew} + r_{obc}} \right), \quad (18)$$

w którym człon $\alpha_t \frac{r_{wew}}{r_{wew} + r_{obc}}$ uwzględnia wpływ obciążenia r_{obc} na wartość ΔU_{H4} i nie występuje przy otwartym obwodzie napięciowym ($r_{obc} = \infty$).

Dla InAs: $\alpha'_t = -0,0005/^\circ\text{C}$; $\alpha_t = 0,001/^\circ\text{C}$, oraz przy założeniu

$$\frac{r_{wew}}{r_{wew} + r_{obc}} \approx \frac{1}{5}, \quad U_{obc\ max} = 500\ \text{mV}; \quad \Theta = 40^\circ; \quad \Delta U_{H4} \approx -7 \cdot 10^{-3}\ \text{V}.$$

Uchyb wywołany przez elementy ferromagnetyczne przetwornika ΔU_{H5} należy rozpatrywać, gdy pole magnetyczne jest wytworzone przez układ magnesów trwałych lub elektromagnesów.

W pierwszym przypadku magnesy trwałe będą przyczyną dwóch uchybów

$$\Delta U_{H5} = \Delta U'_{H5} + \Delta U''_{H5}, \quad (19)$$

gdzie:

$\Delta U'_{H5}$ — uchyb spowodowany starzeniem się magnesów,

$\Delta U''_{H5}$ — uchyb od wpływu temperatury na wartość indukcji magnetycznej w szczelinie magnesów,

$$\Delta U'_{H5} = R_H \frac{I_p \cdot \Delta B_{mt}}{d}, \quad (20)$$

gdzie $\Delta B_{mt} \approx 0,2\% B_{max}$ — przyjęty spadek indukcji magnetycznej spowodowany starzeniem się magnesów po kilkunastu latach.

Przy założeniu $U_{obc\ max} = 500\ \text{mV}$, uchyb

$$\Delta U'_{H5} \approx 1 \cdot 10^{-3}\ \text{V},$$

$$\Delta U''_{H5} = R_H \frac{I_p}{d} B_{max20} \cdot \alpha''(\Theta - 20^\circ), \quad (21)$$

gdzie:

α'' — temperaturowy współczynnik zmian indukcji magnetycznej,

B_{max20} — maksymalna indukcja magnetyczna przy 20°C ,

Θ — temperatura, przy której wykonuje się pomiar.

Przy założeniu $U_{obc\ max} = 500\ \text{mV}$ i $\alpha'' = -0,00016$ (dla Alnico V) oraz $\Theta = 40^\circ$,

$$\Delta U''_{H5} \approx -2,5 \cdot 10^{-3}\ \text{V}.$$

Uchyb $\Delta U'_{H5}$ zazwyczaj można pominąć, gdyż z jednej strony zakłady produkujące magnesy trwałe zapewniają praktycznie stałe wartości włas-

ności stabilizowanych magnesów, z drugiej strony uchyb ten b. łatwo jest skompensować. Kompensacji dokonuje się przy wzorcowaniu bądź przez podregulowanie wzmacniacza napięcia wyjściowego U_H , lub drogą podwyższenia (jak dla rozpatrywanego przypadku) o 0,2% prądu pierwotnego I_p .

Przy stosowaniu elektromagnesów zasilanych prądem stałym lub zmiennym, uchyb może być spowodowany głównie niestabilnością prądu magnesowania.

Uchyb wywołany przez układ mechaniczny ΔU_{H6} jest związany głównie z luzami i tarciami, przede wszystkim w przekładniach mechanicznych, oraz ze zmianą wymiarów lub z naprężeniami mechanicznymi spowodowanymi przez zmianę temperatury.

W przetworniku halotronowym, z uwagi na bardzo małe wymiary i masę (często poniżej 0,5 grama) halotronu oraz minimalne straty na tarcie o powietrze, wystarczą bardzo małe siły dla przesuwania halotronu w polu magnetycznym. Stąd układ mechaniczny może być wykonany z wielką precyzją.

Uchyb układu mechanicznego wyraźniej występuje przy pomiarze przemieszczeń poniżej mikrona.

Oceńić te uchyby można jedynie rozpatrując je na konkretnym przykładzie.

Klasa przetwornika N_p . Halotron przed wstawieniem go do niejednorodnego pola magnetycznego jest cechowany w jednorodnym znanym polu magnetycznym, lecz o różnych wartościach. Warunki te odpowiadają warunkom sprawdzania gausomierza; dlatego początkowo jest określana klasa jak dla gausomierza:

$$N_g = \frac{U_{Ha} + \Delta U_{H12} + \Delta U_{H4}}{U_{obc\ max}}. \quad (22)$$

Do równania powyższego wchodzi niektóre z rozpatrzonych uchybów. Po podstawieniu wynikających z rozważań wielkości oraz założeniu, że halotron jest wykonany z InAs o $U_{obc\ max} = 500$ mV przy $B_{max} = 10.000$ Gs, wynika, że można wykonać gausomierz w klasie 1‰.

Niektóre zakłady produkujące gausomierze z halotronami, po wprowadzeniu układów kompensujących uchyby, zapewniają klasę 0,5‰ z zastrzeżeniem, że halotrony są często sprawdzane na odpowiednich zakresach w specjalnie kalibrowanych magnesach trwałych.

Przy określaniu klasy przetwornika halotronowego N_p należy wziąć pod rozwagę pozostałe przyczyny uchybów:

- U_{Hb} — pomija się przyjmując, że halotron nie drga z wysoką częstotliwością w polu wywołanym przez magnesy trwałe,
- U_{Hc} — pomija się jako niewielkie i łatwe do skompensowania,
- ΔU_{H3} — uwzględnia się,

ΔU_{H5} — przy stosowaniu magnesów trwałych, pomija się łatwą do skompensowania składową wynikającą ze starzenia magnesów $\Delta U'_{H5}$, natomiast uwzględnia się wpływ temperatury na wartość indukcji magnetycznej w szczelinie magnesów, stąd $\Delta U_{H5} = \Delta U''_{H5}$.

ΔU_{H6} — jest uwidoczniiony we wzorze, lecz jego wartość jest uzależniona od konkretnego rozwiązania konstrukcyjnego przyrządu.

Stąd klasę przetwornika można określić równaniem

$$N_p = \frac{U_{Ha} + \Delta U_{H123} + \Delta U_{H4} + \Delta U_{H5} + \Delta U_{H6}}{U_{obc\ max}}. \quad (23)$$

Decydującą różnicę w stosunku do wzoru (22) stanowi obok sprecyzowanego uchybu ΔU_{H5} uchyb ΔU_{H3} , wynikający z nieliniowości funkcji $B = f(x)$ na drodze przemieszczenia halotronu. Uchyby temperaturowe w dużej mierze można skompensować.

Z przeprowadzonych doświadczeń można sądzić, że dla zakresu pomiarowego 1/10 mm, 1/100 mm i 1/1000 mm, funkcje przebiegają liniowo i uchyb $\Delta U_{H3} \leq 1\%$. Stąd po podstawieniu innych wartości na uchyby wynika, że dla wymienionych zakresów pomiarowych można budować przetworniki o klasie N_p rzędu 1,5%.

Przy pomiarach poniżej 1 mikrona zaczynają występować kłopoty częściowo omówione poniżej w części eksperymentalnej.

3. CZUŁOŚĆ NAPIĘCIOWA PRZETWORNIKA HALOTRONOWEGO

Na czułość przetwornika halotronowego składają się dwa podstawowe parametry:

(a) czułość napięciowa halotronu γ_I oraz

(b) gradient indukcji, grad B , dla układu magnetycznego, w którym przemieszcza się halotron.

(a) Czułość halotronu γ_I określa się jako napięcie Halla dla jednostkowej indukcji magnetycznej.

Zagadnienie czułości jest rozpatrywane przez wielu autorów, np. [8], [11], [15].

$$\gamma_I = \frac{U_H}{B} = \frac{R_H I_{p\ max}}{d} 10^{-8} \text{ (V/Gs)}. \quad (24)$$

Jak wynika ze wzoru, czułość halotronu rośnie ze wzrostem prądu pierwotnego $I_{p\ max}$, ten zaś jest ograniczony wymiarami czujnika i przyrostem temperatury płytki spowodowanej efektem cieplnym przez prąd płynący przez próbkę.

Przy założeniu, że wytwarzanie ciepła w płytce jest równomierne, rozmieszczenie temperatury jest symetrycznie paraboliczne w kierunku grubości płytki, skąd otrzymujemy [18]:

$$T_{max} - T(y) = k \cdot y^2, \quad (25)$$

gdzie:

T_{max} — temperatura w środku płytki,

$T(y)$ — temperatura w półprzewodniku, dana jako funkcja położenia w kierunku jej grubości.

Przy $y = \frac{d}{2}$, $T(y) = T_0$, co odpowiada temperaturze na powierzchni płytki; wówczas

$$k = \frac{4}{d^2} (T_{max} - T_0), \quad (26)$$

a temperatura w płytce jest dana przez

$$T(y) = T_{max} - \frac{4}{d^2} (T_{max} - T_0) y^2. \quad (27)$$

Stąd stały stan przepływu przez powierzchnię graniczną jest dany jako:

$$Q = 2\mu_t \cdot b \cdot l \cdot \left. \frac{dT}{dy} \right|_{y=\frac{d}{2}} = - \frac{8\mu_t \cdot b \cdot l}{d} (T_{max} - T_0), \quad (28)$$

gdzie:

μ_t — przewodność termiczna,

$b \cdot l$ — powierzchnia płytki (rys. 1).

Wówczas maksymalny dopuszczalny prąd wynosi:

$$I_{pmax} = \sqrt{\frac{Q \cdot b \cdot d}{\rho \cdot l}}, \quad (29)$$

gdzie ρ — oporność właściwa materiału.

Po podstawieniu Q z równania (28) otrzymamy:

$$I_{pmax} = b \sqrt{\frac{8\mu_t}{\rho} (T_{max} - T_0)}. \quad (30)$$

Dla interesujących nas półprzewodników elektronowych

$$R_H = \frac{3\pi}{8ne} \quad \text{i} \quad \rho = \frac{1}{neu},$$

gdzie u — jest ruchliwością nośników prądu, a inne wielkości są dane zgodnie z (1).

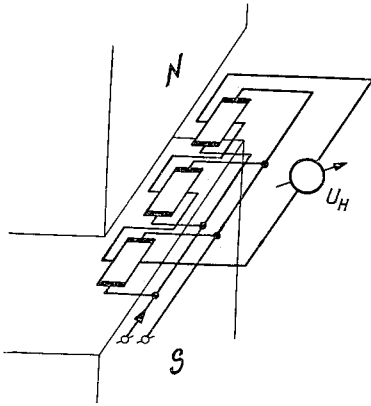
Stąd po podstawieniu do (24) otrzymamy:

$$\gamma_I = 3\pi \frac{b}{d} \sqrt{\frac{\mu_t \cdot u}{8ne} (T_{max} - T_0)}. \quad (31)$$

Na podstawie powyższego równania można wyciągnąć szereg wniosków. Dla danej koncentracji nośników prądu powinny być stosowane materiały o jak najwyższej ruchliwości tychże nośników.

Największą ruchliwość nośników prądu posiada InSb. Jednak wysoka zależność napięcia Halla od temperatury eliminuje generatory z antymonku indu z zastosowań do dokładnych pomiarów. Dlatego obecnie zwykle stosuje się InAs o stosunkowo wysokim u i małych zmianach napięcia Halla w funkcji temperatury.

Dla danego materiału i maksymalnego dopuszczalnego przyrostu temperatury płytka halotronu powinna być jak najcieńsza, podczas gdy jej szerokość powinna być możliwie duża. Można stąd dla zwiększenia czułości przetwornika halotronowego stosować kilka halotronów, w ten sposób, że ich elektrody prądowe są łączone równolegle, zaś napięciowe szeregowo [19]; rys. 9 przedstawia ideowy układ połączeń.



Rys. 9. Zwiększenie czułości przetwornika halotronowego przez równoległe łączenie elektrod prądowych i szeregowe napięciowych (schemat ideowy)

Istotnym czynnikiem wpływającym na wartość napięcia Halla, a zatem i na czułość pomiaru, jest właściwy stosunek wymiarów halotronu $k_d = b/l$, rys. 1, optymalne $k_d = 2 \div 3$ [7], [8], [20].

(b) Dla zakresu przemieszczeń, w którym wartość indukcji magnetycznej jest proporcjonalna do przesunięcia,

$$\text{grad } B = \frac{B}{x} \left(\frac{\text{Gs}}{\text{mm}} \right). \quad (32)$$

Ostatecznie czułość przetwornika halotronowego można określić zgodnie z równaniami (24) i (32) jako:

$$\gamma = \gamma_I \cdot \text{grad } B = \frac{U_H}{x} \quad (\text{V/mm}) \quad (33)$$

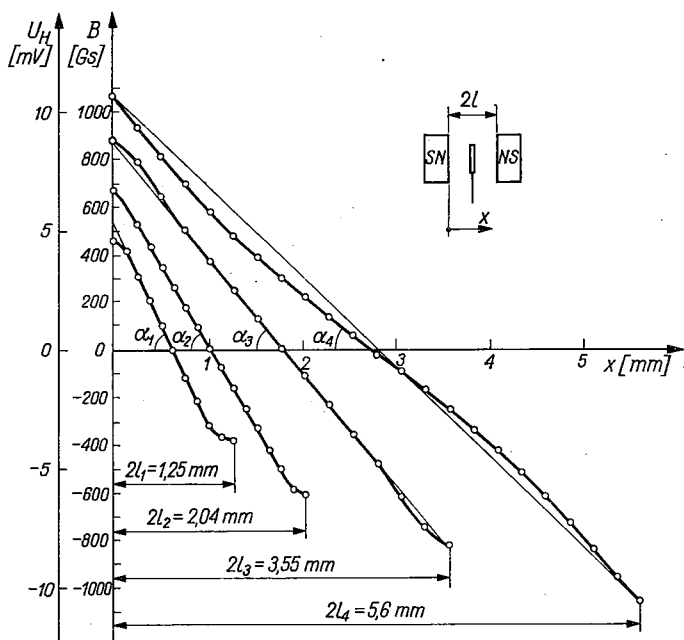
lub

$$\gamma = \text{tg } \alpha, \quad (34)$$

gdzie

α — kąt nachylenia przebiegu $U_H = f(x)$, rys. 3.

Czułość przetwornika halotronowego może być regulowana w prosty sposób przez zmianę prądu pierwotnego I_p (24) lub przez zmianę odległości $2l$ między biegunami magnesów trwałych, rys. 2.



Rys. 10. Funkcje $U_H = f(x)$ dla obwodu magnetycznego z rys. 2b przy różnych $2l$. Czułość przetwornika określają $\tan \alpha_1 > \tan \alpha_2 > \tan \alpha_3 > \tan \alpha_4$. W pomiarach stosowano halotron typu BH 200 firmy Bell Inc.

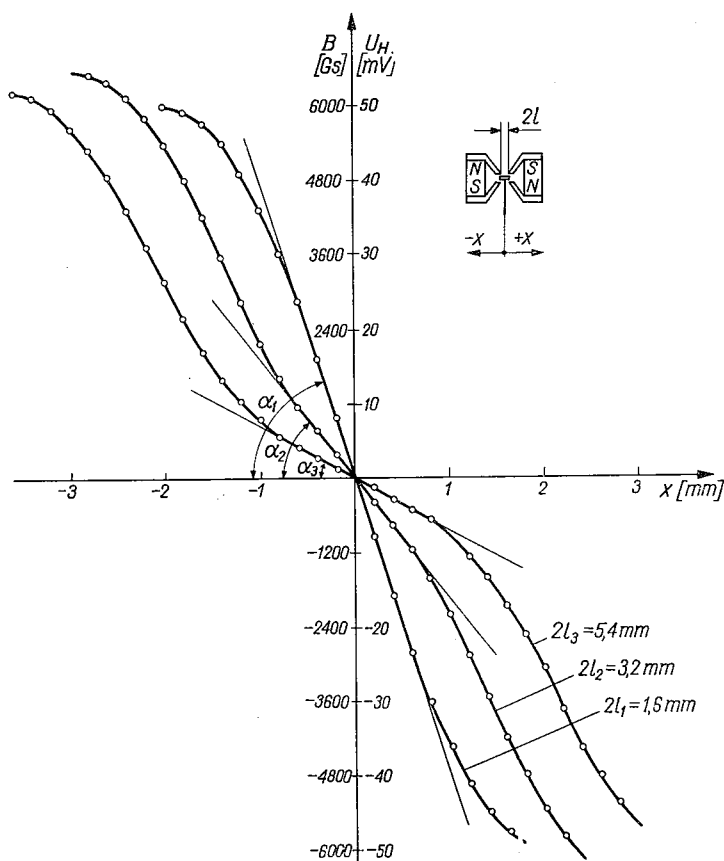
Na rys. 10 przedstawiono krzywe $U_H = f(x)$ dla magnesów ferrytowych firmy D. M. Steward MFG. Co., Chattanooga, Tennessee, USA, typ

o wymiarach $d_m = \frac{1''}{2}$, $l_m = \frac{1''}{4}$ (rys. 2b) w funkcji odległości między jednolitymi biegunami $2l$. Jak wynika z wykresu, $\gamma_{b1} = \tan \alpha_1 = 8,8 \cdot 10^{-3}$ V/mm, $\gamma_{b2} = 6,8 \cdot 10^{-3}$ V/mm, $\gamma_{b3} = 5 \cdot 10^{-3}$ V/mm, $\gamma_{b4} = 3 \cdot 10^{-3}$ V/mm.

Na rys. 11 przedstawiono krzywe $H_H = f(x)$ dla magnesów Alnico produkowanych w Instytucie Metalurgii w Gliwicach o wymiarach $a_m = 30$ mm, $b_m = 30$ mm, $c_m = 45$ mm i szczeliny $\delta = 1,6$ mm (rys. 2c), w funkcji odległości między biegunami $2l$.

Z rys. 11 można odczytać, że $\gamma_{c1} = \tan \alpha_1 = 40 \cdot 10^{-3}$ V/mm, $\gamma_{c2} = 16,5 \cdot 10^{-3}$ V/mm, $\gamma_{c3} = 7 \cdot 10^{-3}$ V/mm.

Obecnie Siemens produkuje halotроны typu SV 240 o czułości $\gamma_I = (1000 \pm 30\%) \cdot 10^{-3}$ V/10 kGs; przy równoczesnym zastosowaniu układu rys. 2c o grad $B = 10$ kGs/mm można otrzymać b. wysoką czułość przetwornika halotronowego $\gamma = (1000 \pm 30\%) 10^{-3}$ V/mm.



Rys. 11. Funkcje $U_H = f(x)$ dla obwodu magnetycznego z rys. 2c przy różnych $2l$. Czułość przetwornika określają $\tan \alpha_1 > \tan \alpha_2 > \tan \alpha_3$. W pomiarach stosowano halotron typu FA 22 firmy Siemens

4. WYNIKI POMIARÓW

Poniżej zostaną przedstawione wyniki pomiarów przemieszczeń mechanicznych, jakie zostały wykonane za pomocą przetwornika halotronowego podczas pobytu autora w Engineering Design Center w Case Institute of Technology (Cleveland, Ohio, USA).

Aby uzyskać największą czułość pomiaru przemieszczeń mechanicznych, przy wykonywaniu doświadczeń korzystano z układu magnesów przedstawionego na rys. 2c. Zastosowano magnesy trwałe, Alnico 5 $\times$ dg General Magnetic Corporation, o wymiarach $a_m = b_m = c_m = 25,4$ mm i szczeliny $\delta = 1$ mm). Układ ten pozwalał na uzyskanie wysokiego gradientu indukcji magnetycznej na drodze przemieszczenia płytki rzędu 10000 Gs/mm. Równocześnie pole magnetyczne było mierzone przy pomo-

cy gausomierza klasy 1 firmy Bell Inc., Columbus OHIO, USA, o 100 kach i następujących 12 zakresach pomiarowych:

0,1	3	100	3000
0,3	10	300	10000
1,0	30	1000	30000 Gs

Układ magnesów był umieszczony na ławie obrabiarki i był przesuwany względem nieruchomego halotronu. Noniusz umieszczony na mechanizmie do przesuwu ławy pozwalał na najmniejsze odczyty 10^{-4} cala. Dla zwiększenia zakresu najmniejszych odczytów do ławy przystawiono przetwornik indukcyjny pracujący na zasadzie zmiany indukcyjności wzajemnej firmy „Cleveland”, pozwalający na najmniejszy odczyt 10^{-5} cala = 0,254 mikrona/działkę. Odczyty te jednak były obciążone dużym uchybem układu mechanicznego ΔU_{H_6} z uwagi na luzy w przesuwie i trudności w dokładnym ustaleniu wielkości przemieszczeń.

Halotron gausomierza typ BH 200, wykonany z InAs, zasilany prądem $I_p = 100$ mA o częstotliwości 1,1 kHz, posiadał czułość $\gamma_I = (100 \pm \pm 25\%) 10^{-3}$ V/10 kGs, stąd biorąc pod uwagę ww. wartość gradientu pola magnetycznego, grad $B = 10$ kGs/mm, czułość pomiaru (rys. 13)

$$\gamma = \gamma_I \cdot \text{grad } B = 10 \alpha_p = (100 \pm 25\%) 10^{-3} \text{ V/mm.}$$

Z wykreślonych krzywych (rys. 12) wynika, że odcinek liniowej zależności $B_1 = f(x_1)$ wynosi około 1,15 mm. Na tej długości przemieszczenia można wykonywać bez większych nieliniowych zniekształceń pomiary przesunięć i drgań mechanicznych. Punkty funkcji $B_1 = f(x_1)$, $B_2 = f(x_2)$ i $B_3 = f(x_3)$ leżą prawie idealnie na prostej.

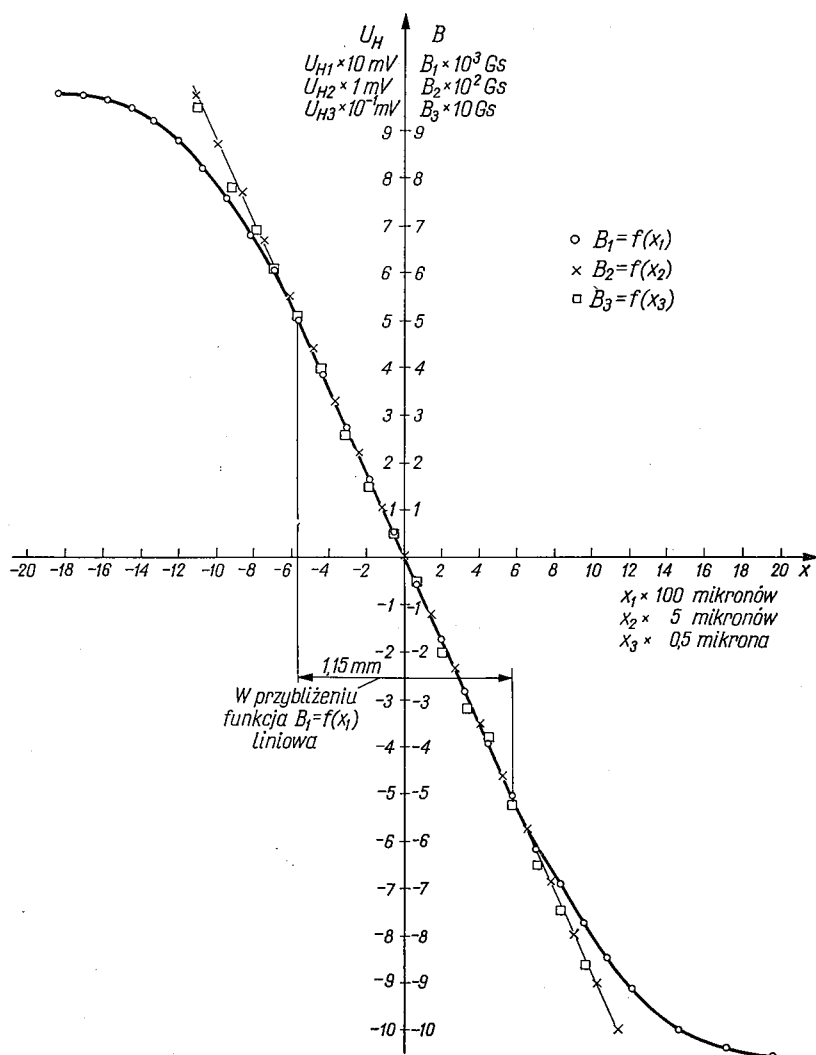
Na rys. 13 przedstawiono wyniki pomiarów na najniższym z badanych zakresów pomiarowych. Jak widać z wykresu, przebiegi są liniowe w granicach dokładności gausomierza i przetwornika indukcyjnego, który służył jako wzorzec.

Stąd wniosek, że przetwornik halotronowy w przedstawionym układzie pozwala na pomiar przesunięcia mechanicznego rzędu 0,1 mikrona.

Przedstawione wyniki pomiarów należy traktować jako wstępne. Pomiary przesunięć mechanicznych rzędu i poniżej 0,1 mikrona wymagają laboratorium wysokiej precyzji. W tym zakresie pomiaru istnieje szereg uchybów niesystematycznych, które wymagają bardziej szczegółowego zbadania.

Szczególnie praca halotronu jak na rys. 4 wymaga bliższej analizy, gdy każda z połówek halotronu znajduje się pod działaniem silnego pola magnetycznego, lecz skierowanego w przeciwnym kierunku. W takim układzie różnicowym halotron jest szczególnie czuły nie tylko na zmiany w przemieszczeniu płytki, lecz również na wpływy zewnętrzne. Na przykład zmiany temperatury, jak to wynika z rozważań uchybu ΔU_{H_2} i ΔU_{H_4} , wpływają tak na zmianę stałej Halla R_H , jak i oporności wewnętrznej

r_{wev} . W omawianym układzie gęstość prądu w środku płytki jest większa niż na krańcach z elektrodami napięciowymi (zjawisko Gaussa), a niejednorodny rozkład temperatury wpływa na niestabilność tego stanu, a stąd i na napięcia Halla. Konieczna jest zatem kompensacja temperatury jak dla układu różnicowego oraz właściwy dobór materiału na halotron.



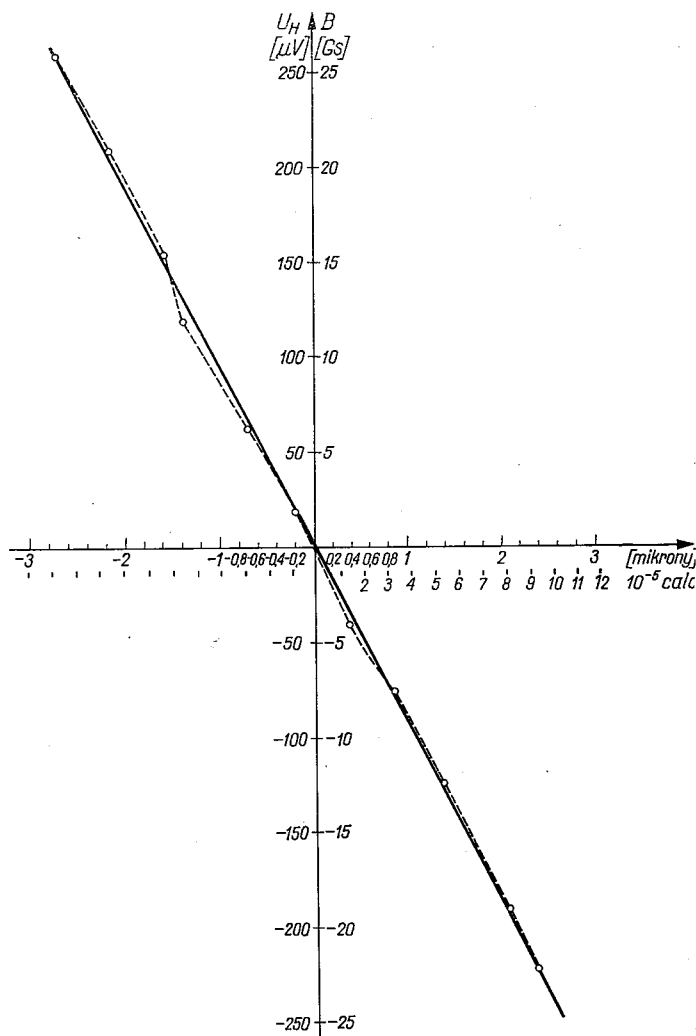
Rys. 12. Funkcje $B = f(x)$ lub $U_H = f(x)$ dla obwodu magnetycznego z rys. 2c przy różnych zakresach pomiaru

Należy się liczyć również z prawdopodobieństwem wpływu szumów w półprzewodniku, w szczególności w silnym polu magnetycznym.

Również gęstość prądu pierwotnego I_p w układzie jak na rys. 4 będzie większa na środku płytki niż bliżej krawędzi, gdzie pole jest większe (zja-

wisko Gaussa). Powodować to może nierównomierne nagrzewanie płytki, a stąd dalsze efekty.

Przy kompensacji wpływów zewnętrznych, bądź wykonaniu pomiarów w pomieszczeniach o stałej temperaturze oraz przy bardziej dokład-



Rys. 13. Funkcje $B = f(x)$ lub $H = f(x)$ dla obwodu magnetycznego z rys. 2c dla najniższego zakresu pomiaru

nym sprawdzeniu liniowości w pobliżu zera napięcia Halla, istnieje możliwość pomiaru bardzo małych przemieszczeń.

Należy zwrócić uwagę na to, że przy korzystaniu z zakresu 30 Gs można zmierzyć 0,1 mikrona. Gausmierz dysponuje ponadto zakresami 10; 3; 1; 0,3 i 0,1 gaussa na 100 działek skali. Istnieje zatem teoretyczna możliwość pomiaru przemieszczeń rzędu 1 Å.

5. WNIOSKI

Przedstawione rozważania i wstępne doświadczenia wykazały, że przetwornik halotronowy można wybudować w klasie 1,5%. Za pomocą przetwornika wykonano pomiar statycznego przemieszczenia mechanicznego rzędu 0,1 mikrona.

Osiągnięto wysoką czułość pomiaru 100 mV/mm, a przy zastosowaniu halotronu Siemens SV 240 czułość tę można zwiększyć do 1 V/mm. Stąd wniosek, że przetwornik halotronowy w wielu przypadkach może być stosowany bez wzmacniacza napięcia wyjściowego.

Jednocześnie b. małe wymiary samego halotronu (np. $0,5 \times 1 \times 0,3$ mm) i praktycznie brak strat na tarcie powodują, że przetwornik tego typu będzie reagował na przyłożenie bardzo małych sił.

Jako pierwsze zostały na zasadzie przetwornika halotronowego zbudowane dwa rodzaje przyrządów:

a) sejsmograf z halotronem, którego maksymalne wzmocnienie jest 100 razy większe niż maksymalne wzmocnienie dla sejsmografu elektrodynamicznego [21], [22] i który posiada zdolność pomiaru kąta nachylenia o wartości 0,001 sekundy kątowej [22], oraz

b) dwa przyspieszeniomierze z halotronami — jeden umożliwiający pomiar 10^{-4} g do 1 g, drugi na zakres 5 g [23].

Wydaje się, że należy prowadzić dalsze prace nad kompensacją rozważonych w pracy uchybów, gdyż istnieje realna możliwość zbliżenia się do pomiaru przemieszczeń mechanicznych rzędu 1 Å.

Instytut Automatyki PAN

WYKAZ LITERATURY

1. Hall E. H.: On a New Action of the Magnet on the Electric Circuit. — *Phil. Magaz.* 10, 1880, str. 225.
2. Pearson G. L.: Magnetic Field Strength Meter Employing the Hall Effect in Ge. — *Rev. Sci. Instr.* 19, 1948, 4, str. 263—265.
3. Nałęcz M., Ziomecki H.: Układy gausmierz z zastosowaniem efektu Halla. — *Przegląd Elektrotechn.*, 10, 1960, str. 469—473.
4. Nałęcz M.: Mechanical Measurement — Use of Hall Effect. — *Electronic Technology*, 1, 1961, str. 15—17.
5. Nałęcz M.: Hall Effect in Measurements of Mechanical Displacement and Vibration. — *Bulletin of the Polish Academy of Science*, 9, 1961, str. 469—475.
6. Giriat W.: Zastosowanie materiałów półprzewodnikowych o dużej ruchliwości nośników prądu. — *Postępy Fizyki*, 10, nr 2, str. 149, 1959.
7. Tuszyński J.: Zjawisko Halla i jego zastosowanie w technice pomiarów i regulacji. — *Pomiary, Automatyka, Kontrola*, 1959, nr 4, str. 144—148, i nr 5, str. 193—199.
8. Kobus A.: Projektowanie i własności halotronów germanowych. — *Przegląd Elektroniki*, 4, 1959, str. 414—435.
9. Kobus A.: Elektryczne własności halotronów CH1 i CH3. — *Pomiary, Automatyka, Kontrola*, 11, 1959, str. 446—449.

10. Welker H.: Neuartige Werkstoffe mit grossen Hall — Effekt. — ETZ-A, 76, 1955, str. 513—517.
11. Żużel W. P., Regel A. R.: Tiechniczieskoje primienienije effiektu Halla. Leningrad, Dom nauczno-tiechniczieskoj propagandy, wyp. 11, 1957.
12. Barlow H. E. M.: The Design of Semi-Conductor Wattmeters for Power-Frequency and Audio-Frequency Applications. — Proc. of I.E.E., 102, B, nr 2, 1955.
13. Pikus G. E., Sorokin O. W.: Nowyj mietod izmierienija naprażennosti magnitnogo polia. — Ž. Tiechn. Fiz., 27, 1957, str. 26.
14. Ross J. M., Saker E. W.: Application of Indium Antimonide. — J. Electronics, 1, 1955, str. 223—231.
15. Wasilczenko L. S., Sientiurina L. W., Sockow B. C.: O tocznisti daczikow ispolżujuszczich effiekt Halla. — Awtomatika i Tielemiechanika, 20, 7, 1959, str. 939—945.
16. Żużel W. P.: Połuprowodniki w nauce i technice, t. I, 1957, str. 368—416.
17. Nałęcz M.: The Hall Generator and Some of its Applications. Case Institute of Technology, Report nr EDC-3-61-TN3, 1961, stron 60.
18. Epstein M., Szulz R. B.: Magnetic-Field Pickup for Low-Frequency Radio-Interference Measuring Sets. — I.R.E. Transactions on Electron Devices, 1, 1961, str. 70—77.
19. Epstein M., Sachs H. M., Greenstein L. J.: Multiple-Element Hall-Effect Sensor. — Proc. I.R.E., vol. 47, 11, 1959, str. 2014.
20. Flanagan W. F., Flinn P. A., Averbach B. L.: Shorting and Field Correction in Hall Measurements. — Rev. Sci. Instr., 25, 6, 1954, str. 593.
21. Nałęcz M., Zawicki I.: Sejsmograf z płytką Halla. — Acta Geophysica Polonica Polskiej Akademii Nauk, nr 1, 1962.
22. Nałęcz M., Zawicki I.: A Seismograph with Hall Transducer. — Bulletin of the Seismological Society of America, kwiecień 1962.
23. Nałęcz M., Ziomecki H.: Hall Effect Accelerometer. — Journal of Franklin Institute, nr 1, 1963, str. 14—25.

M. Nałęcz

ON ACCURACY AND SENSITIVITY OF THE HALL GENERATOR TRANSDUCER IN STATIC MECHANICAL MEASUREMENT

Summary

The Hall effect offers a simple solution for mechanical displacement and vibration measurements. When the Hall Generator is moved mechanically in a non-uniform magnetic field, the Hall voltage output will vary in sympathy. For optimum performance the gradient of magnetic density should be as high as possible and constant.

The work was performed to determine the accuracy and sensitivity of the Hall Generator Transducer in static mechanical measurement.

Several errors were examined like:

- (a) unsymmetrical location of the Hall electrodes,
- (b) voltage induced in wires carrying Hall voltage,
- (c) voltage from earth or outside magnetic field,
- (d) nonlinearity of function $U_H = f(B)$,
- (e) influence of temperature on linearity of the function $U_H = f(B)$,
- (f) nonlinearity of function $B = f(x)$,

- (g) influence of temperature on Hall constant and internal resistance,
- (h) influence of ferromagnetic elements,
- (i) error caused by mechanical arrangement.

Results showed that it is possible to build a Hall effect transducer with an accuracy of better than 1.5 per cent and with sensitivities about 100 times higher than other transducers now in use. Using special magnetic circuits with magnetic field density gradients of 10^4 gauss/mm and Hall generators with sensitivities of 0.1 to 1.0 volts/ 10^4 gauss in special configurations, a sensitivity of 0.1 to 10.0 volts/mm can be reached.

M. NAŁĘCZ

PRÉCISION ET SENSIBILITÉ DES CONVERTISSEURS DE HALL EN CAS DES MESURES MÉCANIQUES STATIQUES

Résumé

L'effet de Hall donne la possibilité d'exécuter simplement les mesures des déplacements et vibrations mécaniques. Si le convertisseur est déplacé mécaniquement dans un champ magnétique irrégulier dans le même temps se changera la tension de Hall. Pour obtenir les conditions de travail le plus avantageuses le gradient du champ magnétique doit être le plus grand que possible et constant.

On a effectué des essais pour déterminer précision et sensibilité des convertisseurs de Hall en cas des mesures mécaniques statiques. On a étudié une série des erreurs comme:

- (a) situation non symétrique des électrodes de Hall,
- (b) tension induite aux conducteurs de la boucle de tension,
- (c) tension causée par le champ magnétique terrestre ou par un autre champ magnétique extérieur,
- (d) non-linéarité de la fonction $U_H = f(B)$,
- (e) influence de la température sur la linéarité de la fonction $U_H = f(B)$,
- (f) non-linéarité de la fonction $B = f(x)$,
- (g) influence de la température sur la constante de Hall et la résistance intérieure,
- (h) influence des éléments ferromagnétiques,
- (i) erreur causée par installation mécanique.

Les résultats montrent, qu'il existe la possibilité de construire les convertisseurs de Hall à précision meilleure que 1,5% et ayant la sensibilité plus grande que d'autres convertisseurs utilisés à présent. Employant les circuits magnétiques spéciaux à un gradient de l'induction magnétique 10^4 Gs/mm et les convertisseurs de Hall à sensibilité 0,1—1,0 V/ 10^4 Gs avec les circuits spéciaux, on peut obtenir la sensibilité d'ordre 0,1—10,0 V/mm.

М. НАЛЭЧ

ТОЧНОСТЬ И ЧУВСТВИТЕЛЬНОСТЬ ПРЕОБРАЗОВАТЕЛЕЙ ХОЛЛА В СТАТИЧЕСКИХ МЕХАНИЧЕСКИХ ИЗМЕРЕНИЯХ

Резюме

Эффект Холла дает возможность простым образом производить измерения механических сдвигов и вибрации. Если преобразователь Холла передвигается

механически в неоднородном магнитном поле, одновременно изменяется напряжение Холла. С целью получения самых выгодных условий работы, градиент магнитного поля должен быть возможно самым большим и постоянным.

Произведено исследование для определения точности и чувствительности преобразователей Холла в статических механических измерениях. Изучено ряд погрешностей, как:

- (а) несимметричное расположение электродов Холла,
- (б) напряжение, наведенное в проводах шлейфа напряжения,
- (в) напряжение, вызванное полем земли или внешним магнитным полем,
- (г) нелинейность функции $U_H = f(B)$,
- (д) влияние температуры на линейность функции $U_H = f(B)$,
- (е) нелинейность функции $B = f(x)$,
- (ж) влияние температуры на постоянную Холла и внутреннее сопротивление,
- (з) влияние ферромагнитных элементов,
- (и) погрешность, вызванная механическими устройствами.

Результаты доказывают, что существует возможность построения преобразователей Холла с точностью лучшей чем 1,5% и чувствительностью большей, чем другие преобразователи, применяемые в настоящее время. Используя специальные магнитные цепи с градиентом магнитной индукции 10^4 гс/мм и преобразователи Холла с чувствительностью 0,1 до 1,0 в/ 10^4 гс/мм в специальных схемах, можно получить чувствительность порядка 0,1 до 10,0 в/мм.

MARIAN DĄBROWSKI

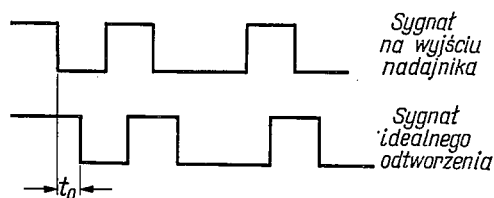
Wpływ zniekształceń opóźnieniowych kanałów teletransmisyjnych na transmisję danych

Rękopis dostarczono 30.4.1964

W pracy omówiono wpływ zniekształceń opóźnieniowych kanałów teletransmisyjnych na transmisję danych. We wstępie podano pojęcia ogólne oraz warunki, jakie powinny spełniać charakterystyki fazowe i opóźnieniowe w funkcji częstotliwości przy wiernym odtwarzaniu przesyłanych przebiegów. W rozdziale drugim rozpatrzono kształt przebiegów przesyłanych kanałami teletransmisyjnymi o różnych kształtach charakterystyk fazowych. Następnie, w rozdziale trzecim, podano zależności między zniekształceniami fazowymi a ekwiwalentnym stosunkiem sygnał-zakłócenie dla charakterystyk fazowych i opóźnieniowych, będących funkcjami potęgowymi i oscylacyjnymi częstotliwości. Zależności końcowe pozwalają ocenić stopień błędów systemu transmisji danych przy zadanych funkcjach przesuwności fazowej lub opóźności grupowej od częstotliwości.

1. WSTĘP

Sygnał przekazywany kanałem teletransmisyjnym od nadajnika do odbiornika zawiera informację w postaci pewnych istotnych cech (np. amplituda, faza, częstotliwość). Wierne odtworzenie tych cech w odbiorniku jest warunkiem wystarczającym dla przekazywania informacji. W transmisji binarnej (np. telegrafii, transmisji danych) odtworzenie jest wierne, jeśli podział czasowy sygnału, powstały w nadajniku, jest zachowany po odtworzeniu w części odbiorczej (rys. 1).



Rys. 1

Przy odtworzeniu wiernym wszystkie momenty charakterystyczne mogą być opóźnione o jednakowy czas t_0 (w porównaniu z sygnałem nadajnika). W technicznej rzeczywistości niemożliwe jest wierne odtworzenie na skutek zniekształceń linearnych i nielinearnych wnoszonych przez ka-

nał teletransmisyjny. Mimo to, możliwe jest wierne przekazywanie informacji, jeżeli łączy teletransmisyjne spełnia dwa postulaty:

- 1) Urządzenie odbiorcze posiada określoną zdolność poprawiania,
- 2) Kanał teletransmisyjny nie wprowadzi zniekształceń cech istotnych sygnału przekraczających zdolność poprawiania urządzenia odbiorczego.

Drugi postulat wymaga określenia zależności zniekształceń sygnału przesyłanego od właściwości kanału. Nowe dziedziny teletransmisji, jak np. transmisja danych i fototelegrafia, zmuszają do rozszerzenia zakresu i wnikliwości badań kanałów. Do podstawowych badań należą zniekształcenia opóźnieniowe poprzez pomiary charakterystyk opóźnieniowych.

Niech na wejście kanału o charakterystyce amplitudowej $A(\omega)$ i charakterystyce fazowej $B(\omega)$ zostanie podane napięcie sinusoidalne:

$$u_n = U_n \cdot \cos \omega_0 t. \quad (1)$$

Funkcja czasowa napięcia wyjściowego kanału ma postać:

$$u_0 = U_n \exp(-A_0) \cdot \cos(\omega_0 t - B_0), \quad (2)$$

gdzie:

A_0 — tłumienność kanału (w neperach) przy pulsacji ω_0 ;

B_0 — przesuwność fazowa kanału (w radianach) przy pulsacji ω_0 .

Określona faza (np. maksimum napięcia), będąca na wejściu kanału w chwili t_1 , osiągnie wyjście kanału w chwili t_2 . Dla rozpatrywanego przypadku słuszna jest równość:

$$\begin{aligned} t_1 &= t_2 - \frac{B_0}{\omega_0} \\ t_2 - t_1 &= \frac{B_0}{\omega_0} = \tau_f. \end{aligned} \quad (3)$$

Równaniem (3) wyraża się fazowy czas przejścia (opóźność fazowa) fali, czyli czas przejścia określonej fazy napięcia od wejścia do wyjścia kanału.

Założmy, że przesyłamy w kanale dwa sinusoidalne napięcia różniące się częstotliwością kątową (pulsacją) o nieskończenie małą wartość $\Delta\omega$.

$$\left. \begin{aligned} u_{n1} &= U_{n1} \cos \omega_0 t, \\ u_{n2} &= U_{n2} \cos[(\omega_0 + \Delta\omega)t]. \end{aligned} \right\} \quad (4)$$

Jeżeli przesuwność fazowa kanału dla pulsacji ω_0 wynosi B_0 , zaś dla pulsacji $(\omega_0 + \Delta\omega)$ wynosi $B_0 + \Delta B$, wówczas odpowiednie napięcia na wyjściu kanału mają następującą postać analityczną:

$$\left. \begin{aligned} u_{o1} &= U_{o1} \cos(\omega_0 t - B), \\ u_{o2} &= U_{o2} \cos[(\omega_0 + \Delta\omega)t - B - \Delta B]. \end{aligned} \right\} \quad (5)$$

Określona faza (np. maksimum napięcia) obwiedni dudnień została opóźniona w kanale o czas wynikający z równości:

$$(\omega_0 t - B) = [(\omega_0 + \Delta\omega)t - B - \Delta B],$$

czyli

$$t = \lim_{\Delta\omega \rightarrow 0} \frac{\Delta B}{\Delta\omega} = \frac{dB}{d\omega} = \tau_g. \quad (6)$$

Równanie (6) określa grupowy czas przejścia, zwany opóźnością grupową.

2. WARUNKI WIERNEGO ODTWARZANIA

W technice teletransmisijnej przesyłane są na ogół przebiegi nieokresowe charakteryzowane funkcją $f_n(t)$ czasu t :

$$f_n(t) = \int_0^{\infty} K(\omega) \cos[\omega t - B(\omega)] d\omega. \quad (7)$$

Jest to granica sumy drgań harmoniczných o pulsacjach od 0 do ∞ . Każde drganie harmoniczne charakteryzuje pulsacja ω , amplituda o gęstości widmowej $K(\omega)$ i faza początkowa $B(\omega)$. Dla wiernego odtwarzania konieczne jest, aby wszystkie momenty charakterystyczne (momenty czasowe) były opóźnione o jednakowy czas t_0 . Wobec tego funkcja czasowa przebiegu po stronie odbiorczej $f_0(t)$ powinna mieć postać:

$$f_0(t) = K_0 \cdot f_n(t - t_0). \quad (8)$$

Dopuszczalne jest zmniejszenie lub zwiększenie amplitudy poszczególnych drgań harmoniczných w stosunku K_0 oraz przesunięcie drgań harmoniczných o czas t_0 . Wymienione warunki wystarczające do wiernego odtworzenia polegają na tym, żeby tłumienność kanału była niezależna od pulsacji, natomiast przesuwność fazowa powinna być do pulsacji proporcjonalna.

Warunek proporcjonalności przesuwności fazowej do pulsacji nie zawsze jest konieczny dla wiernego odtwarzania [2]. W zasadzie odnosi się on do przebiegów przesyłanych w pasmie naturalnym. Przy przesyłaniu przebiegów modulowanych jednowstęgowo, przesuwność fazowa powinna być liniową funkcją pulsacji oraz powinna wynosić $2M\pi$ ($M = 0, 1, 2, 3, \dots$) dla pulsacji nośnej. Analiza tego przypadku znajduje się w rozdz. 4.

Przy przesyłaniu przebiegów modulowanych dwuwstęgowo (obie wstęgi są symetryczne względem pulsacji nośnej) warunek wiernego odtwarzania może być jeszcze bardziej złagodzony. Na przykład przebieg modulowany amplitudowo na wejściu kanału może być zapisany w postaci:

$$u_n = U_n \cos \omega_0 t [1 + m e(t)], \quad (9)$$

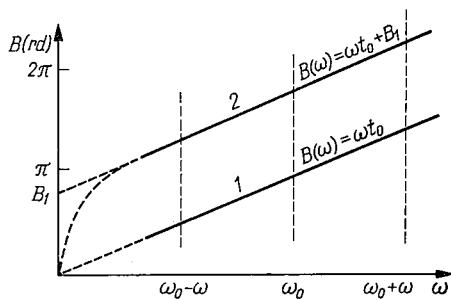
gdzie:

$\frac{\omega_0}{2\pi}$ — częstotliwość nośna,

U_n — amplituda częstotliwości nośnej,

m — współczynnik głębokości modulacji,

$e(t) = \sum_{i=1}^{\infty} a_i \cos i\omega t$ — funkcja czasowa przebiegu modulującego.



Rys. 2

Jeżeli kanał teletransmisyjny ma liniową charakterystykę fazową przecinającą początek układu współrzędnych (krzywa 1 na rys. 2), to każda składowa harmoniczna przebiegu wyrażonego równaniem (9) zostanie odebrana w postaci:

$$u_0 = U_0 \cos \omega_0(t - t_0) + \frac{U_0 a_i \cdot m}{2} \cos[(\omega_0 + i\omega)(t + t_0)] + \frac{U_0 a_i \cdot m}{2} \cos[(\omega_0 - i\omega)(t + t_0)]. \quad (10)$$

W tym przypadku wszystkie modulatory są opóźnione o czas

$$t_0 = \frac{B}{\omega} = \frac{dB}{d\omega}.$$

Wobec tego istotne cechy sygnału wyjściowego z kanału nie będą zniekształcone.

W przypadku gdy charakterystyka fazowa nie przecina osi B w punktach $M\pi$, gdzie M jest zerem lub liczbą całkowitą (krzywa 2 na rys. 2), każda składowa harmoniczna zostanie odebrana w postaci

$$u_0 = U_0 \cdot \cos[\omega_0(t + t_0) + B_1] + \frac{U_0 a_i \cdot m}{2} \cdot \cos[(\omega_0 + i\omega)(t + t_0) + B_1] + \frac{U_0 a_i \cdot m}{2} \cos[(\omega_0 - i\omega)(t + t_0) + B_1], \quad (11)$$

przy czym w postaci obwiedni można równanie (11) zapisać jako:

$$u_0 = U_0 \cos[\omega_0(t + t_0) + B_1] \cdot [1 + m a_i \cos i\omega(t + t_0)]. \quad (12)$$

Modulat o pulsacji nośnej ω_0 jest opóźniony o czas równy opóźności fazowej

$$\tau_f = t_0 + \frac{B_1}{\omega_0}, \quad (13)$$

natomiast obwiednia jest opóźniona o czas równy opóźności grupowej:

$$\tau_g = t_0. \quad (14)$$

Równania (13) i (14) wskazują, że w wyniku zniekształceń fazowych „przecięciowych” przy liniowej charakterystyce fazowej częstotliwość nośna i obwiednia są opóźnione w kanale o niejednakową wartość. Przy tego rodzaju zniekształceniach częstotliwość nośna jest przesunięta w fazie względem obwiedni.

Na ogół cechą istotną przesyłanych sygnałów jest kształt obwiedni. Przebiegi zapisane równaniami (10) i (11) nie wykazują różnic w kształcie obwiedni w porównaniu z przebiegiem nadawanym [równanie (9)]. Zniekształcenia „przecięciowe” przy liniowej charakterystyce fazowej nie powodują pogorszenia transmisji.

Wymieniony poprzednio warunek wystarczający, polegający na tym, żeby przesuwność fazowa była proporcjonalna do pulsacji, można złągodzić; przesuwność fazowa powinna być funkcją liniową pulsacji o dowolnym wyrazie stałym B_1 .

$$B(\omega) = \tau_g \cdot \omega + B_1. \quad (15)$$

Praktycznie użyteczny warunek wiernego odtworzenia można zapisać w postaci:

$$\tau_g = \frac{dB(\omega)}{d\omega} = \text{constans}. \quad (16)$$

Opóźność grupowa powinna być niezależna od pulsacji.

W technicznej rzeczywistości przesyłanie przebiegów elektrycznych odbywa się w pasmie ograniczonym. Warunek stałej wartości opóźności grupowej odnosi się do tego właśnie pasma.

3. WPŁYW ZNIEKSZTAŁCEŃ OPÓŹNIENIOWYCH NA KSZTAŁT PRZEBIEGÓW

Charakterystyki częstotliwościowe (amplitudowa i fazowa) czwórnika, przez który przesyłany jest impuls w postaci skoku jednostkowego, mają wpływ m.in. na dwie cechy przebiegu:

1) stromość zbocza określona np. czasem narostu od 0,1 do 0,9 wartości w stanie ustalonym;

2) oscylacje wokół wartości stanów ustalonych.

Jeżeli funkcja czasowa $f_n(t)$ określa przebieg wejściowy, to funkcja widmowa jest zgodna z transformacją Fouriera:

$$F_n(\omega) = \int_{-\infty}^{+\infty} f_n(t) \exp(-j\omega t) dt. \quad (17)$$

Niech transmitancja czwórnika będzie określona równaniem:

$$T(\omega) = K(\omega) \exp[-jB(\omega)]. \quad (18)$$

Wobec tego widmo przebiegu odbieranego jest określone przez

$$F_0(\omega) = F_n(\omega) \cdot T(\omega). \quad (19)$$

Kształt przebiegu odbieranego, jego funkcję czasową, wyznacza odwrotna transformata Fouriera

$$f_0(t) = \frac{1}{2\pi} \int_{-\infty}^{\infty} F_n(\omega) \cdot T(\omega) \exp(j\omega t) d\omega. \quad (20)$$

Interesująca jest przede wszystkim odpowiedź układu dla stanu nieustalonego, który fizycznie stanowi moment charakterystyczny przebiegu. Po podstawieniu do wzoru (20) w miejsce $F_n(\omega)$ funkcji widmowej dla skoku jednostkowego równanie funkcji $f_0(t)$ dla stanu nieustalonego przyjmie postać (21) [3]:

$$f'_0(t) = \frac{2}{\pi} \int_0^{\infty} \frac{K(\omega)}{\omega} \sin \omega t \cdot \cos[B(\omega)] d\omega. \quad (21)$$

Czynnik $K(\omega) \cdot \cos[B(\omega)]$ jest częścią rzeczywistą transmitancji $T(\omega)$ określonej wzorem (18).

Równanie (21) sugeruje zależność kształtu przebiegu wyjściowego od charakterystyki fazowej czwórnika. Problem ten przeanalizowany zostanie dla kilku funkcji przesuwności fazowej od pulsacji $B(\omega)$.

3.1. Charakterystyka fazowa w ograniczonym zakresie częstotliwości do ω_{gr} o kształcie oscylacyjnym

$$B(\omega) = \omega t_0 + a \sin n\omega, \quad (22)$$

gdzie:

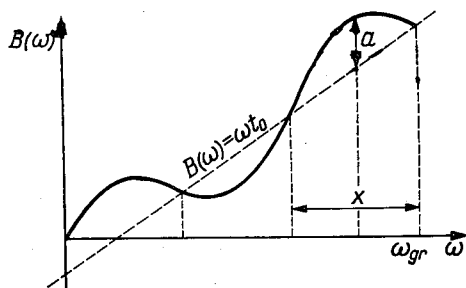
- t_0 — stała wartość opóźności grupowej,
- a — amplituda odchylenia charakterystyki fazowej od prostoliniowej (rys. 3),
- n — częstotliwość oscylacji charakterystyki fazowej w funkcji pulsacji $\left(n = \frac{1}{x}\right)$.

Dla uproszczenia obliczeń zakłada się, że charakterystyka amplitudowa w ograniczonym zakresie do ω_{gr} jest następująca:

$$\left. \begin{aligned} K(\omega) &= 1 & \text{dla} & \quad 0 \leq \omega < \omega_{gr}, \\ K(\omega) &= 0 & \text{dla} & \quad \omega > \omega_{gr}. \end{aligned} \right\} \quad (23)$$

Transmitancja czwórnika przenoszącego impuls na podstawie równań (22) i (23) ma postać:

$$\left. \begin{aligned} T(\omega) &= \exp[-j(\omega t_0 + a \sin n\omega)] & \text{dla } \omega \leq \omega_0, \\ T(\omega) &= 0 & \text{dla } \omega > \omega_0: \end{aligned} \right\} \quad (24)$$



Rys. 3. Przykład charakterystyki fazowej kanału

Po podstawieniu części rzeczywistej transmitancji (24)

$$\operatorname{Re}[T(\omega)] = \cos(\omega t_0 + a \sin n\omega) \quad (25)$$

do ogólnej zależności (21), otrzymamy:

$$\begin{aligned} f_0(t) &= \frac{2}{\pi} \int_0^{\omega_{gr}} \frac{\sin \omega t \cdot \cos[\omega t_0 + a \sin n\omega]}{\omega} d\omega = \\ &= \frac{1}{\pi} \int_0^{\omega_{gr}} \frac{\sin[\omega t + \omega t_0 + a \sin n\omega]}{\omega} d\omega + \frac{1}{\pi} \int_0^{\omega_{gr}} \frac{\sin[\omega t - \omega t_0 - a \sin n\omega]}{\omega} d\omega. \end{aligned} \quad (26)$$

Rozwiązanie ogólne równania (26) jest trudne i nie prowadzi do prostej interpretacji wyników. Jeśli założyć, że czas t_0 jest dostatecznie duży, to pierwsza całka z równania (26) jest prawie równa swojej granicy, to znaczy $\frac{1}{2}$. Wówczas:

$$f'_0(t) \cong \frac{1}{2} + \frac{1}{\pi} \int_0^{\omega_{gr}} \frac{\sin[\omega t - \omega t_0 - a \sin n\omega]}{\omega} d\omega. \quad (26a)$$

Stosując elementarne przekształcenia trygonometryczne

$$\sin(\alpha - \beta) = \sin \alpha \cos \beta - \sin \beta \cos \alpha,$$

gdzie:

$$\alpha = \omega(t - t_0),$$

$$\beta = a \sin n\omega,$$

wzór (26) przyjmie postać:

$$f'_0(t) \approx \frac{1}{2} + \frac{1}{\pi} \int_0^{\omega_{gr}} \frac{\sin \omega(t-t_0) \cos(a \sin n\omega)}{\omega} d\omega + \\ - \frac{1}{\pi} \int_0^{\omega_{gr}} \frac{\cos \omega(t-t_0) \sin(a \sin n\omega)}{\omega} d\omega. \quad (27)$$

Czynniki $\cos(a \sin n\omega)$ i $\sin(a \sin n\omega)$ można przedstawić następującymi wyrażeniami:

$$\left. \begin{aligned} \cos(a \sin n\omega) &= J_0(a) + 2 \sum_{i=1}^{\infty} J_{2i}(a) \cos(2i n\omega), \\ \sin(a \sin n\omega) &= \sum_{i=1}^{\infty} J_{2i-1}(a) \sin[(2i-1)n\omega], \end{aligned} \right\} \quad (28)$$

gdzie: $J_i(a)$ jest funkcją Bessela rzędu i .

Przy założeniu małych amplitud oscylacji charakterystyki fazowej, tzn. że wartości a są małe, można do wzorów (28) podstawić:

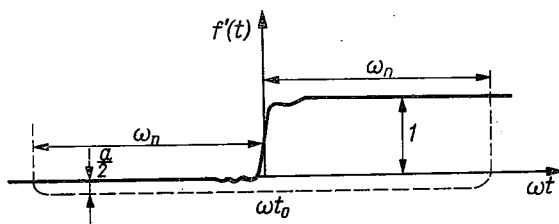
$$\begin{aligned} J_0(a) &\approx 1, \\ J_1(a) &\approx \frac{a}{2}, \\ J_i(a) &= 0 \quad \text{dla} \quad i \geq 2. \end{aligned}$$

Końcowy rezultat wyrażenia (27) będzie następujący:

$$f'_0(t) = \frac{1}{2} + \frac{1}{\pi} \int_0^{\omega_{gr}} \frac{\sin \omega(t-t_0)}{\omega} d\omega - \frac{a}{2\pi} \left[\int_0^{\omega_{gr}} \frac{\sin \omega[(t-t_0)+n]}{\omega} d\omega + \right. \\ \left. - \int_0^{\omega_{gr}} \frac{\sin \omega[(t-t_0)-n]}{\omega} d\omega \right] = \\ = \frac{1}{2} + \frac{1}{\pi} \text{Si } \omega_{gr}(t-t_0) - \frac{a}{2\pi} \text{Si } \omega_{gr}[(t-t_0)+n] - \text{Si } \omega_{gr}[(t-t_0)-n]. \quad (29)$$

Pierwsze dwa składniki (29) charakteryzują przebieg przy liniowym kształcie charakterystyki fazowej, natomiast dwa ostatnie składniki uwzględniają oscylacje charakterystyk. Maksymalna różnica wyrażen całkowitych w nawiasie wynosi π , w związku z tym wartość sygnału zakłócającego nie przekroczy $\frac{a}{2}$ (rys. 4).

Jeżeli przebiegiem wejściowym jest impuls prostokątny, który analitycznie można przedstawić jako superpozycję dwóch skoków jednostkowych przesuniętych w czasie względem siebie, to przebieg wyjściowy posiada kształt pokazany na rys. 5.

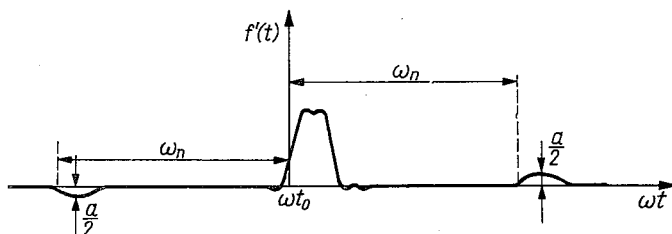


Rys. 4. Odpowiedź na skok jednostkowy czwórnika o charakterystyce fazowej pokazanej na rys. 3

Impulsy zakłócające są przesunięte względem momentu t_0 o czas równy częstotliwości oscylacji charakterystyki fazowej n .

Wyznamy wartość opóźności grupowej dla charakterystyki fazowej określonej równaniem (22). Korzystając z definicji (6) otrzymamy:

$$\tau_g = \frac{dB(\omega)}{d\omega} = t_0 + an \cos n\omega. \quad (30)$$



Rys. 5. Odpowiedź na impuls prostokątny czwórnika o charakterystyce fazowej pokazanej na rys. 3.

Maksymalna różnica opóźności grupowej wynosi:

$$\Delta\tau_{max} = \tau_g - t_0 = [an \cos n\omega]_{max} = an. \quad (31)$$

Stąd

$$a = \frac{\Delta\tau_{max}}{n}.$$

Amplituda zakłócenia jest proporcjonalna do maksymalnej różnicy opóźności grupowej $\Delta\tau$ oraz odwrotnie proporcjonalna do częstotliwości oscylacji charakterystyki fazowej n . Ze zwiększeniem n zmniejsza się amplituda zakłócenia, ale jednocześnie zwiększone zostaje jego przesunięcie czasowe.

3.2. Charakterystyka fazowa w ograniczonym zakresie częstotliwości o kształcie monotonicznym

W tym przypadku można posłużyć się również charakterystyką fazową zapisaną w postaci (22) przy założeniu takich wartości n , że $n\omega_0$ nie przekracza wartości $\frac{\pi}{2}$ w zakresie pulsacji od zera do wartości granicznej ω_{gr} . Praktyczne wykorzystanie w tym przypadku wzoru (26) jest bardzo pracochłonne. Jeżeli:

$$B(\omega) = \omega t_0 + a \sin n\omega,$$

$$\tau_g(\omega) = \frac{dB(\omega)}{d\omega} = t_0 + an \cos n\omega,$$

to

$$\Delta\tau_{max} = \tau_g(0) - \tau_g(\omega_{gr}) = an(1 - \cos n\omega_0). \quad (32)$$

Podstawiając do wzoru (32) następujące wartości:

$$n \cdot \omega_0 = \frac{\pi}{2} \text{ (w zakresie } 0 \div \omega_{gr} \text{ charakterystyka fazowa zmienia się}$$

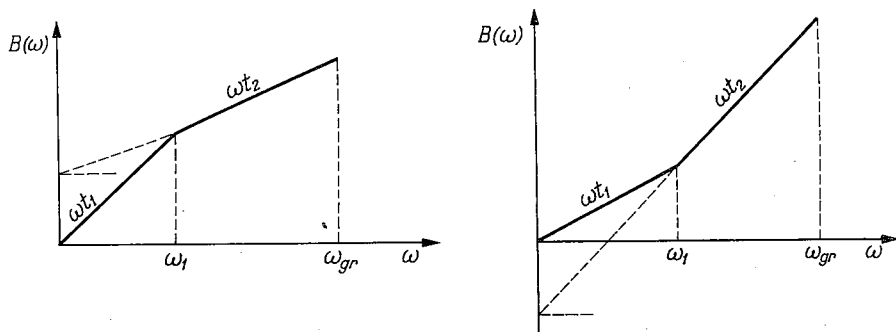
w granicach $\frac{1}{4}$ okresu sinusoidy).

$$\Delta\tau = 200 \text{ } \mu\text{sec},$$

$$\omega_0 = 2\pi \cdot 3000,$$

otrzymamy:

$$a = \frac{\Delta\tau_{max}}{n} = \frac{\Delta\tau_{max}}{\frac{\pi}{2}} \omega_0 = 2,4.$$



Rys. 6. Przykłady aproksymacji charakterystyk fazowych

Tak duża wartość a nie pozwala pominąć w równaniach (28) wartości funkcji Bessela rzędu wyższego od 1, wobec czego zależności końcowe nie dają się w prosty sposób interpretować. Dlatego dla celów praktycznych posługiwać się można aproksymacją charakterystyki fazowej w postaci linii łamanej (dwóch odcinków linii prostych) (rys. 6a i b).

Niech transmitancja czwórnika będzie określona w następujący sposób:

$$\left. \begin{aligned} K(\omega) &= 1 & \text{dla} & \quad 0 \leq \omega \leq \omega_{gr}, \\ K(\omega) &= 0 & \text{dla} & \quad \omega > \omega_{gr}, \\ B(\omega) &= \omega t_1 & \text{dla} & \quad 0 \leq \omega \leq \omega_1 < \omega_{gr}, \\ B(\omega) &= B_0 + \omega t_2 & \text{dla} & \quad \omega_1 < \omega \leq \omega_{gr}. \end{aligned} \right\} \quad (33)$$

Po podstawieniu (33) do ogólnej zależności (21) otrzymamy:

$$f'_0(t) = \frac{2}{\pi} \int_0^{\omega_1} \frac{\sin \omega t \cos \omega t_1}{\omega} d\omega + \frac{2}{\pi} \int_{\omega_1}^{\omega_{gr}} \frac{\sin \omega t \cos (B_0 + \omega t_2)}{\omega} d\omega. \quad (34)$$

Wykorzystując przekształcenia trygonometryczne

$$\sin \alpha \cos \beta = \frac{1}{2} \sin(\alpha + \beta) + \sin(\alpha - \beta),$$

$$\sin(\alpha + \beta) = \sin \alpha \cos \beta + \sin \beta \cos \alpha$$

oraz definicje

$$\text{Si}(x) = - \int_x^\infty \frac{\sin t}{t} dt,$$

$$\text{Ci}(x) = - \int_x^\infty \frac{\cos t}{t} dt,$$

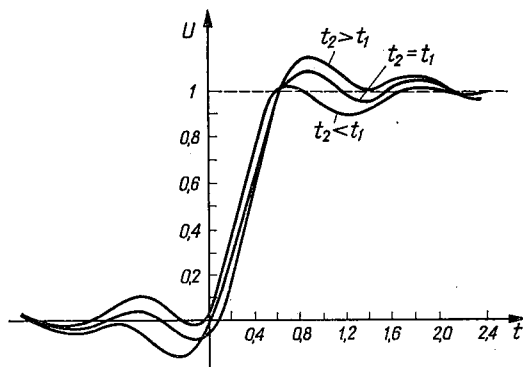
zależność (34) można zapisać w następującej formie:

$$\begin{aligned} f'_0(t) &= \frac{1}{2} + \frac{1}{\pi} \text{Si}[\omega_1(t - t_1)] + \frac{1}{\pi} \cos B_0 \left\{ \text{Si} \left[\omega(t - t_1) - B_0 \frac{\omega_{gr}}{\omega_1} \right] + \right. \\ &\quad \left. - \text{Si}[\omega_1(t - t_1) - B_0] \right\} - \frac{1}{\pi} \sin B_0 \left\{ \text{Ci} \left[\omega_0(t - t_1) - B_0 \frac{\omega_{gr}}{\omega_1} \right] + \right. \\ &\quad \left. - \text{Ci}[\omega_1(t - t_1) - B_0] \right\}. \end{aligned} \quad (35)$$

Interpretacją graficzną (35) jest rys. 7. Załamanie charakterystyki fazowej do góry powoduje zwiększone oscylacje wokół wartości ustalonej 1, załamanie do dołu ma następstwa odwrotne. Gdy zniekształceń fazowych nie ma, oscylacje są z obu stron jednakowe.

Powyżej przytoczono przykłady oddziaływania zniekształceń fazowych na kształt przebiegu odbiorczego przy przesyłaniu skoku jednostkowego. Wnioski powyższe są słuszne również dla przebiegów modulowanych, jeśli pasmo przenoszone jest symetryczne względem pulsacji nośnej Ω . Przy niesymetrycznym ograniczeniu pasma występują zniekształcenia kształtu przebiegu nie tylko związane ze składową synfazową pulsacji nośnej Ω , lecz również zniekształcenia opóźnieńowe powiązane z istnieniem składowej ortogonalnej Ω [3].

Wartość zniekształceń opóźnieniowych zależy również od tego, czy charakterystyka fazowa jest symetryczna lub niesymetryczna względem pulsacji Ω . Mniejsze zniekształcenia występują dla charakterystyk fazowych niesymetrycznych względem pulsacji nośnej Ω .



Rys. 7. Odpowiedź na skok jednostkowy czwórnika o charakterystykach fazowych podanych na rys. 6

W rzeczywistości charakterystyki fazowe są bardziej złożone od tych, które zostały wyżej rozpatrzone. Niektóre z nich nie dadzą się dostatecznie dokładnie aproksymować odcinkami funkcji sinusoidalnych. Często najodpowiedniejszą aproksymacją byłaby funkcja potęgowa lub funkcja sinusoidalna o malejącej lub rosnącej amplitudzie, a także funkcje oscylacyjne o zmieniającym się z pulsacją okresie oscylacji.

Z rozpatrzonego przykładu charakterystyki fazowej w kształcie linii łamanej wynika, że przy takich samych różnicach opóźności grupowej oddziaływanie zniekształceń fazowych może być różne. Dlatego normowanie różnic opóźnienia grupowego jest bardzo ogólnym ujęciem zagadnienia. Wobec powyższego duże znaczenie posiada umiejętność technicznego wyznaczania zniekształceń przebiegu w zależności od kształtu charakterystyki fazowej.

4. WPŁYW ZNIEKSZTAŁCEŃ OPÓŹNIENIOWYCH NA STOPEŃ BŁĘDÓW W TRANSMISJI DANYCH

W poprzednim rozdziale stwierdzono, że niewłaściwe charakterystyki fazowe kanału powodują zmianę czasu narastania impulsów, zwiększenie oscylacji wokół stanów ustalonych oraz powodują dodatkowe impulsy (echo) o amplitudzie i przesunięciu czasowym zależnym od kształtu charakterystyki fazowej. Impulsy echa mogą wystąpić w momencie rozoznania urządzenia odbiorczego i wobec tego można traktować je jak zakłócenia.

Poniższa analiza jest słuszna przy następujących założeniach:

1) Impulsy przesyłane przez kanał mają krótki czas trwania w porównaniu z szerokością pasma kanału. Zakłada się, że:

$$T_0 \cdot \omega_{gr} = \pi, \quad (36)$$

gdzie:

T_0 — czas trwania impulsu (elementu jednostkowego),

ω_{gr} — częstotliwość graniczna kanału.

2) Prawdopodobieństwo pojawienia się impulsów o polaryzacji „0” i „1” na wejściu kanału jest jednakowe, równe 0,5.

3) Sygnały „startu”, „stopu” oraz sygnały synchronizacji nie są przesyłane kanałem teletransmisyjnym, jakkolwiek częstotliwości skali czasu w nadajniku i odbiorniku są utrzymywane z dużą dokładnością.

Wąski impuls jednostkowy nadajnika można analitycznie traktować jako impuls Diraca. Ponieważ charakterystyka amplitudowa nadajnika jest następująca:

$$K_1(\omega) = \frac{1}{\omega_{gr}} \quad \text{dla} \quad 0 < \omega \leq \omega_{gr}$$

$$K_1(\omega) = 0 \quad \text{dla} \quad \omega > \omega_{gr},$$

przeto funkcja czasowa ma postać zgodną z transformatą Fouriera (z dokładnością do stałego współczynnika):

$$S_n(t) = \int_0^{\omega_{gr}} K_1(\omega) \cos \omega t d\omega = \frac{1}{\omega_{gr}} \int_0^{\omega_{gr}} \cos \omega t dt = \frac{\sin \omega_{gr} t}{\omega_{gr} t}. \quad (37)$$

Kształt impulsu na wejściu kanału teletransmisyjnego opisany równaniem (37) nazwiemy znormalizowanym. Jeżeli charakterystyki kanału są zdeterminowane przez funkcję amplitudową $K(\omega)$ i fazową $B(\omega)$, to impuls na wejściu odbiornika ma postać określoną równaniem:

$$S_0(t) = \frac{1}{\omega_{gr}} \int_0^{\omega_{gr}} \exp[-K(\omega)] \cos[\omega t - B(\omega)] d\omega. \quad (38)$$

Rozpatrzone zostanie prawdopodobieństwo błędnego rozeznania w odbiorniku, jeżeli na wyjściu kanału istnieje sygnał właściwy oraz zakłócenia (impulsy echa). Niech impuls w momencie rozeznania posiada wartość chwilową S'_0 , która jest większa od zera, natomiast wartość chwilowa zakłócenia wynosi Z' . Rozeznanie impulsu będzie prawidłowe, jeżeli suma sygnału i zakłócenia przewyższy zero, natomiast rozeznanie będzie błędne, jeżeli:

$$S'_0 - Z' < 0.$$

Niech wartość średniokwadratowa zakłóceń wynosi σ . Prawdopodobieństwo błędnego rozeznania (błędu) jest funkcją stosunku sygnał zakłócenie w momencie rozeznania:

$$P_{bt} = P(Z' > S'_0) = \Phi\left(\frac{S'_0}{\sigma}\right). \quad (39)$$

Zakładając gaussowski rozkład zakłóceń otrzymamy:

$$\Phi\left(\frac{S'_0}{\sigma}\right) = \frac{1}{\sqrt{2\pi}} \int_{\left(\frac{S'_0}{\sigma}\right)}^{\infty} \exp\left[-\frac{\left(\frac{S'_0}{\sigma}\right)^2}{2}\right] d\left[\frac{S'_0}{\sigma}\right]. \quad (40)$$

W dalszym ciągu stosunek $\frac{S'_0}{\sigma}$ będzie traktowany jako różnica poziomów w neperach dla sygnału i zakłócenia:

$$A_z = \ln \frac{S'_0}{\sigma}. \quad (41)$$

Niech impuls określony wzorem (37) zostanie wysłany z nadajnika w momencie $t = 0$, natomiast rozeznanie impulsu w odbiorniku następuje w momencie $t = t'$. Wartość chwilową sygnału w momencie rozeznania można wyrazić następującą zależnością:

$$S'_0 = S_0(t') = \frac{1}{\omega_{gr}} \int_0^{\omega_{gr}} \exp[-K(\omega)] \cdot \cos[\omega t' - B(\omega)] d\omega. \quad (42)$$

Dowolny impuls wysłany poprzednio w momencie $t = nT_0$ wywołuje w momencie rozeznania t' składową zakłóceń:

$$Z_n = \pm S_n(t' + nT_0). \quad (43)$$

Jeżeli impulsy są wysyłane w sposób ciągły, to przy jednakowych prawdopodobieństwach „0” i „1” całkowite (wypadkowe) zakłócenie jest charakteryzowane przez wartość średniokwadratową:

$$\sigma^2 = \sum_{n=-\infty}^{+\infty} S_n^2(t' + nT_0), \quad (44)$$

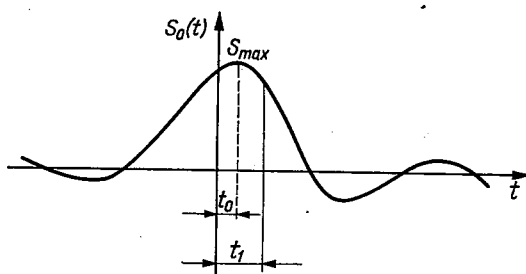
gdzie $n \neq 0$.

Maksymalna różnica poziomów sygnału i zakłócenia wystąpi w momencie czasowym t_0 , dla którego $S'_0 = S_{max}$, to znaczy:

$$\left. \frac{dS'_0(t)}{dt} \right|_{t_0} = 0. \quad (45)$$

Wielkość ta oznaczona zostanie przez A :

$$A = \ln \frac{S_{max}}{\sigma}. \quad (46)$$



Rys. 8. Kształt sygnału w odbiorniku. Optymalny moment rozeznania w punkcie t_0

Według [1] stosunek sygnał-zakłócenie A można z dobrym przybliżeniem (jeśli jest zależny tylko od kształtów charakterystyk amplitudowych i fazowych) obliczyć z następującej zależności:

$$A_p = \ln \frac{S_{max}}{\sigma} \approx \frac{1}{2} \ln \frac{1}{K_\sigma^2 + B_\sigma^2}, \quad (47)$$

gdzie:

K_σ — jest minimalną wartością średniokwadratową charakterystyki $K(\omega)$;

B_σ — jest minimalną wartością średniokwadratową odchylenia rzeczywistej charakterystyki $B(\omega)$ od prostopadłowej wyrażonej zależnością:

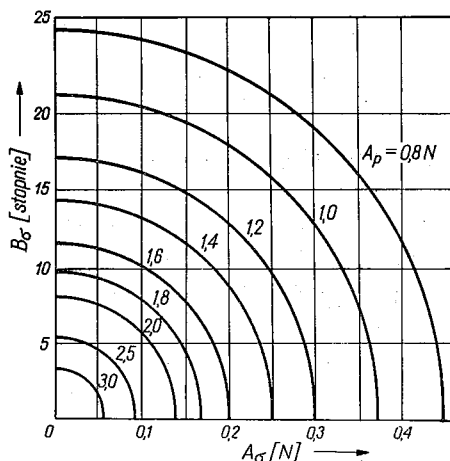
$$B(\omega) = \omega t_0.$$

Przy zadanych funkcjach $K(\omega)$ i $B(\omega)$ wielkości K_σ i B_σ wyznacza się z równań:

$$\left. \begin{aligned} K_\sigma &= \sqrt{\frac{1}{\omega_{gr}} \int_0^{\omega_{gr}} [K(\omega) - K_{sr}]^2 d\omega}, \\ B_\sigma &= \sqrt{\frac{1}{\omega_{gr}} \int_0^{\omega_{gr}} [B(\omega) - \omega t_0]^2 d\omega}, \end{aligned} \right\} \quad (48)$$

$$\left. \begin{aligned} K_{sr} &= \frac{1}{\omega_{gr}} \int_0^{\omega_{gr}} K(\omega) d\omega, \\ t_0 &= \frac{3}{\omega_{gr}^3} \int_0^{\omega_{gr}} B(\omega) d\omega, \end{aligned} \right\}$$

- K_{sr} — jest wartością średnią funkcji amplitudowej w pasmie od zera do ω_{gr} ;
- t_0 — jest współczynnikiem kątowym prostej $B(\omega) = \omega t_0$ na charakterystyce fazowej, obliczonym w taki sposób, aby wartość B_σ była minimalna [pochodną $\frac{dB_\sigma}{dt_0}$ równania (48) należy przyrównać do zera].



Rys. 9.

Rys. 9 jest graficzną ilustracją równania (47). Konstrukcja parametryczna na wykresu pozwala ocenić jednocześnie wpływ zniekształceń amplitudowych i fazowych. W technicznej rzeczywistości stosunek sygnał-zakłócenie A jest determinowany przez urządzenia odbiorcze. Przez badanie odporności odbiornika na zakłócenia można wyznaczyć wartość A , a następnie na podstawie równań (47), (48) i rys. 9 — określić dopuszczalne wartości średniokwadratowe charakterystyk amplitudowej — K_σ i fazowej — B_σ .

Wielkości K_σ i B_σ nie dają się mierzyć bezpośrednio. Na podstawie pomiarów otrzymywane są funkcje $K(\omega)$ i $B(\omega)$ w postaci graficznej. W praktyce mogą powstać trudności ze znalezieniem właściwej aproksymacji charakterystyk, to znaczy ze znalezieniem właściwych funkcji analitycznych, które mogłyby być podstawione do równań (48). Jednak w większości przypadków charakterystyki fazowe kanałów $B(\omega)$ dają się aproksymować funkcjami potęgowymi oraz kosinusoidalnymi. Wówczas obliczenia mogą być przeprowadzone w sposób podany niżej.

4.1. Funkcje potęgowe

Niech charakterystyki fazowe lub opóźnień w funkcji pulsacji będą opisane następującymi równaniami:

$$\left. \begin{aligned} B(\omega) &= \delta B \left(\frac{\omega}{\omega_{gr}} \right)^n, \\ \tau_{gr}(\omega) &= \frac{dB}{d\omega} = \delta \tau \left(\frac{\omega}{\omega_{gr}} \right)^n, \end{aligned} \right\} \quad (49)$$

gdzie:

δB — jest maksymalną różnicą przesuwności fazowej w pasmie;

$\delta \tau$ — jest maksymalną różnicą opóźności grupowej w pasmie;

ω_{gr} — jest pulsacją graniczną pasma;

n — wykładnikiem funkcji potęgowej.

Obliczamy wartość B_σ przy znajomości charakterystyki $B(\omega)$. Podstawiając pierwsze wyrażenie z (49) do wyrażenia czwartego w (48) otrzymuje się:

$$t_0 = \frac{3}{\omega_{gr}^3} \int_0^{\omega_{gr}} \omega B(\omega) d\omega = \frac{3\delta B}{\omega_{gr}^3} \int_0^{\omega_{gr}} \frac{\omega^{n+1}}{\omega_{gr}^n} d\omega = \frac{3\delta B}{\omega_{gr}^{n+3}} \int_0^{\omega_{gr}} \omega^{n+1} d\omega = \frac{3\delta B}{n+2} \frac{1}{\omega_{gr}}. \quad (50)$$

Podstawiając wynik wyrażenia (50) do drugiego wyrażenia (48) otrzymamy:

$$\begin{aligned} B_\sigma &= \sqrt{\frac{1}{\omega_{gr}} \int_0^{\omega_{gr}} \left[\delta B \left(\frac{\omega}{\omega_{gr}} \right)^n - \frac{3\delta B}{n+2} \frac{\omega}{\omega_{gr}} \right]^2 d\omega} = \\ &= \sqrt{\frac{\delta B^2}{\omega_{gr}} \left\{ \int_0^{\omega_{gr}} \left(\frac{\omega}{\omega_{gr}} \right)^{2n} d\omega - \frac{6}{n+2} \int_0^{\omega_{gr}} \left(\frac{\omega}{\omega_{gr}} \right)^{n+1} d\omega + \frac{9}{(n+2)^2} \int_0^{\omega_{gr}} \left(\frac{\omega}{\omega_{gr}} \right)^2 d\omega \right\}} = \\ &= \sqrt{\delta B^2 \left\{ \frac{1}{2n+1} - \frac{6}{(n+2)^2} + \frac{3}{(n+2)^2} \right\}} = \delta B \sqrt{\frac{1}{2n+1} - \frac{3}{(n+2)^2}}. \quad (51) \end{aligned}$$

Wartości średniokwadratowe B_σ są proporcjonalne do maksymalnej różnicy przesuwności fazowej δB oraz do współczynnika m_B , zależnego od wykładnika funkcji potęgowej n .

$$\left. \begin{aligned} B_\sigma &= \delta B \cdot m_B, \\ m_B &= \sqrt{\frac{1}{2n+1} - \frac{3}{(n+2)^2}}. \end{aligned} \right\} \quad (52)$$

W tablicy 1 podano wartości m_B dla kilku wykładników n .

Tablica 1

n	1	2	3	4	5	6
m_B	0	0,112	0,151	0,167	0,172	0,173

Przy ustalonej wartości δB , wartości średniokwadratowe B_σ rosną przy większym wykładniku n . Oczywiście jest, że wartość średniokwadratowa B_σ równa jest zeru dla charakterystyki fazowej $B(\omega)$ prostoliniowej ($m_B = 0$).

Obliczona zostanie wartość B przy znajomości charakterystyki opóźnienia $\tau_{gr}(\omega)$. Ponieważ

$$\tau_{gr}(\omega) = \frac{dB(\omega)}{d\omega} = \delta\tau \left(\frac{\omega}{\omega_{gr}} \right)^n,$$

to charakterystyka fazowa jest określona równaniem:

$$B(\omega) = \int \omega_{gr}(\omega) d\omega = \delta\tau \int \left(\frac{\omega}{\omega_{gr}} \right)^n d\omega = \frac{\delta\tau}{n+1} \left(\frac{\omega}{\omega_{gr}} \right)^{n+1} \cdot \omega. \quad (53)$$

Podstawiając wynik równania (53) do wyrażenia czwartego w równaniu (48) otrzymamy:

$$t_0 = \frac{3\delta\tau}{\omega_{gr}^3(n+1)} \int_0^{\omega_{gr}} \frac{\omega^{n+2}}{\omega_{gr}^n} d\omega = \frac{3\delta\tau}{(n+1)(n+3)}. \quad (54)$$

Z kolei obliczamy B_σ z drugiego wyrażenia równania (48):

$$\begin{aligned} B_\sigma &= \sqrt{\frac{1}{\omega_{gr}} \int_0^{\omega_{gr}} \delta\tau^2 \left[\frac{\omega}{n+1} \left(\frac{\omega}{\omega_{gr}} \right)^n - \frac{3\omega}{(n+1)(n+3)} \right]^2 d\omega} = \\ &= \frac{\delta\tau}{(n+1)} \sqrt{\frac{1}{\omega_{gr}} \left\{ \int_0^{\omega_{gr}} \frac{\omega^{2n+2}}{\omega_{gr}^{2n}} d\omega - \frac{6}{n+3} \int_0^{\omega_{gr}} \frac{\omega^{n+2}}{\omega_{gr}^n} d\omega + \frac{9}{n+3} \int_0^{\omega_{gr}} \omega^2 d\omega \right\}} = \\ &= \frac{\delta\tau}{n+1} \sqrt{\frac{\omega_{gr}^2}{2n+3} - 6 \frac{\omega_{gr}^2}{(n+3)^2} + 3 \frac{\omega_{gr}^2}{(n+3)^2}} = \frac{\delta\tau\omega_{gr}}{n+1} \sqrt{\frac{1}{2n+3} - \frac{3}{(n+3)^2}}. \end{aligned} \quad (55)$$

Podstawiając do powyższego wzoru zależność wynikającą z założeń wymienionych poprzednio

$$\omega_{gr} \cdot T_0 = \pi$$

otrzymamy:

$$B_\sigma = \frac{\delta\tau}{T_0} \cdot \frac{\pi}{n+1} \sqrt{\frac{1}{2n+3} - \frac{3}{(n+3)^2}} \quad (56)$$

lub

$$B_\sigma = \frac{\delta\tau}{T_0} \cdot m_\tau,$$

gdzie:

$$m_\tau = \frac{\pi}{n+1} \sqrt{\frac{1}{2n+3} - \frac{3}{(n+3)^2}}. \quad (57)$$

W tablicy 2 podano wartości m_τ dla kilku wykładników n .

Tablica 2

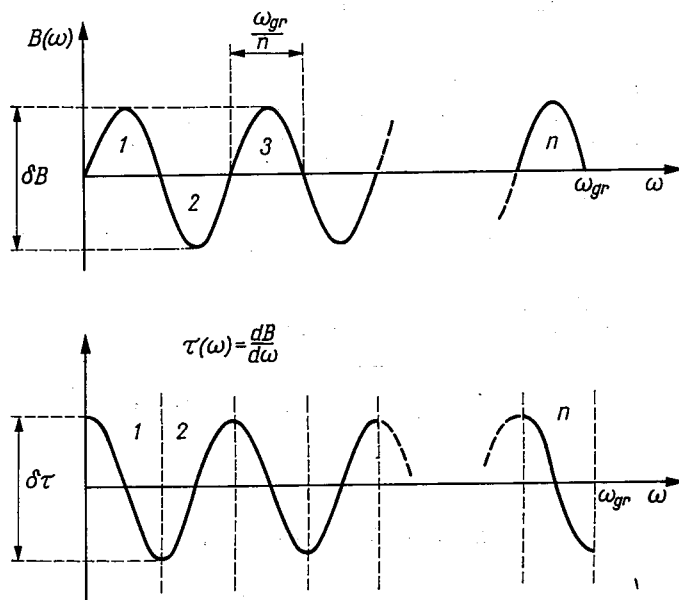
n	1	2	3	4	5	6
m_τ	10,1°	9,1°	7,5°	6,2°	5,0°	4,4°

Wartość średniokwadratową zniekształceń fazowych B_σ należy rozpatrywać przy założonej szybkości modulacji (T_0 jest długością elementu jednostkowego modulacji), ponieważ znajomość maksymalnej różnicy opóźności grupowej w pasmie $\delta\tau$ nie określa zniekształceń fazowych. Przy takich samych wartościach $\frac{\delta\tau}{T_0}$ mniejsze zniekształcenia fazowe występują, gdy funkcja opóźności grupowej od pulsacji jest wyższego rzędu.

4.2. Funkcje oscylacyjne

Niech charakterystyki fazowe i opóźnieniowe będą opisane następującymi funkcjami pulsacji (rys. 10):

$$\left. \begin{aligned} B(\omega) &= \frac{\delta B}{2} \sin\left(\pi n \frac{\omega}{\omega_{gr}}\right), \\ \tau(\omega) &= \frac{\delta\tau}{2} \cos\left(\pi n \frac{\omega}{\omega_{gr}}\right), \end{aligned} \right\} \quad (58)$$



Rys. 10. Przykład charakterystyk oscylacyjnych przesuwności fazowej i opóźności grupowej

gdzie:

δB — maksymalna różnica przesuwności fazowej w pasmie;

$\delta\tau$ — maksymalna różnica opóźności grupowej w pasmie;

n — ilość półokresów charakterystyki fazowej lub opóźnieniowej w zakresie od zera do ω_{gr} .

Podstawiając funkcję $B(\omega)$ zapisaną w postaci (58) do czwartego wyrażenia w równaniu (48) otrzymamy:

$$t_0 = \frac{3\delta B}{2\omega_{gr}^3} \int_0^{\omega_{gr}} \omega \sin\left(\pi n \frac{\omega}{\omega_{gr}}\right) d\omega = \frac{-3\delta B}{2\pi n \omega_{gr}}. \quad (59)$$

Po podstawieniu (59) do drugiego wyrażenia w równaniu (48):

$$B_\sigma = \frac{\delta B^2}{2} \sqrt{\frac{1}{\omega_{gr}} \int_0^{\omega_{gr}} \left[\sin \pi n \frac{\omega}{\omega_{gr}} + \frac{3}{\pi n} \frac{\omega}{\omega_{gr}} \right]^2 d\omega}. \quad (60)$$

Podnosząc wyrażenie w nawiasie do kwadratu, pod pierwiastkiem otrzyma się sumę następujących całek:

$$\left. \begin{aligned} \frac{1}{\omega_{gr}} \int_0^{\omega_{gr}} \sin^2\left(\pi n \frac{\omega}{\omega_{gr}}\right) d\omega &= \frac{1}{\omega_{gr}} \left[\int_0^{\omega_{gr}} \frac{1}{2} d\omega - \frac{1}{2} \int_0^{\omega_{gr}} \sin 2\pi n \frac{\omega}{\omega_{gr}} d\omega \right] = \frac{1}{2} \\ - \frac{6}{\pi n \omega_{gr}^2} \int_0^{\omega_{gr}} \omega \sin\left(\pi n \frac{\omega}{\omega_{gr}}\right) d\omega &= \\ - \frac{6}{\pi n \omega_{gr}^2} \left[\frac{\omega_{gr}^2}{\pi n} - \frac{\omega_{gr}}{\pi n} \int_0^{\omega_{gr}} \cos\left(\pi n \frac{\omega}{\omega_{gr}}\right) d\omega \right] &= - \frac{6}{\pi^2 n^2}, \\ \frac{9}{\pi^2 n^2 \omega_{gr}^3} \int_0^{\omega_{gr}} \omega^2 d\omega &= \frac{3}{\pi^2 n^2}. \end{aligned} \right\} \quad (61)$$

Po podstawieniu (61) do (60) otrzymamy:

$$B_\sigma = \frac{\delta B}{2} \sqrt{\frac{1}{2} - \frac{3}{\pi^2 n^2}} = \delta B \cdot m_B, \quad (62)$$

gdzie:

$$m_B = \frac{1}{2} \sqrt{\frac{1}{2} - \frac{3}{\pi^2 n^2}}.$$

W tablicy 3 podano wartości m_B dla kilku typów charakterystyk oscylacyjnych.

Tablica 3

n	1	2	3	4	5	6
m_B	0,177	0,261	0,353	0,353	0,353	0,353

Z kolei wyznaczmy wartość średniokwadratową B_σ przy znajomości charakterystyki opóźności grupowej $\tau(\omega)$, określonej w postaci (58).

$$B(\omega) = \frac{\delta\tau}{2} \int \cos\left(\pi n \frac{\omega}{\omega_{gr}}\right) d\omega = \frac{\delta\tau}{2} \frac{\omega_{gr}}{\pi n} \sin\left(\pi n \frac{\omega}{\omega_{gr}}\right) = \frac{\delta\tau}{T_0} \frac{1}{2n} \sin\left(\pi n \frac{\omega}{\omega_{gr}}\right). \quad (63)$$

Funkcja $B(\omega)$ określona przez (63) różni się jedynie stałymi czynnikami od funkcji $B(\omega)$ określonej przez (58), wobec czego wynik całkowania wg równań (48) może być zapisany następująco:

$$B_\sigma = \frac{\delta\tau}{T_0} \frac{1}{2n} \sqrt{\frac{1}{2} - \frac{3}{\pi^2 n^2}} = \frac{\delta\tau}{T_0} m_\tau, \quad (64)$$

gdzie:

$$m_\tau = \frac{1}{2n} \sqrt{\frac{1}{2} - \frac{3}{\pi^2 n^2}}.$$

Wartości m_τ (w stopniach) dla kilku typów charakterystyk opóźnieńowych podano w tablicy 4.

Tablica 4

n	1	2	3	4	5	6
m_τ	12,7°	10,1°	6,8°	5,1°	4,1°	3,4°

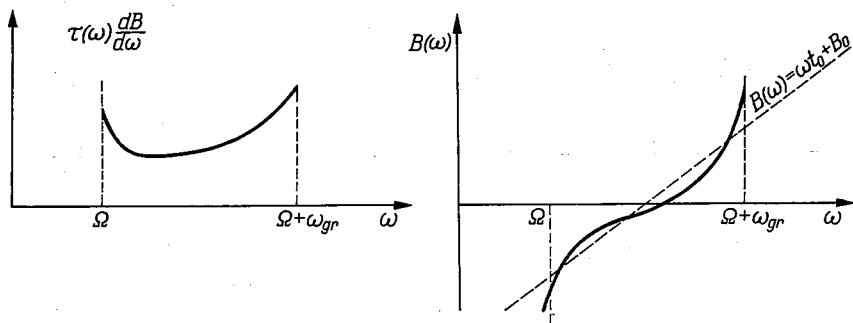
Z tablic 3 i 4, a zwłaszcza z ostatniej, wynika, że mniejszy będzie wpływ zniekształceń fazowych, gdy przy takich samych maksymalnych różnicach opóźności grupowej $\delta\tau$ charakterystyka posiada większą częstotliwość oscylacji.

Porównajmy wyniki w tablicy 2 i 4. Zakładając takie same wartości $\frac{\delta\tau}{T_0}$ widać, że charakterystyki oscylacyjne dają mniejsze wartości B_σ niż charakterystyki potęgowe.

4.3. Wyznaczenie B_σ dla przebiegów modulowanych

Poprzednie zależności pozwalające obliczyć wartości minimalnego odchylenia średniokwadratowego B_σ są słuszne dla przebiegów, których widmo praktycznie zawiera się w granicach od zera do pulsacji ω_{gr} .

Dla przebiegów modulowanych, częstotliwości składowe widma, które mają istotny wpływ na jakość transmisji, zawarte są w przedziale pulsacji Ω , $\Omega + \omega_{gr}$, przy czym Ω nie jest równa zero. Niech charakterystyka opóźnieniowa i fazowa kanału będzie podana jak na rys. 11.



Rys. 11

Wielkość B_σ należy w tym przypadku obliczyć ze wzoru:

$$B_\sigma = \sqrt{\frac{1}{\omega_{gr}} \int_{\Omega}^{\Omega + \omega_{gr}} [B(\omega) - (\omega t_0 + B_0)]^2 d\omega} \quad (65)$$

gdzie:

$\omega_{gr} = (\Omega + \omega_{gr}) - \Omega$ — jest szerokością pasma;

$B(\omega)$ — funkcja przesuwności fazowej od pulsacji w zakresie Ω , $\Omega + \omega_{gr}$;

$\omega t_0 + B_0$ — równanie prostej na charakterystyce fazowej, przy której B_σ osiąga wartość minimalną.

Wzór (65) będzie miał praktyczny użytek, jeżeli podane będą równania pozwalające obliczyć wartość t_0 i B_0 na podstawie znajomości funkcji $B(\omega)$. Traktując B_σ z równania (65) jako funkcję $B_\sigma(t_0, B_0)$, wartości t_0 i B_0 można obliczyć przez przyrównanie pochodnych cząstkowych do zera:

$$\left. \begin{aligned} \frac{\partial B_\sigma}{\partial t_0} &= 0, \\ \frac{\partial B_\sigma}{\partial B_0} &= 0. \end{aligned} \right\} \quad (66)$$

Różniczkując równanie (65), warunki (66) sprowadzą się do następujących zależności:

$$\int_{\Omega}^{\Omega + \omega_{gr}} \frac{\partial}{\partial t_0} [B(\omega) - (\omega t_0 + B_0)]^2 d\omega = 0, \quad (67)$$

$$\int_{\Omega}^{\Omega + \omega_{gr}} \frac{\partial}{\partial B_0} [B(\omega) - (\omega t_0 + B_0)]^2 d\omega = 0. \quad (68)$$

Po zróżniczkowaniu wyrażenia w nawiasie kwadratowym następujące sumy całek powinny być równe zeru:

$$\left. \begin{aligned} t_0 \int_{\Omega}^{\Omega+\omega_{gr}} \omega^2 d\omega - \int_{\Omega}^{\Omega+\omega_{gr}} \omega B(\omega) d\omega + B_0 \int_{\Omega}^{\Omega+\omega_{gr}} \omega d\omega &= 0, \\ B_0 \int_{\Omega}^{\Omega+\omega_{gr}} d\omega - \int_{\Omega}^{\Omega+\omega_{gr}} B(\omega) d\omega + t_0 \int_{\Omega}^{\Omega+\omega_{gr}} \omega d\omega &= 0. \end{aligned} \right\} \quad (69)$$

Wyniki całkowania są następujące:

$$\left. \begin{aligned} \int_{\Omega}^{\Omega+\omega_{gr}} \omega^2 d\omega &= \frac{1}{3} (3\Omega \cdot \omega_{gr}^2 + 3\Omega^2 \cdot \omega_{gr} + \omega_{gr}^3), \\ \int_{\Omega}^{\Omega+\omega_{gr}} \omega d\omega &= \frac{1}{2} (2\Omega \cdot \omega_{gr} + \omega_{gr}^2), \\ \int_{\Omega}^{\Omega+\omega_{gr}} d\omega &= \omega_{gr}. \end{aligned} \right\} \quad (70)$$

Po podstawieniu (70) do (69) otrzymane będą następujące dwa równania:

$$\left. \begin{aligned} \frac{t_0}{3} (3\Omega \cdot \omega_{gr}^2 + 3\Omega^2 \cdot \omega_{gr} + \omega_{gr}^3) + \frac{B_0}{2} (2\Omega \cdot \omega_{gr} + \omega_{gr}^2) &= \int_{\Omega}^{\Omega+\omega_{gr}} \omega B(\omega) d\omega, \\ \frac{t_0}{2} (2\Omega \cdot \omega_{gr} + \omega_{gr}^2) + B_0 \cdot \omega_{gr} &= \int_{\Omega}^{\Omega+\omega_{gr}} B(\omega) d\omega. \end{aligned} \right\} \quad (71)$$

Rozwiązując (71) jako układ dwóch równań z dwiema niewiadomymi t_0 i B_0 , wyniki będą następujące:

$$B_0 = \frac{4}{\omega_{gr}} \left[1 + 3 \left(\frac{\Omega}{\omega_{gr}} \right)^3 + 3 \left(\frac{\Omega}{\omega_{gr}} \right) \right] \int_{\Omega}^{\Omega+\omega_{gr}} B(\omega) d\omega - \frac{6}{\omega_{gr}^2} \left[1 + 2 \left(\frac{\Omega}{\omega_{gr}} \right) \right] \int_{\Omega}^{\Omega+\omega_{gr}} \omega B(\omega) d\omega, \quad (72)$$

$$t_0 = \frac{12}{\omega_{gr}^3} \int_{\Omega}^{\Omega+\omega_{gr}} \omega B(\omega) d\omega + \left[\frac{2}{\omega_{gr}} - 8 \frac{\omega_{gr}^3 + 3\Omega^3 + 3\Omega\omega_{gr}^2}{\omega_{gr}^4 (2\Omega + \omega_{gr})} \right] \int_{\Omega}^{\Omega+\omega_{gr}} B(\omega) d\omega. \quad (73)$$

Wielkości B_0 i T_0 w równaniach (72) i (73) wyznaczają prostą na charakterystyce fazowej $B(\omega)$, przy czym B_0 jest wartością przesuwności fazowej dla pulsacji $\omega = 0$. Dla założonych poprzednio celów potrzebna jest znajomość B_0 jako przesuwności fazowej dla pulsacji $\omega = \Omega$, wobec czego

równania (72) i (73) będą słuszne, jeśli podstawić do nich $\Omega = 0$. Wówczas:

$$\left. \begin{aligned} B_0 &= \frac{4}{\omega_{gr}} \int_{\Omega}^{\Omega+\omega_{gr}} B(\omega) d\omega - \frac{6}{\omega_{gr}^2} \int_{\Omega}^{\Omega+\omega_{gr}} \omega B(\omega) d\omega, \\ t_0 &= \frac{12}{\omega_{gr}^3} \int_{\Omega}^{\Omega+\omega_{gr}} \omega B(\omega) d\omega - \frac{6}{\omega_{gr}^2} \int_{\Omega}^{\Omega+\omega_{gr}} B(\omega) d\omega. \end{aligned} \right\} \quad (74)$$

Przy znajomości funkcji $B(\omega)$ w interesującym pasmie, korzystając z zależności (65) i (74), można obliczyć wartość średniokwadratową B_0 .

5. WNIOSKI

Przy ocenie przydatności kanałów dla transmisji danych należy uwzględniać zniekształcenia opóźnieniowe, które mają bezpośredni wpływ na stopę błędów. Spełnienie warunku wiernego odtworzenia, to znaczy warunku, że opóźność grupowa powinna być proporcjonalna do pulsacji, jest bardzo trudne z dwóch powodów:

1) Kanały teletransmisyjne bez korekcji fazowej nie spełniają tego postulatu w żądanym zakresie częstotliwości.

2) Korekcja fazowa pozwala poprawić charakterystykę $B(\omega)$ lub $\tau(\omega)$ sprowadzając ją do funkcji potęgowej lub oscylacyjnej.

W związku z tym, że zniekształceń opóźnieniowych nie da się wyeliminować, powstają dwa problemy godne uwagi:

1) Jeżeli jest znany system transmisji danych, a ściśle — odporność tego systemu na zakłócenia, to należy umieć wyznaczyć dopuszczalne zniekształcenia opóźnieniowe, jakie może wprowadzać kanał teletransmisyjny.

2) Jeżeli są znane zniekształcenia opóźnieniowe kanału teletransmisyjnego, to należy umieć wyznaczyć minimalną odporność systemu transmisji danych na zakłócenia.

Rozdział 4 tego opracowania jest próbą rozwiązania tych problemów. Przy założeniu gausowskiego rozkładu gęstości prawdopodobieństwa błędu od stosunku sygnał-zakłócenie można podać zależności analityczne między zniekształceniami opóźnieniowymi a ekwiwalentnym stosunkiem sygnał-zakłócenie. W praktyce posiada to duże znaczenie, ponieważ badania odporności układów na zakłócenia są znacznie łatwiejsze niż badania odporności układów na zniekształcenia opóźnieniowe.

Poza tym z analizy przeprowadzonej w rozdziałach 3 i 4 wynika, że normowanie bezwzględnej różnicy opóźności grupowej w wykorzystywa-

nym pasmie jest bardzo ogólnym ujęciem zagadnienia. Analiza różnych kształtów charakterystyk fazowych i opóźnieniowych przy takich samych różnicach opóźności grupowej pozwala sprecyzować wymagania dla korekcji fazowej kanałów.

*Politechnika Warszawska
Katedra Urządzeń Teletransmisyjnych
i Telegraficznych*

WYKAZ LITERATURY

1. CCITT — Dokument N° 21 Kom. Sp. A z 1961 r., „The Influence of Some Transmission Defects on Binary Data Systems”.
2. Dąbrowski M., Pawłowska W.: Studium wymagań technicznych i metod pomiarowych opóźnienia grupowego dla potrzeb transmisji danych. Opracowanie wewnętrzne Zakładu Telegrafii Politechniki Warszawskiej.
3. Dobrowolskij G. W.: Ustanowliwajuszcziesja processy w elektriczeskich cepiach. Moskwa 1948.

M. DĄBROWSKI

THE INFLUENCE OF DELAY-DISTORTIONS IN TELECOMMUNICATION CHANNELS ON DATA-TRANSMISSION SYSTEMS

Summary

In this article the relationship between delay-distortions of telecommunication channels and data-transmission systems is discussed. At the beginning the conditions, which have to be fulfilled by phase and delay-characteristics, are described. Second part to the influence of different channel's phase-characteristics on the electrical phenomena in these channels is devoted. In third part the relations between phase distortions and signal-to-noise value for the different phase and delay characteristics are given. The error-rate of data transmission systems may be estimated from final relations at given values of phase shift and envelope delay.

M. DĄBROWSKI

INFLUENCE DES DISTORSIONS DE RETARD DES CANAUX DES TÉLÉCOMMUNICATIONS SUR LA TRANSMISSION DES DONNÉES

Résumé

Dans l'oeuvre on a analysé l'influence des distorsions de retard des canaux des télécommunications sur la transmission des données. Dans l'introduction on a présenté les idées générales ainsi que les conditions, que les caractéristiques de phase et de retard en fonction de fréquence doivent accomplir afin d'obtenir la reproduction correcte des signaux transmis. Dans la partie suivante on a présenté le problème de la forme des signaux transmis par les canaux des télécommunications

ayant de différentes formes des caractéristiques de phase. Ensuite, dans la troisième partie, on a donné les dépendances entre les distorsions de phase et la relation équivalente signal-bruit en cas des caractéristiques de phase et de retard étant des fonctions exponentielles et oscillatoires de fréquence. Les formules finales permettent d'évaluer le taux des erreurs du système de la transmission des données ayant des caractéristiques de déplacement de phase ou de retard de groupe en fonction de fréquence.

М. ДОМБРОВСКИ

ВЛИЯНИЕ ФАЗОВО-ЧАСТОТНЫХ ИСКАЖЕНИЙ ТЕЛЕФОННЫХ КАНАЛОВ НА ПЕРЕДАЧУ ЦИФРОВЫХ ДАННЫХ

Резюме

В этой статье рассматривается влияние фазово-частотных искажений на качество передачи цифровых данных в телефонных каналах. Сначала автор описывает общие понятия и условия, которым должны удовлетворять фазо-частотные характеристики и характеристики замедления группового времени, при неискаженном приеме передаваемых сигналов. Далее рассматривается формулу принимаемых сигналов в зависимости от вида фазо-частотных характеристик канала. В третьей части представлена связь между фазовыми искажениями и коэффициентом сигнал/шум фазо-частотных характеристик и характеристик замедления группового времени как степенных и колебательных функций частоты. Из конечных формул возможно определить вероятность ошибок системы передачи цифровых данных при заданных характеристиках замедления группового времени и фазового сдвига.

WŁADYSŁAW KOMPAŁO

Technologia oraz podstawowe parametry germanowych tranzystorów stopowych n-p-n

Rękopis dostarczono 2.3.1964

Omówiono technologię oraz podstawowe parametry germanowych tranzystorów stopowych typu TW 6P o częstotliwości granicznej 10–30 MHz i mocy strat kolektora 50 mW oraz typu TW O5P o $f_a = 5\text{--}10$ MHz i P_k rzędu 250 mW.

W pierwszej części artykułu omówiono: a) zagadnienie wyboru i przygotowania materiałów, b) wytwarzanie złącz w próżni, c) montaż tranzystorów, d) elektrolityczne trawienie złącz.

W dalszym ciągu zostały podane podstawowe parametry i ważniejsze charakterystyki tranzystorów n-p-n. Opisano również zasady zastosowanych metod pomiarowych oraz przeprowadzono krótką dyskusję wyników.

W ostatniej części artykułu dokonano oceny tranzystorów n-p-n i p-n-p na bazie porównania maksymalnego wzmocnienia mocy na wielkich częstotliwościach. Omówiono również przebiegi podstawowych parametrów wchodzących w skład wzoru na maksymalne wzmocnienie mocy w funkcji prądu emitera i napięcia kolektora.

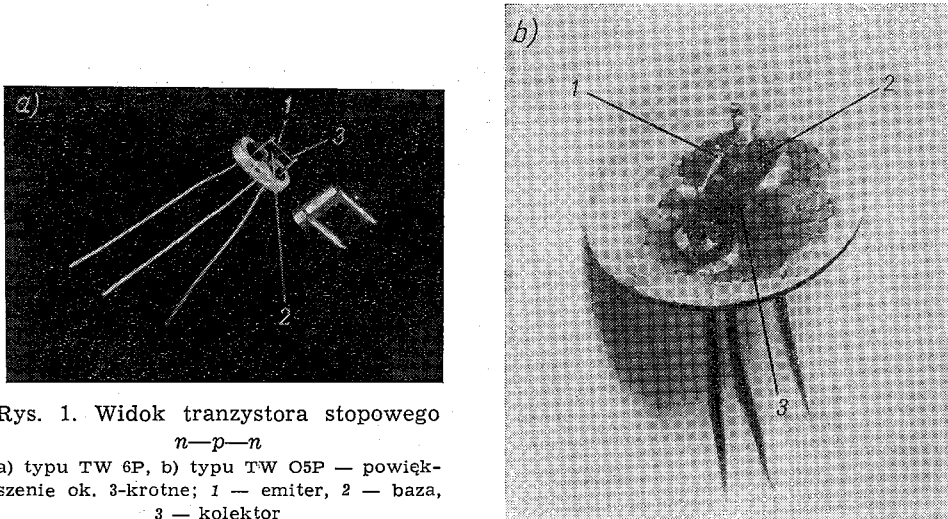
1. WSTĘP

Celem niniejszej pracy jest dokonanie przeglądu procesów technologicznych nad germanowymi tranzystorami stopowymi n-p-n oraz przedstawienie podstawowych parametrów elektrycznych tych tranzystorów. Zakres pracy obejmuje elementy technologii tranzystorów o mocy admysyjnej ok. 50 mW i częstotliwości granicznej rzędu 10–30 MHz (oznaczonych jako TW 6P) oraz wykonanie na tej podstawie tranzystorów o podwyższonej mocy admysyjnej do ok. 250–300 mW i częstotliwości granicznej ponad 5 MHz (oznaczonych jako TW O5P). Z pracami tymi są związane pomiary podstawowych parametrów tranzystorów, zarówno na małych, jak i wysokich częstotliwościach, jako sprawdzian przydatności tych elementów do praktycznych układów. Wszystkie prace zostały wykonane w Zakładzie Elektroniki I.P.P.T. — PAN.

Do wytworzenia obszarów typu n w germanie typu p wykorzystano stosunkowo łatwy pod względem technologicznym proces stopowy. Z procesu tego wynika typowa konstrukcja tranzystorów, polegająca na osiowym wtopieniu dwu kulek stopu o odpowiednim składzie chemicznym

w płytce germanu typu p. Do germanu od strony kulki emitera jest przy-
lutowany pierścieniowy kontakt bazy.

Tranzystory małej mocy są zamykane w standartowych hermetycz-
nych oprawkach, np. takich, jak przedstawione na rys. 1a, albo w opraw-

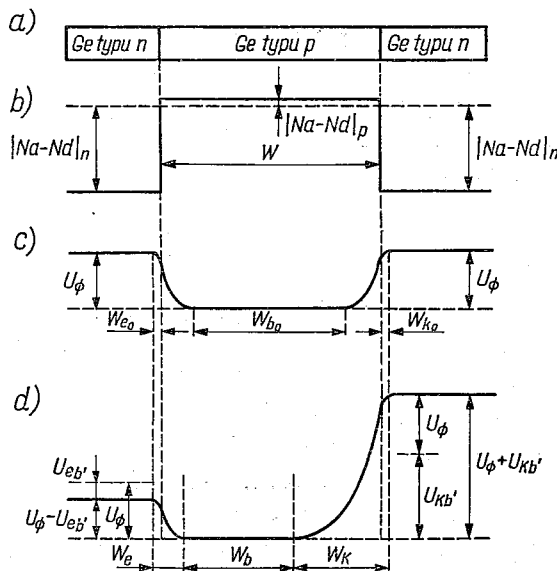


Rys. 1. Widok tranzystora stopowego

$n-p-n$

a) typu TW 6P, b) typu TW 05P — powięk-
szenie ok. 3-krotne; 1 — emiter, 2 — baza,
3 — kolektor

kach tranzystorów TG 10 (20). Tranzystory o podwyższonej mocy są
montowane w oprawkach takich jak tranzystory TG 50 (55) — rys. 1b.



Rys. 2. Schemat tranzystora stopowego $n-p-n$. Uwidoczniono wpływ napięć na
szerokość złącz i grubość bazy

W celu przypomnienia zasad konstrukcji tranzystora stopowego $n-p-n$ oraz w celu ułatwienia zrozumienia zjawisk zachodzących w tranzystorze podczas pracy został podany na rys. 2 schemat ideowy tranzystora stopowego $n-p-n$.

2. MATERIAŁY WYJŚCIOWE I ICH PRZYGOTOWANIE

2.1. German

Fizyczne własności germanu muszą odpowiadać ogólnie przyjętym warunkom tzw. monokryształu tranzystorowego, tzn. oporność właściwa powinna zawierać się w granicach $0,5-4 \Omega\text{cm}$, objętościowy czas życia — w granicach od kilkunastu do kilkudziesięciu μsek , dyslokacje powinny być rozmieszczone równomiernie, o gęstościach rzędu 5000 na cm^2 . Odstępstwa od powyższych wartości spowodują zmianę niektórych parametrów tranzystorów. Tak więc mniejsza oporność właściwa germanu obniży maksymalne napięcie przebiccia, zmniejszy sprawność emitera (γ), może zbyt obniżyć czas życia i ruchliwość nośników mniejszościowych ładunku. Duża oporność właściwa wpłynie niekorzystnie na oporność bazy tranzystora, wzrośnie też wsteczny prąd nasycenia złącza (tzw. prąd zerowy). Zbyt duży czas życia może obniżyć częstotliwość graniczną.

Przygotowanie germanu polega na cięciu monokryształu w płaszczyźnie [111], a następnie na doprowadzeniu powierzchni płytki do wysokiej gładkości przy idealnej czystości. Monokryształ jest najpierw cięty piłami tarczowymi przy użyciu proszków ściernych na płytki grubości ok. 0,3 mm. Płytki te są ścirane przy pomocy drobnoziarnistych proszków karborundu czy alundum do odpowiedniej grubości, a dalej trawione chemicznie do grubości około 50μ dla tranzystorów TW 6P, a około 90μ dla TW O5P. Do trawienia chemicznego używa się mieszanki CP_4 , która daje najlepszą jakość powierzchni. Końcową czynnością jest łamanie i sortowanie płytek w grupy o rozrzucie grubości co 3μ .

2.2. Donorowe stopy domieszkowe

Wybrany metal, ewentualnie stop domieszkowy, musi zapewnić:

- a) dostatecznie dużą przewodność po rekryształizacji stopionego obszaru,
- b) wystarczająco dobrą rozpuszczalność germanu i łatwą do kontroli penetrację stopu w odpowiednich temperaturach,
- c) niepowodowanie powstawania naprężeń mechanicznych w germanie.

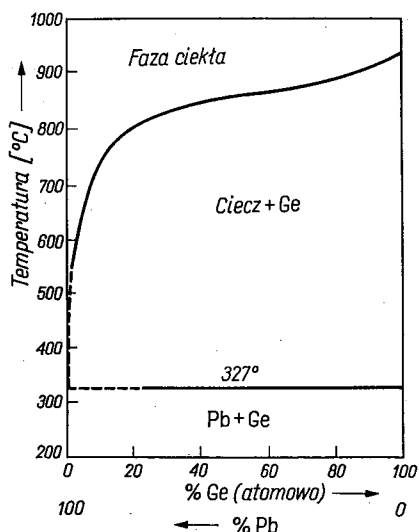
Spośród wielu metali dających przewodnictwo typu n (Bi, Se, Te, P, S) najbardziej odpowiednie są antymon i arsen. Mają one odpowiednią stałą segregacji w germanie (dla Sb: $k = 0,003$, dla As: $k = 0,04$) i dają złącza

o małej oporności w kierunku przewodzenia. Wadami ich są: a) duża twardość i kruchość, inne współczynniki rozszerzalności niż germanu; dlatego stosuje się je łącznie z miękkim metalem nośnym (Pb, Sn), b) duże ciśnienie par w temperaturach wtapienia. Dla orientacji można podać, że dla temperatury około 700°C ciśnienie par As jest rzędu ciśnienia atmosferycznego, dla Pb i Sb wynosi 10^{-2} Tr, a dla cyny jest tylko 10^{-4} Tr.

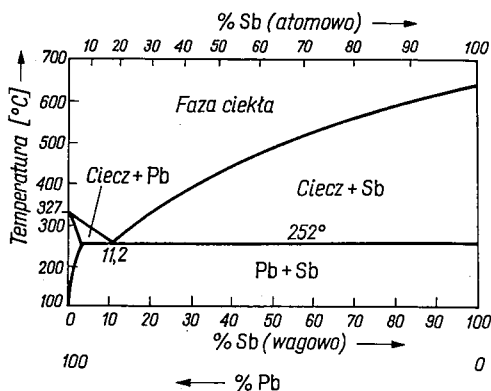
Na metal nośny stopu wybrano ołów, ponieważ mniej roztwarza Ge, choć ciśnienie par ma większe niż cyna. Jest to ważne, gdyż przy wysokich temperaturach procesu stopowego pożądana jest mała zależność głębokości wtopienia od nieznacznych zmian temperatury. Zwilżanie germanu przez stopy cynowe jest natomiast lepsze niż przez stopy z ołowiem.

Jakość zwilżania można ocenić praktycznie obliczając stosunek średnicy kulki po wtopieniu do średnicy kulki przed wtopieniem. Dla stosowanych stopów ołowiowych stosunek ten wynosi 1,2 : 1,3.

Najbardziej odpowiedni skład procentowy stopów jest następujący: 10% Sb — 90% Pb i (1÷3)% As — (99÷97)% Pb. Wybór takiego składu dla antymonu wynika z istnienia eutektyki stopu Sb—Pb przy 11,2% Sb (rys. 4). Przy wyższej procentowości stopy są kruche i twarde. Procentowość As w stopie wynika z większej jego stałej segregacji w germanie



Rys. 3. Wykres podwójny ołów-german



Rys. 4. Wykres podwójny ołów-antymon

(około 13 razy większej niż dla Sb). Przy dalszych pracach stosowano prawie wyłącznie stopy Sb—Pb ze względu na zbyt dużą parowalność arsenu w warunkach procesu próżniowego i z powodu bardziej równomiernej rekryształizacji germanu ze stopów ołowiowych.

Przy wykonywaniu omawianych stopów domieszkowych należy zwrócić uwagę na ich odgazowanie i zachowanie jednorodności.

Kulki otrzymujemy podgrzewając w atmosferze wodoru bryłki pociętego stopu do temperatury ok. 600°C . Następnie, jeżeli kulki mają zanieczyszczenia na powierzchni, są lekko trawione w kilkuprocentowym HCl i ponownie wygrzewane w wodorze. Po wygrzaniu następuje pomiar średnic i segregacja. Postępowanie takie jakkolwiek prymitywne okazało się w praktyce najlepsze, pozwala bowiem uzyskać wystarczająco czystą powierzchnię metaliczną kulek.

Wymiary kulek i złącz przedstawia tablica 1.

Tablica 1

	Średnica kulek		Średnica złącz		$\varnothing$ kulek/ $\varnothing$ złącz	$\varnothing$ kol./ $\varnothing$ em.
	emitera	kolekto- ra	emitera	kolekto- ra		
Tranzystor TW 6P	290—320	400—430	350	450	1,1—1,2	1,3
Tranzystor TW O5P	590—620	810—840	790	1000	1,2—1,3	1,3

Uwaga. Wymiary w mikronach.

2.3. Materiały pomocnicze

Do grupy materiałów pomocniczych zostaną zaliczone wszystkie elementy służące do wykonania połączeń od elektrod tranzystora do przepustów oprawki oraz oprawki z przykrywkami. Jednakże nie będą one tutaj omawiane, ponieważ są zbliżone do analogicznych części stosowanych przy produkcji innych rodzajów tranzystorów, a tym samym są powszechnie znane.

3. WYKONANIE ZŁĄCZ P—N I MONTAŻ TRANZYSTORÓW

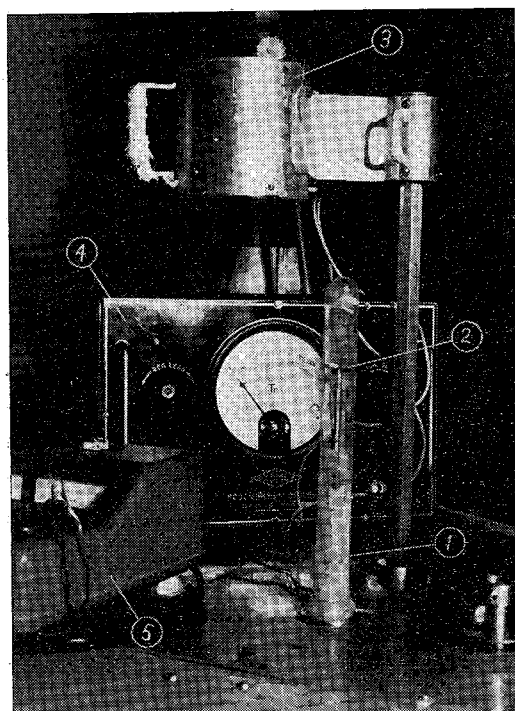
3.1. Opis procesu stopowego

Proces stopowy został przeprowadzony w dwóch etapach. Zadaniem pierwszego etapu jest płytkie wtopienie kulek w german z położeniem nacisku na jak najlepszą ich współosiowość. Właściwe wytwarzanie złącz następuje w drugim etapie.

Pierwsze stopienie kulek z germanem jest przeprowadzane w redukującej atmosferze wodoru w specjalnym przyrządzie grafitowym (w tzw. „kasecie”). Od odpowiedniej konstrukcji kasety zależy, czy po wygrzaniu elektrody będą wtopione współosiowo, czy będą przesunięte. Poprawienie zwilżania germanu przez stopy osiąga się przez wprowadzenie w kasetę nacisku mechanicznego na kulki. Temperatura maksymalna tego procesu

jest rzędu 500°C . Przebieg temperatury w funkcji czasu nie jest istotny dla jakości złącz. Należy jedynie nie dopuścić do powstania naprężeń w germanie wskutek zbyt szybkiego chłodzenia (powstają one przy szybkości spadku temperatury powyżej 100° na minutę).

Etap drugi — to wytwarzanie właściwych złącz. Na płytce grafitowej posiadającej szereg otworów umieszcza się płytki germanu z przyklejonymi w pierwszym etapie kulkami. Dolna kulka stopu zostaje umieszczona w otworze płytki grafitowej i w czasie procesu stapiania może swobodnie rozlewać się po powierzchni germanu. Całość zostaje umieszczona w kloszu stanowiska próżniowego (rys. 5) i wygrzana do temperatury w granicach $750\text{--}850^{\circ}\text{C}$, w zależności od grubości danej grupy płytek germanu.



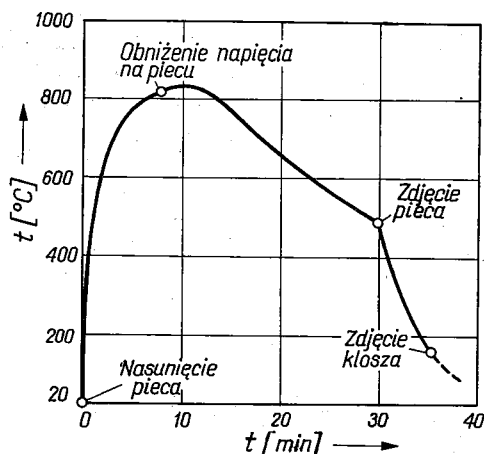
Rys. 5. Fragment stanowiska próżniowego

1 — rura kwarcowa, 2 — talerzyk na kasetki grafitowe, od spodu dotyka termopara, 3 — nasuwany piec oporowy, 4 — próżniomierz, 5 — wskaźnik termopary

Przebieg temperatury w czasie pokazano przykładowo na rys. 6.

Do określenia przybliżonej głębokości powstania złącza stosuje się następującą metodę. Zakłada się mianowicie, że głębokość penetracji stopu o tak dużej zawartości ołowiu można obliczyć z wystarczającą dokładnością z wykresu fazowego Pb—Ge (rys. 3). Jak widać z wykresu, rozpuszczalność germanu w ołowiu w temperaturach do 750°C nie jest duża, wzrasta jednak szybko powyżej 850°C . Zakres stosowanych temperatur przy procesie próżniowym wypada jeszcze w granicach małej zależności głębokości wnikania od temperatury. Wytwarzanie złącz można wykonywać wygrzewając płytki w kasetach od razu do żądanej temperatury.

Stwierdzono jednak, że pozwalając na swobodne rozlewanie się stopu po powierzchni germanu uzyskuje się bardziej powtarzalne głębokości tworzenia się złącz, niż przy ograniczaniu zwilżanej powierzchni, jak to ma miejsce w kasce. Wynika to oczywiście tylko z faktu istnienia kulek



Rys. 6. Przebieg temperatury w czasie próżniowego procesu stopowego

o różnych średnicach wygrzewanych w jednakowych warunkach do jednakowej temperatury. Stwierdzono ponadto, iż w omawianych warunkach laboratoryjnych wygrzewanie w próżni pozwalało na otrzymywanie mniejszych prądów wstecznych i większych napięć maksymalnych złącz niż w przypadku wygrzewania w wodorze.

3.2. Montaż tranzystorów

Wstępną operacją do montażu jest przylutowanie doprowadzenia bazy do płytki germanu. Doprowadzenie bazy musi tworzyć z Ge kontakt ekskluzyjny typu pp^+ o małej oporności. Wymagania te spełniają w stopniu zadowalającym luty cynowe z domieszką obniżającą temperaturę topnienia i z domieszką akceptorową (60% Sn — 30% Pb — 10% In). Niska temperatura topnienia (rzędu 140°C) i miękki lut są konieczne z tego względu, aby nie zepsuć uformowanych już złącz wskutek przetopienia lub przez wprowadzenie dodatkowych naprężeń w germanie.

Lutowanie doprowadzeń bazy wykonujemy na stanowisku wodorowym w temperaturze ok. 350—400°C, chłodzimy niezbyt szybko. Doprowadzenia bazy stanowią odpowiednio wycięte kształtki folii niklowej grubości 0,1 mm. Aby zorientować się o jakości otrzymanych złącz po przylutowaniu bazy, a przed montażem, można wytrawić chemicznie płytki Ge. Ponieważ trawienie to pozostawia jednak nadal dużą szybkość rekombinacji powierzchniowej, a zmniejsza tylko upływność po powierzchni złącz, pozwala więc stwierdzić jedynie, czy tranzystor był właściwie

wygrzany (np. czy obszary emitera i kolektora nie są zwarte ze sobą). Prostszy i bardziej polecanym sposobem jest zmierzenie omomierzem oporności między kulkami emitera i kolektora wprost po wygrzaniu. Przy pewnej wprawie można stąd wnioskować o odległości między złączami.

Montaż polega na przylutowaniu doprowadzenia bazy, emitera i kolektora do drutów przepustu oprawki. Do tych operacji jest używany lut o składzie: Pb 42% — Sn 33% — Bi 25%, posiadający temperaturę topnienia 150°C. Konkretny sposób umocowania płytki i połączenia elektrod zależy od typu oprawki. Tranzystory TW 6P są montowane bez specjalnego uwzględnienia chłodzenia złącz. W tranzystorach TWO 5P realizujemy chłodzenie złącza kolektora lutując kolektor na odpowiednio wygiętym trzpieniu miedzianym umieszczonym bezpośrednio na masie oprawki.

3.3. Trawienie chemiczne i elektrolityczne złącz

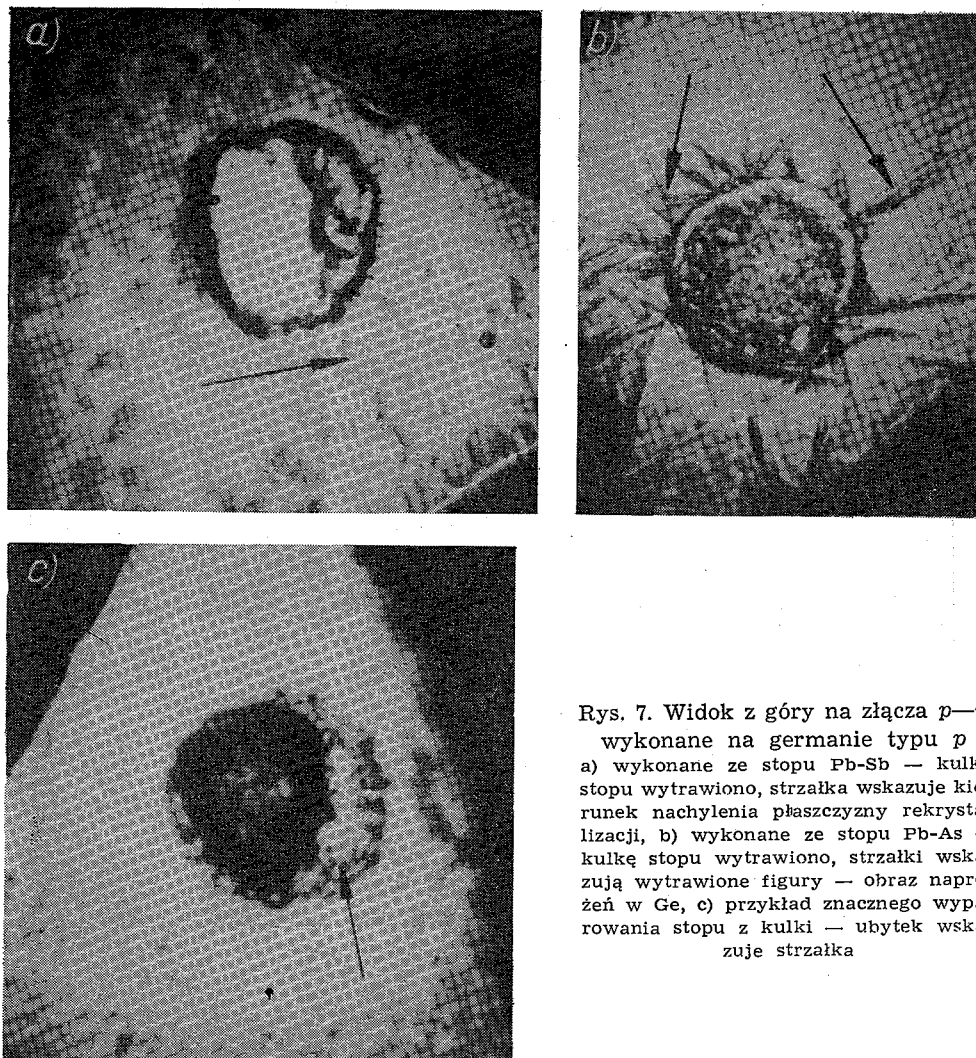
W wyniku dość dużej ilości cząstek parujących składników stopu kulek, w czasie procesu stapiania mamy bezpośrednie bombardowanie nimi germanu na pewnej powierzchni. Cząstki te tworzą cienką płytkę wtopioną w Ge warstwę dookoła złącz. Warstwa ta powoduje bardzo dużą upływność powierzchniową. Trawienie chemiczne mieszką CP₄ czy CPA powoduje rozpuszczenie powierzchniowej warstwy germanu i wytworzenie trudnych do usunięcia związków ołowiu, antymonu czy arsenu. Jakkolwiek upływność powierzchniowa złącz po tym trawieniu maleje, to pozostaje jeszcze bardzo duża szybkość rekombinacji nośników mniejszościowych na powierzchni. Jediną metodą dającą pozytywne wyniki jest trawienie elektrolityczne. Usuwa ono wszystkie zanieczyszczenia dookoła złącz powstałe w czasie stapiania, montażu itp., dając bardzo małe prądy wsteczne złącz i małą prędkość rekombinacji powierzchniowej. Według danych z literatury uzyskiwane minimalne prędkości po trawieniu elektrolitycznym są rzędu $(2 \div 4) \cdot 10^2$ cm/sek. Jako elektrolit stosujemy 5% roztwór wodny KOH. Prąd przepuszczamy od elektrody tranzystora przez elektrolit do elektrody grafitowej. Praktyczną receptę na trawienie każdego rodzaju tranzystora trzeba znaleźć eksperymentalnie. Do dobrego wytrawienia złącza potrzeba kilku jednosekundowych impulsów prądu o wartości 100 ÷ 500 mA na tle podkładu prądu stałego o wartości 5 ÷ 20 mA, przy zamaskowanych lakierem silikonowym częściach metalowych zanurzonych do elektrolitu. Najintensywniej trawią się płytki germanu dookoła kulek emitera i kolektora. Zdjęcia przekrojów tranzystorów wykazują wytrawienia pierścieniowe dookoła wtopionych elektrod, niekiedy bardzo głębokie. Jeżeli teraz po wytrawieniu, wypłukaniu w wodzie destylowanej i wysuszeniu tranzystor posiada dobre parametry elektryczne, to zamyka się oprawkę hermetycznie. Polega to na oblutowaniu lutem niskotopliwym (o temperaturze topnienia ok. 110°C) krawędzi zet-

knęcia pokrywki z zewnętrznym kołnierzem przepustu oprawki. Wewnątrz, pod pokrywką znajduje się przyklejona mała grudka silika-żelu do pochłaniania pary wodnej.

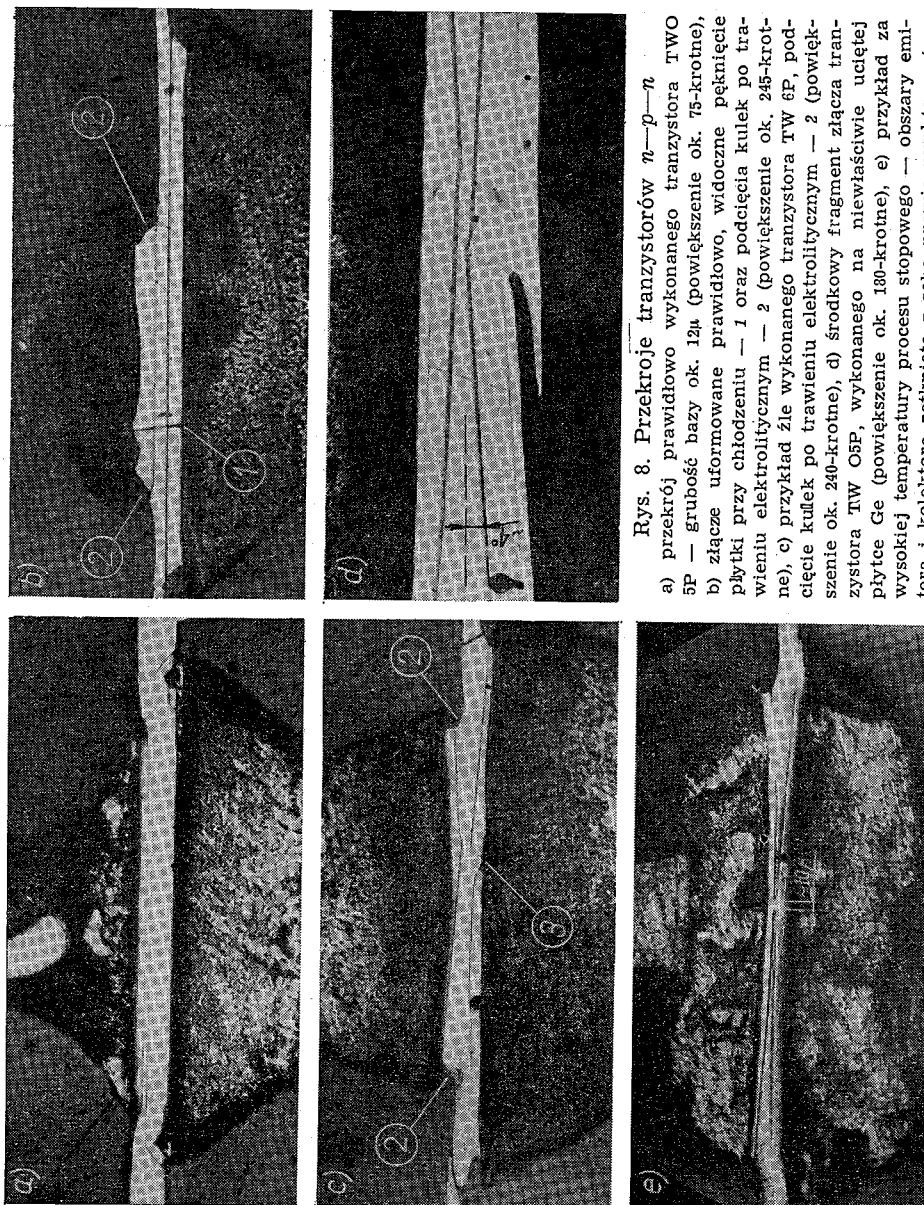
4. ANALIZA WYNIKÓW TECHNOLOGICZNYCH

4.1. Wpływ niedokładności przygotowania materiałów

Podstawowym warunkiem przy produkcji tranzystorów wielkiej częstotliwości jest płasko-równoległość ich złącz. Jednym z powodów uniemożliwiających wytworzenie takich złącz jest nieprawidłowe cięcie płytek germanu. Jeżeli płytki nie są ucięte dokładnie w płaszczyźnie [111]



Rys. 7. Widok z góry na złącza p—n wykonane na germanie typu p
a) wykonane ze stopu Pb-Sb — kulkę stopu wytrawiono, strzałka wskazuje kierunek nachylenia płaszczyzny rekrytalizacji, b) wykonane ze stopu Pb-As — kulkę stopu wytrawiono, strzałki wskazują wytrawione figury — obraz naprężeń w Ge, c) przykład znacznego wyparowania stopu z kulki — ubytek wskazuje strzałka



Rys. 8. Przekroje tranzystorów $n-p-n$

a) przekrój prawidłowo wykonanego tranzystora TWO 5P — grubość bazy ok. 12μ (powiększenie ok. 75-krotne), b) złącze uformowane prawidłowo, widoczne pęknięcie płytki przy chłodzeniu — 1 oraz podcięcie kulek po trawieniu elektrolitycznym — 2 (powiększenie ok. 245-krotne), c) przykład złe wykonanego tranzystora TW 6P, podcięcie kulek po trawieniu elektrolitycznym — 2 (powiększenie ok. 240-krotne), d) środkowy fragment złącza tranzystora TW O5P, wykonanego na niewłaściwie uciętej płytce Ge (powiększenie ok. 130-krotne), e) przykład za wysokiej temperatury procesu stopowego — obszary emitera i kolektora zetknięte z sobą prawie punktowo (powiększenie ok. 130-krotne)

monokryształu, to wtapianie ma kształt „tarasowaty” albo rekrystalizowana warstwa germanu leży pod pewnym kątem do płaszczyzny płytki. Na fotografiach widzimy przykład złącza dobrego, rys. 8a, i złącz nie mających dna płaskiego, rysunki 7a, 8c, 8d. Niedokładności cięcia rzędu paru stopni uniemożliwiają uzyskanie wysokich współczynników α , wysokich częstotliwości granicznych i maksymalnych napięć kolektora. Dopuszczalny błąd w kącie cięcia płytki nie powinien przekraczać $0,5^\circ$.

Dobre przygotowanie powierzchni kulek jest dużo trudniejsze niż przy produkcji tranzystorów *p-n-p* i wymaga szeregu starannych zabiegów. Stosowanie się do podanego w rozdz. 2.2 schematu postępowania przynosi jednak pomyślne wyniki. Dowodem tego jest fakt, że prawie nie obserwuje się odpadu złącz z powodu złego zwilżania powierzchni Ge. Jeżeli jednak w obszarze wtopienia kulki znajdują się zanieczyszczenia, to nie następuje równomierne wtopienie, linia złącza nie jest ciągła i ma wyraźne załamania (rys. 8c—3). Prądy wsteczne takiego złącza są dużo większe od przeciętnie uzyskiwanych.

4.2. Różnice obserwowane przy stosowaniu stopów Sb—Pb i As—Pb

Używając do prób stopów antymonowo-ołowiowych i arsenowo-ołowiowych stwierdzono wyraźną różnicę w jakości rekrystalizowanej warstwy typu *n*. Obrazy rekrystalizowanych obszarów po usunięciu stopu pokazują rysunki 7a i 7b. Są to złącza otrzymane ze stopu Sb 10% — Pb 90% i ze stopu As 3% — Pb 97%. Na fotografiach widać więc, że bardziej równomiernie rekrystalizują stopy z antymonem dając bardziej płaskie dna, podczas gdy dla domieszki arsenu rekrystalizowana powierzchnia jest bardziej nierówna. Stąd wniosek, że dla tranzystorów wielkiej częstotliwości bardziej będzie się nadawał stop z antymonem.

Do usunięcia stopu przykrywającego wyrekrystalizowane obszary typu *n* używano mieszanki o następującym składzie: kwas octowy lodowaty, perhydrol, woda destylowana w stosunku 2 : 2 : 5. Jeżeli po wytrawieniu pozostaje biały osad, to rozpuści się w stężonym HCl w temperaturze 60—70°C. Wymieniona mieszanka rozpuszcza w zadowalający sposób tylko ołów, a nie trawi innych, nawet pokrewnych metali, np. cyny. Ta właściwość ogranicza zastosowanie mieszanki do czystego ołowiu albo ołowiu z bardzo małą zawartością domieszek.

Wykonanie tranzystorów na stopach z arsenem w stanowisku próżniowym napotkało na poważne przeszkody. W temperaturze procesu stapiania ciśnienie par arsenu jest tak duże, że może on nawet całkowicie wyparować przy dłuższym trwającym procesie. Wydaje się nawet, iż obecność w stopie As zwiększa w pewnym stopniu łatwość parowania cząstek ołowiu. W ten sposób możemy tłumaczyć zjawisko uwidocznione na rys. 7c. Ubytek materiału kulki jest tu tak duży, że niecała zwilżona początkowo

powierzchnia germanu jest pokryta stopem. W tych warunkach wytworzenie złącza na żądanej głębokości jest prawie niemożliwe.

4.3. Wpływ temperatury na właściwości germanu i złącz $p-n$

Podczas procesu stopowego german jest podgrzewany do temperatury, w której może zachodzić dyfuzja akceptorów miedziowych. Akceptory te mają tę właściwość, że niewielka ich ilość może znacznie obniżyć czas życia nośników. Pojawienie się w germanie typu p akceptorów wywołuje spadek oporności właściwej Ge . Jeżeli w czasie wygrzewania stworzymy kontakt germanu z ciekłym metalem (mogą to być: Pb , Sn , Zn , In , Au), to wskutek geterującego działania fazy ciekłej na akceptory miedziowe wpływ ich może być nieznaczny. Efekt geterowania wynika z bardzo małej stałej segregacji miedzi w trójskładnikowym stopie przy dużej stałej dyfuzji miedzi w Ge (ok. $0,4 \frac{mm}{min}$ przy $700^{\circ}C$). Doświadczenia przeprowadzone w pracowni ZE—IPPT wydają się potwierdzać powyższe właściwości. Stwierdzono z całą pewnością spadek oporności germanu po wygrzaniu go w próżni do temperatury rzędu $700^{\circ}C$. Podczas wygrzewania bez kontaktów geterujących zmniejszenie oporności właściwej Ge wyniosło ok. 50% zarówno dla płytek o początkowej oporności $\rho = 9,5$ — $10 \Omega cm$, jak i dla $\rho = 2$ — $2,5 \Omega cm$. Przy wygrzewaniu takich samych płytek z kulką ołowiu o średnicy ok. 1 mm spadek ten wynosił ok. 25%.

Podczas wykonywania złącz stopowych wskutek małych wymiarów płytki Ge mamy na pewno geterowanie dzięki stapianym elektrodom, tak że można nie liczyć się z wyraźnym ujemnym wpływem śladów miedzi.

Bardziej istotnym czynnikiem od akceptorów miedziowych, wpływającym na pogorszenie się parametrów germanu, jest sposób chłodzenia płytki ze złączami po procesie wtapiania w próżni. Przy zbyt szybkim tempie chłodzenia mamy „zamrożenie” naprężeń, w wyniku czego powstają trwałe defekty struktury. Ujawniają się one dopiero po wytrawieniu chemicznym (rys. 7b). Zbyt duże naprężenia mogą powodować mikropeknięcia (rys. 8b). Uszkodzenia struktury powodują tak znaczne zwiększenie szybkości rekombinacji nośników, że uniemożliwiają pracę tranzystora. Praktycznie daje to bardzo duże prądy zerowe i bardzo mały współczynnik α . Ustalono doświadczalnie, że maksymalna dopuszczalna szybkość chłodzenia nie przekracza $100^{\circ}C$ na minutę, a wystarczająco wolna — rzędu $20^{\circ}C$ na minutę.

Innym przykładem termicznego uszkodzenia tranzystorów jest przetopienie obszaru bazy. Dążenie do maksymalnej głębokości wtopienia elektrod w celu uzyskania jak najmniejszej grubości bazy prowadzi czasami do przetopienia płytki na wskroś i do zwarcia ze sobą obszarów emitera i kolektora. Na rys. 8e mamy przykład zwarcia spowodowanego zbyt wy-

soką temperaturą przy niedokładnym cięciu płytki germanu w płaszczyźnie [111]. Linie łącz obu obszarów biegną do siebie zbieżnie, a zetknięcie w środku jest prawie punktowe na powierzchni o średnicy ok. 10 μ .

4.4. Wpływ trawienia na parametry tranzystora

Oceną jakości trawienia jest wartość parametrów tranzystorów uzyskana po tej operacji.

Jak już wspomniano w rozdz. 3.3, rezultaty trawienia elektrolitycznego są bez porównania lepsze od wyników otrzymanych po mieszkankach chemicznych, takich jak CP_4 , CPA, HNO_3 + perhydrol. Najlepiej przekona o tym porównanie odpowiednich rubryk w zamieszczonej tablicy:

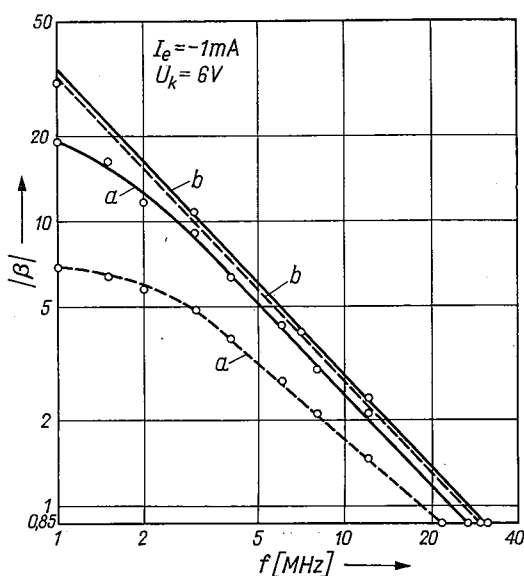
Tablica 2

	$-h_{21}$ [mA/mA]	h_{11} [Ω]	$h_{22} \times 10^8$ [S]	$h_{12} \times 10^5$ [V/V]	I_{e0} [μ A]	I_{k0} [μ A]
po trawieniu chemicznym	0,87	60	200	100	10	15
po trawieniu elektrolit.	0,995	30	30	60	1	1

Interesująco przedstawia się również porównanie przebiegu modułu współczynnika wzmocnienia prądowego w układzie OE w funkcji częstotliwości [$|\beta| = f(f)$] przed i po trawieniu elektrolitycznym, rys. 9. Obserwujemy tu dwa zjawiska: 1) uzyskanie prawidłowego nachylenia krzywych — 6 dB/oktawę, 2) wzrost częstotliwości granicznej.

Prawidłowe nachylenie krzywej osiągnięto dzięki zmniejszeniu rekombinacji powierzchniowej nośników (wzrost β dla małych częstotliwości), a wzrost f_a uzyskano dzięki usunięciu, wytrawieniu brzegowych, najbardziej niepożądanych obszarów złącz.

Ponieważ przykłady tranzystorów z rys. 9 oraz tablicy 2 są bardzo typowe, można więc stwierdzić, że trawienie chemiczne nie pozwoli z regu-



Rys. 9. Wykres współczynnika wzmocnienia prądowego układu OE w funkcji częstotliwości

a) przed trawieniem elektrolitycznym,
b) po trawieniu

ły na uzyskanie wystarczająco dobrych parametrów tranzystorów n - p - n . Trawienie elektrolityczne ma też i ujemne strony. W wyniku przepływu prądu powstaje dookoła złącza pierścieniowy ubytek materiału płytki. Na fotografiach (rys. 8b i 8c—2) zaznaczono miejsca największego podcięcia płytki Ge. Oprócz osłabienia mechanicznego płytki zwiększa to oporność bazy tranzystora. Sprawdzone doświadczalnie dla tranzystorów TW 6P wykonanych z Ge o $\rho = 3 \div 5 \Omega\text{cm}$, że oporność bazy wzrasta średnio o ok. 10—20%, podczas gdy dla $\rho = 0,5 \Omega\text{cm}$ wzrost oporności może wynieść 60—80%.

Ostatnią czynnością w procesie technologicznym, która może mieć ujemny wpływ na parametry tranzystorów, jest zamknięcie oprawek. Stwierdzono, że prawidłowo wykonane zamknięcie nie ma praktycznie żadnego wpływu na tranzystory. Zauważono jedynie w ciągu kilkudziesięciu godzin niewielki spadek α (o ok. 0,1%) i spadek oporności wyjściowej (o ok. 10%). Spowodowane to jest jednak stabilizowaniem się powierzchni germanu w czasie.

Na zakończenie należy zaznaczyć, że opisane procesy technologiczne zostały opracowane i sprawdzone w warunkach laboratoryjnych.

5. PARAMETRY TRANZYSTORÓW N-P-N

Zebrane poniżej w tablicy parametry tranzystorów TW 6P i TW O5P pozwalają zorientować się w wartościach podstawowych parametrów dobrych egzemplarzy tranzystorów n - p - n w typowych i maksymalnych punktach pracy.

Tablica 3

Parametr		TW 6P		TW O5P		Warunki pomiaru
$-h_{21b}$	mA/mA	0,970 — 0,999		0,900 — 0,998		dla $I_e = -1 \text{ mA}$, $U_k = 6 \text{ V}$ $f_{pom} = 1 \text{ kHz}$
h_{11b}	Ω	30 — 35		25 — 30		
$h_{22b} \times 10^8$	S	10 — 50		10 — 50		
$h_{12b} \times 10^5$	V/V	20 — 150		20 — 100		
I_{EBo}, I_{KBo}	μA	0,1 — 2,0*	1 — 5**	1 — 5*	5 — 10**	mierzono przy 15 V
I_{KEo}	μA	1 — 6*	5 — 10**	10 — 80*	30 — 100**	dla $R_b = 0$
U_{kmax}	V	20 — 25*	30 — 45**	20 — 25*	30 — 45**	dla ukł. WB i $OER_b = 0$
poj. kolekt. $C_b'k$	pF	$\sim 20^*$	$\sim 10^{**}$	$\sim 80^*$	$\sim 40^{**}$	dla $I_e = -1 \text{ mA}, U_k = 6 \text{ V}$
oporność bazy r_b	Ω	20 — 150*	200 — 350**	50 — 200*	70 — 250**	dla $I_e = -10 \text{ mA}, U_k = 6 \text{ V}$
częstotł. gran. f_a	MHz	8 — 30		3 — 15		dla $I_e = -5 \text{ mA}, U_k = 6 \text{ V}$
maks. temp. otocz. T_a	$^{\circ}\text{C}$	65		65		
P_{kmax} (w 20°C)	mW	50		200 — 250		

* — $\rho_{Ge} = 0,5 \div 1,5 \Omega\text{cm}$, ** — $\rho_{Ge} = 3 \div 5 \Omega\text{cm}$

Analizując poszczególne rubryki tablicy można stwierdzić, że w zasadzie wartości liczbowe podanych parametrów nie różnią się od odpowiednich wartości liczbowych parametrów tranzystorów stopowych p-n-p.

5.1. Częstotliwość graniczna (f_a)

Z powodu dużych trudności występujących na ogół przy pomiarach wysokich częstotliwości granicznych, a jednocześnie dla dużej wagi tego parametru zostanie pokrótce opisana stosowana w tej pracy metoda pomiaru f_a tranzystorów. Wykorzystano tu metodę graficznego przedstawienia przebiegu modułu współczynnika wzmocnienia prądowego w układzie OE w funkcji częstotliwości [$|\beta| = f(f)$]. Do określenia f_a wykorzystano efekt liniowego opadania $|\beta| = f(f)$ wg prawa 6 dB/oktawę dla częstotliwości powyżej częstotl. gran. tranzystora w ukł. OE:

$$f > f_\beta \approx (1 - \alpha_0) f_a. \quad (1)$$

Pritchard udowodnił, że dla tranzystorów stopowych f_a równa jest częstotliwości, przy której $|\beta| = 0,85$. W podobny sposób można określić częstotliwość f_T . Jest to częstotliwość, przy której $|\beta| = 1$. Zależność $|\beta| = f(f)$ dla tranzystorów TW 6P jest pokazana na rys. 6. Dla tranzystorów, które spełniają zależność:

$$\frac{|\beta_1| \text{ dla } f_1}{|\beta_0| \text{ dla } f_2} = 2, \quad (2)$$

gdzie częstotliwość pomiarowa $f_1 = 1/2 f_2$ jest dowolna i większa od f_β , dokładność odczytu f_a wynosi $+10\%$ w całym zakresie pomiaru. Dla tranzystorów, które mają stosunek $\frac{|\beta_1|}{|\beta_2|}$ nieco większy czy mniejszy od 2, dotychczasowa interpretacja wyników może nie być słuszna. Dla tych przypadków należałoby przeprowadzić nieco szerszą analizę teoretyczną w celu usunięcia wątpliwości co do poprawności wyników. Inną metodą określenia częstotliwości granicznej tranzystorów, jednak mniej użyteczną, jest obliczenie f_a ze zmierzonej pojemności dyfuzyjnej złącza emiter-baza ($C_{b'e}$). Wzór na pojemność $C_{b'e}$ otrzymujemy z wyrażenia na część urojoną admitancji wejściowej macierzy $[Y]$ dla tranzystora w układzie OE:

$$C_{b'e} = \Lambda I_e \frac{W_b^2}{2D_n}, \quad (3)$$

gdzie:

$$\Lambda = \frac{q}{kT}, \quad \text{dla } 300^\circ\text{K: } \Lambda = -38,6[\text{V}^{-1}],$$

W_b — grubość bazy,

D_n — stała dyfuzji elektronów w bazie.

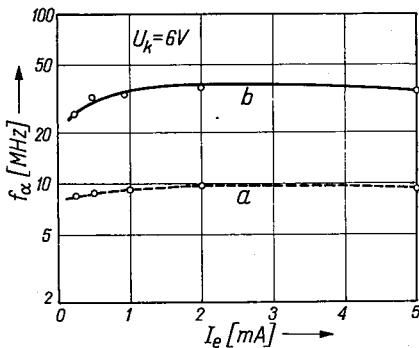
Z obliczenia czasu przejścia nośników przez bazę otrzymano przybliżony wzór:

$$f_a = \frac{D_n}{\pi W_b^2} \quad (4)$$

Po prostym podstawieniu z (3) i (4) otrzymujemy wyrażenie na f_a w zależności od $C_{b'e}$ i prądu polaryzacji emitera (I_e):

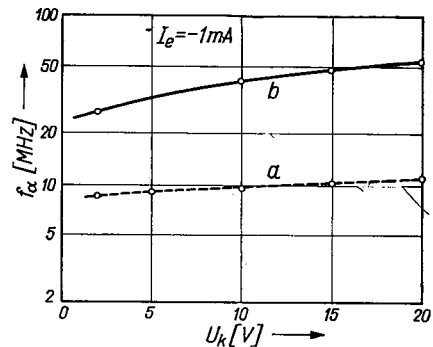
$$f_a = \frac{A \cdot I_e}{2 \cdot \pi \cdot C_{b'e}} \quad (5)$$

Jakkolwiek powyższe wzory są przybliżone, to jednak mogą być w pewnym stopniu przydatne do obliczenia f_a ze zmierzonej pojemności $C_{b'e}$ oraz do dyskusji f_a w funkcji prądu emitera. Przykład zależności częstotliwości granicznej od prądu emitera jest przedstawiony na rys. 10. Z podanych wzorów wynika, że gdyby współczynnik α i stała dyfuzji D_n nie ulegały zmianom w funkcji prądu emitera, to częstotliwość graniczna w funkcji I_e byłaby const. Stwierdzono jednak, że D_n i oczywiście α nie są wielkościami stałymi w funkcji I_e . Wynik działania zmian tych czynników jest taki, że krzywa $f_a = f(I_e)$ osiąga pewne niewielkie maksimum, a opada nieco zarówno dla małych, jak i dla większych prądów emitera.



Rys. 10. Zależność częstotliwości granicznej tranzystorów od prądu emitera

a, b — dwa różne egzemplarze tranzystorów TW 6P



Rys. 11. Zależność częstotliwości granicznej tranzystorów od napięcia kolektora

Wzrost częstotliwości granicznej w funkcji napięcia kolektora (rys. 11) jest spowodowany głównie zjawiskiem zmniejszania się efektywnej grubości bazy, ponieważ w miarę wzrostu napięcia wstecznego rośnie szerokość złącza p—n (tzw. efekt Early'ego).

W celu porównania wyników pomiaru częstotliwości granicznej tranzystorów dokonano zestawienia uzyskanych rezultatów: a) na drodze bezpośredniego pomiaru a, b) z krzywej $|\beta| = f(f)$ oraz c) z pomiaru $C_{b'e}$:

Tablica 4

a	f_a — bezpośredn.	MHz	6,7	10,2	16,8	—
b	f_a — z $ \beta = f(f)$	MHz	6,9	12,5	20,0	28,8
c	f — z obl. $C_{b'e}$	MHz	7,8	14,5	27,4	50,5

Z tablicy tej wynika, że im większą mierzy się częstotliwość graniczną, tym rozbieżności liczbowe stają się coraz większe. Spowodowane jest to głównie błędami metod pomiarowych. Przy bezpośrednich pomiarach $\alpha = f(f)$ (rubryka a tabl. 4). Wskutek niedostosowania przystawki pomiarowej do pracy na wyższych częstotliwościach, powyżej 12—15 MHz, zdejmujemy właściwie charakterystykę częstotliwościową układu, a nie tranzystora, dlatego też otrzymuje się wyniki dużo niższe od rzeczywistych. Wręcz przeciwne mamy zjawisko przy obliczeniach f_a z pojemności $C_{b'e}$ (rubr. c). Obliczone wartości f_a są zbyt duże w stosunku do odpowiednich liczb w rubr. a i b, a różnica między tymi liczbami rośnie szybko wraz ze wzrostem mierzonej f_a . Jest to spowodowane korzystaniem z uproszczonych wzorów, nie uwzględniających szeregu czynników, jak np. stałej czasu kolektora czy oporności bazy. Czynniki te nie odgrywają roli przy częstotliwościach rzędu kilku MHz, natomiast na pewno nie są do pominięcia przy częstotliwościach 20—50 MHz.

Wypada w tym miejscu podkreślić, że metoda określania częstotl. gran. (czy f_T) z krzywych $|\beta| = f(f)$ jest niezmiernie prosta i łatwa, o zadowalającej w praktyce dokładności i jest przydatna do określania częstotliwości granicznych (albo f_T) nawet rzędu 100—200 MHz.

5.2. Pojemność wyjściowa (C_{22})

Na wielkość pojemności C_{22} składa się pojemność bariery kolektora C_{Tk} i pojemność dyfuzyjna złącza kolektora C_{dk} . Pojemność C_{dk} można wyrazić następującym wzorem:

$$C_{dk} = I_e \frac{W_b}{D_n} \cdot \frac{\partial W_b}{\partial U_{b'k}} \quad (6)$$

przy zachowaniu warunku: $\left(\frac{W_b}{L_b}\right)^2 \ll 1$ i $\alpha_0 \approx 1$.

Pojemność barierowa złącza kolektora jest określona przybliżonym wzorem:

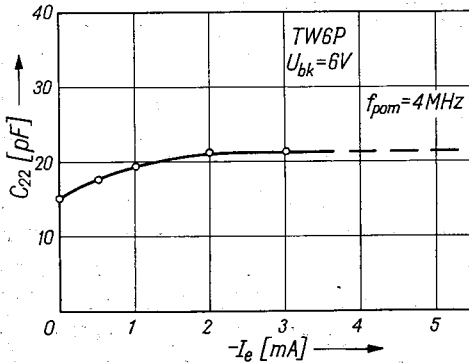
$$C_{Tk} = 3,4 \cdot 10^{-4} \sqrt{\frac{N_a}{U_{b'k}}} \text{ [pF/cm}^2\text{]}, \quad (7)$$

gdzie:

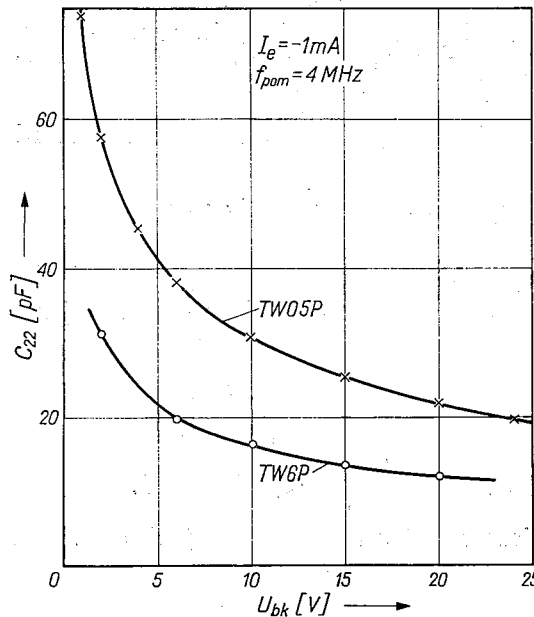
N_a — koncentracja akceptorów w bazie,

$U_{b'k}$ — napięcie baza — kolektor.

Zależność powyższa jest słuszna tylko dla tzw. złącz „skokowych” (do których zaliczamy złącza stopowe) wykonanych na germanie i dla napięć zaporowych większych od 1 V. W zakresie normalnej pracy tranzystora stopowego pojemność dyfuzyjna C_{dk} jest mniejsza o przeszło rząd wielkości od pojemności barierowej C_{Tk} . Można więc przyjąć, że pojemność wyjściowa tranzystora C_{22} równa jest pojemności C_{Tk} , zatem wykres pojemności C_{22} w funkcji I_e (rys. 12) i U_{bk} (rys. 13) należy omawiać w oparciu o wzór (7). Z zależności tej wynika, że pojemność kolektora jest propor-



Rys. 12. Przebieg pojemności wyjściowej C_{22} w funkcji prądu emitera



Rys. 13. Przebieg pojemności wyjściowej C_{22} w funkcji napięcia kolektora

cjonalna do pierwiastka z koncentracji domieszek bazy. Oznacza to z kolei, że dla coraz mniejszych oporności właściwych germanu pojemność C_{22} tranzystorów będzie się powiększać.

Teoretycznie pojemność wyjściowa tranzystora nie zależy od prądu emitera. W rzeczywistości istnieje jednak pewna różnica w wielkości pojemności dla różnych polaryzacji emitera (rys. 12). Wy tłumaczyć to można jedynie zmianą oporności bazy (tzw. modulacja oporności) w wyniku przepływu prądu przez obszar bazy. Ze względu na zastosowaną metodę pomiaru pojemności (rozstrojenie obwodu rezonansowego) nie można zachować dokładności pomiaru dla prądów $I_e > 3-4$ mA — stąd na rys. 12 linia przerywana dla prawdopodobnego przebiegu C_{22} .

Zależność pojemności kolektora od napięcia U_{bk} (rys. 13) przebiega odwrotnie proporcjonalnie do pierwiastka z tego napięcia, zgodnie z teoretyczną zależnością (7).

6. PORÓWNANIE TRANZYSTORÓW N—P—N Z TRANZYSTORAMI P—N—P

6.1. Wpływ ruchliwości elektronów i dziur na współczynnik dobroci tranzystora

Ogólnie panuje przekonanie, że o przydatności tranzystora na wyższych częstotliwościach decyduje ruchliwość nośników mniejszościowych bazy, a nie uwzględnia się wpływu ruchliwości nośników większościowych. Obecnie zostanie przeprowadzona krótka analiza porównawcza obu typów tranzystorów na bazie współczynnika dobroci z uwzględnieniem obu rodzajów nośników prądu.

Pojęcie współczynnika dobroci powstało przy obliczaniu maksymalnego wzmocnienia tranzystora na częstotliwościach pracy rzędu częstotliwości granicznej. Wzmocnienie tranzystora definiujemy oczywiście jako stosunek mocy na wyjściu do mocy na wejściu tranzystora przy założeniu dopasowania generatora doporności wejściowej i obciążenia do oporności wyjściowej tranzystora:

$$G = \frac{P_{wyj}}{P_{wej}}. \quad (8)$$

Stosunek ten obliczono w funkcji parametrów tranzystora przy następujących założeniach:

1. $\alpha_0 \approx 1$ dla małych częstotliwości, a $\gamma_0 \approx 1$ dla całego przedziału częstotliwości.

2. O wielkości oporności wejściowej tranzystora decyduje oporność bazy.

3. Pojemność dyfuzyjna kolektor — baza jest do pominięcia. Wtedy, w rezultacie stosunkowo prostych obliczeń, otrzymuje się przybliżoną zależność:

$$G = \frac{\alpha \cdot f_a}{8\pi \cdot f^2 \cdot r_b'' \cdot C_{b'k}} \approx \frac{f_a}{25 \cdot f^2 \cdot r_b'' \cdot C_{b'k}}. \quad (9)$$

Chcąc obliczyć teraz maksymalną możliwą częstotliwość oscylacji tranzystora zakładamy $G = 1$ i ze wzoru (9) otrzymujemy:

$$f_{max} = \sqrt{\frac{f_a}{25 \cdot r_b'' \cdot C_{b'k}}} = \sqrt{M}. \quad (10)$$

Wyrażenie oznaczone przez M nazywamy współczynnikiem dobroci tranzystora.

Obliczając f_{max} dla różnych układów i stosując różne metody obliczeń okazuje się, że współczynnik dobroci jest niezmiennikiem, wielkość jego zależy tylko od konstrukcji tranzystora. Teraz możemy zapisać wzór (9) w skróconej postaci:

$$G = \frac{M}{f^2}. \quad (11)$$

Należy jednak pamiętać, że jest on słuszny w przedziale częstotliwości:

$$(0,05-0,1) < \frac{f}{f_a} < 2.$$

Otrzymana we wzorze (10) postać współczynnika M jest dogodna do dyskusji o charakterze układowym. Do dyskusji nad projektowaniem tranzystorów natomiast nadaje się lepiej inna, nieco skomplikowana postać, do której można dojść po pewnych przekształceniach:

$$M = Gf^2 \approx \frac{K \cdot \mu_p \cdot \mu_n}{4 \cdot \pi^2 \cdot \left(\frac{2 \cdot \epsilon_0 \cdot \epsilon_r}{q \cdot N_a \cdot U_{b'k}} \right)^{1/2} \cdot W_b^2 \cdot A_k \cdot A}, \quad (12)$$

gdzie:

K — stała zależna od kształtu i wymiarów geometrycznych,

μ_p, μ_n — ruchliwości dziur i elektronów,

ϵ_0 — przenikalność względna,

ϵ_r — przenikalność germanu,

W_b — grubość bazy,

A_k — powierzchnia kolektora,

N_a — koncentracja akceptorów w bazie,

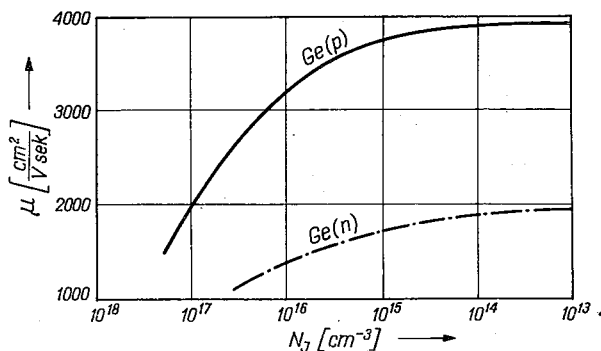
$$A = \frac{q}{k \cdot T}$$

$U_{b'k}$ — napięcie baza — kolektor.

Dążenie do wysokiego współczynnika M prowadzi często do sprzecznych warunków, których jednoczesne spełnienie jest niemożliwe. Np. podniesienie N_a obniża napięcie przebicia germanu (U_B), a więc i możliwość stosowania dużego $U_{b'k}$. Zależność U_B od N_a jest typu: $U_B = N_a^{-0.7}$. Ponadto wzrost N_a obniża ruchliwość μ_p i μ_n .

Pozostałe czynniki wchodzące w skład wzoru (12), jak współczynnik K , A_k czy W_b zależą od kształtu złącz i procesów technologicznych.

Podstawowym kryterium porównawczym tranzystorów p—n—p i n—p—n przy pracy na wielkich częstotliwościach jest stosunek maksymalnego wzmocnienia mocy G_{npn}/G_{pnp} . Obliczamy go właśnie przy pomocy omówionego wzoru (12) zakładając identyczne wymiary geometryczne



Rys. 14. Ruchliwość nośników mniejszościowych w Ge w funkcji koncentracji domieszek

tranzystorów, jednakowe napięcia i częstotliwości pracy. Przy dodatkowym założeniu jednakowej koncentracji domieszek bazy otrzymujemy:

$$\frac{G_{npn}}{G_{pnp}} = \frac{\mu_p \cdot \mu_n}{\mu_n \cdot \mu_p} \approx 1. \quad (13)$$

Przybliżenie wynika z mało dokładnych równań opisujących zależność ruchliwości nośników od koncentracji.

Dla równych przewodności (albo oporności) właściwych bazy otrzymujemy:

$$\frac{G_{npn}}{G_{pnp}} = \frac{(\mu_p)^{1/2} \cdot \mu_n}{(\mu_n)^{1/2} \cdot \mu_p} \approx \left(-\frac{\mu_n}{\mu_p} \right)^{1/2} = b^{1/2}. \quad (14)$$

Dla germanu $b^{1/2} = 2,05^{1/2} = 1,43$ (dla małych przewodności). Dla większych przewodności ok. 1[S/cm]: $b^{1/2} = 1,34^{1/2} = 1,16$.

Po znalezieniu pochodnej cząstkowej równania (12) względem przewodności właściwej bazy i ruchliwości nośników obliczono optymalną wartość σ_b dla dostatecznie dużej jeszcze ruchliwości. Dla germanu typu p optymalna przewodność właściwa bazy wynosi 8 [S/cm], a dla typu n 22 [S/cm]. Jeżeli teraz obliczymy stosunek $\frac{G_{npn}}{G_{pnp}}$ dla tych wartości σ_b , to wyniesie on tylko 0,95.

Jak założono na wstępie, wywody te są słuszne w przypadku równych napięć kolektora. Dla stałego stosunku $\frac{U_{bk}}{U_B}$ zarówno dla struktury n—p—n, jak i p—n—p we wszystkich powyższych przypadkach otrzymamy stosunek $\frac{G_{npn}}{G_{pnp}} \approx 1$.

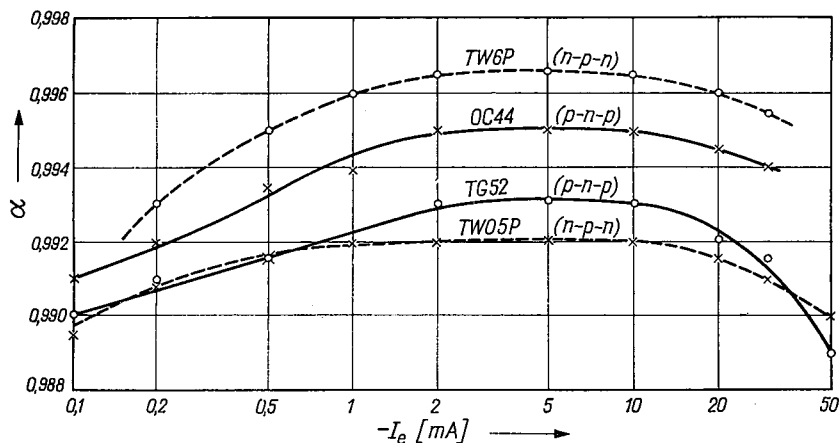
6.2. Porównanie podstawowych parametrów

W oparciu o wzory teoretyczne i wyniki pomiarów rzeczywistych tranzystorów zostaną wykazane różnice między najważniejszymi parametrami tranzystorów $n-p-n$ i $p-n-p$.

Przed wszystkim trzeba założyć, że porównywane mogą być jedynie tranzystory o podobnych wymiarach geometrycznych i wykonane tą samą techniką (w tym przypadku stopową).

Najbardziej interesujące jest porównanie przebiegu współczynnika wzmocnienia prądowego, $\alpha = f(I_e)$. Według ogólnie znanych zależności na wartość α wpływają trzy czynniki: tzw. współczynnik wprowadzania γ (inaczej sprawność emitera), rekombinacja objętościowa nośników w bazie i rekombinacja powierzchniowa nośników. Rekombinacja powierzchniowa wpływa głównie na wielkość α przy małych prądach I_e , przy dużych I_e decyduje sprawność emitera.

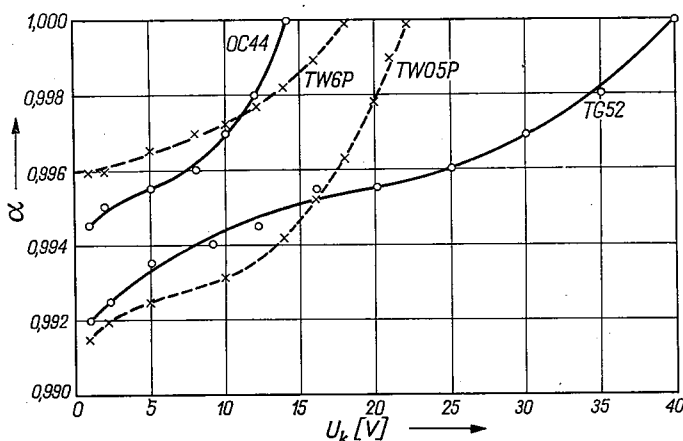
Zostało udowodnione, że wyżej wspomniane czynniki wskutek różnic ruchliwości nośników mniejszościowych bazy powodują mniejszą zależność współczynnika α od prądu emitera dla tranzystorów $n-p-n$ niż dla tranzystorów $p-n-p$. Porównanie przebiegów $\alpha = f(I_e)$ dla tranzystorów większej mocy na rys. 15 wykazuje zgodność z przewidywaniami teoretycznymi, natomiast krzywe odnoszące się do tranzystorów małej mocy odbiegają od założeń teoretycznych. Podstawowym tego powodem jest dowolność wyboru egzemplarzy bez zachowania warunków wymienionych na wstępie. Nie bez znaczenia jest też fakt zestawienia na jednym



Rys. 15. Porównanie zależności $\alpha = f(I_e)$ tranzystorów klasy 50 mW (OC44 — TW 6P) i klasy 200 mW (TG51 — TWO 5P)

wykreście tranzystorów różnych producentów (różne materiały wyjściowe, różne technologie). Takie same zestawienie dla tranzystorów $p-n-p$ wykonanych w ZE-IPPT (tranzystory TW 12) wypadło zgodnie z przewidywaniami.

Omówienie przebiegu $\alpha = f(U_{bk})$ zostanie poprzedzone krótkim przypomnieniem zjawisk lawinowych w złączu. W wyniku lawinowego powielania nośników w złączu kolektora oraz zmniejszania się grubości ba-



Rys. 16. Porównanie przebiegu $\alpha = f(U_{bk})$ tranzystorów klasy 50 mW (OC44 — TW 6P) i klasy 200 mW (TG51 — TWO 5P)

zy przez wzrost szerokości bariery (efekt Earlyego) mamy znaczny wzrost α przy zwiększaniu napięcia kolektora. Dla tranzystorów stopowych słuszne jest następujące empiryczne wyrażenie uwzględniające zjawiska lawinowe w złączu:

$$\alpha(U_{bk}) = \alpha_0 \alpha^* = \frac{\alpha_0}{1 - \left(\frac{U_{bk}}{U_B}\right)^n}, \quad (15)$$

gdzie:

U_{bk} — napięcie kolektora,

U_B — napięcie przebicia wewnętrznego w germanie,

n — wykładnik zależny od typu germanu,

dla typu p: $n = 4,5 \div 6,5$, dla typu n: $n = 3$.

Zmiany szerokości bariery można obliczyć przy pomocy następującego wzoru:

$$\frac{\partial W_b}{\partial U_{bk}} \approx \frac{1}{2} \left(\frac{2 \varepsilon_0 \varepsilon_r \mu_w}{\sigma_b \cdot U_{bk}} \right)^{1/2}, \quad (16)$$

gdzie μ_w ruchliwość nośników większościowych bazy: dla elektronów

$$\mu_n \approx 3800 \frac{\text{cm}^2}{\text{V sek}}, \text{ dla dziur } \mu_p \approx 1700 \frac{\text{cm}^2}{\text{V sek}}.$$

Z obydwu powyższych wzorów wynika, że tranzystory n—p—n powinny mieć mniejszą zależność α od napięcia kolektora niż tranzystory

$p-n-p$, przy założeniu jednakowego stosunku $\frac{U_{bk}}{U_B}$ i jednakowej σ_b . Gdyby założenia te, jak również warunek identycznych wymiarów geometrycznych zostały spełnione, można byłoby przeprowadzić analizę krzywych na rys. 16 w oparciu o zależności (15) i (16). Ponieważ do pomiarów wybór egzemplarzy był przypadkowy, dopuszcza to powstanie poważnych nieprawidłowości w przebiegu krzywych, które zresztą mają miejsce na rys. 15.

Inne interesujące nas zagadnienie to częstotliwość graniczna tranzystorów $p-n-p$ i $n-p-n$. Korzystając ze wzorów podanych w rozdziałach poprzednich można stwierdzić, że w wyniku dwukrotnie większej stałej dyfuzji elektronów pojemność dyfuzyjna $n-p-n$ ów jest ok. dwa razy mniejsza. To powoduje z kolei, iż częstotliwość graniczna jest ok. dwa razy większa. Wyniki porównania f_a obu typów tranzystorów dla tej samej grubości bazy podaje tablica 5.

Tablica 5

W_b	μ	22,0	18,8	17,5	15,5	13,5	11,0	9,6
dla $n-p-n$ f_a	MHz	7,2	9,4	10,1	15,0	17,0	24,0	33,0
dla $p-n-p$ f_a	MHz	3,5	5,0	5,5	7,1	9,5	14,3	18,0
$\frac{f_a \text{ dla } n-p-n}{f_a \text{ dla } p-n-p}$		2,18	1,88	1,84	2,11	1,79	1,68	1,83

Z powyższej tablicy wynika, że zmierzone wartości częstotliwości granicznej dla tej samej grubości bazy (obliczonej z pomiarów $C_{b'e}$), w wystarczającym stopniu są zgodne z rozważaniami teoretycznymi.

Pojemności barierowe i szerokości złącz w tranzystorach zależą od pierwiastka kwadratowego z ruchliwości nośników większościowych bazy. Stąd pojemność złącz tranzystorów $n-p-n$ powinna być około $b^{1/2}$ razy większa niż dla analogicznych tranzystorów $p-n-p$, a szerokość bariery $b^{1/2}$ razy mniejsza w stosunku do tranzystorów $p-n-p$. Dalszą konsekwencją tego faktu jest mniejsza zależność grubości bazy od napięcia kolektora (tzw. efekt Earlyego), przez co osiągnięcie zjawiska „punch-through” w tranzystorach $n-p-n$ jest prawie niemożliwe. Mniejszy efekt Earlyego potwierdzają również doświadczalnie zdjęte krzywe $r_b''(U_{bk})$.

6.3. Symetria dopełniająca tranzystorów

Korzystnymi cechami tranzystorów, których nie posiadają inne przyrządy wzmacniające, np. lampy próżniowe, są ich właściwości symetryczne. Wyróżniono symetrię własną każdego egzemplarza (zamiana roli emi-

tera z kolektorem) i tzw. symetrię dopełniającą, wynikającą ze stosowania elementów $n-p-n$ i $p-n-p$.

Określenie „symetria dopełniająca” pochodzi z faktu uzupełniania się charakterystyk jako wynik przeciwnych znaków napięć i prądów występujących w tych dwóch typach przyrządów. Można więc mówić, że jeden typ tranzystora jest współpartnerem drugiego, ale nie na zasadzie superpozycji. Istnienie symetrii dopełniającej wykorzystano w układach tranzystorowych między innymi np. do skonstruowania przeciwsobnych wzmacniaczy bez transformatorów albo do podwójnego stopnia wzmacniającego o bezpośrednim sprzężeniu tranzystorów. Zainteresowani sprawami wykorzystania par $n-p-n$ i $p-n-p$ w układach znajdą wyczerpujące informacje w bogatej już na ten temat literaturze. Tranzystory o wysokich symetriach zarówno własnych, jak i dopełniających są łatwe do uzyskania na drodze techniki stopowej.

6.4. Uwagi końcowe

Na zakończenie wydaje się celowe podkreślić w skrócie cechy wspólne oraz różnice zarówno w technologii, jak i w wartościach parametrów tranzystorów stopowych $n-p-n$ i $p-n-p$. W procesie technologicznym wspólna jest zasada wytwarzania złącz oraz częściowo aparatura, przyrządy, przygotowanie materiałów. Różnice występują w typach germanu, w rodzajach stopów domieszkowych, maksymalnych temperaturach stapiania, częściowo w sposobie trawienia. Otrzymywanie dobrych tranzystorów $n-p-n$ jest uwarunkowane tak samo przestrzeganiem tych samych ogólnych zasad, jakie obowiązują przy wytwarzaniu każdego rodzaju elementów półprzewodnikowych.

Odnosnie parametrów elektrycznych podstawowa różnica polega na przeciwnych kierunkach polaryzacji tranzystorów $n-p-n$ i $p-n-p$, wielkości częstotliwości granicznej, pojemności barierowej złącz. Dla tej samej oporności właściwej Ge tranzystory $n-p-n$ mają teoretycznie mniejsze napięcie przebicia.

Wszystkie opisane dotąd właściwości tranzystorów $n-p-n$ pozwalają na wyciągnięcie wniosku o możliwości szerokiego zastosowania ich jako nowych elementów układów. Tranzystory $n-p-n$ można stosować we wszelkiego typu wzmacniaczach, pojedynczo lub jako współpartnerów tranzystorów $p-n-p$. Duża oporność wyjściowa i częstotliwość graniczna pozwala na zastosowanie $n-p-n$ ów w układach przełącznikowych szybkiego działania oraz w układach impulsowych. Duża stałość współczynnika α w funkcji prądu pozwala na pracę przy dużym wysterowaniu bez nadmiernych zniekształceń sygnałów.

Praca niniejsza nie omawia pomiarów szeregu parametrów zarówno na małych, jak i wielkich częstotliwościach. Wydaje się jednak, iż wyniki tych pomiarów mogą zainteresować tylko ścisły krąg użytkowników,

a nie wniosą już wiele nowego do ogólnej charakterystyki tego typu tranzystorów.

*Polska Akademia Nauk
Instytut Maszyn Matematycznych*

WYKAZ LITERATURY

1. Early J. M.: Design Theory of Junction Transistors. — Bell Syst. Techn. Journ., Nov., 1953.
2. Ebers I. I., Miller S. L.: Design of Alloyed-Junction Germanium Transistors for High-Speed Switching. — B.S.T.J., July, 1955.
3. Giacoletto L. J.: Study of P—N—P Alloy Junction Transistor from D—C Through Medium Frequencies. — RCA Review, Dec., 1954.
4. Giacoletto L. J.: Comparative High-Frequency Operation of Junction Transistors Made of Different Semiconductor Materials. — RCA Review, March, 1955.
5. Jenny D. A.: A Germanium N—P—N Alloy Junction Transistor. — PIRE, Dec., 1953.
6. Jonscher A. K.: Podstawy działania przyrządów półprzewodnikowych. WNT, Warszawa 1962.
7. Kompało W.: Tranzystory stopowe $n-p-n$. — Przegl. Elektroniki, nr 9, 1962.
8. Lederhandler S., Giacoletto L. J.: Measurements of Minority-Carrier Lifetime and Surface Effects in Junction Devices. — PIRE, April, 1955.
9. Lohman R. D.: Complementary Symmetry Circuits. — Electronics, Sept., 1953.
10. Mueller C. W., Ditrick N. H.: Uniform Planar Alloy Junctions for Germanium Transistors. — RCA Review, March, 1956.
11. Pritchard R. L.: Frequency Variations of Junction-Transistor Parameters. — PIRE, May, 1954.
12. Pritchard R. L.: High Frequency Power Gain of Junction Transistors. — PIRE, Sept., 1955.
13. Pritchard R. L.: Transistor Tests Predict High Frequency Performance. — Electronic Industries and Tele-Techn., March, 1957.
14. Rosiński W.: Zasady technologii tranzystorów. WNT, Warszawa 1960.
15. Rosiński W.: Zasady działania tranzystora i jego parametry, WNT, Warszawa 1961.
16. Sziklai G. C.: Symmetrical Properties of Transistors and their Applications. — PIRE, June, 1953.
17. Slade B. N.: N—P—N Power Transistors. Transistors I. Princeton 1956.
18. Transistor Technology (praca zbiorowa), tom III. D. Van Nostrand Co., Princeton 1958.

W. KOMPAŁO

TECHNOLOGY AND BASIC PARAMETERS OF N—P—N GERMANIUM ALLOY TRANSISTORS

Summary

In the present paper the technology and some basic parameters of $n-p-n$ germanium alloy transistors are discussed. These transistors have the following

cut-off frequency and collector power dissipation respectively: type TW 6P: 10—30 Mc/s and 50 mW, type TW O5P: 5—10 Mc/s and about 250 mW.

The first part is concerned with technological problems:

- a) selection and preparing of materials,
- b) alloying in vacuum of $p-n$ junctions,
- c) soldering of contacts,
- d) electrolytical etching of $p-n$ junctions.

Next sections bring data of more important parameters and some characteristics of these $n-p-n$ transistors. One may find also the description and principles of the measurement method applied and brief discussion of resulting curves.

The comparative evaluation of $n-p-n$ and $p-n-p$ alloy transistors is accomplished on the basis of H.F. maximum available power gain for both types. Basic transistor parameters which are the factors of the expression for maximum available gain are briefly discussed as functions of emitter current and collector voltage.

W. KOMPALO

LA TECHNOLOGIE ET LES PARAMÈTRES FONDAMENTAUX DES TRANSISTORS N—P—N AU GERMANIUM PAR ALLIAGE

Résumé

L'article présente la technologie et les paramètres fondamentaux des transistors $n-p-n$ au germanium par alliage, type TW 6P, fréquence de coupure 10—30 MHz, puissance dissipée du collecteur 50 mW et type TW O5P: $f_a = 5-10$ MHz P_c ordre de 250 mW.

La première partie présente:

- 1) le choix et la préparation des matériaux,
- 2) la production des jonctions dans le vide,
- 3) le montage des transistors,
- 4) le traitement électrolytique des jonctions.

Ensuite on a donné les paramètres fondamentaux ainsi que les caractéristiques plus importantes des transistors $n-p-n$. Les principes des techniques de mesure appliquées sont donnés et les résultats discutés. L'article donne aussi l'estimation des transistors $n-p-n$ et $p-n-p$ sur la base de la comparaison de l'amplification de puissance maximum aux fréquences élevées. Le comportement des paramètres composants la formule de l'amplification de puissance maximum en fonction du courant de l'émetteur et de la tension du collecteur sont discutés.

W. KOMPALO

TECHNOLOGIE UND GRUNDKENNGRÖßEN DER N—P—N GERMANIUM — LEGIERUNGSTRANSISTOREN

Zusammenfassung

In dem vorliegenden Artikel wurden Technologie und Grundparametern der Germanium-Legierungstransistoren besprochen, die die folgenden Grenzfrequenzen und Verlustleistungen haben: Type TW 6P: 10—30 MHz und 50 mW, Type TW O5P: 5—10 MHz und etwa 250 mW.

Es wurde in dem ersten Teil beschrieben:

- a) die Materialauswahlweise und die Werkstoffvorbereitung,
- b) die Herstellung den $p-n$ Übergängen im Vakuum,
- c) die Montage des Transistors,
- d) das elektrolytische Abätzen von $p-n$ Übergängen.

Weiter sind die Grundparameter und wichtigste Kennlinien von $n-p-n$ Transistoren gegeben worden. Es gibt hier auch die Beschreibung des Meßprinzips von gegebenen Transistorkenngrößen und eine Diskussion über die erhaltenen Ergebnissen.

In dem letzten Abschnitt hat man die Vergleichung von $n-p-n$ und $p-n-p$ Transistoren auf Grund des Vergleiches der maximalen H.F.-Leistungsverstärkung von beiden Transistortypen ausgeführt. Es wurde auch der Verlauf der Transistorkenngrößen, die in der maximalen Leistungsverstärkung enthalten sind in Abhängigkeit vom Emitterstrom und Kollektorspannung kurz besprochen.

В. КОМПАЛО

ТЕХНОЛОГИЯ И ОСНОВНЫЕ ПАРАМЕТРЫ ГЕРМАНИЕВЫХ СПЛАВЛЕННЫХ ТРАНЗИСТОРОВ ТИПА $n-p-n$

Резюме

В работе описаны технология и свойства германиевых высоко-частотных транзисторов типа $n-p-n$ полученных методом сплавления. Были разработаны два варианты этих транзисторов — транзисторов с мощностью рассеяния коллектором 50 мвт и предельной частотой 10—30 Мгц и транзистор с мощностью 250 мвт и частотой 5—10 Мгц.

В первой технологической части работы описаны вопросы подбора исходных материалов, получение $p-n$ переходов в вакуумном процессе, а также монтаж и электрохимическая обработка транзистора.

Во второй части даны сведения об применяемых методах измерений электрических и технологических параметров этих кристаллических триодов.

В заключение проведено сравнение сплавленных германиевых транзисторов $n-p-n$ типа с общеприменяемыми сплавленными транзисторами $p-n-p$. Основным критерием этого сравнения является максимальное усиление мощности на высоких частотах. Анализируется также влияние режима работы транзистора на параметры, от которых зависит максимальное усиление мощности.

IRENEUSZ WÓJCIK

Germanowy tranzystor mocy o wielkiej częstotliwości granicznej wykonany techniką stopowo-dyfuzyjną

Rękopis dostarczono 2 3.1964

W artykule rozpatrzono zagadnienia technologii i konstrukcji germanowych tranzystorów mocy wielkiej częstotliwości wykonanych techniką stopowo-dyfuzyjną.

Zbadano zależność parametrów tranzystora od doboru materiałów wyjściowych oraz od przebiegu procesów technologicznych. Omówiono i przeanalizowano kolejno wszystkie procesy potrzebne do wykonania tranzystora ze szczególnym uwzględnieniem wpływu gęstości dyslokacji w germanie na prawidłowość formowania się złącz podczas wtapiania-dyfuzji. Opisano również stosowaną aparaturę.

Pomiary międzyoperacyjne jak też pomiary gotowych elementów były przeprowadzone przede wszystkim pod kątem zbadania prawidłowości procesów technologicznych oraz ich wpływu na właściwości tranzystorów.

W oparciu o przeprowadzone w ten sposób badania podano optymalne warunki technologiczne na uzyskanie różnych (ze względu na zastosowania praktyczne) wariantów stopowo-dyfuzyjnego tranzystora mocy.

WSTĘP

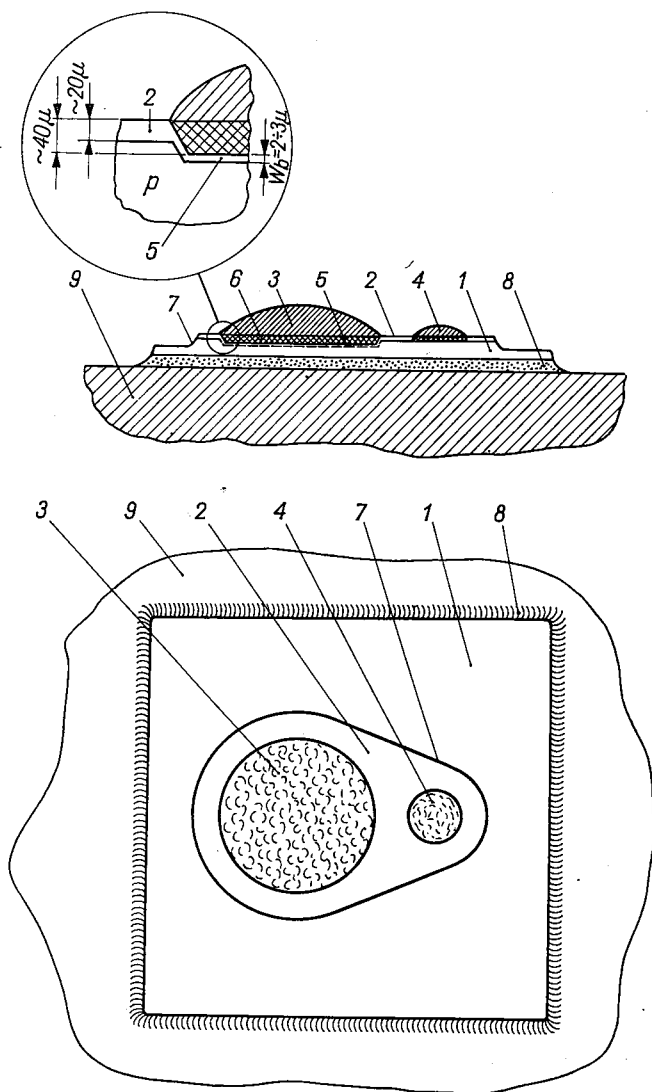
Tematem niniejszej pracy są podstawowe zagadnienia projektowania i technologii tranzystora wielkiej częstotliwości o mocy admisyjnej kolektora do 4 W. Opracowanie to powstało w trakcie kontynuacji prac nad ogólniejszym problemem dotyczącym możliwości zastosowań techniki stopowo-dyfuzyjnej do wykonania różnego rodzaju tranzystorów. Poprzednie prace w tej dziedzinie wykonane w Zakładzie Elektroniki IPPT PAN są podane w spisie literatury (poz. [8], [11], [15]).

I. KONSTRUKCJA I TECHNOLOGIA

1. ZASADA KONSTRUKCJI GERMANOWEGO TRANZYSTORA STOPOWO-DYFUZYJNEGO

Zasada konstrukcji tranzystora stopowo-dyfuzyjnego została podana przez Beale'a [1]. Jest tu wykorzystane zjawisko dużej różnicy w szybkościach dyfuzji pomiędzy niektórymi akceptorami a donorami oraz fakt, że domieszki powoli dyfundujące posiadają zwykle dużą stałą segregacji.

Większość donorów dyfunduje w germanie ok. 1000-krotnie szybciej niż akceptory. Jednocześnie stała segregacji akceptorów jest od 10 do 100 razy większa niż donorów. Wynika stąd, że gdy do germanu typu *p* będziemy wtapiali stop zawierający zarówno donory, jak i akceptory, to dobierając odpowiednio czas i temperaturę uzyskamy strukturę *p-n-p*. Baza typu *n* utworzy się dzięki dyfuzji donorów z fazy ciekłej, złącze zaś



Rys. 1. Konstrukcja stopowo-dyfuzyjnego tranzystora mocy. Powiększenie ok. 13,3-krotne

1 — german typu *p*, 2 — warstwa typu *n* uzyskana po wstępnej dyfuzji, 3 — kontakt emitera, 4 — kontakt bazy, 5 — baza typu *n* powstała w procesie stopowo-dyfuzyjnym, 6 — rekrystalizowany obszar emitera, 7 — brzeg wyspy kolektora, 8 — lut łączący płytkę Ge z podstawką, 9 — oprawka miedziana