

ności magnetycznej. Dokładne wykonanie takich indukcyjności w produkcji masowej jest niemożliwe, a dostrojenie w czasie montażu — trudne; cena ich jest znaczna, jak również ciężar duży. Najważniejszą wadą jest jednak nieliniowość dużych indukcyjności wykluczająca ich zastosowanie w wielu urządzeniach. Dlatego zwiększone zapotrzebowanie na filtry dolnoprzepustowe o niskiej częstotliwości ω_p (głównie ze strony automatyki) zmusiło do poszukiwania filtrów nowego typu. Znalaziono je wśród biernych i czynnych czwórników złożonych z pojemności C i rzeczywistych oporności R .

W celu dokonania syntezy pewnego typu układów elektrycznych, np. filtrów, należy zawsze rozwiązać trzy następujące problemy:

1) Określić, do jakiej klasy funkcji należą interesujące nas charakterystyki tego typu układów (w naszym przypadku charakterystykami tymi są: funkcja przenoszenia $h(s)$ i charakterystyka tłumieniowa $w(\omega^2) = |h(s)|$).

2) Znaleźć sposób wyboru pożądanej charakterystyki spośród wszystkich funkcji określonej klasy.

3) Znaleźć sposób wyznaczenia struktury i elementów układu realizującego wybraną charakterystykę.

Pierwszy z wymienionych trzech problemów jest w naszym przypadku dość zawiły, ale dla biernych czwórników RC został on całkowicie rozwiązany w pracach Fialkova i Gersta [1], [2]. Wyniki uzyskane przez tych autorów podamy w następnym rozdziale jako podstawę do dalszych rozważań.

Drugi problem syntezy jest w zasadzie problemem aproksymacyjnym. Chodzi w nim o przybliżenie idealnej charakterystyki (nie dającej się zazwyczaj fizycznie zrealizować) przy pomocy funkcji określonych w rozwiązaniu problemu pierwszego, przy czym spełnić należy pewne dodatkowe warunki. W naszym przypadku idealną charakterystyką tłumieniową jest funkcja

$$\begin{aligned} w(\omega^2) &= \text{const} > 0, & 0 \leq \omega^2 \leq \omega_p^2, \\ w(\omega^2) &= 0, & \omega_p^2 < \omega^2 \leq \infty. \end{aligned}$$

Charakterystyka $w(\omega^2)$ spełniająca warunki Fialkova i Gersta powinna dostatecznie mało różnić się od tej funkcji. Dodatkowym warunkiem jest, aby filtr składał się z możliwie małej liczby elementów. Po wprowadzeniu paru upraszczających założeń będziemy mogli wykorzystać znane rozwiązania pewnych zagadnień aproksymacyjnych do skonstruowania pożądanej charakterystyki $w(\omega^2)$. Niniejsza praca poświęcona jest właśnie adaptacji tych rozwiązań do celów syntezy filtrów RC oraz zbadaniu niektórych konsekwencji takiej adaptacji.

Kompletne rozwiązanie trzeciego problemu syntezy w przypadku czwórników RC nie jest dotychczas znane. Istnieje szereg metod obliczania takich czwórników realizujących zadaną funkcję przenoszenia, ale nie każda z nich może być stosowana do wszystkich funkcji $h(s)$ spełniających warunki Fialkova i Gersta, ponadto nie ma sposobu określenia dla danej

funkcji $h(s)$ takiej metody, która dawałaby wynik optymalny np. ze względu na liczbę elementów. W celu zilustrowania zagadnień związanych z realizacją zadanej funkcji przenoszenia zamieszczamy w trzecim podrozdziale dwa przykłady liczbowe syntezy czwórników RC najdogodniejszą w praktyce metodą Darlingtona-Dashera. Metoda ta nadaje się wyłącznie do projektowania biernych układów RC, ale takimi będziemy się głównie zajmować w niniejszym artykule, gdyż dla układów czynnych trudno jest podać warunki realizowalności analogiczne do warunków Fialkova i Gersta. Tym niemniej wyniki uzyskane w rozdziałach 1 i 2 można na ogół wykorzystać również przy projektowaniu czwórników czynnych.

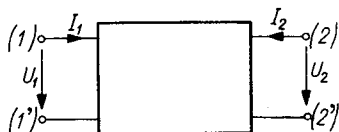
1.2. Charakterystyki częstotliwościowe czwórników RC

W niniejszym rozdziale podamy w postaci szeregu twierdzeń konieczne i dostateczne warunki realizowalności biernego czwórnika RC, a następnie konieczne i dostateczne warunki realizowalności funkcji $h(s)$ oraz $w(\omega^2) = |h(s)|$ jako charakterystyk częstotliwościowych takiego czwórnika. Tezy tych twierdzeń weźmiemy następnie pod uwagę w przykładach liczbowych w rozdziale 1.3 oraz przy formułowaniu zagadnień aproksymacyjnych w rozdziale 1.4.

Wiadomo, że własności biernego czwórnika liniowego określają się jednoznacznie jedną z tzw. macierzy charakterystycznych, czyli zespolem 3 funkcji zmiennej $s = i\omega$ wiążących z sobą zespolone przebiegi na wejściu i wyjściu czwórnika. W przypadku tzw. macierzy przewodnościowej odpowiednie związki mają postać:

$$\begin{aligned} I_1 &= Y_{11} U_1 + Y_{12} U_2, \\ I_2 &= Y_{12} U_1 + Y_{22} U_2, \end{aligned} \quad (1.1)$$

gdzie symbole I_1, I_2, U_1, U_2 są przebiegami określonymi na rys. 2, zaś $Y_{\mu\nu}$ — funkcjami zmiennej zespolonej $s = i\omega$. Macierz o elementach $Y_{\mu\nu}$ oznacza



Rys. 2. Czwórnik elektryczny i wielkości określające jego macierz charakterystyczną

się symbolem $\|Y\|$. Jej własności określa następujące twierdzenie ([3], rozdz. V).

Twierdzenie 1. Jeżeli $\|Y\|$ jest macierzą przewodnościową czwórnika RC, to jej elementy $Y_{11}(s), Y_{12}(s), Y_{22}(s)$ spełniają następujące warunki:

1) $Y_{11}(s), Y_{22}(s)$ są rzeczywistymi funkcjami wymiernymi (tzn. ilorazami dwóch wielomianów o rzeczywistych współczynnikach) o pojedyn-

czych pierwiastkach i biegunach leżących na przemian na ujemnej, domkniętej półosi zmiennej s . Stąd wynika możliwość zapisania tych funkcji w postaci

$$Y_{vv}(s) = \frac{h_0^{(vv)}}{s} + \sum_{i=1}^m \frac{h_i^{(vv)}}{s + k_i^{(vv)}} + h_\infty^{(vv)} s, \quad (1.2)$$

gdzie

$$k_i^{(vv)} > 0 \quad (i = 1, 2, \dots, m) \quad h_i^{(vv)} \geq 0 \quad (i = 0, 1, 2, \dots, m, \infty);$$

2) $Y_{12}(s)$ jest wymierną funkcją rzeczywistą o pojedynczych biegunach leżących na domkniętej ujemnej półosi zmiennej s . Stąd wynika możliwość zapisania tej funkcji w postaci

$$Y_{12}(s) = \frac{h_0^{(12)}}{s} + \sum_{i=1}^n \frac{h_i^{(12)}}{s + k_i^{(12)}} + h_\infty^{(12)} s, \quad (1.3)$$

gdzie

$$k_i^{(12)} > 0 \quad (i = 1, 2, \dots, n) \quad \infty < h_i^{(12)} < \infty \quad (i = 0, 1, 2, \dots, n, \infty);$$

$$3) \quad h_i^{(11)} h_i^{(22)} - (h_i^{(12)})^2 \geq 0 \quad (i = 0, 1, 2, \dots, n, \infty).$$

Warunek 3 jest istotny. Wynika z niego w szczególności, że funkcja $Y_{12}(s)$ nie może mieć takiego bieguna $k_i^{(12)}$, który nie byłby jednocześnie biegunem funkcji $Y_{11}(s)$ oraz $Y_{22}(s)$. Twierdzenie 1 głosi, że jeżeli którykolwiek z warunków 1÷3 nie jest spełniony, to nie istnieje czwórnik RC o danej macierzy przewodnościowej $\|Y\|$, 1÷3 są zatem koniecznymi warunkami fizycznej realizowalności macierzy $\|Y\|$. Okazuje się, że warunki 1÷3 są również wystarczające, tzn. że gwarantują one realizację czwórnika RC o spełniającej je macierzy $\|Y\|$, ale jedynie wówczas, jeżeli dopuścimy możliwość zastosowania w tym czwórniku idealnych transformatorów.

Aby wyrazić funkcję przenoszenia $h(s)$ przy pomocy elementów macierzy przewodnościowej, podstawiamy do drugiego z równań (1.1) $I_2 = 0$ oraz $U_1 = E$. Zgodnie z rys. 1 i definicją 1 dostaniemy

$$h(s) = \frac{U_2}{E} = - \frac{Y_{12}(s)}{Y_{22}(s)}. \quad (1.4)$$

Powyższe równanie oraz znajomość własności funkcji $Y_{12}(s)$, $Y_{22}(s)$ pozwalają z łatwością określić warunki, jakie musi spełniać funkcja przenoszenia $h(s)$ czwórnika RC. Warunki te nie gwarantowałyby jednak realizowalności $h(s)$ przy pomocy czwórnika beztransformatorowego, a tylko takie czwórnik możemy brać w praktyce pod uwagę, jeżeli mamy do czynienia z niskimi częstotliwościami. Fialkow i Gerst [1], [2] udowodnili następujące twierdzenie o funkcji $h(s)$ beztransformatorowych czwórników RC.

Twierdzenie 2. Funkcja $h(s)$ daje się zrealizować jako funkcja przenoszenia czwórników RC wtedy i tylko wtedy, gdy:

1) $h(s)$ jest rzeczywistą funkcją wymierną:

$$h(s) = A \frac{a(s)}{b(s)} = A \frac{s^m + a_1 s^{m-1} + a_2 s^{m-2} + \dots + a_{m-1} s + a_m}{s^n + b_1 s^{n-1} + b_2 s^{n-2} + \dots + b_{n-1} s + b_n}; \quad (1.5)$$

2) pierwiastki wielomianu $b(s)$ są pojedyncze, rzeczywiste, ujemne;

$$3) \quad |A| \leq \min_{0 \leq s < \infty} \left| \frac{b(s)}{a(s)} \right|,$$

przy czym znak równości może mieć miejsce jedynie wówczas, gdy funkcja $|b(s)/a(s)|$ osiąga swe minimum na krańcach przedziału $[0, \infty]$.

Należy zwrócić uwagę, że z warunku 3 wynika $m \leq n$, gdyż w przeciwnym razie otrzymalibyśmy $\min |b(s)/a(s)| = 0$, czyli $A = 0$. Konieczności warunków 1÷3 można dowieść drogą analizy równań Kirchhoffa dla czwórników RC, ich dostateczność udowodnił Fialkow i Gerst podając metodę syntezy odpowiedniego czwórnika RC. Metoda ta ma znaczenie jedynie teoretyczne, gdyż jest uciążliwa pod względem obliczeniowym i prowadzi do układu o nadmiernej liczbie elementów. W praktyce stosuje się metody prostsze, przy których nie można jednak z góry zadać współczynnika A , ustala się on niejako automatycznie w procesie syntezy i nie osiąga na ogół liczby $\min |b(s)/a(s)|$. Dlatego spośród trzech warunków Fialkova i Gersta istotne są jedynie dwa pierwsze oraz zastrzeżenie, że $m \leq n$.

Konieczne i dostateczne warunki realizowalności charakterystyki tłumieniowej $w(\omega^2)$ nietrudno jest określić znając własności funkcji przenoszenia $h(s)$. W tym celu spostrzegamy, że dla $s = i\omega$

$$w^2(\omega^2) = |h(s)|_{s=i\omega}^2 = h(s) \cdot h(-s)|_{s=i\omega}, \quad (1.6)$$

gdyż na urojonej osi s funkcje $h(s)$ i $h(-s)$ są zespolone sprzężone. Na podstawie twierdzenia 2 i równania (1.6) można o funkcji $w(\omega^2)$ twierdzić, co następuje:

Twierdzenie 3. Aby funkcja $w(\omega^2)$ dała się zrealizować jako charakterystyka tłumieniowa beztransformatorowego czwórnika RC muszą być spełnione następujące warunki:

1)

$$\begin{aligned} w(\omega^2) = |h(s)|_{s=i\omega} &= \sqrt{A^2 \frac{\alpha(\omega^2)}{\beta(\omega^2)}} = \\ &= \sqrt{A^2 \frac{\omega^{2m} + a_1 \omega^{2(m-1)} + \dots + a_{m-1} \omega^2 + a_m}{\omega^{2n} + b_1 \omega^{2(n-1)} + \dots + b_{n-1} \omega^2 + b_n}}; \end{aligned} \quad (1.7)$$

gdzie liczby A, a_ν, b_ν są rzeczywiste.

2)

$$[w(\omega^2)]^2 \geq 0 \text{ dla } 0 \leq \omega^2 \leq \infty; \quad m \leq n;$$

3) Pierwiastki wielomianu $\beta(\omega^2)$ są pojedyncze, ujemne.

Jasną jest rzeczą, że warunki 1÷3 są również wystarczającymi warunkami realizacji beztransformatorowego czwórnika RC o charakterystyce tłumieniowej różniącej się co najwyżej stałym mnożnikiem od funkcji $w(\omega^2)$. Istotnie, warunki 1 i 2 pozwalają przy pomocy równania (1.6) wyznaczyć funkcję przenoszenia $h(s)$, której licznikiem i mianownikiem będą wielomiany rzeczywiste stopnia m i n , przy czym $m \leq n$, natomiast warunek 3 zapewnia możliwość uzyskania funkcji $h(s)$ o pojedynczych, ujemnych biegunach. Wyznaczoną w ten sposób funkcję $h(s)$ po ewentualnym pomnożeniu jej przez odpowiednią liczbę $c < 1$ potrafimy fizycznie zrealizować, o czym była mowa w związku z twierdzeniem 2.

Obecnie zilustrujemy na dwóch przykładach liczbowych zagadnienie syntezy czwórników RC o zadanych charakterystykach tłumieniowych. Będą to funkcje typu (1.7) o współczynniku $A^2 = 1$. Charakterystyki o takiej szczególnej postaci oznaczać będziemy symbolem $w_1(\omega)$.

1.3. Przykłady realizacji zadanej charakterystyki tłumieniowej

Przykład 1. Skonstruować czwórnik RC, którego charakterystyką tłumieniową byłaby, z dokładnością do stałego mnożnika, funkcja

$$w_1(\omega^2) = \sqrt{\frac{\omega^2 + 0,023276}{(\omega^2 + 0,088803)(\omega^2 + 0,38676)}}.$$

Przed wszystkim stwierdzamy, że wyżej zadana funkcja $w_1(\omega^2)$ spełnia warunki 1÷3 twierdzenia 3. Istotnie, jej kwadrat jest rzeczywistą funkcją wymierną większą od 0 w przedziale $0 \leq \omega^2 \leq \infty$, o ujemnych, pojedynczych pierwiastkach mianownika. Wobec tego przystępujemy do skonstruowania funkcji przenoszenia $h(s)$ projektowanego czwórnika. W tym celu zapisujemy $w^2(\omega^2)$ w postaci

$$\begin{aligned} w_1^2(\omega^2) &= h_1(s) h_1(-s) \Big|_{s=i\omega} = \\ &= \frac{(0,15257 + i\omega)(0,15257 - i\omega)}{(0,29800 + i\omega)(0,29800 - i\omega)(0,62190 + i\omega)(0,62190 - i\omega)}, \end{aligned}$$

stąd

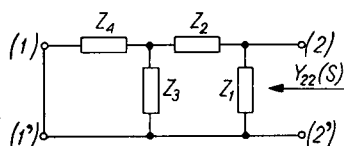
$$h_1(s) = \frac{(s + 0,15257)}{(s + 0,29800)(s + 0,62190)}. \quad (1.8)$$

Otrzymana funkcja przenoszenia spełnia warunki 1, 2 twierdzenia 2, a po

opatrzeniu jej odpowiednim mnożnikiem spełni również warunek 3, a więc będzie ją można wówczas fizycznie zrealizować. Dokonamy tego metodą Guillemina ([7] str. 207), która jest wersją procesu Darlingtona [4] stosowanego w syntezie układów czwórnikowych LC. W tym celu wykorzystujemy równanie (1.4) zapisując znaną funkcję $h(s)$ w postaci

$$h(s) = -\frac{Y_{12}(s)}{Y_{22}(s)} = \frac{s + 0,15257}{(s + 0,29800)(s + 0,62190)} \cdot \frac{s + 0,45}{s + 0,45} \quad (1.9)$$

Skonstruowane w ten sposób elementy Y_{12} i Y_{22} macierzy $\|Y\|$ spełniają warunki 1, 2 twierdzenia 1, a po dobraniu odpowiedniej funkcji $Y_{11}(s)$



Rys. 3. Struktura czwórnika RC o jednym rzeczywistym pierwiastku funkcji przenoszenia

spełnią również warunek 3. Wobec tego można będzie skonstruować tzw. kanoniczny czwórnik RC z idealnymi transformatorami ([3], rozdz. V), którego wtórna zwarciowa przewodność wyniesie $Y_{22}(s)$, a wzajemna $Y_{12}(s)$. Czwórnik ten będzie miał oczywiście funkcję przenoszenia określoną przez (1.9). Jeżeli do realizacji Y_{22} , Y_{12} nie zastosujemy kanonicznego układu Cauera, lecz wykorzystamy metodę Darlingtona-Guillemina, to uzyskamy czwórnik beztransformatorowy o funkcji przenoszenia (1.9), przemnożonej jednak przez pewien czynnik mniejszy od jedności. Metoda Darlingtona-Guillemina polega na skonstruowaniu czwórnika RC, którego wtórna przewodność zwarciowa równa się funkcji $Y_{22}(s)$ w równaniu (1.9), a pierwiastki przewodności wzajemnej pokrywają się z pierwiastkami funkcji $Y_{12}(s)$. Ponieważ na mocy warunku 3 w twierdzeniu 1 przewodność wzajemna nie może mieć innych biegunów niż bieguny funkcji $Y_{22}(s)$, a więc z dokładnością do stałego mnożnika musi być ona funkcją $Y_{12}(s)$ określoną w (1.9).

Przewodność $Y_{22}(s)$ realizujemy w postaci drabinkowego połączenia dwójników przedstawionego na rys. 3. Dwójnik Z_1 dobieramy tak, aby oporność $Z_2(s) = \infty$ dla $s = -0,15257$. Dzięki temu przewodność wzajemna $Y_{12}(s)$ będzie miała pierwiastek w tym właśnie punkcie płaszczyzny s . Dwójnik Z_4 będzie opornością rzeczywistą i nie wniesie żadnych innych pierwiastków funkcji Y_{12} . W celu wyznaczenia Z_1 obliczamy

$$Y_{22}(-0,15257) = \frac{(-0,15257 + 0,29800)(-0,15257 + 0,62190)}{(-0,15257 + 0,45)} = 0,2945,$$

Z_1 jest zatem opornością rzeczywistą, równą

$$Z_1 = R_1 = \frac{1}{0,22945} = 4,3576.$$

W dalszym ciągu obliczamy przewodność

$$\begin{aligned} Y_{22}^{(1)}(s) &= Y_{22}(s) - \frac{1}{Z_1} = \frac{s^2 + 0,69041s + 0,082057}{s + 0,45} = \\ &= \frac{(s + 0,15257)(s + 0,53785)}{s + 0,45} \end{aligned}$$

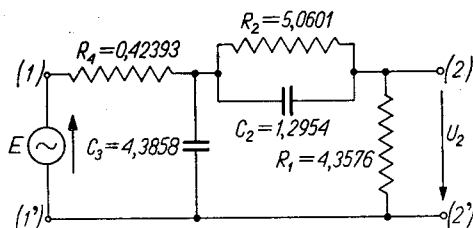
oraz oporność $Z^{(1)}(s) = 1/Y_{22}^{(1)}(s)$ w postaci sumy ułamków częściowych

$$Z^{(1)}(s) = \frac{A_1}{s + k_1} + \frac{A_2}{s + k_2} = \frac{0,77199}{s + 0,15257} + \frac{0,22801}{s + 0,53785}.$$

Łatwo sprawdzić, że oporność ta jest sumą oporności dwóch równoległych obwodów rezonansowych RC. Rezonans pierwszego z tych obwodów wypada dla $s = -0,15257$, mamy tu więc do czynienia z dwójnikiem Z_2 na rys. 3. Drugi obwód rezonansowy — to oporność Z_4 oraz połączony z nią równolegle kondensator C_3 realizujący dwójnik Z_3 . Wartości C_2 , R_2 , C_3 , R_4 obliczamy przy pomocy następujących, elementarnych wzorów

$$\begin{aligned} C_2 &= \frac{1}{A_1} = \frac{1}{0,77199} = 1,2954; & R_2 &= \frac{1}{C_2 k_1} = \frac{1}{1,2954 \cdot 0,15257} = 5,0601; \\ C_3 &= \frac{1}{A_2} = \frac{1}{0,22801} = 4,3858; & R_4 &= \frac{1}{C_3 k_2} = \frac{1}{4,3858 \cdot 0,53785} = 0,42393. \end{aligned}$$

Na rysunku 4 podane są struktura i wartości elementów zaprojektowanego czwórnika.



Rys. 4. Czwórnik RC o jednym rzeczywistym pierwiastku funkcji przenoszenia

Na zakończenie należy wyznaczyć stały mnożnik A funkcji przenoszenia $h(s)$ zrealizowanej przy pomocy tego czwórnika. W tym celu podstawiamy do (1.8) $s = 0$ otrzymując

$$h_1(0) = \frac{0,15257}{0,29800 \cdot 0,62190} = 0,82323.$$

Natomiast z rysunku 4 dostajemy

$$h(0) = A \cdot h_1(0) = \frac{R_1}{R_1 + R_2 + R_4} = 0,44277.$$

Wobec tego

$$A = \frac{0,44277}{0,82323} = 0,53785.$$

Ostatecznie zamiast funkcji $h_1(0)$ danej wzorem (1.8) zrealizowaliśmy

$$h(s) = 0,53785 \frac{s + 0,15257}{(s + 0,29800)(s + 0,62190)}.$$

Przykład 2. Zrealizować czwórnik RC, którego charakterystyką przenoszenia byłaby, z dokładnością do stałego mnożnika, funkcja

$$w_1(\omega^2) = \sqrt{\frac{\omega^4 + \omega^2 0,050056 + 0,0053506}{(\omega^2 + 0,057720)(\omega^2 + 0,24421)(\omega^2 + 0,60251)}}.$$

Jak w poprzednim przykładzie sprawdzamy, że powyższa charakterystyka spełnia warunki 1÷3 twierdzenia 3, a więc może być zrealizowana po ewentualnym przemnożeniu przez dostatecznie mały czynnik. W celu dokonania syntezy odpowiedniego czwornika RC wyznaczamy jego funkcję przenoszenia $h_1(s)$ wychodząc z równania (1.6)

$$\begin{aligned} [w_1(\omega^2)]^2 &= h_1(s)h_1(-s)|_{s=i\omega} = \\ &= \frac{(\omega + 0,15511 + i0,22156)(\omega + 0,15511 - i0,22156)(\omega - 0,15511 + \\ &\quad + i0,22156)(\omega - 0,15511 - i0,22156)}{(0,24025 + i\omega)(0,24025 - i\omega)(0,49417 + i\omega)(0,49417 + \\ &\quad - i\omega)(0,77622 + i\omega)(0,77622 - i\omega)} \end{aligned}$$

Stąd

$$\begin{aligned} h_1(s) &= \frac{(s + 0,22156 + i0,15511)(s + 0,22156 - i0,15511)}{(s + 0,24025)(s + 0,49417)(s + 0,77622)} = \\ &= \frac{s^2 + 0,44312s + 0,073148}{(s + 0,24025)(s + 0,49417)(s + 0,77622)}. \end{aligned}$$

W celu dokonania fizycznej realizacji powyższej funkcji przenoszenia zastosujemy metodę Dashera ([7], str. 220—231). Jest to wersja metody Darlingtona-Guillemina stosowana wówczas, gdy licznik $h(s)$ ma zespolone pierwiastki. Przede wszystkim zapisujemy $h_1(s)$ w postaci

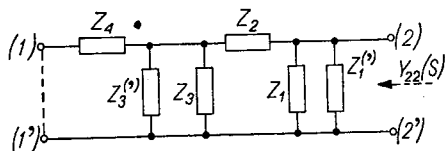
$$h_1(s) = - \frac{Y_{12}(s)}{Y_{22}(s)} = \frac{\frac{s^2 + 0,44312s + 0,073148}{(s + 0,35)(s + 0,6)}}{\frac{(s + 0,24025)(s + 0,49417)(s + 0,77622)}{(s + 0,35)(s + 0,6)}}$$

określając tym sposobem przewodności $Y_{22}(s)$ oraz $Y_{12}(s)$ realizowalne w sensie podanym w rozdziale 1.2. W dalszym ciągu realizujemy czwórnik RC, którego wyjściowa oporność zwarciova wyniesie $Y_{22}(s)$, a pierwiastki przewodności wzajemnej pokryją się z pierwiastkami funkcji $Y_{12}(s)$.

Ze względów omówionych w przykładzie 1 funkcja przenoszenia takiego układu jest identyczna z dokładnością do stałego mnożnika z funkcją $h_1(s)$.

Ogólny schemat czwórnika RC realizującego powyższe wymagania przedstawiony jest na rys. 5.

Dwójniki $Z_1^{(1)}$, Z_4 na tym rysunku są opornościami rzeczywistymi, Z_1 oraz Z_3 szeregowymi obwodami rezonansowymi RC, $Z_3^{(1)}$ — kondensato-



Rys. 5. Struktura czwórnika RC o jednym zespolonym pierwiastku funkcji przenoszenia

rem, Z_2 — dwójnikiem o oporności równej ∞ dla $s = s_0 = -0,22156 \pm \pm i 0,1511$, czyli dla pierwiastków funkcji $Y_{12}(s)$. Proces syntezy rozpoczynamy od obliczenia wartości funkcji $Y_{22}(s)$ i jej pochodnej w punkcie $s_0 = -\alpha + i\beta = -0,22156 + i 0,11511$ otrzymując

$$Y_{22}(s_0) = g + ib = 0,19294 + i 0,28317,$$

$$Y'_{22}(s_0) = g' + ib' = 0,99026 - i 0,766049.$$

Następnie ze wzoru

$$\frac{1}{R'_1} = g - b \frac{b - \beta g' + \alpha b'}{\left(\frac{\alpha}{\beta}\right) b - \alpha g' - \beta b'}$$

obliczamy pomocniczą oporność $Z_1^{(1)} = R'_1 = 4,3411$. Rola tej oporności zostanie wyjaśniona później. Po odjęciu przewodności $(R'_1)^{-1}$ do zrealizowania pozostanie przewodność

$$Y_{22}^{(1)}(s) = Y_{22}(s) - \frac{1}{R'_1} = \frac{s^2 + 1,2803 s^2 + 0,46996 s + 0,043781}{(s + 0,35)(s + 0,6)}.$$

Dla $s = s_0$ otrzymujemy oczywiście

$$Y_{22}^{(1)}(s_0) = g^{(1)} + ib^{(1)} = g - \frac{1}{R'_1} + ib = -0,037420 + i 0,28317.$$

Od przewodności $Y_{22}^{(1)}(s)$ należy w dalszym ciągu odjąć taką przewodność $Y_1(s) = Z_1^{-1}(s)$ (zob. rys. 5), która dla $s = s_0$ wyniesie $g^{(1)} + ib^{(1)}$. Wówczas dwójnik $Z_2(s)$ można będzie skonstruować w ten sposób, aby jego oporność wyniosła ∞ dla $s = s_0$. Przewodność $Y_1(s)$ realizujemy w postaci szeregowego połączenia oporności R_1 i pojemności C_1 , zatem

$$Y_1(s) = \frac{1/R_1}{s + 1/R_1 C_1} = \frac{A_1}{s + \sigma}.$$

Łatwo sprawdzić, że w celu uzyskania $Y_1(s_0) = Y_1(-\alpha + i\beta) = g^{(1)} + i b^{(1)}$ musimy wziąć

$$A_1 = \frac{1}{b^{(1)}} \frac{g^{(1)2} + b^{(1)2}}{g^{(1)}/b^{(1)} + \alpha/\beta} = 0,22227,$$

$$\sigma = \frac{a^2/\beta + \beta}{g/b + \alpha/\beta} = 0,36381.$$

W dalszym ciągu obliczamy przewodność $Y_{22}^{(2)}(s)$, jaka pozostanie do zrealizowania po odjęciu $Y_1(s)$ od $Y_{22}^{(1)}(s)$

$$Y_{22}^{(1)}(s) - Y_1(s) = Y_{22}^{(2)}(s) - \frac{A_1}{s + \sigma} =$$

$$= \frac{s^4 + 1,4218 s^3 + 0,72458 s^2 + 0,16808 s + 0,015928}{(s + 0,35)(s + 0,6)(s + 0,36381)}.$$

Przewodność dwójnika Z_2 na rys. 5 zakładamy w postaci

$$Y_2(s) = B \frac{(s + \alpha + i\beta)(s + \alpha - i\beta)}{s + \sigma} = B \frac{s^2 + 2\alpha s + \alpha^2 + \beta^2}{s + \sigma},$$

dzięki czemu jego oporność wyniesie ∞ dla $s = s_0 = -\alpha \pm i\beta$. Stałą B wyznaczamy z warunku, aby oporność

$$\frac{1}{Y_{22}^{(3)}} = \frac{1}{Y_{22}^{(2)}} - \frac{1}{Y_2}$$

nie miała bieguna w s_0 (dzięki temu redukuje się jej stopień). Warunek ten ma postać

$$B = \lim_{s \rightarrow s_0} \left[\frac{s + \sigma}{s^2 + 2\alpha s + \alpha^2 + \beta^2} \cdot Y_{22}^{(2)}(s) \right],$$

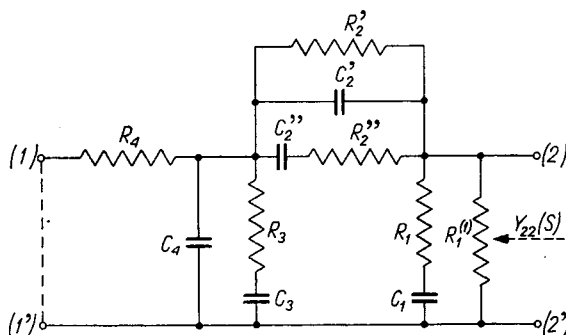
skąd $B = 1,0566$. Po uwzględnieniu tej wartości B i rozbiciu $Y_2(s)$ na ułamki częściowe dostaniemy

$$Y_2(s) = \frac{A_2}{s + \sigma} + \frac{1}{R_2'} + C_2' s = -\frac{0,12865}{s + 0,36381} + 0,21245 + 1,0566 s.$$

Dwójnik $Z_2(s)$ na rys. 5 może być więc zrealizowany w postaci równoległego połączenia szeregowego obwodu rezonansowego $R_2'' C_2''$ o ujemnych elementach z rzeczywistą opornością R_2' i pojemnością C_2' , jak to przedstawia rys. 6.

W dalszym ciągu obliczamy przewodność $Y_{22}^{(3)}(s)$, jaką otrzymamy po wydzieleniu dwójnika $Y_2(s)$ z przewodności $Y_{22}^{(2)}(s)$

$$\frac{1}{Y_{22}^{(3)}(s)} = \frac{1}{Y_{22}^{(2)}(s)} - \frac{1}{Y_2(s)} = \frac{0,053604(s + \sigma)(s^2 + 2\alpha s + \alpha^2 + \beta^2)}{(s^2 + 0,97871 s + 0,21775)(s^2 + 2\alpha s + \alpha^2 + \beta^2)},$$



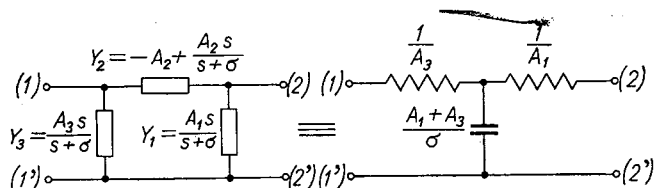
Rys. 6. Czwórnik RC o jednym zespolonym pierwiastku funkcji przenoszenia zrealizowany przy pomocy ujemnych elementów

czyli

$$Y_{22}^{(3)}(s) = 18,655 \frac{s^2 + 0,97871s + 0,21775}{s + 0,36381} = \frac{A_1}{s + \sigma} + \frac{1}{R_4} + sC_4 =$$

$$= \frac{0,30551}{s + 0,36381} + \frac{1}{0,089559} + 18,655s.$$

Otrzymaną przewodność $Y_{22}^{(3)}(s)$ można więc zrealizować w postaci równoległego połączenia szeregowego obwodu rezonansowego R_3C_3 , kondensatora C_4 i rzeczywistej oporności R_4 otrzymując ostatecznie czwórnik na



Rys. 7. Przekształcenie czwórnik RC eliminujące ujemne elementy

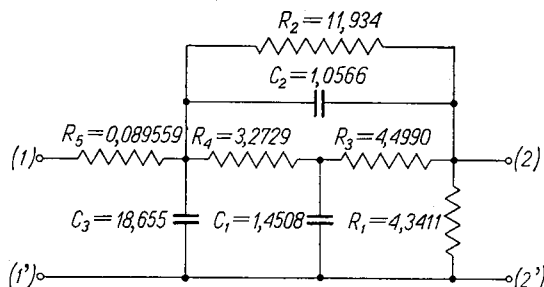
rys. 6. Czwórnik ten nie jest fizycznie realizowalny, gdyż zawiera ujemne elementy w dwójniku Z_2 (por. rys. 5). Zrealizujemy go jednak po zastosowaniu równoważności podanej na rys. 7.

Równoważność ta zachodzi jedynie wówczas, gdy spełniony jest warunek

$$A_1 A_2 + A_1 A_3 + A_2 A_3 = 0.$$

W naszym wypadku, jak łatwo sprawdzić, jest on spełniony. Stało się to dzięki odjęciu przewodności $1/R_1'$ od $Y_{22}(s)$ obliczonej według wzoru podanego na początku naszego przykładu. Po zastosowaniu tożsamości z rys. 7 dostaniemy ostatecznie czwórnik o strukturze i elementach podanych na rys. 8. Stały mnożnik funkcji przenoszenia wyznaczamy identycznie, jak w przykładzie 1. Obliczamy mianowicie

$$h_1(0) = \frac{0,073148}{0,24025 \cdot 0,49417 \cdot 0,77622} = 0,79374.$$



Rys. 8. Czwórnik RC o jednym zespolonym pierwiastku funkcji przenoszenia zrealizowany przy pomocy dodatnich elementów

Natomiast z rysunku 8 mamy

$$h(0) = A \cdot h_1(0) = \frac{R_1}{R_1 + R_5 + \frac{R_2(R_3 + R_4)}{R_2 + R_3 + R_4}} = 0,47510,$$

czyli

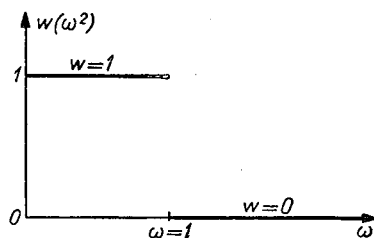
$$A = \frac{0,475098}{0,79374} = 0,59856.$$

Otrzymaliśmy zatem czwórnik o funkcji przenoszenia

$$h(s) = 0,59856 \frac{s^2 + 0,44312s + 0,073148}{(s + 0,24025)(s + 0,49417)(s + 0,77622)}.$$

1.4. Wymagania dotyczące charakterystyki tłumieniowej dolnoprzepustowego filtra RC oraz jej parametry

W poprzednich rozdziałach poznaliśmy własności charakterystyki tłumieniowej czwórnika RC oraz najważniejszą w praktyce metodę jej urze-



Rys 9. Charakterystyka tłumieniowa idealnego filtra

czywistniania. Obecnie zajmiemy się problemem skonstruowania tej charakterystyki w taki sposób, aby zaprojektowany w oparciu o nią filtr dolnoprzepustowy miał dobre własności przy możliwie małej liczbie elementów. Wspominaliśmy już, że problem ten jest w zasadzie problemem

aproksymacji przebiegu idealnego funkcją realizowalną jako charakterystyka przenoszenia. Dla filtru dolnoprzepustowego przebieg idealny przedstawiony jest na rys. 9.

Filtr o takim przebiegu charakterystyki $w(\omega^2)$ przenosiłby bez zniekształceń amplitudy wszystkie częstotliwości od 0 do 1 i nie przepuszczałby żadnej większej od 1. Warto przy tym podkreślić, że założenie 1 jako górnej granicy pasma przepustowego nie jest żadnym ograniczeniem, gdyż po zaprojektowaniu filtru można granicę tę przesunąć do dowolnego punktu f_0 na osi częstotliwości mnożąc w tym celu wszystkie pojemności filtru przez $1/2\pi f_0$.

Wracając do idealnej charakterystyki na rys. 9 stwierdzamy, że nie może ona być zrealizowana przy pomocy jakiegokolwiek czwórnika elektrycznego. Istotnie, charakterystyki takich czwórników jako pierwiastki kwadratowe ze stosunku dwóch wielomianów rzeczywistych nie mogą być stałe w żadnym przedziale zmiennej ω , a tym bardziej nie mogą mieć skoku dla $\omega = 1$. Można je uczynić co najwyżej dowolnie bliskimi funkcji na rys. 9, ale tylko przy dostatecznie wysokim stopniu ich licznika i mianownika. Zwiększenie stopnia charakterystyki powoduje jednak zwiększenie liczby elementów realizującego je czwórnika, czego przykładem mogą być przeliczenia przeprowadzone w poprzednim rozdziale. Definiując jako stopień funkcji (1.7) sumę stopni wielomianów w jej liczniku i mianowniku należałoby problem skonstruowania pożądanej charakterystyki $w(\omega^2)$ sformułować następująco:

Wśród wszystkich funkcji rzeczywistych $w(\omega^2)$ określonych wzorem (1.7) stopnia $s = m + n$ nie większego niż s_0 znaleźć taką $W(\omega^2)$, która w pasmie zaporowym $\omega_z \leq \omega \leq \infty$, gdzie $\omega_z > 1$ jest nie większa od danej liczby ε i dla której różnica $|1 - W(\omega^2)|$ w pasmie przepustowym $0 \leq \omega \leq 1$ jest minimalna.

Rozwiązanie tak sformułowanego problemu podał Darlington [4], ([3] rozdz. VIII) w postaci pierwiastka kwadratowego z funkcji wymiernej od ω^2 o specyficznym rozkładzie pierwiastków i biegunów. Podał on również sposób realizowania tych charakterystyk przy pomocy drabinkowych czwórników LC zakończonych rzeczywistymi opornościami. Nasuwa się od razu pytanie, czy charakterystyki Darlingtona nie mogłyby być zastosowane również w teorii dolnoprzepustowych filtrów RC. Otóż nie jest to możliwe, ponieważ bieguny charakterystyk Darlingtona są zespolone, czyli nie spełniają warunku 3 w twierdzeniu 3, rozdział 1.2. Również rozmieszczenie pierwiastków charakterystyki Darlingtona nie jest dla nas dogodne. Leżą one w pasmie zaporowym filtru, a więc są rzeczywiste. Istnieją wprawdzie czwórniki RC, które tłumią rzeczywiste częstotliwości (np. mostek Wiena), ale zawierają wiele elementów i trudno je dostosować do współpracy z lampami i tranzystorami. Dlatego problem skonstruowania pożądanej charakterystyki $W(\omega^2)$ sformułujemy w inny sposób niż Darlington dla filtrów LC.

Pierwsza różnica polega na obraniu z góry biegunów funkcji $w(\omega^2)$ jako pojedynczych, ujemnych, a więc na rezygnacji z ustalenia ich w wyniku rozwiązania zagadnienia aproksymacyjnego. Druga różnica polega na tym, że nie minimalizujemy funkcji $w(\omega^2)$ w przedziale $\omega_2^2 \leq \omega^2 \leq \infty$, co prowadzioby do pojawienia się w nim rzeczywistych pierwiastków. Żądamy jedynie, aby dla $\omega > 1$ charakterystyka $w(\omega^2)$ malała monotonicznie i dostatecznie szybko do 0, gdy $\omega \rightarrow \infty$. Stąd wniosek, że różnica $2N = 2n - 2m$ między stopniami mianownika i licznika funkcji $w^2(\omega^2)$ (zob. wzór 1.7), jako miara szybkości jej malenia, należy do ważnych parametrów. Parametr ten określamy przyjmując następującą definicję:

Definicja 2. Klasę n charakterystyki tłumieniowej $w(\omega^2)$ nazywać będziemy połowę liczby biegunów funkcji $w^2(\omega^2)$. Parametrem odcięcia N charakterystyki tłumieniowej $w(\omega^2)$ nazywać będziemy połowę różnicy między liczbą biegunów i pierwiastków funkcji $w^2(\omega^2)$. Parę liczb (n, N) nazwiemy typem charakterystyki tłumieniowej.

Po zdefiniowaniu parametrów $n, N, (n, N)$ formułujemy problem aproksymacyjny dla dolnoprzepustowych filtrów RC, który w dalszym ciągu nazywać będziemy *problemem filtracji*.

Wśród wszystkich charakterystyk tłumieniowych

$$w(\omega_1^2, \sigma_v^2) = \sqrt{\frac{A_0 \omega^{2m} + A_1 \omega^{2(m-1)} + \dots + A_{(m-1)} \omega^2 + A_m}{(\omega^2 + \sigma_1^2)(\omega^2 + \sigma_2^2) \dots (\omega^2 + \sigma_n^2)}} \quad (1.10)$$

typu (n, N) , gdzie σ_v są liczbami rzeczywistymi, znaleźć charakterystykę $W(\omega_1^2, \sigma_v^2)$, która w przedziale $0 \leq \omega \leq 1$ aproksymuje w sensie Czebyszewa stałą $C > 0$, czyli dla której liczba

$$L = \max_{0 \leq \omega \leq 1} |C - W(\omega_1^2, \sigma_v^2)|$$

osiąga minimalną wartość.

W rozdziale 2 podamy twierdzenie o istnieniu rozwiązania $W(\omega^2, \sigma_v)$ tak sformułowanego problemu filtracji. Dla rozmaitych stałych C użyjemy oczywiście różne funkcje W , okaże się jednak, że funkcje te różnią się jedynie stałym mnożnikiem. W związku z tym rozwiązaniem zagadnienia filtracji nazwiemy całą rodzinę funkcji $W(C, \omega^2, \sigma_v^2)$, a poszczególne funkcje należące do tej rodziny — *postaciami rozwiązania zagadnienia filtracji*.

W rozdziale 2 pokażemy, że rodzina $W(C, \omega_1^2, \sigma_v^2)$ jest jednoznacznie określona przez swój typ (n, N) oraz przez zbiór biegunów σ_v . Oprócz liczb n, N oraz σ_v zachodzi jednak konieczność wprowadzenia jeszcze dwóch parametrów charakteryzujących praktyczną przydatność rozwiązania $W(C, \omega^2, \sigma_v)$. Konieczność ta wynika z następującego rozumowania.

Założmy, że zrealizowaliśmy dwa filtry RC o charakterystykach $W(\omega^2)$ oraz $kW(\omega^2)$, gdzie $k > 1$, będących oczywiście dwiema postaciami rozwiązania zagadnienia filtracji. Założmy dalej, że na wejścia obu filtrów załączaliśmy siłę elektromotoryczną $E = 1$ V o częstotliwości ω zmienia-

jącej się od 0 do 1. Przy takich zmianach ω amplitudy napięć wyjściowych $U_2 = E \cdot W(\omega^2)$ nie będą stałe, filtry będą więc wnosić zniekształcenia tłumieniowe. Różnica między maksymalną a minimalną wartością U_2 pierwszego filtru będzie oczywiście mniejsza niż drugiego, ale stąd nie wynika, że drugi będzie gorszy. Istotnie, zniekształcenia w mierze logarytmicznej

$$\Delta = 20 \log U_{2max} - 20 \log U_{2min} \Big|_{0 \leq \omega \leq 1}$$

lub zniekształcenia względne

$$\frac{U_{2max} - U_{2min}}{U_{2min}} \Big|_{0 \leq \omega \leq 1}$$

dla obu filtrów będą jednakowe. Zachodzi zatem konieczność wprowadzenia parametru, który pozwoliłby ocenić pod względem zniekształceń tłumieniowych całą rodzinę charakterystyk $W(\omega^2, \sigma_v)$. Dokonamy tego przyjmując następujące dwie definicje.

Definicja 3. Postacią znormalizowaną $W_0(\omega^2, \sigma_v^2)$ rozwiązania $W(C, \omega^2, \sigma_v^2)$ zagadnienia filtracji nazwiemy tę jego postać, której minimum w przedziale $0 \leq \omega^2 \leq 1$ równa się jedności.

Definicja 4. Falistością δ rozwiązania $W(C, \omega^2, \sigma_v^2)$ zagadnienia filtracji nazywamy różnicę między maksymalną i minimalną wartością kwadratu jego znormalizowanej postaci w przedziale $0 \leq \omega^2 \leq 1$

$$\delta = \max_{0 \leq \omega^2 \leq 1} W_0^2(\omega^2, \sigma_v^2) - 1.$$

Jeżeli znamy jakąkolwiek postać $W(\omega^2, \sigma_v^2)$ rozwiązania $W(C, \omega^2, \sigma_v^2)$, to mnożąc ją przez czynnik liczbowy

$$k = 1 / \min_{0 \leq \omega^2 \leq 1} W(\omega^2, \sigma_v^2)$$

znajdujemy postać znormalizowaną $W_0(\omega^2, \sigma_v^2) = kW(\omega^2, \sigma_v^2)$ tego rozwiązania. Stąd wniosek, że falistość wyraża się wzorem

$$\delta = \frac{\max_{0 \leq \omega^2 \leq 1} W(C, \omega^2, \sigma_v^2) - \min_{0 \leq \omega^2 \leq 1} W(C, \omega^2, \sigma_v^2)}{\min_{0 \leq \omega^2 \leq 1} W(C, \omega^2, \sigma_v^2)} \quad (1.11)$$

dla każdego C .

Oprócz szczególnej charakterystyki $W_0(\omega^2, \sigma_v^2)$ dużą rolę odegra w dalszych rozważaniach jeszcze jedna funkcja z rodziny $W(C, \omega^2, \sigma_v^2)$ określona w następujący sposób.

Definicja 5. Postacią wyjściową $W_1(\omega^2, \sigma_v^2)$ rozwiązania $W(C, \omega^2, \sigma_v^2)$, zagadnienia filtracji nazwiemy tę funkcję z rodziny $W(C, \omega^2, \sigma_v^2)$, której współczynniki przy najwyższych potęgach ω w liczniku i mianowniku równają się jedności.

Jeżeli znamy jakąkolwiek postać $W(\omega^2, \sigma_v^2)$ rozwiązania $W(C, \omega^2, \sigma_v^2)$, to z łatwością możemy wyznaczyć jego postać wyjściową. W tym celu należy

funkcję $W(\omega^2, \sigma_v^2)$ zapisać w formie (1.10), a następnie podzielić przez $\sqrt{A_0}$. Aby ocenić znaczenie postaci $W_1(\omega^2, \sigma_v^2)$, wróćmy do przykładów syntezy podanych w rozdziale 1.3. Otóż zrealizowane tam za pomocą czwórnika RC charakterystyki tłumieniowe miały postać wyjściową. W wyniku syntezy uległy one przemnożeniu przez czynnik stały równy w pierwszym przykładzie 0,53785, w drugim 0,59856. Przyjmijmy nieco sztuczne założenie, że czynnik ten zależy jedynie od typu charakterystyki, nie zależy natomiast od zbioru biegunów $\sigma_1^2, \sigma_2^2 \dots \sigma_n^2$, oraz przypuśćmy, że mamy dwa różne rozwiązania zagadnienia filtracji tego samego typu i o tej samej falistości δ (rozwiązania tego samego typu mogą się różnić jedynie rozmieszczeniem biegunów σ_v^2). W celu zrealizowania filtru wybieramy oczywiście to z dwu rozwiązań, które w swej postaci wyjściowej posiada większe maksimum w przedziale $0 \leq \omega^2 \leq 1$. Istotnie, ponieważ obie postaci wyjściowe ulegają w procesie syntezy przemnożeniu przez ten sam czynnik, więc filtr o tak dobranym rozwiązaniu $W(\omega^2, \sigma_v^2)$ będzie dawał większe napięcie na wyjściu. Aby scharakteryzować rozwiązanie zagadnienia filtracji z punktu widzenia napięcia wyjściowego, zdefiniujemy jeszcze jeden parametr tego rozwiązania.

Definicja 6. *Zyskiem K rozwiązania $W(C, \omega^2, \sigma_v^2)$ zagadnienia filtracji nazywamy maksymalną wartość jego postaci wyjściowej $W_1(\omega^2, \sigma_v^2)$ w przedziale $0 \leq \omega^2 \leq 1$.*

Jeżeli znamy znormalizowaną postać $W_0(\omega^2, \sigma_v^2)$ rozwiązania $W(C, \omega^2, \sigma_v^2)$ oraz jego falistość δ , to jego zysk możemy obliczyć ze wzoru

$$k = \sqrt{\frac{1 + \delta}{A_0^{(0)}}}, \quad (1.12)$$

gdzie $A_0^{(0)}$ jest współczynnikiem przy najwyższej potędze ω w liczniku funkcji $W_0(\omega^2, \sigma_v^2)$ zapisanej w postaci (1.10).

Aby na zakończenie dokonać przeglądu wszystkich wprowadzonych dotychczas parametrów rozpatrzmy konkretne rozwiązanie problemu filtracji znalezione w rozdziale (2.5) i wykorzystane w rozdziale (1.3), przykład 2 jako punkt wyjścia w projekcie czwórnika RC. Rozwiązanie to można zapisać w postaci

$$\begin{aligned} W(C, \omega_2^2, \sigma_1^2, \sigma_2^2, \sigma_3^2) &= k(C) \sqrt{\frac{A_0^{(0)} \omega^4 + A_1^{(0)} \omega^2 + A_2^{(0)}}{(\omega^2 + \sigma_1^2)(\omega^2 + \sigma_2^2)(\omega^2 + \sigma_3^2)}} = \\ &= k(C) \sqrt{\frac{1,9981 \omega^4 + 0,10002 \omega^2 + 0,010692}{(\omega^2 + 0,057720)(\omega^2 + 0,24421)(\omega^2 + 0,60251)}}. \end{aligned} \quad (1.13)$$

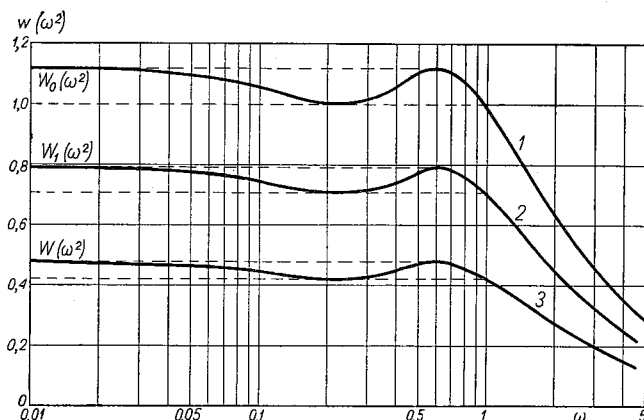
Otóż wartości parametrów tego rozwiązania wynoszą

Klasa $n = 3$ (6 biegunów urojonych, zob. definicja 2)

Parametr odcięcia $N = n - m = 3 - 2 = 1$ (zob. definicja 2)

Typ $(n, N) = (3, 1)$ (definicja 2).

W celu znalezienia falistości δ zakładamy pewną wartość $k(C)$, np.



Rys. 10. Różne postaci rozwiązania zagadnienia filtracji

$k(C) = 1$ i badamy funkcję $k(C) \cdot W$, której przebieg przedstawia krzywa 1 na rys. 10. Ponieważ minimum tej funkcji w przedziale $0 \leq \omega^2 \leq 1$ wynosi 1, więc zgodnie z definicją 3

$$W_0(\omega^2, \sigma_1^2, \sigma_2^2, \sigma_3^2) = \sqrt{\frac{1,9982 \omega^4 + 0,10002 \omega^2 + 0,010692}{(\omega^2 + 0,057720)(\omega^2 + 0,24421)(\omega^2 + 0,60251)}}$$

jest postacią *znormalizowaną* rozwiązania (1.13). Stosując definicję 4 znajdujemy *falistość* tego rozwiązania (por. krzywa 1 na rys. 10).

$$\delta = \max_{0 < \omega < 1} W_0^2(\omega^2, \sigma_1^2, \sigma_2^2, \sigma_3^2) - 1 = 1,1220^2 - 1 = 0,25893.$$

Funkcja $W_0(\omega^2)$ oscyluje w przedziale $0 \leq \omega^2 \leq 1$ między wartościami 1 oraz $\sqrt{1 + \delta} = 1,1220$. W rozdziale 2 udowodnimy, że funkcja ta odchyła się od stałej $C_0 = \frac{1}{2}(1 + 1,1220) = 1,0610$ najmniej spośród wszystkich funkcji typu (3,1) o biegunach

$$\sigma_1^2 = 0,057720, \quad \sigma_2^2 = 0,24421, \quad \sigma_3^2 = 0,60251.$$

Można zatem zapisać

$$W_0(\omega^2, \sigma_1^2, \sigma_2^2, \sigma_3^2) = W(1,0610, \omega^2, \sigma_1^2, \sigma_2^2, \sigma_3^2).$$

Aby znaleźć wyjściową postać rozwiązania, dzielimy $W_0^2(\omega^2)$ przez współczynnik przy najwyższej potędze ω w jej liczniku, czyli przez $A_0^{(0)} = 1,9982$. Dostaniemy w ten sposób funkcję

$$W_1(\omega^2, \sigma_1^2, \sigma_2^2, \sigma_3^2) = \sqrt{\frac{\omega^4 + 0,050056 \omega^2 + 0,0053506}{(\omega^2 + 0,057720)(\omega^2 + 0,24421)(\omega^2 + 0,60251)}}$$

przedstawioną za pomocą krzywej 2 na rys. 10. Przybliży ona oczywiście stałą $C_1 = C_0 / \sqrt{A_0^{(0)}} = 1,0610 / \sqrt{1,9982} = 0,75056$,

a więc

$$W_1(\omega^2, \sigma_1^2, \sigma_2^2, \sigma_3^2) = W(0,75056, \omega^2, \sigma_1^2, \sigma_2^2, \sigma_3^2).$$

Zysk K rozwiązania (1.13) obliczamy zgodnie z definicją 6 jako $\max W_1(\omega^2, \sigma_1^2, \sigma_2^2, \sigma_3^2)$ dla $0 \leq \omega^2 \leq 1$, czyli według wzoru (1.12)

$$K = \sqrt{\frac{1+\delta}{A_0^{(0)}}} = \sqrt{\frac{1,2589}{1,9982}} = 0,79374.$$

Krzywa 3 na rys. 10 przedstawia jeszcze jedną postać rozwiązania (1.13), mianowicie tę, która przybliży stałą $C_w = 0,59856 \cdot C_1 = 0,44926$, czyli tę, która została zrealizowana przy pomocy czwórnika na rys. 8 w przykładzie 2 rozdział 1.3.

2. CZEBYSZEWOWSKIE CHARAKTERYSTYKI TŁUMIENIOWE DOLNOPRZEPUSTOWEGO FILTRU RC

2.1. Twierdzenie Achiezera

Rozpatrzmy następujące zagadnienie aproksymacyjne: Wśród wszystkich funkcji $q(x)$ typu

$$q(x) = s(x) \frac{a(x)}{b(x)} = s(x) \frac{a_0 x^m + a_1 x^{m-1} + \dots + a_m}{b_0 x^n + b_1 x^{n-1} + \dots + b_n} \quad (2.1)$$

znaleźć taką funkcję $Q(x)$, dla której wartość

$$L = \max |f(x) - Q(x)|$$

w pewnym przedziale $[a, b]$ jest najmniejsza z możliwych. Jest to zagadnienie aproksymacji w sensie Czebyszewa funkcji $f(x)$ w przedziale $a \leq x \leq b$ za pomocą funkcji wymiernej $a(x)/b(x)$ z „wagą” $s(x)$. O zagadnieniu tym twierdzenie Achiezera mówi, co następuje ([10] str. 60):

Twierdzenie 4. Jeżeli funkcje $f(x)$ oraz $s(x)$ są ciągłe w przedziale $[a, b]$ i ponadto $s(x) \neq 0$ gdy $a \leq x \leq b$, to funkcja $Q(x)$ istnieje i jest jedyna. Funkcja ta charakteryzuje się całkowicie następującą własnością: Jeżeli zapiszemy ją w postaci

$$Q(x) = s(x) \frac{a(x)}{\beta(x)} = s(x) \frac{a_0 x^{m-\mu} + a_1 x^{m-\mu-1} + \dots + a_{m-\mu}}{\beta_0 x^{n-\nu} + \beta_1 x^{n-\nu-1} + \dots + \beta_{n-\nu}},$$

gdzie $a(x)/\beta(x)$ jest ułamkiem nieskracalnym oraz $a_0, \beta_0 \neq 0, 0 \leq \mu \leq m, 0 \leq \nu \leq n$, wówczas liczba N kolejnych punktów przedziału $[a, b]$, w których różnica $f(x) - Q(x)$ osiąga bezwzględnie maksymalną wartość L ze znakiem na przemian $+$ i $-$, jest nie mniejsza niż $m + n + 2 - \min(\mu, \nu)$.

Zagadnienie filtracji sformułowane w rozdziale 1.4 sprowadza się łatwo do zagadnienia aproksymacyjnego, o którym mówi twierdzenie 4. W tym celu wystarczy przyjąć $\omega^2 = x$ i podnieść charakterystykę (1.10) do kwadratu. Dostaniemy wówczas

$$f(x) = C^2 = \text{const},$$

$$s(x) = 1/(x + \sigma_1^2)(x + \sigma_2^2) \dots (x + \sigma_n^2), \quad a(x) = A_0 x^n + A_1 x^{n-1} + \dots + A_n,$$

$$b(x) = 1, \quad [a, b] = [0, 1]$$

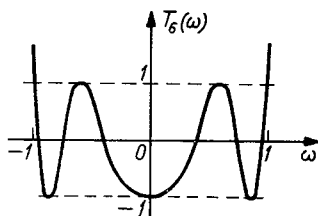
Charakterystyka $W(\omega^2, \sigma_v^2)$ istnieje więc na pewno. Twierdzenie 4 nie mówi wprawdzie, jak ją znaleźć, ale pozwoli ono stwierdzić, czy dana funkcja $W(\omega^2, \sigma_v^2)$ jest rozwiązaniem zagadnienia filtracji.

2.2. Wielomianowa charakterystyka tłumieniowa

Wielomianowymi będziemy nazywali charakterystyki typu $[n, n]$, czyli

$$w(\omega^2) = \sqrt{\frac{1}{B(\omega^2)}} = \sqrt{\frac{A_0}{B_0 \omega^{2n} + B_1 \omega^{2(n-1)} + \dots + B_n}}.$$

Postaramy się skonstruować taką charakterystykę wielomianową klasy n , która w swej postaci wyjściowej (tzn. dla $A_0 = 1, B_0 = 1$) aproksymuje



Rys. 11. Wykres wielomianu Czebyszewa

w sensie Czebyszewa pewną stałą $f(x) = C$ w przedziale $0 \leq \omega^2 \leq 1$ (a więc w przedziale $-1 \leq \omega \leq 1$ z uwagi na symetrię funkcji $w(\omega^2)$). W tym celu bierzemy pod uwagę tę gałąź wieloznacznej funkcji

$$\Psi(\omega) = \arccos(2\omega^2 - 1), \quad (2.2)$$

która w przedziale $-1 \leq \omega \leq 1$ wzrasta monotonicznie od 0 do 2π . Dla gałęzi tej mamy $\cos \Psi = 2\omega^2 - 1$, $\sin \Psi = 2\omega \sqrt{1 - \omega^2}$, gdzie $\sqrt{1 - \omega^2}$ jest ujemny dla $\omega = 0$. Wprowadźmy w dalszym ciągu funkcję

$$T_{2n}(\omega) = \cos[n\Psi(\omega)],$$

która n -krotnie oscyluje między wartościami 1 i -1 , gdy ω wzrasta od -1 do 1. Stosując wzór na kosinus wielokrotności kąta dostajemy

$$T_{2n}(\omega) = \frac{1}{2} [(\omega + \sqrt{\omega^2 - 1})^{2n} + (\omega - \sqrt{\omega^2 - 1})^{2n}]. \quad (2.3)$$

Podnosząc nawiasy do potęgi $2n$ i redukując wyrazy podobne nietrudno się przekonać, że funkcja $T_{2n}(\omega)$ jest wielomianem parzystym stopnia $2n$ o współczynniku przy najwyższej potędze równym 2^{2n-1} . Na rys. 11 naszkicowano przebieg tego wielomianu dla $n = 3$.

Jeżeli wielomian $T_{2n}(\omega)$ podzielimy przez 2^{2n-1} , to otrzymamy wielomian $\tilde{T}_{2n}(\omega)$ o współczynniku przy najwyższej potędze równym 1. $\tilde{T}_{2n}(\omega)$ różni

się od zera mniej niż o jedność, a wartość tej różnicy $|\tilde{T}_{2n}(\omega)| = 1/2^{2n-1}$ maleje ze wzrostem n . Jeżeli $\tilde{T}_{2n}$ zapiszemy w postaci $\omega^2 - p_{2n-2}(\omega)$, gdzie $p_{2n-2}(\omega)$ jest wielomianem stopnia $2n-2$, to łatwo stwierdzimy, że różnica

$$L = f(\omega) - Q(\omega) = (\omega^{2n} - p_{2n-2}(\omega))$$

przyjmuje w przedziale $[-1, 1]$ $N = 2n + 1$ krotnie swą bezwzględnie maksymalną wartość $1/2^{-n+1}$ na przemian ze znakami $+$ i $-$. Stosując twierdzenie 4 wyprowadzamy stąd wniosek, że wielomian $p_{2n-2}(\omega)$ spośród wszystkich wielomianów stopnia nie wyższego niż $n' = 2n - 1$ aproksymuje najlepiej funkcję ω^{2n} w przedziale $[-1, 1]$, gdyż spełniony jest warunek

$$N = 2n + 1 = n' + 2$$

Można zatem twierdzić, że wielomian $\tilde{T}_{2n}(\omega)$ aproksymuje najlepiej zero spośród wszystkich wielomianów stopnia nie wyższego, niż $2n$ o współczynniku przy najwyższej potędze ω równym 1.

Wielomiany $T_{2n}(\omega)$, dany wzorem (2.3), oraz $\tilde{T}_{2n}(\omega) = T_{2n}(\omega)/2^{2n-1}$, nazywają się wielomianami Czebyszewa $2n$ -tego stopnia. Przy pomocy wielomianu $T_{2n}(\omega)$ z łatwością konstruujemy czebyszewowską charakterystykę wielomianową

$$W(\omega^2) = \sqrt{\frac{1}{\tau + T_{2n}(\omega^2)}}, \quad \tau > 1, \quad (2.4)$$

która aproksymuje stałą

$$C = \sqrt{\frac{1}{2} \left[\frac{1}{\tau + 1} + \frac{1}{\tau - 1} \right]}$$

najlepiej spośród wszystkich charakterystyk wielomianowych klasy nie wyższej niż n , o współczynniku przy najwyższej potędze równym 2^{2n-1} . Istotnie, gdyby istniała charakterystyka wielomianowa różna od (2.4) klasy n , o współczynniku przy ω^{2n} równym 2^{2n-1} , aproksymująca stałą C lepiej niż (2.4), to przy jej pomocy można by skonstruować wielomian $t_{2n}(\omega)$ o współczynniku przy ω^{2n} równym 1 aproksymujący zero lepiej niż $\tilde{T}_{2n}(\omega)$. To jest jednak niemożliwe.

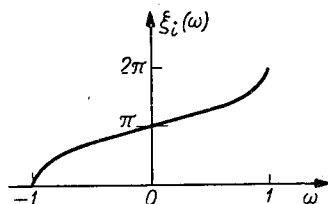
Filtr RC o wielomianowej charakterystyce tłumieniowej miałby dobre własności techniczne: wyrównany przebieg tłumienia w pasmie przepustowym, maksymalny parametr odcięcia w klasie n . Filtru takiego nie można jednak zrealizować, ponieważ czebyszewowska charakterystyka wielomianowa ma zespolone bieguny. Zajęliśmy się nią w celu zilustrowania metody, która w dalszym ciągu posłuży do skonstruowania czebyszewowskich charakterystyk o zadanych biegunach. Poza tym o charakterystyce wielomianowej będzie jeszcze mowa w rozdziale 3.

2.3. Czebyszewska charakterystyka typu $(n, 1)$

Aby skonstruować czebyszewską charakterystykę $W(\omega^2)$ o $2n$ zadanych biegunach $\pm i\sigma_i$, wprowadzimy pomocniczą funkcję $M(\omega^2)$, która odegra podobną rolę, jak wielomian $T(\omega^2)$ w poprzednim zagadnieniu. Rozpatrzmy zależność

$$\cos \xi_i = \frac{(2\sigma_i^2 + 1)\omega^2 - \sigma_i^2}{\omega^2 + \sigma_i^2}. \quad (2.5)$$

Łatwo stwierdzić, że gdy zmienna ω dąży od -1 do 1 , to prawa strona tej zależności maleje monotonicznie od 1 do -1 , a następnie wzrasta mono-



Rys. 12. Wykres jednej z gałęzi funkcji $\xi_i = \arccos \{[(2\sigma_i^2 + 1)\omega^2 - \sigma_i^2]/(\omega^2 + \sigma_i^2)\}$

tonicznie z powrotem do 1 . Zależność $\xi_i = \xi_i(\omega)$ określona wzorem (2.5) jest oczywiście funkcją wieloznaczną. Zajmiemy się jedynie tą jej gałęzią, której wykres przedstawiony jest na rys. 12.

Pochodną uwikłanej funkcji (2.5) wyraża się wzorem

$$\frac{d\xi_i}{d\omega} = -\frac{2\sigma_i\sqrt{1+\sigma_i^2}}{(\omega^2 + \sigma_i^2)\sqrt{1-\omega^2}},$$

skąd wniosek, że gałęzie pierwiastków $\sqrt{1+\sigma_i^2}$ oraz $\sqrt{1-\omega^2}$ muszą mieć różne znaki, jeżeli mają odpowiadać rosnącej gałęzi funkcji $\xi_i(\omega)$ przedstawionej na rys. 12. Dodajmy teraz n funkcji $\xi_i(\omega)$ odpowiadających n różnym liczbom σ_i i oznaczmy

$$M(\omega^2) = \cos [\xi_1(\omega) + \xi_2(\omega) + \dots + \xi_n(\omega)].$$

Jeżeli ω zmienia się od -1 do 1 , to kąt w nawiasie kwadratowym wzrasta od 0 do $2\pi n$, a więc w przedziale $[-1, 1]$ $M(\omega^2)$ oscyluje n -krotnie między wartościami -1 a 1 . Aby zbadać jakiego typu funkcją jest $M(\omega^2)$, zapisujemy ją w postaci

$$\begin{aligned} M(\omega^2) &= \frac{1}{2} [e^{i(\xi_1 + \xi_2 + \dots + \xi_n)} + e^{-i(\xi_1 + \xi_2 + \dots + \xi_n)}] = \frac{1}{2} \prod_{\nu=1}^n (\cos \xi_\nu + i \sin \xi_\nu) + \\ &+ \frac{1}{2} \prod_{\nu=1}^n (\cos \xi_\nu - i \sin \xi_\nu). \end{aligned}$$

Podstawiając w nawiasach za $\cos \xi_\nu$ wyrażenie (2.5), a za $\sin \xi_\nu$ wynikającą z tego wyrażenia funkcję

$$\sin \xi_\nu = \frac{2 \sigma_\nu \omega \sqrt{(1 - \omega^2)(1 + \sigma_\nu^2)}}{\omega^2 + \sigma_\nu^2}$$

z właściwie dobranymi gałęziami $\sqrt{1 - \omega^2}$ oraz $\sqrt{1 + \sigma_\nu^2}$ otrzymamy

$$M(\omega^2) = \left\{ \prod_{\nu=1}^n [(2 \sigma_\nu^2 + 1) \omega^2 - \sigma_\nu^2 + 2 \sigma_\nu \omega \sqrt{(\omega^2 - 1)(\sigma_\nu^2 + 1)}] + \right. \\ \left. + \prod_{\nu=1}^n [(2 \sigma_\nu^2 + 1) \omega^2 - \sigma_\nu^2 - 2 \sigma_\nu \omega \sqrt{(\omega^2 - 1)(\sigma_\nu^2 + 1)}] \right\} : \left\{ 2 \prod_{\nu=1}^n (\omega^2 + \sigma_\nu^2) \right\}. \quad (2.6)$$

Choć w liczniku funkcji $M(\omega^2)$ występują wyrażenia pierwiastkowe, to nie jest on jednak funkcją przestępną, gdyż po wymnożeniu kwadratów nawiasów wyrażenia te ulegają redukcji. $M(\omega^2)$ jest więc funkcją wymierną o $2n$ urojonych biegunach $\pm i \sigma_\nu$ oscylującą w przedziale $[-1, 1]$ między 1 i -1 .

$$M(\omega^2) = \frac{B_0 \omega^{2n} + B_1 \omega^{2n-2} + \dots + B_n}{\prod_{\nu=1}^n (\omega^2 + \sigma_\nu^2)}.$$

Posługując się zależnością (2.6) można dla danego n obliczyć współczynniki B_ν . Otrzymujemy w ten sposób

$$B_0 = \frac{1}{2} \left[\prod_{\nu=1}^n (\sigma_\nu + \sqrt{\sigma_\nu^2 + 1})^2 + \prod_{\nu=1}^n (\sigma_\nu - \sqrt{\sigma_\nu^2 + 1})^2 \right], \quad (2.7)$$

$$B_2 = -\frac{1}{2} \left[\prod_{\nu=1}^n (\sigma_\nu + \sqrt{\sigma_\nu^2 + 1})^2 \cdot \sum_{\mu=1}^n \frac{\sigma_\mu}{\sigma_\mu + \sqrt{\sigma_\mu^2 + 1}} + \right. \\ \left. + \prod_{\nu=1}^n (\sigma_\nu - \sqrt{\sigma_\nu^2 + 1})^2 \sum_{\mu=1}^n \frac{\sigma_\mu}{\sigma_\mu - \sqrt{\sigma_\mu^2 + 1}} \right] \quad (2.8)$$

itd. oraz dla B_n

$$B_n = (-1)^n \prod_{\nu=1}^n \sigma_\nu^n. \quad (2.9)$$

Przy pomocy funkcji $M(\omega^2)$ można skonstruować następująco kwadrat charakterystyki tłumieniowej typu $(n, 1)$:

$$w^2(\omega^2) = \left[1 - \frac{M(\omega^2)}{B_0} \right]. \quad (2.10)$$

Ponieważ $M(\omega^2)$ oscyluje n -krotnie między wartościami -1 oraz 1 dla $-1 \leq \omega \leq 1$, więc funkcja (2.10) przyjmuje $2n + 1$ krotnie na przemian

swą maksymalną i minimalną wartość $1 + 1/B_0$ oraz $1 - 1/B_0$. Z twierdzenia 4 wynika, że funkcja (2.10) aproksymuje stałą $C = 1$ najlepiej spośród wszystkich funkcji typu

$$w^2(\omega^2) = \frac{a_0 \omega^{2(n-1)} + a_1 \omega^{2(n-2)} + \dots + a_{2(n-1)}}{(\omega^2 + \sigma_1^2)(\omega^2 + \sigma_2^2) \dots (\omega^2 + \sigma_n^2)}.$$

Pierwiastek kwadratowy z (2.10) reprezentuje zatem poszukiwaną rodzinę $W(\omega^2)$ rozwiązań typu $(n, 1)$ zagadnienia filtracji. Minimalna wartość funkcji (2.10) w przedziale $[0, 1]$ wynosi $1 - 1/B_0 = (B_0 - 1)/B_0$. Wobec tego kwadrat znormalizowanej postaci znalezionej rozwiązania $W(\omega^2)$ wyraża się wzorem

$$W_0^2(\omega^2) = \frac{B_0}{B_0 - 1} \left(1 - \frac{M(\omega^2)}{B_0} \right). \quad (2.11)$$

Przy pomocy wyrażenia (2.11) i definicji 4, rozdział 1,4, wyznaczamy falistość znalezionej rozwiązania $W(\omega^2)$

$$\delta = \max_{-1 \leq \omega \leq 1} W_0^2(\omega^2) - 1 = \frac{2}{B_0 - 1}. \quad (2.12)$$

Współczynnik przy najwyższej potędze w funkcji (2.11) wyraża się jako

$$A_0^{(0)} = \frac{B_0}{B_0 - 1} \left(\sigma_1^2 + \sigma_2^2 + \dots + \sigma_n^2 - \frac{B_2}{B_0} \right), \quad (2.13)$$

wobec czego kwadrat wyjściowej postaci rozwiązania $W(\omega^2)$ obliczamy z zależności

$$W_1^2(\omega^2) = \frac{W_0^2(\omega^2)}{A_0^{(0)}} = \frac{1}{\sigma_1^2 + \sigma_2^2 + \dots + \sigma_n^2 - \frac{B_2}{B_0}} \left(1 - \frac{M(\omega^2)}{B_0} \right) \quad (2.14)$$

Znając związek między postacią wyjściową i znormalizowaną rozwiązaniem $W(\omega^2)$ obliczamy jego zysk określony w definicji 6, rozdział 1.4

$$K = \max_{-1 \leq \omega \leq 1} W_1(\omega) = \max_{-1 \leq \omega \leq 1} \frac{W_0(\omega)}{A_0^{(0)}} = \sqrt{\frac{1 + \sigma}{A_0^{(0)}}}. \quad (2.15)$$

Jeżeli chodzi o bieguny σ_i , to wybieramy je stosując rozmaite kryteria. Można zażądać np. minimalnej falistości δ dla danego typu (n, N) charakterystyki $W(\omega, \sigma_i)$ lub maksymalnego zysku. Tutaj podamy metodę doboru biegunów ułatwiającą obliczenie parametrów rozwiązania $W(\omega^2)$. Polega ona na założeniu

$$\sigma = \text{sh } \nu \gamma \quad (\nu = 1, 2, \dots, n), \quad (2.16)$$

gdzie γ jest stałą dodatnią. Ponieważ $\text{sh}^2 x + 1 = \text{ch}^2 x$, wobec tego

$$\sqrt{\sigma_\nu^2 + 1} = \text{ch } \nu \gamma. \quad (2.17)$$

Wstawiając (2.16) i (2.17) do (2.6) otrzymamy wzory (2.7), (2.8) i (2.9) w postaci

$$B_0 = \operatorname{ch} n(n+1)\gamma, \quad (2.18)$$

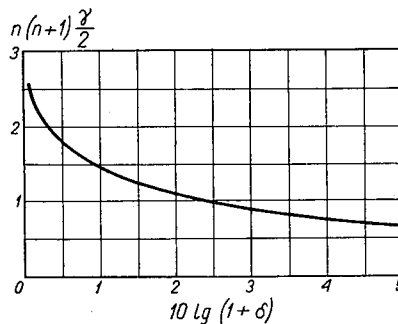
$$B_2 = -\frac{n}{2} \operatorname{ch} n(n+1)\gamma + \frac{1}{2} \{ \operatorname{ch}[n(n+1)-2]\gamma + \operatorname{ch}[n(n+1)-4]\gamma + \dots \\ + \operatorname{ch}[n(n-1)\gamma] \}, \quad (2.19)$$

$$B_{2n} = (-1)^n \prod_{\nu=1}^n \operatorname{sh}^2 \nu \gamma. \quad (2.20)$$

Najważniejszy jest dla nas jednak wzór wiążący klasę n charakterystyki typu $(n,1)$ z jej falistością. Podstawiając mianowicie (2.18) do (2.12) dostajemy

$$\delta = \frac{1}{\operatorname{th}^2 \left[n(n+1) \frac{\gamma}{2} \right]} - 1. \quad (2.21)$$

Wykres na rys. 13 przedstawia wynikającą z powyższego równania zależność między iloczynem $n(n+1)\gamma/2$ oraz wartością $1 + \delta$ w decybelach.



Rys. 13. Związek między falistością i klasą czhebyszewowskiego filtra RC

Przy jego pomocy możemy wyznaczyć klasę n lub stałą γ , jeżeli dana jest falistość δ . Należy przy tym zwrócić uwagę na następującą sprawę. Otóż z rys. 13 wynika, że dowolnie małą falistość δ można osiągnąć już przy pomocy charakterystyki typu (1,1) (jednobiegunowej) obierając w tym celu dostatecznie dużą wartość γ . Z (2.16) wynika, że bieguny σ_ν znajdują się wówczas daleko od punktu $\omega = 0$, co z kolei spowoduje wolny spadek charakterystyki bezpośrednio za pasmem przepustowym, mimo że w okolicy punktu $\omega = \infty$ maleć ona będzie jak $1/\omega$. Dlatego w konkretnych wypadkach należy zbadać szybkość spadku dla częstotliwości bliskich 1 i ewentualnie zwiększyć klasę, gdy szybkość ta jest zbyt mała.

Na zakończenie tego rozdziału podamy pewne uogólnienie rozpatrzonej

funkcji $M(\omega^2)$. Otrzymamy je zakładając

$$M(\omega^2) = \cos(m\Psi + \xi_1 + \xi_2 + \dots + \xi_n),$$

gdzie kąty Ψ i ξ_v wynikają ze wzorów (2.2) i (2.5). Stosując wzór na kosinus kątów stwierdzamy, że $M(\omega^2)$ jest funkcją wymierną o zadanych $2n$ biegunach $\pm i\sigma_v$ i o liczniku stopnia wyższego o $2m$ niż stopień mianownika

$$M(\omega^2) = \frac{B_0 \omega^{2(m+n)} + B_2 \omega^{2(m+n-1)} + \dots + B_{2(m+n)}}{(\omega^2 + \sigma_1^2)(\omega^2 + \sigma_2^2) \dots (\omega^2 + \sigma_n^2)}.$$

Gdy ω rośnie od -1 do 1 , to funkcja $M(\omega^2)$ oscyluje $(m+n)$ -krotnie między liczbami 1 i -1 . Zapisując funkcję $\tilde{M}(\omega^2) = M(\omega^2)/B_0$ w postaci różnicy

$$\tilde{M}(\omega^2) = f(\omega) - Q(\omega) = \frac{\omega^{2(m+n)}}{(\omega^2 + \sigma_1^2) \dots (\omega^2 + \sigma_n^2)} - \frac{B_2 \omega^{2(m+n-2)} + \dots + \tilde{B}_{2(m+n)}}{(\omega^2 + \sigma_1^2) \dots (\omega^2 + \sigma_n^2)}$$

stwierdzamy, że w przedziale $[-1, 1]$ przyjmuje ta różnica $2(m+n) + 1$ krotnie i na przemian swą maksymalną i minimalną wartość. Z twierdzenia 4 wynika zatem, że $Q(\omega)$ aproksymuje w sensie Czebyszewa funkcję $f(\omega)$, $M(\omega^2)$, przybliża zatem stałą $C = 0$ najlepiej spośród wszystkich funkcji wymiernych o $2n$ zadanych biegunach $\pm i\sigma_v$, o stopniu licznika nie większym od $2(m+n)$ i o współczynnikach przy najwyższej potęgze ω w liczniku i mianowniku równych jedności.

Przy pomocy funkcji $\tilde{M}(\omega^2)$ łatwo stworzyć czebyszewowską charakterystykę tłumieniową $W(\omega^2)$ typu $(m+n, n)$ o $2n$ zadanych pierwiastkach $\pm i\sigma_v$, a mianowicie

$$W(\omega^2) = \frac{1}{1 + \tilde{M}(\omega^2)}.$$

Z uwagi na zespolone bieguny tej charakterystyki nie da się ona zrealizować przy pomocy czwórnika RC, ale można ją wykorzystać w innych dziedzinach syntezy, np. do projektowania filtrów LC metodą Darlingtona.

2.4. Czwórniki RC o czebyszewowskim opóźnieniu

Niech będzie dany czwórnik RC o funkcji przenoszenia

$$h(s) = \frac{1}{(s - \sigma_1)(s - \sigma_2) \dots (s - \sigma_n)}, \quad (2.22)$$

gdzie σ_v są liczbami rzeczywistymi, dodatnimi. Przesunięciem fazowym $\beta(\omega)$ tego czwórnika nazwiemy argument funkcji $h(s)$ na osi rzeczywistych częstotliwości. Ponieważ argument liczby zespolonej równy jest części urojonej jej logarytmu, więc możemy zapisać:

$$\ln h(s) = \alpha(\omega) + i\beta(\omega).$$

Opóźnienie D czwórnika jest, jak wiadomo, pochodną jego przesunięcia $\beta(\omega)$. Ze wzoru (2.22) wynika, że argument $\beta(\omega)$ jest nieparzystą funkcją

częstotliwości, jego pochodna jest więc parzysta. W związku z tym zapisujemy

$$D = D(\omega^2) = \frac{d\beta(\omega)}{d\omega}.$$

Oznaczając przez $\vartheta(s)$ pochodną logarytmiczną funkcji (2.22) otrzymamy

$$\vartheta(s) \Big|_{s=i\omega} = \frac{d}{ds} \ln h(s) \Big|_{s=i\omega} = -i \frac{d\alpha(\omega)}{d\omega} + \frac{d\beta(\omega)}{d\omega}, \quad (2.23)$$

a stąd wzór na opóźnienie

$$D(\omega^2) = \operatorname{Rel} \vartheta(s) \Big|_{s=i\omega} = \frac{1}{2} [\vartheta(s) - \vartheta(-s)] \Big|_{s=i\omega}$$

wynikający z faktu, że na urojonej osi s funkcje $\vartheta(s)$ i $\vartheta(-s)$ są zespolone, sprzężone. Ze wzorów (2.22) i (2.23) wynika, że

$$\vartheta(s) = \frac{d}{ds} \ln h(s) = -\frac{1}{s - \sigma_1} - \frac{1}{s - \sigma_2} - \dots - \frac{1}{s - \sigma_n},$$

i wobec tego

$$\begin{aligned} D(\omega^2) &= \frac{\sigma_1}{\omega^2 + \sigma_1^2} + \frac{\sigma_2}{\omega^2 + \sigma_2^2} + \dots + \frac{\sigma_n}{\omega^2 + \sigma_n^2} = \\ &= \frac{d_0 \omega^{2(n-1)} + d_1 \omega^{2(n-2)} + \dots + d_{n-1}}{(\omega^2 + \sigma_1^2)(\omega^2 + \sigma_2^2) \dots (\omega^2 + \sigma_n^2)}. \end{aligned}$$

Widzimy, że opóźnienie czwórnika RC o funkcji przenoszenia określonej wzorem (2.22) wówczas zbliża się najbardziej do zadanej stałej Δ , gdy zależne od σ_v współczynniki d_μ równają się współczynnikom B_μ czebyszewowskiej charakterystyki tłumieniowej $W(\omega^2)$ o biegunach σ_v , aproksymującej stałą Δ . Współczynniki te potrafimy obliczyć przy pomocy wzorów (2.7), (2.8), (2.9). Porównując d_μ i B_μ oraz zadając falistość δ otrzymamy n równań na n biegunów σ_v . Po rozwiązaniu tych równań (dość skomplikowanych, nieliniowych względem σ_v) możemy zapisać funkcję przenoszenia (2.22) i przy jej pomocy skonstruować czwórnik RC, którego opóźnienie przybliży zadaną stałą Δ z zadaną falistością δ .

2.5. Przykłady projektowania dolnoprzepustowych filtrów RC typu (n, 1)

Przykład 1. Zaprojektować dolnoprzepustowy filtr typu (2.1) o zniekształceniach $10 \log(1 + \delta)$ nie większych niż 3 dB.

Z warunków projektu wynika, że falistość δ filtru powinna równać się 1. Istotnie

$$10 \log(1 + \delta) = 10 \log 2 = 10 \cdot 0,3010300 \cong 3 \text{ dB}.$$

Ponieważ parametr odcięcia $N = 1$, stosujemy więc wzory podane w rozdziale 2.3. Z (2.21) obliczamy pomocniczy parametr γ

$$\gamma = \frac{2}{n(n+1)} \operatorname{arth} \sqrt{\frac{1}{1+\delta}} = \frac{2}{2(2+1)} \operatorname{arth} \sqrt{\frac{1}{1+1}} = 0,29376.$$

Przy pomocy zależności (2.16) i tablic funkcji hiperbolicznych wyznaczamy bieguny funkcji przenoszenia

$$\sigma_1 = \operatorname{sh} \gamma = \operatorname{sh} 0,29376 = 0,29800,$$

$$\sigma_2 = \operatorname{sh} 2\gamma = \operatorname{sh} 2 \cdot 0,29376 = 0,62190.$$

Ze wzorów (2.18), (2.19) i (2.20) obliczamy współczynniki B_μ pomocniczej funkcji $M(\omega^2)$. Dostajemy

$$B_0 = \operatorname{ch} n(n+1)\gamma = \operatorname{ch} 2(2+1)0,29376 = 2,9994,$$

$$B_2 = -\frac{n}{2} \operatorname{ch} n(n+1) + \frac{1}{2} \{ \operatorname{ch} [n(n+1) - 2]\gamma + \operatorname{ch} n(n+1)\gamma \} = -1,5239,$$

$$B_4 = (-1)^n \operatorname{sh}^2 \gamma \operatorname{sh}^2 2\gamma = 0,034345.$$

Stąd

$$M(\omega^2) = \frac{B_0 \omega^4 + B_2 \omega^2 + B_4}{(\omega^2 + \sigma_1^2)(\omega^2 + \sigma_2^2)} = \frac{2,9994 \omega^4 - 1,5239 \omega^2 + 0,034345}{(\omega^2 + 0,088803)(\omega^2 + 0,38676)}.$$

Wzór (2.11) daje nam kwadrat czebyszewowskiej charakterystyki tłumieniowej w postaci znormalizowanej

$$\begin{aligned} W_0^2(\omega^2) &= \frac{B_0}{B_0 - 1} \left[1 - \frac{M(\omega^2)}{B_0} \right] = \frac{B_0}{B_0 - 1} \frac{\left(\omega^2(\sigma_1^2 + \sigma_2^2 - \frac{B_4}{B_0}) + \left(\sigma_1 \sigma_2 - \frac{B_4}{B_0} \right) \right)}{(\omega^2 + \sigma_1^2)(\omega^2 + \sigma_2^2)} = \\ &= 1,4756 \frac{\omega^2 + 0,0232761}{(\omega^2 + 0,088803)(\omega^2 + 0,38676)}. \end{aligned}$$

Ponieważ współczynnik $A_0^{(0)} = 1,4756$ przy najwyższej potędze w liczniku funkcji $W(\omega^2)$ został wyciągnięty przed ułamek, więc ułamek ten jest postacią wyjściową charakterystyki tłumieniowej projektowanego filtra. Zysk tej charakterystyki obliczamy według wzoru (2.15)

$$K = \sqrt{\frac{1+\delta}{A_0^{(0)}}} = \sqrt{\frac{1+1}{1,4756}} = 1,1642.$$

W dalszym ciągu należy dokonać syntezy czwórnika RC o charakterystyce

$$W_1(\omega^2) = \sqrt{\frac{\omega^2 + 0,023261}{(\omega^2 + 0,088803)(\omega^2 + 0,38676)}}.$$

Dokonaliśmy już tego metodą Darlingtona-Guillemina w rozdziale 1.3, przykład 1. W procesie syntezy charakterystyka wyjściowa uległa prze-

mnożeniu przez czynnik 0,53785. Tak więc rys. 4 przedstawia filtr RC o charakterystyce typu (2.1), o falistości $\delta = 1$ i zysku wynikowym

$$K_w = 0,53785 \cdot 1,1642 = 0,62616.$$

Przykład 2. Zaprojektować filtr dolnoprzepustowy RC typu (3.1) o zniekształceniach tłumieniowych $10 \log(1 + \delta)$ nie większych niż 1 dB.

Analogicznie, jak w poprzednim przykładzie znajdujemy falistość δ oraz parametr γ . Mamy

$$\log(1 + \delta) = 0,1 \quad \delta = 10^{0,1} - 1 = 0,25893$$

$$\gamma = \frac{2}{n(n+1)} \operatorname{arth} \sqrt{\frac{1}{1+\delta}} = \frac{2}{3(3+1)} \operatorname{arth} \sqrt{\frac{1}{1+0,25893}} = 0,23800.$$

Jak poprzednio wyznaczamy bieguny funkcji przenoszenia

$$\sigma_1 = \operatorname{sh} \gamma = 0,24025,$$

$$\sigma_2 = \operatorname{sh} 2\gamma = 0,49417,$$

$$\sigma_3 = \operatorname{sh} 3\gamma = 0,77622.$$

Przy pomocy wzorów (2.18), (2.19) i (2.20) obliczamy współczynniki pomocniczej funkcji $M(\omega^2)$:

$$B_0 = 8,7243, \quad B_2 = -7,5443, \quad B_6 = -0,0084926.$$

Natomiast wzór na B_4 wyprowadzamy wykorzystując zależność (2.6) i otrzymując

$$B_4 = (2\sigma_1^2 + 1)\sigma_2^2\sigma_3^2 + (2\sigma_2 + 1)\sigma_1^2\sigma_3^2 + (2\sigma_3^2 + 1)\sigma_1^2\sigma_2^2 + 4\sigma_1^2\sigma_2\sigma_3\mathcal{I}_2\mathcal{I}_3 + \\ + 4\sigma_2^2\sigma_1\sigma_3\mathcal{I}_1\mathcal{I}_3 + 4\sigma_3^2\sigma_1\sigma_2\mathcal{I}_1\mathcal{I}_2,$$

gdzie $\tau_v = \sqrt{\sigma_v^2 + 1} = \operatorname{ch} v\gamma$. Po wstawieniu do tego wzoru wartości na σ_v dostaniemy

$$B_4 = 0,93742.$$

Znając parametry B_v funkcji pomocniczej $M(\omega^2)$ obliczamy wg wzoru (2.11) znormalizowaną postać czebyszewowskiej charakterystyki tłumieniowej

$$W_0^2(\omega^2) = \frac{B_0}{B_0 - 1} \left[1 - \frac{1}{B_0} M(\omega^2) \right] = \\ = \frac{B_0}{B_0 - 1} \left\{ \omega^4 \left(\sigma_1^2 + \sigma_2^2 + \sigma_3^2 - \frac{B_2}{B_0} \right) + \omega^2 \left(\sigma_1^2\sigma_2^2 + \sigma_1^2\sigma_3^2 + \sigma_2^2\sigma_3^2 - \frac{B_4}{B_0} \right) + \right. \\ \left. + \left(\sigma_1^2\sigma_2\sigma_3^2 - \frac{B_6}{B_0} \right) : [(\omega^2 + \sigma_1^2)(\omega^2 + \sigma_2^2)(\omega^2 + \sigma_3^2)] \right\} = \\ = 1,9982 \frac{\omega^4 + \omega^2 0,050056 + 0,0053506}{(\omega^2 + 0,057720)(\omega^2 + 0,24421)(\omega^2 + 0,60251)}.$$

Jak w poprzednim projekcie wyciągnęliśmy współczynnik $A_0^{(0)} = 1,9982$ przy potędze ω^4 przed ułamek, dzięki czemu ułamek ten przedstawia

kwadrat wyjściowej postaci charakterystyki projektowanego filtru. Czwórnik RC realizujący tę charakterystykę przemnożoną przez 0,59856 został obliczony w przykładzie 2, rozdział 1.3. W związku z tym układ na rys. 8 jest filtrem RC typu (3.1) o falistości 0,25893. Zysk K_w zrealizowany przez ten filtr obliczamy mnożąc zysk charakterystyki $W(\omega^2)$ przez czynnik wprowadzony w procesie syntezy. Zgodnie ze wzorem (2.15) mamy

$$K_w = 0,59856 \sqrt{\frac{1+\delta}{A_0^{(0)}}} = 0,59856 \sqrt{\frac{1+0,25893}{1,9982}} = \\ = 0,59856 \cdot 0,79374 = 0,47510.$$

Na rysunku 10 przedstawiono tłumieniową charakterystykę filtru w jej postaci znormalizowanej, wyjściowej i zrealizowanej przez układ na rys. 8. Otrzymane krzywe pozwalają zorientować się, czy spadek charakterystyki bezpośrednio za pasmem przepustowym jest dostatecznie szybki.

3. OPTYMALNE CHARAKTERYSTYKI TŁUMIENIOWE

3.1. Definicja optymalnej charakterystyki tłumieniowej

W podrozdziale (2.3.) była mowa o tym, że typ charakterystyki czebyszewowskiej (n, N) nie określa jednoznacznie tej charakterystyki. W szczególności stwierdziliśmy, że dowolnie małą falistość δ można uzyskać już w pierwszej klasie charakterystyk czebyszewowskich, odsuwając w tym celu ich biegun dostatecznie daleko od początku układu współrzędnych s . Jeżeli stosujemy filtry klasy n wyższej niż 1, to głównie w celu zapewnienia dostatecznie szybkiego spadku charakterystyki bezpośrednio za pasmem przepustowym. Szybkość ta mogłaby więc, obok typu (n, N) i falistości δ , być czynnikiem określającym z dokładnością do stałego mnożnika pewną pożądaną charakterystykę tłumieniową filtru RC. Jednak wobec luk w teorii aproksymacji obliczenie takiej charakterystyki natrafiłoby na poważne trudności. Dlatego w niniejszym rozdziale oberzemy zysk K za czynnik definiujący pożądaną charakterystykę i pokażemy na paru przykładach, że charakterystyka cechująca się maksymalnym zyskiem jest określona z dokładnością do stałego mnożnika.

Aby móc wprowadzić kryterium zysku, należy przede wszystkim zdefiniować zysk i falistość dla ogólniejszych rodzin funkcji typu (n, N) , niż to miało miejsce w podrozdziale 1.4, gdzie skoncentrowaliśmy uwagę na czebyszewowskich rozwiązaniach zagadnienia filtracji. Istotnie, nie jest oczywiste, że charakterystyka typu (n, N) o danych zniekształceniach i maksymalnym zysku musi być charakterystyką czebyszewowską. Dlatego rozszerzymy rodzinę rozpatrywanych funkcji obejmując nią pierwiastki kwadratowe ze wszystkich funkcji wymiernych od ω^2 , dodatnich na dodatniej osi ω^2 . Charakterystyki te nazwiemy *niezerowymi*.

Łatwo spostrzec, że każda niezerowa charakterystyka $w(\omega^2)$ typu (n, N) , gdzie $N > 0$, może być wykorzystana do zaprojektowania filtru dolno-przepustowego o pasmie przepustowym $[0, 1]$. Istotnie, ponieważ stopień mianownika tej charakterystyki jest o N większy od stopnia licznika, więc od pewnej częstotliwości odcięcia ω_0 począwszy będzie ona monotonicznie maleć. Wystarczy zatem dokonać podstawienia

$$\omega = \tilde{\omega} \omega_0, \quad (3.1)$$

aby uzyskać charakterystykę $w(\tilde{\omega}^2)$ opadającą monotonicznie dla $\tilde{\omega} > 1$. Po dokonaniu podstawienia (3.1) nic nie przeszkadza w przemnożeniu charakterystyki $w(\omega^2)$ przez taką stałą, aby jej minimum w przedziale $0 \leq \omega^2 \leq 1$ równało się 1. Tak przemnożoną funkcję $w(\omega^2)$ oznaczymy przez $w_0(\tilde{\omega}^2)$ i nazwiemy charakterystyką *znormalizowaną*, zachowując tym sposobem w mocy definicję 3 rozdziału 1.4. Dzięki temu nie ulegają również zmianie określenia: *falistości* δ jako

$$\delta = \max_{0 \leq \tilde{\omega}^2 \leq 1} w_0^2(\tilde{\omega}^2) - 1$$

(por. definicja 4), postaci wyjściowej charakterystyki $w(\tilde{\omega}^2)$ jako

$$w_1(\tilde{\omega}^2) = \frac{w_0(\tilde{\omega}^2)}{A_0^{(0)}}$$

(por. definicja 5), gdzie $A_0^{(0)}$ jest ilorazem współczynników przy najwyższej potędze ω w liczniku i mianowniku funkcji $w_0(\tilde{\omega}^2)$ oraz zysku jako

$$K = \max_{0 \leq \tilde{\omega}^2 \leq 1} w_1(\tilde{\omega}^2).$$

Podstawienie (3.1) oraz przemnożenie charakterystyki $w(\tilde{\omega}^2)$ przez czynnik sprowadzający ją do postaci znormalizowanej $w_0(\tilde{\omega}^2)$ nazywać będziemy *dopuszczalną normalizacją* charakterystyki $w(\omega^2)$ i oznaczać będziemy przez r_w . Dla każdej niezerowej charakterystyki $w(\omega^2)$ istnieje oczywiście nieskończenie wiele takich normalizacji. Po zdefiniowaniu dopuszczalnej normalizacji określamy następująco optymalną charakterystykę tłumieniową.

Definicja 7. *Optymalną charakterystyką tłumieniową typu (n, N) o falistości δ nazywamy taką niezerową charakterystykę $w(\omega^2)$, która poddana pewnej normalizacji dopuszczalnej daje falistość nie większą niż δ i maksymalny zysk.*

3.2. Analiza charakterystyk czebyszewowskich z punktu widzenia maksymalnego zysku

Poszukajmy optymalnej charakterystyki tłumieniowej wśród charakterystyk $w(\omega^2)$ typu (n, n) o falistości δ , czyli wśród charakterystyk wielomianowych, których maksimum w ich postaci znormalizowanej jest dla

$0 \leq \omega^2 \leq 1$ nie większe niż $\sqrt{1+\delta}$. W rozdziale 2.2 rozpatrzyliśmy czebyszewowską postać takiej charakterystyki.

$$W(\omega^2) = \frac{1}{\tau + T_{2n}(\omega)}; \quad \tau > 1. \quad (3.2)$$

Liczba τ jest tu jednoznacznie określona falistością δ . Istotnie, ponieważ wielomian $T_{2n}(\omega)$ różni się w przedziale $-1 \leq \omega \leq 1$ od zera co najwyżej o ± 1 , więc stosując wzór (1.11) dostajemy

$$\delta = \frac{\max W^2(\omega^2) - \min W^2(\omega^2)}{\min W^2(\omega^2)} \bigg|_{0 \leq \omega^2 \leq 1} = \frac{\frac{1}{\tau-1} - \frac{1}{\tau+1}}{\frac{1}{\tau+1}} = \frac{2}{\tau-1}. \quad (3.3)$$

Nietrudno pokazać, że wśród wszystkich charakterystyk wielomianowych klasy nie wyższej niż n o falistości (3.3) optymalną jest charakterystyka (3.2). Aby się o tym przekonać, zapiszmy czebyszewowską charakterystykę wielomianową w postaci wyjściowej

$$W_1(\omega^2) = \sqrt{\frac{1}{\frac{\tau}{2^{2n-1}} + \tilde{T}_{2n}(\omega)}} = \sqrt{\frac{1}{\vartheta + \tilde{T}_{2n}(\omega)}} \quad (3.4)$$

i załóżmy, że istnieje różna od niej charakterystyka wielomianowa $w(\omega^2)$ o falistości nie większej i nie mniejszym zysku. W postaci wyjściowej charakterystykę tę można zapisać jako

$$w_1(\omega^2) = \sqrt{\frac{1}{\vartheta + t(\omega)}},$$

gdzie $t(\omega)$ wielomian stopnia co najwyżej $2n$, o współczynniku przy najwyższej potędze równym 1. Z porównania zysków mamy

$$\max_{-1 \leq \omega \leq 1} W_1^2(\omega^2) = \frac{1}{\vartheta + \min_{-1 \leq \omega \leq 1} \tilde{T}_{2n}(\omega)} \leq \max_{-1 \leq \omega \leq 1} w^2(\omega^2) = \frac{1}{\vartheta + \min_{-1 \leq \omega \leq 1} t(\omega)},$$

stąd

$$\min_{-1 \leq \omega \leq 1} t(\omega) \geq \min_{-1 \leq \omega \leq 1} \tilde{T}_{2n}(\omega). \quad (3.5)$$

Z drugiej strony mamy jednak

$$\min_{-1 \leq \omega \leq 1} W_2^2(\omega^2) = \frac{1}{\vartheta + \max_{-1 \leq \omega \leq 1} \tilde{T}_{2n}(\omega)} \leq \min_{-1 \leq \omega \leq 1} w^2(\omega^2) = \frac{1}{\vartheta + \max_{-1 \leq \omega \leq 1} t(\omega)},$$

gdyż w przeciwnym razie falistość charakterystyki $w(\omega^2)$ byłaby większa, niż $W_1(\omega^2)$ [por. wzór (1.11)]. Zatem

$$\max_{-1 \leq \omega \leq 1} \tilde{T}_{2n}(\omega) \geq \max_{-1 \leq \omega \leq 1} t(\omega). \quad (3.6)$$

Jednoczesne zachodzenie nierówności (3.5) oraz (3.6) jest jednak niemożliwe, jeżeli wielomiany $t(\omega)$ i $\tilde{T}_{2n}(\omega)$ są różne i $\tilde{T}_{2n}(\omega)$ jest wielomianem Czebyszewa.

Przejdźmy teraz do charakterystyk tłumieniowych typu $(n, 1)$. Wyniki badania charakterystyki wielomianowej sugerują założenie, że optymalna charakterystyka typu $(n, 1)$ jest również charakterystyką czebyszewowską. Ale charakterystyk czebyszewowskich typu $(n, 1)$ o danej falistości δ jest nieskończenie wiele, gdyż istnieje nieskończenie wiele możliwości doboru biegunów $\pm i\sigma_v$, $v = 1, 2, \dots, n$. Poszukamy takiego zbioru liczb $\sigma_1, \sigma_2, \dots, \sigma_n$, który zapewni nam maksymalny zysk przy danej falistości δ i klasie n .

Wzór na zysk K został wyprowadzony w podrozdziale 2.3.

$$K = \sqrt{\frac{1+\delta}{A_0^{(0)}}}.$$

Wynika z tego, że przy zadanym δ K osiąga maksimum dla takiego zbioru liczb $\sigma_1, \sigma_2, \dots, \sigma_n$, dla którego współczynnik $A_0^{(0)} = A_0^{(0)}(\sigma_1, \sigma_2, \dots, \sigma_n)$ przy najwyższej potędze ω w liczniku charakterystyki $W(\omega^2)$ w jej znormalizowanej postaci osiąga minimum. To minimum funkcji $A_0^{(0)}(\sigma_1, \sigma_2, \dots, \sigma_n)$ nie jest w naszym wypadku bezwzględne, lecz warunkowe, gdyż między zmiennymi σ_v musi być zachowany dodatkowy związek $\delta = \delta(\sigma_1, \sigma_2, \dots, \sigma_n) = \text{const}$. Zagadnienia na minimum warunkowe rozwiązuje się zazwyczaj metodą mnożników Lagrange'a ([12], str. 146). Stosując w naszym wypadku tę metodę zapisujemy związek między zmiennymi σ_v

$$g_1(\sigma_1, \sigma_2, \dots, \sigma_n) = 0$$

w postaci

$$g_1(\sigma_1, \sigma_2, \dots, \sigma_n) = \delta(\sigma_1, \sigma_2, \dots, \sigma_n) - \text{const} = 0,$$

a następnie tworzymy funkcję pomocniczą

$$F(\sigma_1, \sigma_2, \dots, \sigma_n, \lambda_0, \lambda_1) = \lambda_0 A_0^{(0)}(\sigma_1, \sigma_2, \dots, \sigma_n) + \lambda_1 g_1(\sigma_1, \sigma_2, \dots, \sigma_n).$$

Aby $A_0^{(0)}$ osiągnęło minimum warunkowe, muszą być spełnione następujące równania:

$$\left. \begin{aligned} \frac{\partial F}{\partial \sigma_v} &= \lambda_1 \frac{\partial g_1}{\partial \sigma_v} + \lambda_0 \frac{\partial A_0^{(0)}}{\partial \sigma_v} = \lambda_1 \frac{\partial \delta}{\partial \sigma_v} + \lambda_0 \frac{\partial A_0^{(0)}}{\partial \sigma_v} = 0, \\ \frac{\partial F}{\partial \lambda_1} &= \delta(\sigma_1, \sigma_2, \dots, \sigma_n) - \text{const} = 0. \end{aligned} \right\} \quad (v = 1, 2, \dots, n), \quad (3.7)$$

Znając wymaganą wartość $\delta = \text{const}$ i zakładając $\lambda_0 = 1$ możemy z równań tych wyznaczyć wszystkie niewiadome $\sigma_1^{(0)}, \dots, \sigma_n^{(0)}$ oraz λ_1 . Następnie należy zbadać, czy $A_0^{(0)}(\sigma_1^{(0)}, \sigma_2^{(0)}, \dots, \sigma_n^{(0)})$ jest istotnie minimalną wartością funkcji $A_0^{(0)}(\sigma_1, \sigma_2, \dots, \sigma_n)$, a nie jej wartością w punkcie stacjonarnym innego typu (jak maksimum, punkt przegięcia, punkt hiperboliczny).

Podstawmy wyprowadzone w rozdziale 2.3 zależności

$$\delta = \frac{2}{B_0 - 1}, \quad (2.12)$$

$$A_0^{(0)} = \frac{B_0}{B_0 - 1} \left[\sigma_1^2 + \sigma_2^2 + \dots + \sigma_n^2 - \frac{B_2}{B_0} \right] \quad (2.13)$$

do równania (3.7). Otrzymamy w ten sposób

$$\lambda_1 \frac{\partial B_0}{\partial \sigma_\nu} + \left[2\sigma_\nu B_0 - \frac{\partial B_2}{\partial \sigma_\nu} \right] = 0, \quad (\nu = 1, 2, \dots, n). \quad (3.8)$$

Przy pomocy dalszych, wyprowadzonych w rozdziale 2.3 wzorów

$$B_0(\sigma, \sigma_2, \dots, \sigma_n) = \frac{1}{2} \left[\prod_{\nu=1}^n (\sigma_\nu + \sqrt{\sigma_\nu^2 + 1})^2 + \prod_{\nu=1}^n (\sigma_\nu - \sqrt{\sigma_\nu^2 + 1})^2 \right], \quad (2.7)$$

$$B_2(\sigma_1, \sigma_2, \dots, \sigma_n) = -\frac{1}{2} \left[\prod_{\mu=1}^n (\sigma_\mu + \sqrt{\sigma_\mu^2 + 1})^2 \sum_{\nu=1}^n \frac{\sigma_\nu}{\sigma_\nu + \sqrt{\sigma_\nu^2 + 1}} + \right. \\ \left. + \prod_{\mu=1}^n (\sigma_\mu - \sqrt{\sigma_\mu^2 + 1})^2 \sum_{\nu=1}^n \frac{\sigma_\nu}{\sigma_\nu - \sqrt{\sigma_\nu^2 + 1}} \right] \quad (2.8)$$

wyznaczamy

$$\frac{\partial B_0}{\partial \sigma_\mu} = \frac{\prod_{\nu=1}^n (\sigma_\nu + \sqrt{\sigma_\nu^2 + 1})^2}{\sqrt{\sigma_\mu^2 + 1}} - \frac{\prod_{\nu=1}^n (\sigma_\nu - \sqrt{\sigma_\nu^2 + 1})^2}{\sqrt{\sigma_\mu^2 + 1}},$$

$$\frac{\partial B_2}{\partial \sigma_\mu} = \frac{1}{2} \left\{ \frac{\prod_{\nu=1}^n (\sigma_\nu - \sqrt{\sigma_\nu^2 + 1})^2}{\sqrt{\sigma_\mu^2 + 1}} \left[(\sigma_\mu + \sqrt{\sigma_\mu^2 + 1})^2 + 2 \sum_{\nu=1}^n \frac{\sigma_\nu}{\sigma_\nu - \sqrt{\sigma_\nu^2 + 1}} \right] + \right. \\ \left. - \frac{\prod_{\nu=1}^n (\sigma_\nu + \sqrt{\sigma_\nu^2 + 1})^2}{\sqrt{\sigma_\mu^2 + 1}} \left[(\sigma_\mu - \sqrt{\sigma_\mu^2 + 1})^2 + 2 \sum_{\nu=1}^n \frac{\sigma_\nu}{\sigma_\nu + \sqrt{\sigma_\nu^2 + 1}} \right] \right\}.$$

Wstawiając powyższe pochodne do (3.8) dostajemy

$$\lambda_1 \left[\prod_{\nu=1}^n (\sigma_\nu + \sqrt{\sigma_\nu^2 + 1})^2 - \prod_{\nu=1}^n (\sigma_\nu - \sqrt{\sigma_\nu^2 + 1})^2 \right] + \\ + \sigma_\mu \sqrt{\sigma_\mu^2 + 1} \left[\prod_{\nu=1}^n (\sigma_\nu + \sqrt{\sigma_\nu^2 + 1})^2 + \prod_{\nu=1}^n (\sigma_\nu - \sqrt{\sigma_\nu^2 + 1})^2 \right] + \\ + \frac{1}{2} \prod_{\nu=1}^n (\sigma_\nu + \sqrt{\sigma_\nu^2 + 1})^2 \left[(\sigma_\mu - \sqrt{\sigma_\mu^2 + 1})^2 + 2 \sum_{\nu=1}^n \frac{\sigma_\nu}{\sigma_\nu + \sqrt{\sigma_\nu^2 + 1}} \right] +$$

$$-\frac{1}{2} \prod_{\nu=1}^n (\sigma_{\nu} - \sqrt{\sigma_{\nu}^2 + 1})^2 \left[(\sigma_{\mu} + \sqrt{\sigma_{\mu}^2 + 1})^2 + 2 \sum_{\nu=1}^n \frac{\sigma_{\nu}}{\sigma_{\nu} - \sqrt{\sigma_{\nu}^2 + 1}} \right] = 0, \\ (\mu = 1, 2 \dots n)$$

Odejmując i -te z powyższych równań od j -tego otrzymamy

$$(\sigma_i^2 - \sigma_j^2) \left[\prod_{\nu=1}^n (\sigma_{\nu} + \sqrt{\sigma_{\nu}^2 + 1})^2 - \prod_{\nu=1}^n (\sigma_{\nu} - \sqrt{\sigma_{\nu}^2 + 1}) \right] = 0. \quad (3.9)$$

Widzimy, że równanie to jest spełnione jedynie gdy $\sigma_i = \sigma_j$ dla każdego i oraz j . Punkt stacjonarny funkcji $A_0^{(0)}(\sigma_1, \sigma_2, \dots, \sigma_n)$ ma więc jednakowe współrzędne w układzie kartezjańskim $\sigma_1, \sigma_2, \dots, \sigma_n$. Wiedząc o tym nie-trudno dla danej falistości δ znaleźć związek $\sigma_1 = \sigma_2 = \dots = \sigma_n = \sigma(\delta, n)$. Otóż wstawiając $\sigma_1 = \sigma_2 = \dots = \sigma_n = \sigma$ do wzorów (2.7) i (2.8) dostajemy

$$B_0 = \frac{1}{2} [(\sigma + \sqrt{\sigma^2 + 1})^{2n} + (\sigma - \sqrt{\sigma^2 + 1})^{2n}], \quad (3.10)$$

$$B_2 = -\frac{n\sigma}{2} [(\sigma + \sqrt{\sigma^2 + 1})^{2n+1} + (\sigma - \sqrt{\sigma^2 + 1})^{2n+1}]. \quad (3.11)$$

Ponieważ zgodnie z (2.12) mamy

$$1 + \delta = \frac{B_0 + 1}{B_0 - 1}, \quad (3.12)$$

więc przy pomocy (3.10) możemy obliczyć stąd $\sigma = \sigma(\delta, n)$ dla każdego δ i n . Zależiony w ten sposób stacjonarny punkt funkcji $A_0^{(0)}(\sigma_1, \sigma_2, \dots, \sigma_n)$ jest punktem, w którym osiąga ona swe minimum, o czym można się przekonać obliczając $A_0^{(0)}$ dla innego zbioru liczb $\sigma_1 \dots \sigma_n$ spełniających związek $\delta(\sigma_1, \dots, \sigma_n) = \text{const}$. Stwierdziliśmy tym sposobem, że charakterystyka optymalna wśród chebyszewowskich charakterystyk typu $(n, 1)$, jest granicznym przypadkiem charakterystyk realizowalnych. Istotnie, ze względu na n -krotny biegun nie daje się ona zrealizować przy pomocy czwórnika RC, można się jednak do niej dowolnie zbliżyć obierając charakterystykę o biegunach różniących się odpowiednio mało od $\omega = \pm i\sigma(\delta, n)$. Dlatego warto obliczyć wartość maksymalnego zysku i zbadać jego zależność od n i δ . Otóż korzystając z tożsamości

$$\frac{-1}{\sigma + \sqrt{\sigma^2 + 1}} = \sigma - \sqrt{\sigma^2 + 1}$$

obliczamy przy pomocy (3.10)

$$B_0 - 1 = \frac{1}{2} [(\sigma + \sqrt{\sigma^2 + 1})^n - (-1)^n (\sigma - \sqrt{\sigma^2 + 1})^n]^2; \quad (3.13)$$

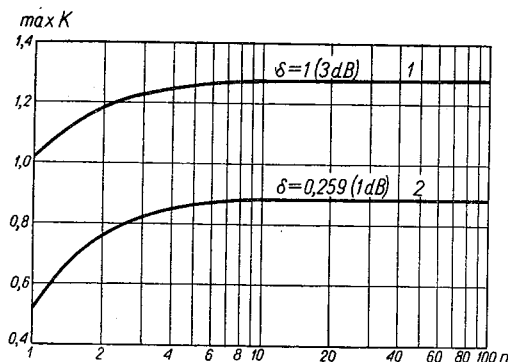
wstawiając w dalszym ciągu (3.10), (3.11) i (3.13) do (2.13) dostajemy

$$A_0^{(0)} = \frac{B_0}{B_0 - 1} \left[n\sigma - \frac{B_2}{B_0} \right] = n\sigma \sqrt{\sigma^2 + 1} \frac{(\sigma + \sqrt{\sigma^2 + 1})^n + (-1)^n (\sigma - \sqrt{\sigma^2 + 1})^n}{(\sigma + \sqrt{\sigma^2 + 1})^n - (-1)^n (\sigma - \sqrt{\sigma^2 + 1})^n}. \quad (3.14)$$

Ponieważ

$$K = \sqrt{\frac{1 + \delta}{A_0^{(0)}}}, \quad (2.15)$$

więc możemy stąd, po wyznaczeniu przy pomocy (3.12) wartości $\sigma = \sigma(\delta)$ znaleźć maksymalny zysk w każdej klasie n charakterystyk czhebyszewow-



Rys. 14. Zależność maksymalnego zysku filtru RC od jego klasy

skich typu $(n, 1)$ i dla każdej falistości δ . Rozpatrzmy tytułem przykładu przypadek zniekształceń tłumieniowych $\Delta = 10 \log(1 + \delta) = 3$ dB, czyli $\delta = 1$. Z (3.12) wynika, że wówczas $B_0 = 3$. Rozwiązując w dalszym ciągu (3.10) względem σ i podstawiając $B_0 = 3$ dostaniemy

$$\sigma = \frac{(1 + \sqrt{2})^{\frac{2}{n}} - 1}{2(1 + \sqrt{2})^{\frac{1}{n}}},$$

a stąd, przy pomocy (3.14)

$$A_0^{(0)} = \frac{n}{2\sqrt{2}} \frac{(1 + \sqrt{2})^{\frac{4}{n}} - 1}{(1 + \sqrt{2})^{\frac{2}{n}}}.$$

Zatem z (2.15)

$$\max K_{\delta=1} = \max K_1 = 2 \sqrt{\frac{\sqrt{2}}{n} \frac{(1 + \sqrt{2})^{\frac{2}{n}}}{(1 + \sqrt{2})^{\frac{4}{n}} - 1}}.$$

Powyższa funkcja została przedstawiona na rys. 14 w postaci krzywej 1. Funkcja ta wzrasta monotonicznie wraz z n dążąc do granicy

$$\lim_{n \rightarrow \infty} \max K_1 = \sqrt{\frac{\sqrt{2}}{\ln(1 + \sqrt{2})}} = 1,267.$$

Zakładając zniekształcenia $\Delta = 10 \log(1 + \delta) = 1$ dB, czyli $\delta = 0,259$,

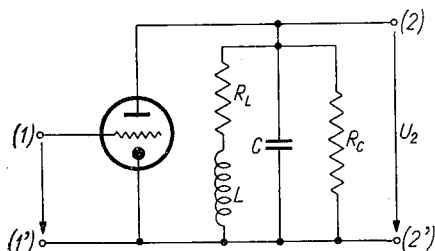
dostajemy w identyczny sposób

$$\max K_{0,259} = 2,12 \sqrt{\frac{1}{n} \frac{4,17^{\frac{2}{n}}}{4,17^{\frac{4}{n}} - 1}},$$

a stąd dolną krzywą na rys. 14. Widzimy zatem, że maksymalny zysk maleje wraz ze zniekształceniami tłumieniowymi. Ponadto stwierdzamy, że krzywa dla 1 dB zbliża się do swej asymptoty wolniej niż dla 3 dB. Jeżeli chodzi nam więc o zbliżenie się do maksymalnego zysku, to należy zakładać mniejszą falistość dla układów o większej liczbie biegunów, natomiast w układach prostszych falistość można przyjąć mniejszą.

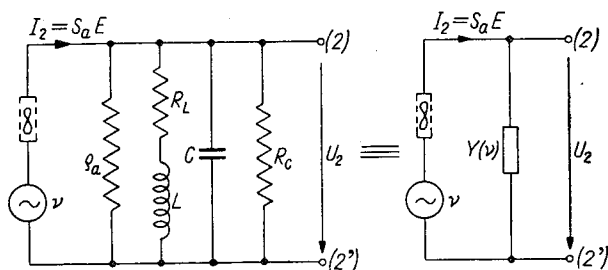
3.3. Porównanie charakterystyk tłumieniowych różnego typu filtrów i wzmacniaczo-filtrów

Rozpatrzmy stopień wzmacniacza rezonansowego na rys. 15 oraz jego schemat zastępczy na rys. 16.



Rys. 15. Stopień wzmacniaczo-filtru

Rys. 16. Schematy zastępcze stopnia wzmacniaczo-filtru



Na rysunkach tych ν oznacza częstotliwość źródła prądu, S_a i q_a są nachyleniem i opornością lampy, R_L , R_C — stratami w cewce i kondensatorze, $Y(\nu)$ — przewodnością, na jaką pracuje zastępcze źródło o wydajności $I = S_a E$. Okazuje się ([13], str. 71—79), że przy odpowiednio dużej oporności q_a i małych stratach w cewce i kondensatorze, przewodność $Y(\nu)$ można obliczyć z następującego wzoru ([13], str. 78, 1.4.4.III):

$$Y(\nu) = G \frac{R}{d} \left[\frac{d}{R} - i\omega_r + i\omega \right]. \quad (3.15)$$

We wzorze tym G oznacza całkowitą przewodność strat kondensatora:

$$G = \frac{1}{R_c} + \frac{1}{q_a},$$

d — całkowity współczynnik strat obwodu rezonansowego dla częstotliwości rezonansowej $\nu_r = 1/\sqrt{LC}$

$$d = \frac{R_L}{\nu_r L} + \frac{G}{\nu_r C},$$

ω — znormalizowaną częstotliwość bieżącą:

$$\omega = \frac{2(\nu - \nu_0)}{\nu_0 R}, \quad (3.16)$$

ω_r — znormalizowaną częstotliwość rezonansową

$$\omega_r = \frac{2(\nu_r - \nu_0)}{\nu_0 R},$$

R — pewien czynnik normalizujący, którego znaczenie wyjaśni się później, ν_0 — środek pasma przepustowego filtru środkowoprzepustowego, który można zbudować łącząc łańcuchowo czwórniki jak na rys. 15. Przewodność $Y(\nu)$ można obliczyć ze wzoru (3.15) jedynie wówczas, gdy częstotliwości ν oraz ν_0 nie są zbyt odległe od częstotliwości rezonansowej ν_r

$$\frac{2(\nu - \nu_0)}{\nu_0} < 0,1 \quad \frac{2(\nu_r - \nu_0)}{\nu_0} < 0,1. \quad (3.17)$$

Napięcie wyjściowe U_2 czwórnika na rys. 15 wynosi $I_2/Y(\nu) = ES_a/Y(\nu)$ wobec czego funkcja przenoszenia tego czwórnika

$$h = \frac{U_2}{E} = \frac{S_a}{G \frac{R}{d} \left[\frac{d}{R} - i\omega_r + i\omega \right]} \bigg|_{\omega = \omega(\nu)}$$

Funkcja h rozpatrywana w zależności od zmiennej $s = i\omega$ posiada biegun w punkcie $-\sigma_i = d/R - i\omega_r$. Jeżeli kilka wzmacniaczy połączymy w łańcuch, to funkcja przenoszenia $h(s)$ tego łańcucha będzie iloczynem funkcji przenoszenia poszczególnych wzmacniaczy, a więc wyrazi się wzorem

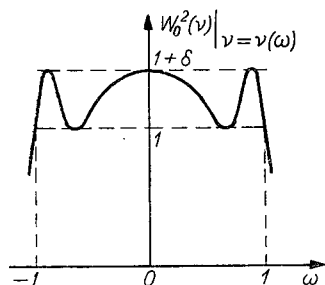
$$h(s) = h_1(s) \cdot h_2(s) \dots h_n(s) = \frac{C}{(s - \sigma_1)(s - \sigma_2) \dots (s - \sigma_n)}.$$

Widać stąd, że stratności d_r obwodów oraz ich częstotliwości rezonansowe ν_r można tak dobrać, aby charakterystyka tłumieniowa łańcucha rozpatrywana jako funkcja zmiennej ω

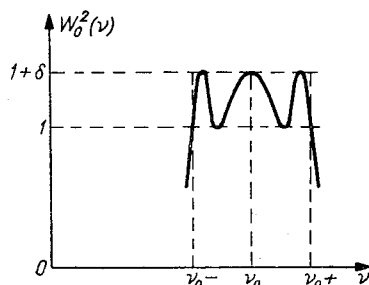
$$w^2(\omega^2) = h(s)h(-s) \bigg|_{s=i\omega} = \frac{C^2}{(\omega^2 + \sigma_1^2)(\omega^2 + \sigma_2^2) \dots (\omega^2 + \sigma_n^2)}$$

była czebyszewowską charakterystyką wielomianową omówioną w rozdziale 2.2, czyli aby częstotliwości $\pm i\sigma$, były pierwiastkami wielomianu $\tau + T_{2n}(\omega)$. Na rys. 17 naszkicowano przebieg funkcji $W(\omega^2)$ w postaci znormalizowanej dla $n = 3$.

Jeżeli chodzi o przebieg funkcji $W_0(\omega)$ na osi częstotliwości ν , to ze wzoru (3.16) wynika, że będzie on identyczny, jak na rys. 17 z tym, że środek pasma przepustowego przesunie się z punktu $\omega = 0$ do punktu $\nu = \nu_0$, a samo pasmo ulegnie $\frac{1}{2} R \nu_0$ krotnemu rozszerzeniu.



Rys. 17. Znormalizowana charakterystyka tłumieniowa chebyszewowskiego wzmacniacza-filtru



Rys. 18. Charakterystyka tłumieniowa chebyszewowskiego wzmacniacza-filtru

To wszystko, co powiedzieliśmy dotąd o układzie na rys. 15, jest więc szkicem teorii czynnych, wąskopasmowych filtrów RLC .

Istotnie, jeżeli chcemy zaprojektować filtr środkowoprzepustowy o środku pasma ν_0 i o jego szerokości $R\nu_0$, złożony z obwodów rezonansowych, to wystarczy w tym celu nastroić te obwody na takie częstotliwości ν_{ri} i dobrać dla nich takie stratności d_i , aby kwadraty liczb

$$i \sigma_i = -i \frac{d_i}{R} - \frac{2(\nu_{ri} - \nu_0)}{R \nu_0}$$

były pierwiastkami wielomianu $\tau + T_{2n}(\omega)$. Po połączeniu odseparowanych lampami, obliczonych w ten sposób rezonansowych obwodów w łańcuch otrzymamy czwórnik, którego charakterystyka tłumieniowa będzie miała chebyszewowski przebieg w przedziale $\left[\nu_0 + \frac{1}{2} R \nu_0, \nu_0 - \frac{1}{2} R \nu_0 \right]$.

Będzie to miało miejsce oczywiście pod warunkiem spełnienia nierówności (3.17), tj. pod warunkiem, że stosunek szerokości pasma do jego środkowej częstotliwości ν_0 nie będzie zbyt duży. Falistość charakterystyki $W(\nu^2)$ zależęć będzie od liczby obwodów i parametru τ .

Z tego, co powiedzieliśmy o wąskopasmowych filtrach RLC , wynika wielkie znaczenie chebyszewowskich charakterystyk wielomianowych w teorii filtrów. Istnieją jednak względy, dla których bieguny $\omega_p^2 = -\sigma_p^2$ charakterystyki $\omega(\nu)$ obiera się jako różne od pierwiastków wielomianu $\tau + T_{2n}(\omega)$. Rozpatrzmy dwie szczególne z tych charakterystyk i porównamy pod względem zysku z charakterystykami chebyszewowskimi tego samego stopnia.

Budując uproszczony środkowoprzepustowy filtr RLC można obrać jednakowe rezonanse i stratności obwodów. Zrealizujemy tym sposobem

synchroniczną charakterystykę tłumieniową. Na osi ω jej kwadrat przybiera następującą postać znormalizowaną

$$w_0^2(\omega^2) = A_0^{(0)} \frac{1}{(\omega^2 + \sigma^2)^n}. \quad (3.18)$$

Powyższą charakterystykę można również otrzymać sprzęgając lampy przy pomocy równoległych obwodów rezonansowych RC. Ponieważ prawa strona (3.18) przybiera w pasmie przepustowym $[-1, 1]$ minimalną wartość $w_0^2(\omega^2) = 1$ dla $\omega^2 = 1$, więc

$$A_0^{(0)} = (1 + \sigma^2)^n.$$

Swoje maksimum, $1 + \delta$, osiąga funkcja (3.18) dla $\omega = 0$, wobec czego z równości

$$1 + \delta = \frac{(1 + \sigma^2)^n}{\sigma^{2n}} \quad (3.19)$$

można wyznaczyć falistość δ , gdy dany jest n -krotny biegun σ lub na odwrót. Ponieważ

$$K = \sqrt{\frac{1 + \delta}{A_0^{(0)}}},$$

więc dla charakterystyk synchronicznych mamy

$$K = \frac{1}{\sigma^n}.$$

Rozpatrzmy specjalny przypadek $\delta = 1$, tzn. przypadek, gdy funkcja tłumieniowa filtru na krańcach pasma przepustowego opada o 3 dB. W teorii wzmacniaczy pasmo częstotliwości, w którym wzmocnienie nie jest mniejsze niż 3 dB od maksymalnego, nazywa się wstęgą wzmacniacza, a jego dobroć charakteryzuje się parametrem zysk $\times$ wstęga lub zysk $\times$ wstęga na stopień wzmocnienia. Dla $\delta = 1$ „wstęga” znormalizowanego filtru dolnoprzepustowego zawiera się między 0 a 1, a więc równa się 1. W tym przypadku zysk $K(\delta) = K(1)$ jest równocześnie parametrem zysk $\times$ wstęga. Rozwiązując (3.19) dla $\delta = 1$ dostajemy

$$\sigma^2 = \frac{1}{\sqrt[n]{2} - 1},$$

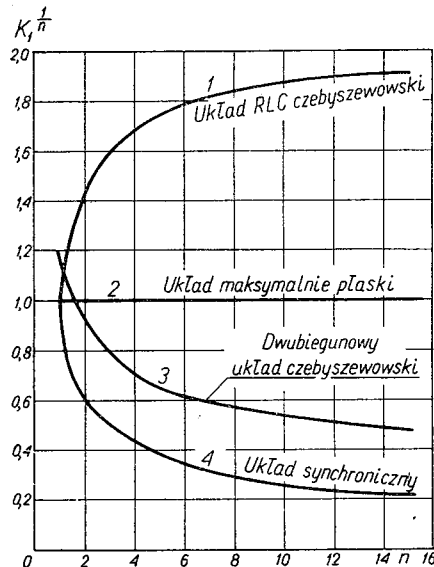
zatem

zysk = zysk $\times$ wstęga = $K(1) = (\sqrt[n]{2} - 1)^n$,
natomiast

$$\text{zysk} \times \text{wstęga na stopień} = \sqrt[n]{K(1)} = \sqrt[n]{2} - 1.$$

Przebieg tego ostatniego parametru w zależności od n przedstawiono na rys. 19, krzywa 4.

Jeżeli filtr środkowoprzepustowy ma przenosić sygnał sinusoidalny modulowany impulsami 0,1, to ważne są nie tyle zniekształcenia tłumieniowe filtru, co liniowa zmiana argumentu jego funkcji przenoszenia roz-



Rys. 19. Zysk jednego stopnia filtru w zależności od klasy n dla różnego typu charakterystyk tłumieniowych

patrywanego w zależności od częstotliwości ω lub ν . Okazuje się, że stosunkowo liniowy przebieg tego argumentu otrzymujemy w przypadku funkcji $h(\omega)$ odpowiadających tzw. *maksymalnie płaskiej* charakterystyce tłumieniowej. W swej postaci znormalizowanej charakterystyka ta wyraża się wzorem

$$w_0(\omega^2) = \sqrt{\frac{A_0^{(0)}}{\omega^{2n} + \sigma^{2n}}}.$$

Wartość stałej $A_0^{(0)}$ wyznaczamy spostrzegając, że charakterystyka znormalizowana osiąga swą wartość minimalną $w_0 = 1$ w pasmie przepustowym dla $\omega = 1$. Zatem

$$\frac{A_0^{(0)}}{1 + \sigma^{2n}} = 1; \quad A_0^{(0)} = 1 + \sigma^{2n}.$$

Ponieważ wartość maksymalną $w_0 = 1 + \delta$ charakterystyka ta przybiera dla $\omega = 0$, więc dostajemy następujący związek między parametrami δ , σ oraz n :

$$1 + \delta = \frac{1 + \sigma^{2n}}{\sigma^{2n}}. \quad (3.20)$$

Zysk maksymalnie płaskiej charakterystyki obliczamy stosując znany nam wzór

$$K = \sqrt[n]{\frac{1+\delta}{A_0^{(0)}}} = \frac{1}{\sigma^n}.$$

Stąd zysk na stopień wzmocnienia

$${}^n\sqrt{K} = \frac{1}{\sigma}.$$

Zakładając 3-decybelowe zniekształcenia tłumieniowe, czyli $\delta = 1$, otrzymamy z równania (3.20) $\sigma = 1$. Stąd wynika, że dla maksymalnie płaskiej charakterystyki tłumieniowej o falistości $\delta = 1$

$$\text{zysk} \times \text{wstęga na stopień} = \text{zysk na stopień} = {}^n\sqrt{K_1} = 1/\sigma_1 = 1.$$

Parametr ten nie zależy więc od liczby stopni, co zobrazowano na rys. 19 w postaci krzywej 2.

Aby dokonać oceny wyników otrzymanych dla dwu specjalnych charakterystyk wielomianowych, wyznaczmy zysk i zysk na stopień wielomianowej charakterystyki czebyszewowskiej. W postaci znormalizowanej charakterystykę tę zapisujemy jako

$$W_o(\omega) = \frac{\tau + 1}{\tau + T_{2n}(\omega)}, \quad \tau > 1,$$

gdzie $T_{2n}(\omega)$ jest znanym nam z rozdziału 2.2 wielomianem Czebyszewa oscylującym między wartościami -1 i 1 , gdy ω dąży od -1 do 1 . Ponieważ współczynnik przy najwyższej potędze ω wielomianu $T_{2n}(\omega)$ równa się 2^{2n-1} , więc parametr $A_0^{(0)}$ wyraża się wzorem

$$A_0^{(0)} = \frac{\tau + 1}{2^{2n-1}}. \quad (3.21)$$

Równanie wiążące falistość z parametrem τ otrzymujemy pamiętając, że minimalna wartość $T_{2n}(\omega)$ dla $-1 \leq \omega \leq 1$ wynosi -1 . Zatem

$$\max_{-1 \leq \omega \leq 1} W_o(\omega) = 1 + \delta = \frac{\tau + 1}{\tau - 1}.$$

Dla $\delta = 1$ otrzymujemy stąd $\tau = 3$, a więc po wstawieniu tej wartości do (3.21) obliczymy

$$K_1 = \sqrt[n]{\frac{1+\delta}{A_0^{(0)}}} = 2^{n-1}.$$

Tak więc zysk na stopień dla czebyszewowskiej charakterystyki o falistości $\delta = 1$, czyli zysk $\times$ wstęga na stopień wynosi

$${}^n\sqrt{K_1} = \frac{2}{{}^n\sqrt{2}} \xrightarrow{n \rightarrow \infty} 2.$$

Powyższą zależność przedstawiono na rys. 19 w postaci krzywej 1.

Omawiając wielomianową charakterystykę synchroniczną wspominaliśmy, że jedynie ona spośród trzech rozpatrzonych charakterystyk wielomianowych może być zrealizowana przy pomocy równoległych obwodów rezonansowych RC zastosowanych między lampami jako dwójniki sprzęgające. Każda inna wymaga zastosowania elementów RLC . Z rysunku 19 wynika jednak, że charakterystyka synchroniczna cechuje się mniejszym zyskiem niż pozostałe. Powstaje pytanie, czy nie można by zwiększyć tego zysku choćby kosztem skomplikowania dwójnika sprzęgającego. Otóż nie jest to możliwe, ponieważ charakterystyki o większym zysku muszą mieć zespolone bieguny, a pierwiastki przewodności RC leżą wszystkie na ujemnej półosi s . W celu zwiększenia zysku w czynnym filtrze RC można jednak zastosować jako sprzężenie czwórnik np. dwubiegunowe typu (2.1), takie jak zaprojektowany w przykładzie 1, rozdział 1.3. Funkcją przenoszenia łańcucha n lamp sprzężonych takimi czwórnikiem będzie iloczyn n czebyszewowskich charakterystyk typu (2.1), różniących się ewentualnie biegunami od obliczonej w przykładzie 1, rozdział 2.5. Jeżeli falistość łańcucha ma wynieść δ , to falistość każdego z n czwórników sprzęgających powinna równać się $\delta_n = \sqrt[n]{1+\delta} - 1$, aby spełniony był związek

$$(1+\delta) = (1+\delta_n)^n.$$

W szczególności dla $\delta = 1$ $\delta_n = 2^{1/n} - 1$. Załóżmy, że każdy z n czwórników sprzęgających zaprojektowany został na falistość $\delta_n = 2^{1/n} - 1$ i na maksymalny graniczny zysk K_n . Wzory w rozdziale (3.2) pozwalają wyznaczyć podwójny w tym przypadku biegun $\sigma = \sigma(\delta_n)$ charakterystyki tłumieniowej oraz zysk K_n dla takiego czwórnika. Ponieważ rozpatrywany przez nas filtr jest łańcuchowym połączeniem n tych czwórników odseparowanych elementami czynnymi, więc K_n jest parametrem zysk $\times$ wstęga na stopień tego filtru. Wykres funkcji K_n podany został na rys. 19 w postaci krzywej 3. Jak widać, przebiega ona dla każdego n powyżej analogicznej krzywej dla charakterystyki synchronicznej.

*Instytut Podstawowych Problemów Techniki PAN
Zakład Teorii Łączności*

WYKAZ LITERATURY

1. Fialkow A., Gerst J.: The Transfer Function of General Two Terminal-Pair Networks. — Quart. Appl. Math., tom 10, lipiec 1952.
2. Fialkow A.: Two Terminal-Pair Networks Containing Two Kinds of Elements Only. — Symposium of Modern Networks Synthesis, New York 1952.
3. Cauer W.: Theorie der linearen Wechselstromschaltungen. Berlin 1954.
4. Darlington S.: Synthesis of Reactance 4-Poles. — Journ. of Math. Phys., 1939.
5. Guillemin E. A.: Synthesis of Passive Networks. — Wiley, New York 1957.

6. Guillemin E. A.: A Summary of Modern Methods of Networks Synthesis Advances in Electronics, том III. New York 1951.
7. Storer J.: Passive Network Synthesis. Mac Graw-Hill, New York 1957.
8. Balabanian N.: Network Synthesis. Prentice Hall, New York 1958.
9. Natanson J. P.: Konstruktiwnaja teorija funkcij. Moskwa 1949.
10. Achiezer N. J.: Teoria aproksymacji. PWN, Warszawa 1957.
11. Golde W., Norek C., Paszkowski S.: Zarys teorii aproksymacji. PWN, Warszawa 1958.
12. Courant R., Hilbert D.: Metody matematycznej fizyki, том I, Moskwa 1951.
13. Lenkowski J.: Teoria wąskopasmowych filtrów wielkiej częstotliwości. PWN, Warszawa 1959.
14. Valley G. E., Wallman H.: Vacuum Tube Amplifiers. — Rad. Lab. Series, том 18, New York 1948.
15. Bosy J. N. D.: Elektrieskije filtry. Kiew 1955.
16. Ragan G. L.: Microwave Transmission Circuits. — Rad. Lab. Series, том 9, New York 1948.

П. ШУЛЬКИН, Ч. НОРЕК

ВОПРОСЫ АПРОКСИМАЦИИ В ТЕОРИИ НИСКОЧАСТОТНЫХ ФИЛЬТРОВ RC

Резюме

Главной темой работы является применение несколько эффективно решенных вопросов аппроксимации в смысле Чебышева при проектировании низкочастотных фильтров RC. В первой главе приведены элементы теории четырехполосников RC, обсуждены свойства их частотных характеристик и способы их реализации. Там же даны два численные примеры синтеза пассивных четырехполосников RC с предопределенными чебышевскими характеристиками затухания. После такого ознакомления читателей с характеристическими особенностями синтеза схем RC приведены дефиниции ряда параметров и понятий позволяющих в дальнейшем на точную формулировку аппроксимационных проблем и определение оптимальных характеристик фильтра RC.

Аппроксимационные проблемы совместно с их формулировкой и решением обсуждены во второй главе. Полученные результаты иллюстрированы двумя численными примерами при конструировании чебышевских характеристик затухания, которые в первой главе являлись исходными для примеров синтеза четырехполосников RC. Решения некоторых аппроксимационных проблем, рассмотренных во второй главе, можно использовать также во многих других случаях синтеза, например при конструировании схем автоматики при проектировании замедляющих схем, в теории узкополосных фильтров RC. и т. п.

Глава 3 имеет дискуссионный характер. Предложен в ней критерий — постулат максимального выигрыша — позволяющий на избрание оптимальной характеристики затухания из целого семейства этих характеристик. Вопрос избрания для семейства характеристик не имеющих конечных корней решен путем доказательства, что оптимальнейшей среди них является чебышевская характеристика. Определен также выигрыш для нескольких часто встречаемых в практике видов характеристик и оптимальные результаты сравнены между собой.

APPROXIMATION PROBLEMS IN THE THEORY OF LOW-PASS RC FILTERS

Summary

The main theme of the work is the application of several efficiently solved approximation problems (in the Tchebyshev sense) to designing of low-pass RC filters. In the first chapter a sketch of the theory of quadripoles has been given; properties of their frequency characteristic and conditions of their realization have been discussed. Two numerical examples of synthesis of passive RC quadripoles with prefixed Tchebyshev damping characteristics have also been given. Having thus acquainted the reader with the specificness of the synthesis of RC systems a number of parameters has been defined as well as notions enabling to formulate precisely the approximation problems and to define optimum characteristics of the RC filter.

The approximation problems as well as their terms and solutions have been discussed in chapter 2. The attained results have been illustrated by two numerical examples, constructing these Tchebyshev damping characteristics, which served as a starting point in the examples of the synthesis of RC quadripoles in the first chapter. Solutions of some approximation problems considered in chapter 2 may be also made use of in lot of other fields of synthesis e.g. when constructing automatic systems, when designing delay systems, in the theory of narrow-band RCL filters etc.

Chapter 3 is a discussible one. A criterion has been suggested there — a postulate of maximum gain — which enable to choose the optimum damping characteristic from among a whole group of those characteristics. The problem of the choice for the set of characteristics without roots in finiteness has been solved by proving that the Tchebyshev characteristic is the optimum one. The gain for some kinds of characteristics often met with in practice has been determined and the results have been compared.

ZBIGNIEW KACZKOWSKI

Równania magnetostrykcyjne i ich współczynniki

Rękopis dostarczono 4.4.1960

Podano podstawowe zależności magnetyczne i mechaniczne stosowane przy rozpatrywaniu zjawisk magnetomechanicznych zachodzących w materiałach magnetostrykcyjnych. Następnie — przy założeniu, że przebiegi są adiabatyczne (co w większości przypadków spotykanych w praktyce jest słuszne) — podano wszystkie możliwe układy równań magnetostrykcyjnych typu B i J . Równania typu B przedstawiono w układzie jednostek MKSA zracjonalizowanym (tablica 2) oraz elektromagnetycznym CGS niezracjonalizowanym (oznaczonym przez CGSM) (tablica 5). W dalszej części pracy omówiono zależności między wielkościami wchodzącymi w skład równań typu B a wielkościami równań typu J w układzie jednostek MKSA zr (tablica 6). Podano także związki między wielkościami typu B w układach MKSA zr i CGSM (tablica 6).

Na zakończenie podano wykaz najczęściej stosowanych wielkości magnetycznych (tablica 7), mechanicznych (tablica 8) i magnetomechanicznych (tablica 9) i ich oznaczeń oraz ich związki z wielkościami wchodzącymi w skład równań typu B w układzie MKSA zr.

1. WSTĘP

Ostatnio w przetwornikach ultradźwiękowych i podzespołach telekomunikacyjnych coraz częściej znajdują zastosowanie materiały magnetostrykcyjne. Zjawisko magnetostrykcji polega na zmianie wymiarów geometrycznych pod wpływem natężenia pola magnetycznego lub na zmianie własności magnetycznych pod wpływem sił przyłożonych na próbkę. Zjawiska te obok doniosłego znaczenia w teorii magnetyzmu mają również znaczenie praktyczne. Zjawisko pierwsze zostało wykorzystane w nadajnikach ultradźwiękowych, stabilizatorach, filtrach i liniach opóźniających w elementach pobudzających do drgań ultradźwiękowych, natomiast zjawisko drugie znalazło zastosowanie w odbiornikach ultradźwiękowych oraz elementach filtrów czy linii opóźniających, przekształcających wytworzoną przez przetworniki nadawcze energię akustyczną z powrotem na energię magnetyczną. Bardziej szczegółowe omówienie tych zjawisk jest tematem kilku publikacji polskich [15], [22], [23], [24], [30], [40], [41], [43], [44] oraz kilkuset zagranicznych, np. [1], [2], [3], [4], [25], [31], [33], [34]. Kompletnie ujęcie matematyczne tych zjawisk z uwzględnieniem zjawisk cieplnych [30], [33] jest sprawą skomplikowaną i dla praktyków, mających zwykle do czynienia jednocześnie tylko z niektórymi przejawami

mi tych zjawisk, jest zwykle mało przydatne. W technice spotyka się zazwyczaj układy równań magnetostrykcyjnych [4], [25], [42], w których zaniebano sprawy cieplne, co upraszcza znacznie zagadnienie, a z punktu widzenia zastosowań jest w większości przypadków słuszne. Jednak dowolność w przyjmowaniu zmiennych niezależnych i układów jednostek spowodowała pojawienie się wielu wersji równań i wielkiej liczby współczynników określających magnetyczne, mechaniczne i magnetostrykcyjne własności materiałów. Dlatego często bardzo trudne jest porównywanie wyników podawanych przez różnych autorów i niekiedy czytelnicy interesujący się magnetostrykcją marginesowo zmuszeni są do historycznych studiów prac z tej dziedziny.

Autor, obok wysuniętych propozycji stosowania równań magnetostrykcyjnych o zmiennych niezależnych: H (natężenie pola magnetycznego) i σ (naprężenia normalne) i zmiennych zależnych: B (indukcja magnetyczna) i ε (względna zmiana długości — współczynnik magnetostrykcji), podaje wszystkie przeliczenia do pozostałych wersji układów równań. Należy zaznaczyć, że proponowany przez autora układ równań magnetostrykcyjnych oparty jest na obowiązującym od 1953 r. [38] w Polsce i powszechnie już stosowanym za granicą układzie jednostek MKSA zracjonalizowanym. Stosowanie proponowanego układu nie pociąga za sobą żadnych istotnych przeliczeń, jeśli chodzi o główne dotychczas stosowane współczynniki magnetostrykcyjne, a ma tę zaletę, że układ obejmuje w zasadzie najważniejsze wielkości magnetomechaniczne, jak naprężenia σ , moduł sprężystości E , magnetostrykcję ε , czułość naprężeń d , indukcję magnetyczną B , przenikalność μ i natężenie pola magnetycznego H , a pozostałe wielkości, jak współczynnik sprzężenia magnetomechanicznego k i stałą magnetostrykcji h łatwo z nich wyznaczyć. Ponadto wydaje się słuszniejsze przyjęcie jako wielkości pierwotnych, niezależnych: σ i H , a wtórnych zależnych: ε i B (przyjęcie magnetyzacji J , słuszniejsze z punktu widzenia fizyki, w technice jest mniej dogodne, gdyż B jest bezpośrednio wyznaczalne z pomiarów i nie pociąga za sobą dodatkowych przeliczeń, co miałoby miejsce w przypadku J). Stosowane przez autora oznaczenia literowe omawianych wielkości oparte są na oznaczeniach tradycyjnych lub zalecanych przez normy krajów o największych osiągnięciach w dziedzinie magnetostrykcji. Oznaczenie np. współczynnika względnych zmian długości, jakim jest magnetostrykcja, przez ε ma usprawiedliwienie w powszechnym stosowaniu tego współczynnika w mechanice, natomiast stosowane w fizyce oznaczenie współczynnika magnetostrykcji statycznej przez λ jest w dziedzinie magnetostrykcji bardziej wieloznaczne i może być mylone ze stałą magnetostrykcji, długością fali dźwiękowej, współczynnikami cieplnymi, stałymi Lamego itp.

Współczynniki magnetostrykcyjne są oznaczane podobnie jak odpowiadające im współczynniki piezoelektryczne [18], [25], co ułatwi konstruktorom przetworników ultradźwiękowych, interesujących się tylko magne-

tostrykcją ze strony użytkowej, przeprowadzanie bezpośrednich porównań współczynników, bez głębszego wnikania w treść artykułów, w celu rozszyfrowania znaczenia danego oznaczenia.

Równania magnetostrykcyjne, przy założeniu, że przebiegi są adiabatyczne (tzn. zachodzą bez wymiany ciepła z otoczeniem), stanowią układ dwu równań zawierających dwie zmienne magnetyczne: natężenie pola magnetycznego H i indukcję magnetyczną B (lub magnetyzację J) oraz dwie zmienne mechaniczne: naprężenia normalne σ (przy drganiach wzdlużnych) i odkształcenia względne ε .

Przed przystąpieniem do omówienia tych równań rozpatrzone zostaną podstawowe zależności między wielkościami magnetycznymi oraz między wielkościami mechanicznymi występującymi w tych równaniach.

2. WIELKOŚCI MAGNETYCZNE

Pole magnetyczne scharakteryzowane jest w każdym punkcie przez dwie podstawowe wielkości wektorowe: natężenie pola magnetycznego H i indukcję magnetyczną B . Wielkości te powiązane są ze sobą następującą zależnością znaną na drodze eksperymentalnej:

$$B = \mu H, \quad (1)$$

gdzie współczynnik μ jest wielkością tensorową, zależną od ośrodka, i nosi nazwę przenikalności magnetycznej. Wartość i jednostki przenikalności zależą od przyjętego układu jednostek.

W niniejszej pracy przyjęto układ jednostek MKSA zracjonalizowany (oznaczany w dalszych częściach przez MKSA *zr*).

Na skutek tego, że w dziedzinie magnetyzmu, zarówno w fizyce, jak i naukach technicznych, jest jeszcze w powszechnym użyciu układ CGSM, niezracjonalizowany, równolegle będą podawane (drobnym drukiem) zależności również i w tym ostatnim układzie, przy czym w częściach 4 i 5 zostaną podane wzajemne zależności i wzory przeliczeniowe.

Przyjmując, że indukcja B jest mierzona w weberach na metr kwadratowy (Wb/m^2), a natężenie pola magnetycznego H — w amperach na metr (A/m), przenikalność magnetyczna próżni w układzie MKSA *zr* będzie:

$$\mu_0 = \frac{B_0}{H} = 4\pi 10^{-7} [\text{H/m}] \quad (2)$$

gdzie henr:

$$1 \text{ H} = 1 \Omega \text{ s} = 1 \frac{\text{Vs}}{\text{A}}. \quad (3)$$

Przenikalność innych materiałów μ podaje się najczęściej w odniesieniu do przenikalności próżni μ_0 :

$$\mu' = \frac{\mu}{\mu_0}. \quad (4)$$

W układzie CGSM przenikalność próżni wynosi

$$\mu_0 = \frac{B_0}{H} = 1 \text{ [Gs/Oe]}, \quad (5)$$

gdzie:

B_0 — w gausach (Gs),

H — w erstedach (Oe).

Stąd wynika, że wartość liczbową przenikalności względnej materiału w układzie MKSA μ_r równa jest wartości liczbowej przenikalności magnetycznej materiału w układzie CGSM niezracjonalizowanym.

Jeżeli w danym miejscu w próżni występuje indukcja magnetyczna B_0 , to po wprowadzeniu w to miejsce jakiegoś ciała indukcja ta zmieni swą wartość i będzie:

$$B = B_0 + J, \quad (6)$$

gdzie J nazywa się natężeniem magnesowania (lub magnetyzacją) i jest związana z danym ciałem. Na podstawie (1), (2), (4), (6) można napisać:

$$J = B - B_0 = \mu' \mu_0 H - \mu_0 H = \mu_0 H (\mu' - 1). \quad (7)$$

Po oznaczeniu:

$$\mu' - 1 = \kappa' \quad (8)$$

będzie:

$$J = \kappa' B_0 = \kappa' \mu_0 H = \kappa H, \quad (9)$$

gdzie:

κ' — podatność magnetyczna względna,

κ — podatność magnetyczna bezwzględna w H/m.

W układzie CGSM wzory (6), (7), (8), (9) będą odpowiednio:

$$B = B_0 + 4\pi J, \quad (10)$$

$$J = \frac{B - B_0}{4\pi} = \frac{\mu H - \mu_0 H}{4\pi} = \frac{\mu - \mu_0}{4\pi} H, \quad (11)$$

$$\frac{\mu' - 1}{4\pi} \mu_0 = \kappa, \quad (12)$$

więc:

$$J = \kappa H \quad (13)$$

gdzie:

κ — podatność w Gs/Oe,

μ — przenikalność w Gs/Oe,

μ_0 — przenikalność próżni = 1 Gs/Oe,

μ' — przenikalność względna,

B — indukcja w Gs,

J — magnetyzacja w Gs,

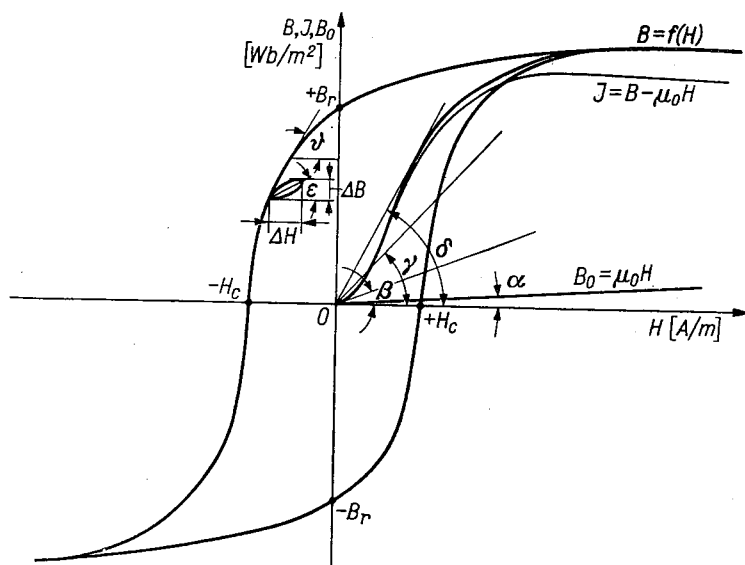
H — natężenie pola w Oe.

Indukcja nasycenia B_s jest to wartość indukcji, przy której natężenie magnesowania J przestaje się zmieniać. Remanencja (pozostałość magnetyczna) B_r określa, jaką wartość indukcji będzie miał materiał magnetycznie nasycony po usunięciu pola zewnętrznego. Koercja (natężenie powścią-

gające) H_c jest to wartość natężenia pola, jaką należy przyłożyć, by wartość indukcji magnetycznej sprowadzić z remanencji do zera.

Definicji przenikalności magnetycznych jest bardzo wiele. Tu omówione zostaną tylko te, które będą występowały w dalszej części pracy i w ogóle w dziedzinie magnetostrykcji.

Przenikalność magnetyczna przy danym polu (lub indukcji) μ określona jest ze wzoru (1) jako stosunek indukcji magnetycznej do odpowiadającego



Rys. 1. Ilustracja graficzna definicji magnetycznych

$$\operatorname{tg} \alpha = \frac{B_0}{H} = \mu_0$$

$$\operatorname{tg} \beta = \lim_{H \rightarrow 0} \frac{B}{H} = \mu_p$$

$$\operatorname{tg} \gamma = \frac{B}{H}$$

$$\operatorname{tg} \delta = \max \operatorname{tg} \gamma = \mu_m$$

$$\operatorname{tg} \epsilon = \left(\frac{\Delta B}{\Delta H} \right)_{H_0} = \mu_{\Delta}$$

$$\operatorname{tg} \varphi = \lim_{\Delta H \rightarrow \Delta H_r} \left(\frac{\Delta B}{\Delta H} \right)_{H_0} = \mu_r$$

$$\operatorname{tg} \vartheta = \frac{dB}{dH} = \mu_d$$

jej natężenia pola, przy czym można określić ją jako tangens nachylenia stycznej poprowadzonej z początku układu do odpowiedniego punktu na krzywej pierwotnej magnesowania ($\operatorname{tg} \gamma$ na rys. 1). Wyróżnia się tu dwie

wartości: przenikalność początkową μ_p , która określana jest jako tangens nachylenia stycznej do krzywej dziewiczej rozpatrywanego materiału poprowadzonej w początku układu współrzędnych B i H ($\text{tg } \beta$ na rys. 1) i przenikalność maksymalną μ_m określaną jako tangens nachylenia stycznej do pierwotnej krzywej magnesowania poprowadzonej z początku układu współrzędnych ($\text{tg } \delta$ na rys. 1). Obok tych definicji dotyczących krzywej pierwotnej magnesowania występują jeszcze inne definicje przenikalności związane zarówno z krzywą pierwotną, jak i z pętlą histerezy magnetycznej. I tak przenikalność przyrostowa dotyczy niewielkich zmian pola na tle stałego podmagnesowania i jest stosunkiem przyrostu indukcji do przyrostu pola przy stałym natężeniu pola H_0 :

$$\mu_\Delta = \left(\frac{\Delta B}{\Delta H} \right)_{H_0}, \quad (14)$$

przy czym μ_Δ jest tangensem nachylenia osi pętliki do osi H ($\text{tg } \varepsilon$ na rys. 1). Ogólnie biorąc zmiany te są niepowtarzalne, tzn. gdy natężenie H po obiegnięciu małego cyklu wróci do swej pierwotnej wartości, wartość indukcji będzie inna niż przed obiegiem i następna pętlika będzie obiegać inną drogą. W szczególnym przypadku, gdy zmiany ΔH są tak małe, że za każdym razem po powrocie do wartości wyjściowej H , otrzyma się wartość wyjściową B , tzn. cykliczne zmiany B i H przebiegają zawsze po tych samych torach, przenikalność przyrostowa nosi nazwę przenikalności odwracalnej μ_r ($\text{tg } \varphi$ na rys. 1).

Pochodna indukcji względem pola, tj. tangens nachylenia stycznej w dowolnym punkcie krzywej magnesowania nosi nazwę przenikalności różniczkowej μ_d ($\text{tg } \vartheta$ na rys. 1).

Należy zaznaczyć, że przed porównywaniem wartości tak zdefiniowanych wielkości z odpowiednimi przenikalnościami w układzie CGSM należy wartości z układu MKSA zr dzielić przez $4\pi 10^{-7}$.

W konkretnym zastosowaniu do magnetostrykcji powyższe definicje nie wystarczają. Należy wprowadzić oznaczenia dodatkowe. Z doświadczeń wiadomo, że inną wartość przenikalności posiada próbka swobodna, tzn. bez naprężeń zewnętrznych, a inną próbka zamocowana, tzn. o stałej długości. Dlatego dla potrzeb magnetostrykcji wprowadzono definicje przenikalności magnetycznych materiału (zwykle odwracalnych lub przyrostowych) przy stałych naprężeniach μ_σ i przy stałej długości próbki μ_ε . A więc przenikalność odwracalna próbki swobodnej będzie:

$$\mu_\sigma = \lim_{\Delta H \rightarrow \Delta H_r} \left(\frac{\Delta B}{\Delta H} \right)_{\sigma=0}, \quad (15)$$

a przenikalność odwracalna próbki zamocowanej:

$$\mu_\varepsilon = \lim_{\Delta H \rightarrow \Delta H_r} \left(\frac{\Delta B}{\Delta H} \right)_{\varepsilon=0}. \quad (16)$$

Znając różnicę tych wielkości można określić w sposób następujący stosunek energii magnetycznej przetwarzanej na mechaniczną. Jeżeli na stałe pole magnetyczne H_0 nałożone zostanie pole zmienne o amplitudzie $\bar{H}$, to zmagazynowana energia właściwa, tzn. energia na jednostkę objętości próbki swobodnej W_σ równa jest iloczynowi amplitud indukcji $\bar{B}_\sigma$ i natężenia pola $\bar{H}$:

$$W_\sigma = \frac{1}{2} \bar{B}_\sigma \bar{H} = \frac{1}{2} \mu_\sigma \bar{H}^2, \quad (17)$$

natomiast zmagazynowana energia magnetyczna próbki zamocowanej W_ϵ :

$$W_\epsilon = \frac{1}{2} \bar{B}_\epsilon \bar{H} = \frac{1}{2} \mu_\epsilon \bar{H}^2. \quad (18)$$

Wynika stąd, że wewnętrzna energia sprężystości W_m w przecie swobodnym będzie różnicą tych energii i może być przetwarzana na energię mechaniczną

$$W_m = W_\sigma - W_\epsilon = \frac{1}{2} (\mu_\sigma - \mu_\epsilon) \bar{H}^2, \quad (19)$$

gdzie:

$W_m, W_\sigma, W_\epsilon$ — w dżulach (J),

$\bar{B}_\sigma, \bar{B}_\epsilon$ — w Wb/m²,

$\bar{H}$ — w A/m²,

μ_σ, μ_ϵ — w H/m.

W układzie CGSM we wzorach tych należy wprowadzić mnożnik $\frac{1}{4\pi}$ [1] i odpowiednie jednostki.

Stosunek części energii magnetycznej, która jest przetwarzana na energię mechaniczną W_m do całkowitej energii magnetycznej zmagazynowanej w próbce swobodnej W określany jest jako kwadrat współczynnika sprzężenia magnetomechanicznego:

$$k^2 = \frac{\mu_\sigma - \mu_\epsilon}{\mu_\sigma} = 1 - \frac{\mu_\epsilon}{\mu_\sigma}. \quad (20)$$

3. WIELKOŚCI MECHANICZNE

Obecnie zostaną rozpatrzone zależności między wielkościami mechanicznymi. Próbka poddana działaniu sił ulega odkształceniu. Siły odkształcające próbkę wykonują pewną pracę, która zostaje zmagazynowana w próbce w postaci energii mechanicznej. Po usunięciu sił zewnętrznych próbka wracając do swego pierwotnego kształtu może częściowo zwrócić nagromadzoną w niej energię. Próbka jest doskonale sprężysta, jeżeli po usunięciu sił zewnętrznych przyjmuje swój pierwotny kształt i wymiary. Większość materiałów w pewnych granicach obciążeń można uważać za

ciała sprężyste. Do materiałów tych stosuje się prawo Hooke'a mówiące, że odkształcenia są proporcjonalne do naprężeń:

$$\varepsilon = \frac{\sigma}{E}, \quad (21)$$

przy czym:

$$\varepsilon = \frac{\Delta l}{l}, \quad (22)$$

$$\sigma = \frac{P}{A}, \quad (23)$$

gdzie:

- ε — wydłużenie względne,
- Δl — przyrost długości w m,
- l — długość początkowa w m,
- σ — naprężenia mechaniczne w N/m^2 ,
- P — siła (obciążenie) w zakresie do granicy sprężystości w N,
- A — powierzchnia w m^2 ,
- E — moduł sprężystości podłużnej (Younga) w N/m^2 .

Odkształcenia mogą być postaciowe i objętościowe. Odkształcenie postaciowe — to takie, w którym następuje zmiana kształtu przy niezmięnionej objętości, a objętościowe — to takie, w którym przy tym samym kształcie następuje zmiana objętości. Oznaczając przez ε_x , ε_y , ε_z wydłużenia względne wzdłuż kierunków osi współrzędnych można napisać, że dla odkształceń postaciowych:

$$\varepsilon_x + \varepsilon_y + \varepsilon_z = 0, \quad (24)$$

a dla odkształceń objętościowych:

$$\varepsilon_x = \varepsilon_y = \varepsilon_z \neq 0, \quad (25)$$

przy czym kąty między ścianami w przypadku wielościanów nie ulegają zmianie.

Względny przyrost objętości będzie:

$$\omega = \frac{\Delta V}{V} = \varepsilon_x + \varepsilon_y + \varepsilon_z. \quad (26)$$

W przetwornikach magnetostrykcyjnych występują przeważnie odkształcenia postaciowe. Odkształcenia objętościowe nałożone na odkształcenia postaciowe zaczynają występować w pobliżu nasycenia magnetycznego.

Wymiary poprzeczne próbki zmieniają się zgodnie z równaniem (21). Stwierdzono, że stosunek względnego odkształcenia w kierunku poprzecznym ε' do względnego odkształcenia wzdłużnego ε jest wielkością stałą dla danego materiału.

Wielkość tę nazwano liczbą Poissona:

$$\nu = -\frac{\varepsilon'}{\varepsilon} \quad (27)$$

W przypadkach ścinania słuszna jest następująca zależność:

$$\gamma = \frac{\tau}{G}, \quad (28)$$

gdzie:

γ — kąt odkształcenia postaciowego w radianach,

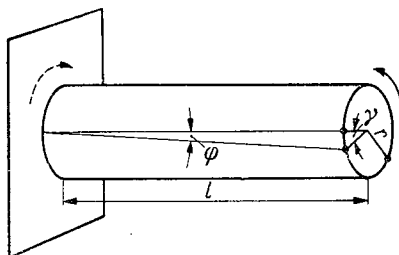
τ — naprężenia ścinające w N/m^2 ,

G — moduł ścinania lub skręcania, zwany też modułem sprężystości postaciowej w N/m^2 .

Moduł sprężystości postaciowej G powiązany jest z modułem sprężystości wzdłużnej E i z liczbą Poissona ν następującym wzorem:

$$G = \frac{E}{2(1+\nu)}. \quad (29)$$

W przypadku skręceń (rys. 2) odkształcenie pręta o przekroju okrągłym



Rys. 2. Skręcanie walca

określone jest wartością kąta skręcenia odniesionego do jednostki długości:

$$\varphi' = \frac{\varphi}{l} \quad (30)$$

i jest określone z zależności (28)

$$\varphi' = \frac{\gamma}{r} = \frac{\tau}{Gr}, \quad (31)$$

gdzie:

φ — kąt skręcenia w radianach,

l — długość pręta w m,

r — promień przekroju pręta w m,

γ — kąt odkształcenia postaciowego w radianach,

τ — naprężenia ścinające w N/m^2 ,

G — moduł sprężystości postaciowej w N/m^2 .

Kąty skręceń próbek o innych kształtach są określane bardziej skomplikowanymi wzorami, które można znaleźć w pracach z dziedziny mechaniki [20] lub publikacjach dotyczących przetworników ultradźwiękowych.

Zgodnie z zasadą zachowania energii praca wykonana przy odkształceniu sprężystym zamienia się w całości (w warunkach idealnych) na energię potencjalną wewnętrznych sił sprężystości lub krócej energię sprężystą próbki. Wykonana praca L wyraża się zależnością:

$$L = \frac{1}{2} \Delta l P, \quad (32)$$

gdzie:

Δl — wydłużenie w m,

P — siła w N,

L — praca w J.

Oznaczając przez W energię sprężystą właściwą, tj. energię sprężystości na jednostkę objętości, po odpowiednich przekształceniach dla przebiegów dynamicznych otrzymuje się:

$$W = \frac{1}{2} \bar{\sigma} \bar{\varepsilon}. \quad (33)$$

Uwzględniając zależność (21):

$$W = \frac{1}{2} \bar{\varepsilon}^2 E = \frac{1}{2} \frac{\bar{\sigma}^2}{E}, \quad (34)$$

gdzie:

W — energia właściwa sprężystości w J,

$\bar{\sigma}$ — amplituda naprężeń normalnych w N/m²,

$\bar{\varepsilon}$ — amplituda wydłużeń względnych,

E — moduł sprężystości w N/m².

Po umieszczeniu próbki o własnościach magnetostrykcyjnych w polu magnetycznym na skutek zmian modułu sprężystości ulegnie zmianie energia sprężystości. Wartość modułu sprężystości E_H pomierzona przy stałym natężeniu pola H będzie nieco się różnić od wartości modułu E_B pomierzonego przy stałej indukcji. Odpowiednio będą się różnić wzory na energię. I tak przy stałym natężeniu pola energia sprężysta W_H będzie:

$$W_H = \frac{1}{2} \frac{\bar{\sigma}^2}{E_H}, \quad (35)$$

a przy stałej indukcji W_B :

$$W_B = \frac{1}{2} \frac{\bar{\sigma}^2}{E_B}. \quad (36)$$

Różnica tych energii może być przetworzona na energię magnetyczną,

$$W_{mag} = W_H - W_B = \frac{1}{2} \bar{\sigma}^2 \left(\frac{1}{E_H} - \frac{1}{E_B} \right). \quad (37)$$

Podobnie jak w przypadku zamiany energii magnetycznej na mechaniczną również i teraz, tzn. przy zamianie energii mechanicznej na magnetyczną można wprowadzić definicję współczynnika sprzężenia mechanomagnetycznego. A więc stosunek części energii mechanicznej przetwarzanej na magnetyczną do całkowitej energii mechanicznej zmagazynowanej w próbce umieszczonej w stałym polu magnetycznym określony jest jako kwadrat współczynnika sprzężenia mechanomagnetycznego:

$$k^2 = \frac{\frac{1}{E_H} - \frac{1}{E_B}}{\frac{1}{E_H}} = 1 - \frac{E_H}{E_B}. \quad (38)$$

Dla przebiegów odwracalnych wartości k wyznaczone z równań (20) i (38) są sobie równe.

4. UKŁADY JEDNOSTEK

W pracach z dziedziny magnetostrykcji można spotkać jednostki magnetyczne, elektryczne i mechaniczne, pochodzące ze wszystkich stosowanych układów jednostek.

Celem niniejszej pracy jest, między innymi, uporządkowanie tych zagadnień i podanie odpowiednich zestawień i wzorów przeliczeniowych. Przy omawianiu zjawisk i własności mechanicznych stosowane są układy CGS, MKS i układ techniczny. W dziedzinie magnetyzmu obok zalecanego i obowiązującego w wielu krajach układu MKSA zrationalizowanego spotyka się często układy: zrationalizowane i niezrationalizowane CGSM i GGSE, układ niezrationalizowany MKSA i układ Gaussa. Autor ograniczy się do podania zestawień obejmujących wielkości występujące w niniejszej pracy, a podstawowe w dziedzinie magnetostrykcji, zaznaczając, że zagadnienia układów jednostek zostały wyczerpująco omówione w kilku pracach polskich [16], [28], [29], [38], [45].

Jednostkami podstawowymi w układzie MKSA są: metr (m), kilogram (kg) (jednostka masy), sekunda (s) i amper (A). W tym układzie główną jednostką siły jest niuton (N), pracy — dżul (J), mocy — wat (W), gdzie:

$$1 \text{ W} = 1 \text{ Js}^{-1} = 1 \text{ Nms}^{-1} = 1 \text{ kg m}^2 \text{ s}^{-3}, \quad (39)$$

napięcia — wolt (V), ładunku — kulomb (C), indukcyjności — henr (H), pojemności — farad (F), i oporności — om (Ω).

Główne jednostki magnetyczne w układzie MKSA to: weber (Wb) — jednostka strumienia magnetycznego Φ , Wb/m^2 — jednostka indukcji B , A/m — jednostka natężenia pola magnetycznego H i H/m — jednostka przenikalności magnetycznej (bezwzględnej). Powiązania między jednostkami elektrycznymi i magnetycznymi na przykładzie jednostki mocy są następujące:

$$1 \text{ W} = 1 \text{ A V} = 1 \text{ C s V} = 1 \text{ H A}^2 \text{ s}^{-1} = \text{F V}^2 \text{ s} = \text{V}^2 \Omega^{-1} = \text{A W b s}^{-1}. \quad (40)$$

Jednostki MKS i CGS różnią się wykładnikiem potęgowym przy liczbie 10. Podobnie różnią się jednostki MKSA i CGSM (gdy oba układy są równocześnie zrationalizowane lub niezrationalizowane). Stosunek jednostek układu CGSE i CGSM określa współczynnik, którym jest prędkość światła ($\sim 3 \cdot 10^{10}$ cm/s) lub jej kwadrat. Układ zrationalizowany MKSA powszechnie stosowany w Niemczech, Związku Radzieckim i Stanach Zjednoczonych, różni się od układu CGSM niezrationalizowanego współczyn-

Tablica 1

Zależności między niektórymi jednostkami zrationalizowanego układu praktycznego MKSA a jednostkami niezrationalizowanych układów CGSM oraz jednostkami układu technicznego

L. p.	Wielkość	Sym-bol	Jedn. w MKSA zr	Równowartość w jednostkach		
				CGSM	CGSE	układ techniczny
1	2	3	4	5	6	7
1	długość	l	1 m	100 cm	100 cm	1 m
2	czas	t	1 s	1 s	1 s	1 s
3	masa	m	1 kg	100 g	100 g	0,102 kG s ² /m
4	natężenie prądu elektrycznego	i	1 A	0,1 jem	$\sim 3 \cdot 10^9$ jes	
5	powierzchnia	s	1 m ²	10 ⁴ cm ²	10 ⁴ cm ²	1 m ²
6	objętość	v	1 m ³	10 ⁶ cm ³	10 ⁶ cm ³	1 m ³
7	siła	F	1 N	10 ⁵ dyn	10 ⁵ dyn	0,102 kG
8	naprężenia	σ, τ	1 N/m ²	10 dyn/cm ²	10 dyn/cm ²	0,102 kG/m ² 1,02 · 10 ⁻⁷ kG/mm ²
9	moduł sprężystości	E, G	1 N/m ²	10 dyn/cm ²	10 dyn/cm ²	0,102 kG/m ² 1,02 · 10 ⁻⁷ kG/mm ²
10	praca energia	A W	1 J	10 ⁷ erg	10 ⁷ erg	0,102 kGm
11	moc	P	1 W	10 ⁷ erg/s	10 ⁷ erg/s	0,102 kG/s
12	indukcja magnetyczna	B	1 Wb/m ²	10 ⁴ Gs	$\sim \frac{1}{3} \cdot 10^{-6}$ jes	
13	natężenie pola magnetycznego	H	1 A/m	$4 \pi \cdot 10^{-3}$ Oe	$\sim 12 \pi \cdot 10^7$ jes	
14	bezwzględna przenikalność magnetyczna	μ	1 H/m	$\frac{1}{4\pi} 10^7$ Gs/Oe	$\sim \frac{1}{36\pi} 10^{-13}$ s ² /cm ²	
15	względna przenikalność magnetyczna	μ'	1	1	1	

nikiem 4π lub jego odwrotnością, pomnożonym przez 10^9 przy takich wielkościach, jak przenikalność magnetyczna μ , natężenie pola magnetycznego H , masa magnetyczna m , oporność R_μ i przewodność magnetyczna Λ , moment magnetyczny M_μ , magnetyzacja (namagnesowanie) J , przenikalność dielektryczna ε i przesunięcie dielektryczne D .

Należy również zwrócić uwagę na to, że w układzie technicznym przyjęto jako jednostkę podstawową jednostkę siły (kG) w miejsce występującej we wszystkich pozostałych układach jednostki masy, co wprowadza przy przeliczeniach współczynnik równy przyspieszeniu ziemskiemu g . Pozostałe jednostki podstawowe tego układu to metr i sekunda.

Odpowiednie zależności między niektórymi jednostkami układu MKSA a jednostkami układów: CGSM, CGSE i technicznego podano w tablicy 1.

5. UKŁADY RÓWNAŃ MAGNETOSTRYKCYJNYCH

Ogólne równania magnetostrykcyjne wiążą ze sobą wielkości mechaniczne, magnetyczne i cieplne. Rozpatrzenie jednoczesne wszystkich zależności jest praktycznie niemożliwe i w obecnej chwili ma jedynie znaczenie teoretyczne.

W przetwornikach ultradźwiękowych zamiana energii magnetycznej na mechaniczną (lub odwrotnie) zachodzi bardzo szybko, tak, że nie ma czasu na wymianę ciepła i można przyjąć, że proces przebiega adiabatycznie. Przy założeniu, że przebiegi są adiabatyczne, zjawisko magnetostrykcji można opisać przy pomocy układu równań różniczkowych zawierających 4 zmienne, przy czym dwie z nich to wielkości magnetyczne H i B (lub J), a dwie pozostałe — to dynamiczne wielkości mechaniczne σ i ε . Przyjmując jako zmienne niezależne jedną wielkość magnetyczną, a drugą mechaniczną, a dwie pozostałe wielkości jako zmienne zależne, można napisać 4 układy równań magnetostrykcyjnych, przy czym zależnie od tego, czy przyjęto indukcję B , czy magnetyzację J , każdy z tych układów będzie w dwóch wersjach (B i J). Poniżej podano wszystkie możliwe zależności funkcjonalne wielkości magnetycznych i mechanicznych dla drgań wzdłużnych:

$$\left. \begin{array}{l} \varepsilon = f_1(\sigma, H) \\ B = f_2(\sigma, H) \end{array} \right\} \text{I B} \quad \left. \begin{array}{l} \varepsilon = f_1(\sigma, H) \\ J = f_2(\sigma, H) \end{array} \right\} \text{I J} \quad (41) \quad (42)$$

$$\left. \begin{array}{l} \varepsilon = f_1(\sigma, B) \\ H = f_2(\sigma, B) \end{array} \right\} \text{II B} \quad \left. \begin{array}{l} \varepsilon = f_1(\sigma, J) \\ H = f_2(\sigma, J) \end{array} \right\} \text{II J} \quad (43) \quad (44)$$

$$\left. \begin{array}{l} \sigma = f_1(\varepsilon, B) \\ H = f_2(\varepsilon, B) \end{array} \right\} \text{III B} \quad \left. \begin{array}{l} \sigma = f_1(\varepsilon, J) \\ H = f_2(\varepsilon, J) \end{array} \right\} \text{III J} \quad (45) \quad (46)$$

$$\left. \begin{array}{l} \sigma = f_1(\varepsilon, H) \\ B = f_2(\varepsilon, H) \end{array} \right\} \text{IV B} \quad \left. \begin{array}{l} \sigma = f_1(\varepsilon, H) \\ J = f_2(\varepsilon, H) \end{array} \right\} \text{IV J} \quad (47) \quad (48)$$

Oczywiście w postaci najogólniejszej równania magnetostrykcyjne będą miały postać macierzy. W celu uproszczenia rozważań przeprowadzona zostanie analiza na równaniach prostszych, słusznych przy pewnych dodatkowych założeniach.

Warunki, w jakich słuszne są omawiane równania magnetostrykcyjne, są następujące: 1) amplitudy okresowo (a przy tym założeniu sinusoidalnie) zmiennych wielkości magnetycznych i mechanicznych są bardzo małe i praktycznie można przyjąć zależności liniowe i zaniedbać histerezę magneto mechaniczną, 2) materiał polikrystaliczny (gdy nie jest magnesowany ani odkształcany), z którego wykonana jest badana próbka, jest magnetycznie i mechanicznie jednorodny (izotropowy), 3) odkształcenia dynamiczne zgodne są z prawem Hooke'a, 4) pręty są dostatecznie długie, a przekrój odpowiednio mały, aby można było zaniedbać współczynnik od magnesowania, 5) zmienne naprężenia i pola, zależnie od rodzaju przetwornika, są tylko równoległe lub tylko prostopadłe do stałego pola pod magnesującego, 6) praca odbywa się przy pod magnesowaniu niższym od nasycającego, w celu wyeliminowania w rozważaniach wpływu zjawiska magnetokalorycznego [30], [4], które wprowadzie niewiele zmienia wielkości mechaniczne i magnetyczne (dużo poniżej temperatury Curie), ale ma duży wpływ na zmiany wielkości magneto mechanicznych.

Zjawisko magnetostrykcji, jak już wyżej powiedziano, opisane jest przez jeden z układów równań różniczkowych uzyskanych przez różniczkowanie funkcji (41) ÷ (48). Należy zaznaczyć, że dla drgań skrętnych lub ścinających można napisać analogiczny zespół funkcji i równań jak dla drgań wzdłużnych, przyjmując zamiast naprężeń normalnych (rozciągających lub ściskających) σ naprężenia ścinające τ , zamiast względnych wydłużeń ε — odkształcenia kątowe γ , a zamiast dynamicznych modułów sprężystości wzdłużnej E_B i E_H — moduły sprężystości postaciowej G_B i G_H .

5.1. Równania magnetostrykcyjne o zmiennych niezależnych σ i H

Ustalenie, która z wielkości fizycznych jest pierwotna, a która wtórna, jest często zagadnieniem dyskusyjnym, gdyż przy pewnych założeniach początkowych można w wielu przypadkach zagadnienie odwrócić. Wiadomo, że w nadajniku ultradźwiękowym przetwarza się energię magnetyczną (lub elektryczną) na mechaniczną i akustyczną, ale również wiadomo, że energia akustyczna może być w odbiorniku ultradźwiękowym zamieniona na energię magnetyczną i elektryczną. O tym, która energia jest pierwotna, decydują warunki, w jakich rozpatrywane jest dane zagadnienie. Dlatego nie jest intencją autora rozstrzygać metafizycznie, co jest pierwotne. Z punktu widzenia magnetyzmu i mechaniki wydaje się słuszne założenie, że natężenie pola H i naprężenia σ są wielkościami pierwotnymi,

wywołującymi odkształcenia ε i zmiany indukcji B . Dlatego bardziej szczegółowa analiza zostanie przeprowadzona dla tych założeń [układ równań I B (41), (42)], a dla pozostałych układów zostaną podane tylko niektóre przeliczenia i wyniki końcowe.

Różniczki zupełne równań I B mają następującą postać:

$$d\varepsilon = \left(\frac{\partial \varepsilon}{\partial \sigma} \right)_H d\sigma + \left(\frac{\partial \varepsilon}{\partial H} \right)_\sigma dH \quad (49)$$

$$dB = \left(\frac{\partial B}{\partial \sigma} \right)_H d\sigma + \left(\frac{\partial B}{\partial H} \right)_\sigma dH \quad (50)$$

Pochodne cząstkowe zawarte w tych równaniach są zdefiniowane w sposób następujący:

$$\left(\frac{\partial \varepsilon}{\partial \sigma} \right)_H = \frac{1}{E_H} \text{ — współczynnik sprężystości (odwrotność dynamicz-} \quad (51)$$

nego modułu sprężystości wzdłużnej) mierzony przy stałym natężeniu pola,

$$\left(\frac{\partial \varepsilon}{\partial H} \right)_\sigma = d' \text{ — dynamiczny współczynnik magnetostrykcji odwra-} \quad (52)$$

calnej przy stałych naprężeniach,

$$\left(\frac{\partial B}{\partial \sigma} \right)_H = d'' \text{ — czułość naprężeń,} \quad (53)$$

$$\left(\frac{\partial B}{\partial H} \right)_\sigma = \mu_\sigma \text{ — przenikalność odwracalna mierzona przy stałych} \quad (54)$$

naprężeniach.

Otrzymane 4 wielkości po uwzględnieniu zależności termodynamicznych można zredukować do 3.

Zgodnie z pierwszą zasadą termodynamiki zmiana energii właściwej dW , tj. energii na jednostkę objętości, jest sumą ciepła dQ pochłoniętego przez ciało i pracy dL dostarczanej z zewnątrz:

$$dW = dQ + dL. \quad (55)$$

Druga zasada termodynamiki mówi, że zmiana entropii układu dS jest stosunkiem zmiany ciepła dQ do temperatury bezwzględnej T :

$$dS = \frac{dQ}{T}. \quad (56)$$

Praca dostarczona z zewnątrz może mieć różny charakter. W przypadku rozważanego ciała ferromagnetycznego praca dostarczona z zewnątrz będzie określona przez energię magnesowania materiału:

$$dL = H dJ = H dB - \mu_0 H dH, \quad (57)$$

gdzie:

$H dB$ — praca zużyta na magnesowanie,

$H dJ$ — energia magnesowania materiału ferromagnetycznego,

$\mu_0 H dH$ — praca zużyta przy magnesowaniu, gdyby nie było materiału ferromagnetycznego.

Uwzględniając zależności (56) i (57), równanie (55) można napisać w następującej postaci:

$$dW = T dS + H dJ. \quad (58)$$

Równanie (58) jest słuszne, gdy zmiany ograniczone są do ciepła i magnetyzacji. W ciałach magnetostrykcyjnych występuje ponadto zmiana wymiarów, a więc dochodzi człon wyrażający energię dostarczoną na zmiany sprężyste. W założeniach wstępnych (p. 5) przyjęto, że odkształcenia (współczynnik magnetostrykcji) pod wpływem naprężeń normalnych rozciągających (σ dodatnie) występują w kierunku wzdłużnym, przy czym wzrastająca długość próbki jest znacznie większa od wymiarów poprzecznych i praktycznie przekrój poprzeczny nie ulega zmianie. Na podstawie tego założenia można napisać:

$$dW = T dS + H dJ + \sigma d\varepsilon, \quad (59)$$

przy czym wzór (59) słuszny jest również przy ściskaniu (σ ujemne), gdy długość próbki ulega skróceniu (ε ujemne). Potencjał termodynamiczny φ określony jest zależnością:

$$\varphi = F - HJ - \sigma\varepsilon, \quad (60)$$

gdzie:

HJ — energia wymiany między polem a materiałem,

$\sigma\varepsilon$ — energia mechaniczna sprężystości,

F — energia swobodna, zdefiniowana wzorem:

$$F = W - TS, \quad (61)$$

a więc

$$\varphi = W - TS - HJ - \sigma\varepsilon. \quad (62)$$

Różniczkę zupełną potencjału termodynamicznego po uwzględnieniu zależności (59) można napisać:

$$\begin{aligned} d\varphi = T dS + H dJ + \sigma d\varepsilon - S dT - T dS - H dJ - J dH - \\ - \sigma d\varepsilon - \varepsilon d\sigma = -S dT - J dH - \varepsilon d\sigma. \end{aligned} \quad (63)$$

Pochodna cząstkowa potencjału φ względem natężenia pola przy stałych naprężeniach i temperaturze będzie:

$$\left(\frac{\partial \varphi}{\partial H} \right)_{\sigma, T} = -J. \quad (64)$$

Natomiast pochodna cząstkowa potencjału φ względem naprężeń normalnych przy stałych: T i H będzie:

$$\left(\frac{\partial \varphi}{\partial \sigma} \right)_{T, H} = -\varepsilon. \quad (65)$$

Obecnie wykorzystuje się znaną zależność matematyczną, że jeżeli wyrażenie jest różniczką zupełną pewnych zmiennych, to jej drugie pochodne mieszane względem tych wielkości są sobie równe:

$$\frac{\partial^2 \varphi}{\partial H \partial \sigma} = - \left(\frac{\partial J}{\partial \sigma} \right)_{H,T} = - \left(\frac{\partial B}{\partial \sigma} \right)_{H,T} \quad (66)$$

$$\frac{\partial^2 \varphi}{\partial \sigma \partial H} = - \left(\frac{\partial \varepsilon}{\partial H} \right)_{\sigma,T} = - \frac{1}{l} \left(\frac{\partial (\Delta l)}{\partial H} \right)_{\sigma,T}, \quad (67)$$

a więc

$$\left(\frac{\partial B}{\partial \sigma} \right)_{H,T} = \left(\frac{\partial \varepsilon}{\partial H} \right)_{\sigma,T}. \quad (68)$$

Wynika z tego, że

$$d' = \left(\frac{\partial \varepsilon}{\partial H} \right)_{\sigma} = \left(\frac{\partial B}{\partial \sigma} \right)_H = d'' = d, \quad (69)$$

gdzie d nazywany jest współczynnikiem magnetostrykcji odwracalnej lub czułością naprężeń. Należy zaznaczyć, że zależności te odnoszą się do zmian odwracalnych i są słuszne w ograniczonym zakresie σ i H .

W układzie CGSM

$$\left(\frac{\partial \varepsilon}{\partial H} \right)_{\sigma} = \left(\frac{\partial J}{\partial \sigma} \right)_H = \frac{1}{4\pi} \left(\frac{\partial B}{\partial \sigma} \right)_H = \frac{1}{\pi 4} d, \quad (70)$$

Przy rozważaniach ogólniejszych zamiast naprężeń normalnych σ i odkształceń wzdlużnych ε należy rozpatrywać zmiany objętości V i ciśnienia p i wtedy zależność (68) przyjmie postać:

$$\left(\frac{\partial B}{\partial p} \right)_{H,T} = - \frac{1}{V} \left(\frac{\partial V}{\partial H} \right)_{p,T}. \quad (71)$$

Uwzględniając określenia (51÷54) i zależność (69) oraz w oparciu o założenia, że rozpatrywane zmiany są bardzo małe, można pomijając znak różniczek napisać ostatecznie układ równań I B (49 i 50) w sposób następujący:

$$\left. \begin{aligned} \varepsilon &= \frac{\sigma}{E_H} + d H \\ B &= d \sigma + \mu_0 H \end{aligned} \right\} \text{ I B,} \quad (72)$$

$$(73)$$

gdzie:

- ε — magnetostrykcja (względne zmiany długości),
- σ — naprężenia normalne w N/m²,
- E_H — dynamiczny moduł sprężystości przy stałym natężeniu pola w N/m²,
- d — dynamiczny współczynnik magnetostrykcji odwracalnej (czułość naprężeń) w Wb/N,
- H — natężenie pola magnetycznego w A/m,

B — indukcja magnetyczna w Wb/m^2 ,
 μ_σ — przenikalność magnetyczna przy stałych naprężeniach (próbki swobodnej) w H/m .

Współczynniki układu równań I B są ściśle związane ze współczynnikiem sprzężenia magnetomechanicznego k . Przyjmując, że próbka jest sztywno zamocowana ($\varepsilon = 0$ równanie (72) będzie równe zero:

$$\varepsilon = \frac{\sigma}{E_H} + dH = 0, \quad (74)$$

stąd

$$\sigma = -dE_H H, \quad (75)$$

a

$$B_\varepsilon = (-d^2 E_H + \mu_\sigma) H, \quad (76)$$

więc

$$\mu_\sigma - \mu_\varepsilon = d^2 E_H. \quad (77)$$

Ponieważ współczynnik sprzężenia magnetomechanicznego

$$k^2 = 1 - \frac{\mu_\varepsilon}{\mu_\sigma}, \quad (78)$$

ostatecznie można napisać, że:

$$k^2 = \frac{d^2 E_H}{\mu_\sigma}. \quad (79)$$

Z równań I B (72) i (73) wynikają następujące wnioski: odkształcony materiał, dla którego słuszne jest prawo Hooke'a [wzór (21)], po umieszczeniu w polu magnetycznym przy niezmiennych naprężeniach zewnętrznych (rozciągających lub ściskających) ulegnie dalszemu odkształceniu, przy czym kierunek tych odkształceń zależny jest jedynie od znaku dynamicznego współczynnika magnetostrykcji d , gdyż współczynnik magnetostrykcji ε jest funkcją parzystą wielkości magnetycznych (tzn. zmiana znaku B lub H nie powoduje zmiany znaku ε). Znak współczynnika d określa, czy dany materiał ma współczynnik magnetostrykcji lub krócej magnetostrykcję ε dodatnią czy ujemną. Jeżeli nie ma naprężeń ($\sigma = 0$), to w przypadku wzrostu pola wartość bezwzględna magnetostrykcji rośnie, przy czym ε jest ujemna, gdy czułość d jest ujemna, lub dodatnia w przypadku, gdy czułość d jest dodatnia.

Z równania (73) wynika, że próbka magnesowana, dla której jest słuszna zależność (1), poddana naprężeniom w przypadku rozciągania i magnetostrykcji dodatniej ($d > 0$), np. permalój 50% lub magnetyt), będzie miała większą indukcję, a więc i przenikalność przy niezmiennym natężeniu pola H . To samo będzie w przypadku ściskania i magnetostrykcji ujemnej.

W układzie jednostek GGSM układ równań I B można napisać w dwóch wersjach:

$$\varepsilon = \frac{\sigma}{E_H} + \frac{A}{4\pi} H, \quad (80)$$

$$B = A \sigma + \mu_\sigma H, \quad (81)$$

gdzie:

$$A = \left(\frac{\partial B}{\partial \sigma} \right)_H = 4\pi \left(\frac{\partial \varepsilon}{\partial H} \right)_\sigma, \quad (82)$$

$$k^2 = \frac{A^2 E_H}{4\pi \mu_\sigma}, \quad (83)$$

lub

$$\varepsilon = \frac{\sigma}{E_H} + d_o H, \quad (84)$$

$$B = 4\pi d_o \sigma + \mu_\sigma H, \quad (85)$$

gdzie:

$$d_o = \frac{1}{4\pi} \left(\frac{\partial B}{\partial \sigma} \right)_H = \left(\frac{\partial \varepsilon}{\partial H} \right)_\sigma, \quad (86)$$

$$k^2 = \frac{4\pi d_o^2 E_H}{\mu_\sigma}, \quad (87)$$

przy czym wersją pierwszą (rów. 80—83) posłużono się w pracach [22], [23], [24], [35], [42], a wersją drugą (rów. 84—87) w pracach [40], [41]. Współczynnik A w wersji 1 wyrażony w 10^{-6} Gs cm²/dyn odpowiada współczynnikowi d w układzie MKSA z wyrażenemu w 10^{-9} Wb/N.

5.2. Inne układy równań magnetostrykcyjnych

Obok układu równań magnetostrykcyjnych I B można spotkać w literaturze [4] układ I J, który po podobnych przekształceniach jak w przypadku równań I B i po opuszczeniu znaków różniczek, przy założeniu, że rozpatrywane wielkości zmienne są dostatecznie małe, aby procesy były odwracalne, można napisać w ostatecznej postaci:

$$\varepsilon = \frac{1}{E_H} \sigma + d H, \quad \left. \varepsilon = \frac{1}{E_H} \sigma + d H, \right\} \text{ I J} \quad (88)$$

$$J = d \sigma + \kappa_\sigma H, \quad (89)$$

gdzie:

$J = B - \mu_0 H$ — magnetyzacja w Wb/m²,

$\kappa_\sigma = \mu_\sigma - \mu_0$ — podatność magnetyczna próbki swobodnej ($\sigma_{zew} = 0$) w H/m,

pozostałe wielkości zdefiniowane jak w układzie I B.

Współczynnik k określa zależność:

$$k^2 = \frac{d^2 E_H}{\kappa_\sigma}. \quad (90)$$

Często w literaturze można spotkać współczynniki wyprowadzone z układu III B, np. [25]. Układ ten po uwzględnieniu zależności termodynamicznej

$$-\left(\frac{\partial \sigma}{\partial B}\right)_\varepsilon = -\left(\frac{\partial H}{\partial \varepsilon}\right)_B = h \quad (91)$$

i po opuszczeniu znaków różniczek przyjmie postać:

$$\sigma = E_B \varepsilon - h B, \quad \left. \begin{array}{l} \\ \\ \end{array} \right\} \text{III B} \quad (92)$$

$$H = -h \varepsilon + \frac{1}{\mu_e} B. \quad (93)$$

Współczynnik k określony będzie wzorem:

$$k^2 = \frac{h^2 \mu_e}{E_B}, \quad (94)$$

gdzie między modulem sprężystości przy stałym natężeniu pola E_H , znanym z układów I B i I J, i występującym tu modulem sprężystości przy stałej indukcji magnetycznej E_B istnieje związek:

$$E_H = E_B - h^2 \mu_e, \quad (95)$$

gdzie nowe wielkości to:

h — stała magnetostrykcji w N/Wb,

μ_e — odwracalna przenikalność magnetyczna przy stałych odkształceniach (magnetostrykcji) w H/m.

Układ równań magnetostrykcyjnych III B spotyka się najczęściej w literaturze w układzie jednostek CGSM [10], [35], [42] i wtedy:

$$\sigma = E_B \varepsilon - h B, \quad \left. \begin{array}{l} \\ \\ \end{array} \right\} \text{III Bc} \quad (96)$$

$$H = -4\pi h \varepsilon + \frac{1}{\mu_e} B \quad (97)$$

i

$$k^2 = \frac{4\pi h^2 \mu_e}{E_B}, \quad (98)$$

przy czym h jest najczęściej w tym przypadku oznaczane przez λ .

W podobny sposób można rozpisać pozostałe układy równań magnetostrykcyjnych i ustalić związki współczynników zawartych w równaniach ze współczynnikiem sprzężenia magnetostrykcyjnego k , np. układ II B zastosowano w publikacji [33].

W tablicy 2 podano odpowiednio przekształcone równania typu B w układzie MKSA zr oraz podano definicje matematyczne stosowanych współczynników.

W tablicy 3 podano wzory na wyznaczanie współczynników magnetostrykcyjnych w zależności od pozostałych współczynników pochodzących z różnych układów równań magnetostrykcyjnych typu B zebranych w ta-

Tablica 2

Układy równań magnetostrykcyjnych typu B dla drgań wzdłużnych, w których jako zmienne występują naprężenia normalne σ , odkształcenia względne (magnetostrykcja) ε , indukcja magnetyczna B i natężenie pola magnetycznego H. Układ jednostek MKSAzr. Dla drgań skrętnych w miejsce σ , ε , E_H i E_B należy wstawić odpowiednio τ , γ , G_H i G_B

Układ	Równania magnetostrykcyjne	Definicje współczynników	Związki między współczynnikami
I B	$d\varepsilon = \frac{1}{E_H} d\sigma + d d H \quad (1)$	$\frac{1}{E_H} = \left(\frac{\partial \varepsilon}{\partial \sigma} \right)_H \quad (3)$	$d = \left(\frac{\partial \varepsilon}{\partial H} \right)_\sigma = \left(\frac{\partial B}{\partial \sigma} \right)_H \quad (5)$
	$d B = d d \sigma + \mu_\sigma d H \quad (2)$	$\mu_\sigma = \left(\frac{\partial B}{\partial H} \right)_\sigma \quad (4)$	$k = \frac{d^2 E_H}{\mu_\sigma} \quad (6)$
II B	$d\varepsilon = \frac{1}{E_B} d\sigma + g d B \quad (7)$	$\frac{1}{E_B} = \left(\frac{\partial \varepsilon}{\partial \sigma} \right)_B \quad (9)$	$g = \left(\frac{\partial \varepsilon}{\partial B} \right)_\sigma = - \left(\frac{\partial H}{\partial \sigma} \right)_B \quad (11)$
	$d H = -g d \sigma + \frac{1}{\mu_\sigma} d B \quad (8)$	$\frac{1}{\mu_\sigma} = \left(\frac{\partial H}{\partial B} \right)_\sigma \quad (10)$	$k^2 = \frac{g^2 \mu_\sigma E_B}{1 + g^2 \mu_\sigma E_B} \quad (12)$
III B	$d\sigma = E_B d\varepsilon - h d B \quad (13)$	$E_B = \left(\frac{\partial \sigma}{\partial \varepsilon} \right)_B \quad (15)$	$h = - \left(\frac{\partial \sigma}{\partial B} \right)_\varepsilon = - \left(\frac{\partial H}{\partial \varepsilon} \right)_B \quad (17)$
	$d H = -h d \varepsilon + \frac{1}{\mu_\varepsilon} d B \quad (14)$	$\frac{1}{\mu_\varepsilon} = \left(\frac{\partial H}{\partial B} \right)_\varepsilon \quad (16)$	$k^2 = \frac{h^2 \mu_\varepsilon}{E_B} \quad (18)$
IV B	$d\sigma = E_H d\varepsilon - e d H \quad (19)$	$E_H = \left(\frac{\partial \sigma}{\partial \varepsilon} \right)_H \quad (21)$	$e = - \left(\frac{\partial \sigma}{\partial H} \right)_\varepsilon = \left(\frac{\partial B}{\partial \varepsilon} \right)_H \quad (23)$
	$d B = e d \varepsilon + \mu_\varepsilon d H \quad (20)$	$\mu_\varepsilon = \left(\frac{\partial B}{\partial H} \right)_\varepsilon \quad (22)$	$k^2 = \frac{e^2}{e^2 + \mu_\varepsilon E_H} \quad (24)$

blicy 2. Podano tu tylko określenia danego współczynnika magnetostrykcyjnego w zależności od dwóch innych, a jedynie w przypadku współczynnika sprzężenia magnetomechanicznego k podano dodatkowo niektóre związki określające k w zależności od trzech innych współczynników. Ponadto podano określenia iloczynów dwóch współczynników magnetostrykcyjnych za pomocą pozostałych wielkości. Z zależności podanych w tablicach 2 i 3 można w razie potrzeby wyprowadzić pozostałe związki. Wystarczy znać komplet współczynników występujących w danym układzie równań, aby z zależności podanych w tablicy 3 można było wyznaczyć wszystkie pozostałe współczynniki magnetostrykcyjne występujące w innych układach równań.

Obok równań typu B spotyka się w literaturze [4] równania typu J, które są dogodniejsze przy uwzględnianiu zjawiska odmagnesowania występującego w przetwornikach nie stanowiących zamkniętego obwodu magnetycznego. Warianty tych równań wraz z odpowiednimi definicjami matematycznymi podano w tablicy 4.

Moduł sprężystości przy stałym natężeniu pola E_H jest identyczny z E_H z układu równań typu B, natomiast E_J różni się nieco od E_B .

Przenikalność względna z podatnością względną związana jest za pomocą wzoru (9). Współczynniki magnetostrykcyjne d i e są identyczne,

Tablica 3

Zależności między współczynnikami równań magnetostrykcyjnych typu B podanych w tablicy 2

Nazwa współczynnika	Sym- bol	Jednost- ka	Wzory określające		
1	2	3	4	5	6
Współczynnik indukcji naprężeniowej (czułość naprężeń, współcz. magnetostrykcji odwracalnej)	d	Wb/N	$\frac{e}{E_H}$	$\mu_\sigma g$	$\frac{k^2}{(1-k^2)h}$
Współczynnik naprężeń polowych (indukcji odkształceniowej)	e	Wb/m ²	$d E_H$	$\mu_\epsilon h$	$\frac{k^2}{e}$
Współczynnik odkształceń indukcyjnych	g	m ² /Wb	$\frac{h}{E_B}$	$\frac{d}{\mu_\sigma}$	$\frac{k^2}{e}$
Stała magnetostrykcji	h	N/Wb	$g E_B$	$\frac{e}{\mu_\epsilon}$	$\frac{k^2}{(1-k^2)d}$
Moduł sprężystości wzdłużnej przy stałej indukcji	E_B	N/m ²	$\frac{h}{g}$	$\frac{E_H}{1-k^2}$	
Moduł sprężystości wzdłużnej przy stałym natężeniu pola	E_H	N/m ²	$\frac{e}{d}$	$(1-k^2)E_B$	
Magnetyczna przenikalność odwracalna przy stałych odkształceniach	μ	H/m	$\frac{e}{h}$	$(1-k^2)\mu_\sigma$	
Magnetyczna przenikalność odwracalna przy stałych naprężeniach	μ_σ	H/m	$\frac{d}{g}$	$\frac{\mu_\epsilon}{1-k^2}$	
Współczynnik sprzężenia magneto- mechanicznego	k^2	—	$1 - \frac{\mu_\epsilon}{\mu_\sigma}$	$1 - \frac{E_H}{E_B}$	eg
			$\frac{dh}{1+dh}$	$\frac{eh}{E_B}$	$\frac{e \cdot d}{\mu_\sigma}$
			$gh \mu_\sigma$	$g^2 \mu_\epsilon E_B$	$\frac{e^2}{\mu_\sigma E_H}$
			oraz wzory: 6, 12, 18 i 24 z tablicy 2		
$de = \mu_\sigma - \mu_\epsilon$	$dh = \frac{k^2}{1-k^2}$		$eh = E_B - E_H$		
$dg = \frac{1}{E_H} - \frac{1}{E_B}$	$eg = k^2$		$gh = \frac{1}{\mu_\epsilon} - \frac{1}{\mu_\sigma}$		

natomiast współczynnik sprzężenia magnetomechaniczne g i h są odmiennie zdefiniowane, przy czym wzajemny związek między nimi jest następujący:

$$f B(B) = f(J) \cdot \frac{1}{1 + \frac{\mu_0}{\kappa}} \quad (99)$$

Tablica 4

Układy równań magnetostrykcyjnych typu J dla drgań wzdłużnych, w których jako zmienne występują: naprężenia normalne σ , odkształcenia względne (magnetostrykcja) ε , magnetyzacja J i natężenie pola magnetycznego H . Układ jednostek MKSA_{zr}. Dla drgań skrętnych w miejscach ε i σ należy podstawić γ i τ , a w miejsce E_H i E_J — G_H i G_J

Układ	Równania magnetostrykcyjne		Definicje współczynników		Związki między współczynnikami										
I J	$d\varepsilon = \frac{1}{E_H} d\sigma + d dH$	(1)	$\frac{1}{E_H} = \left(\frac{\partial \varepsilon}{\partial \sigma}\right)_H$	(3)	$d = \left(\frac{\partial \varepsilon}{\partial H}\right)_\sigma = \left(\frac{\partial J}{\partial \sigma}\right)_H$	(5)									
	$dJ = d d\sigma + \kappa_\sigma dH$	(2)	$\kappa_\sigma = \left(\frac{\partial J}{\partial H}\right)_\sigma$	(4)	$k^2 = \frac{d^2 E_H}{\kappa_\sigma}$	(6)									
II J	$d\varepsilon = \frac{1}{E_J} d\sigma + g dJ$	(7)	$\frac{1}{E_J} = \left(\frac{\partial \varepsilon}{\partial \sigma}\right)_J$	(9)	$g = \left(\frac{\partial \varepsilon}{\partial J}\right)_\sigma = -\left(\frac{\partial H}{\partial \sigma}\right)_J$	(11)									
	$dH = -g d\sigma + \frac{1}{\kappa_\sigma} dJ$	(8)	$\frac{1}{\kappa_\sigma} = \left(\frac{\partial H}{\partial J}\right)_\sigma$	(10)	$k^2 = \frac{g^2 \kappa_\sigma E_J}{1 + g^2 \kappa_\sigma E_J}$	(12)									
III J	$d\sigma = E_J d\varepsilon - h dJ$	(13)	$E_J = \left(\frac{\partial \sigma}{\partial \varepsilon}\right)_J$	(15)	$h = -\left(\frac{\partial \sigma}{\partial J}\right)_\varepsilon = -\left(\frac{\partial H}{\partial \varepsilon}\right)_J$	(17)									
	$dH = -h d\varepsilon + \frac{1}{\kappa_\varepsilon} dJ$	(14)	$\frac{1}{\kappa_\varepsilon} = \left(\frac{\partial H}{\partial J}\right)_\varepsilon$	(16)	$k^2 = \frac{h^2 \kappa_\varepsilon}{E_J}$	(18)									
IV J	$d\sigma = E_H d\varepsilon - e d h$	(19)	$E_H = \left(\frac{\partial \sigma}{\partial \varepsilon}\right)_H$	(21)	$e = -\left(\frac{\partial \sigma}{\partial H}\right)_\varepsilon = \left(\frac{\partial J}{\partial \varepsilon}\right)_H$	(23)									
	$dJ = e d\varepsilon + \kappa_\varepsilon dH$	(20)	$\kappa_\varepsilon = \left(\frac{\partial J}{\partial H}\right)_\varepsilon$	(22)	$k^2 = \frac{e^2}{e^2 + \kappa_\varepsilon E_H}$	(24)									
Związki między współczynnikami magnetostrykcyjnymi różnych układów															
$h = g E_J = \frac{e}{\kappa_\varepsilon} = \frac{d E_H}{\kappa_\sigma}$			(25)				$k^2 = 1 - \frac{E_H}{E_J} = 1 - \frac{\kappa_\varepsilon}{\kappa_\sigma} = e g = \frac{d h}{1 + d h}$			(26)					
$d e = \kappa_\sigma - \kappa_\varepsilon$		(27)		$d g = \frac{1}{E_H} - \frac{1}{E_J}$		(28)		$e h = E_B - E_H$		(29)		$g h = \frac{1}{\mu_\varepsilon} - \frac{1}{\mu_\sigma}$		(30)	
Jednostki															
Wielkość	H	J	$\kappa_\sigma, \kappa_\varepsilon$	σ, E_H, E_B	ε, k	d	e	g	h						
Jednostka	A/m	Wb/m ²	H/m	N/m ²	—	Wb/N	Wb/m ²	m ² /Wb	N/Wb						

gdzie:

$f(B)$ — odpowiednie współczynniki (k , g lub h) w równaniach typu B ,

$f(J)$ — odpowiednie współczynniki (k , g lub h) w równaniach typu J ,

μ_0 — przenikalność próżni = 10^{-7} H/m,

κ — podatność: przy stałych naprężeniach κ_σ lub przy stałych wydłużeniach κ_ε , zależnie od określanej wielkości.

Szczegółowe wzory będą podane poniżej (tablica 6).

W tablicy 5 podano odpowiednio równania typu B w układzie CGSM (oznaczone przez B_c), przy czym współczynniki są tak zdefiniowane, że

Tablica 5

Układy równań magnetostrykcyjnych typu B_c dla drgań wzdłużnych, w których jako zmienne występują: naprężenia normalne σ , odkształcenia względne (magnetostrykcja) ε , indukcja magnetyczna B i natężenie pola magnetycznego H . Układ jednostek CGSM. Dla drgań skrętnych w miejsce σ , ε , E_H i E_B należy wstawić odpowiednio τ , γ , G_H i G_B

Układ	Równania magnetostrykcyjne	Definicje współczynników	Związki między współczynnikami						
I B _c	$d\varepsilon = \frac{1}{E_H} d\sigma + \frac{d}{4\pi} dH \quad (1)$	$\frac{1}{E_H} = \left(\frac{\partial \varepsilon}{\partial \sigma} \right)_H \quad (3)$	$d = 4\pi \left(\frac{\partial \varepsilon}{\partial H} \right)_\sigma = \left(\frac{\partial B}{\partial \sigma} \right)_H \quad (5)$						
	$dB = d\varepsilon + \mu_\sigma dH \quad (2)$	$\mu_\sigma = \left(\frac{\partial B}{\partial H} \right)_\sigma \quad (4)$	$k^2 = \frac{d^2 E_H}{4\pi \mu_\sigma} \quad (6)$						
II B _c	$d\varepsilon = \frac{1}{E_B} d\sigma + g dB \quad (7)$	$\frac{1}{E_B} = \left(\frac{\partial \varepsilon}{\partial \sigma} \right)_B \quad (9)$	$g = \left(\frac{\partial \varepsilon}{\partial B} \right)_\sigma = -\frac{1}{4\pi} \left(\frac{\partial H}{\partial \sigma} \right)_B \quad (11)$						
	$dH = -4\pi g d\sigma + \frac{1}{\mu_\sigma} dB \quad (8)$	$\frac{1}{\mu_\sigma} = \left(\frac{\partial H}{\partial B} \right)_\sigma \quad (10)$	$k^2 = \frac{4\pi \mu_\sigma g^2 E_B}{1 + 4\pi \mu_\sigma g^2 E_B} \quad (12)$						
III B _c	$d\sigma = E_B d\varepsilon - h dB \quad (13)$	$E_B = \left(\frac{\partial \sigma}{\partial \varepsilon} \right)_B \quad (15)$	$h = -\left(\frac{\partial \sigma}{\partial B} \right)_\varepsilon = -\frac{1}{4\pi} \left(\frac{\partial H}{\partial \varepsilon} \right)_B \quad (17)$						
	$dH = -4\pi h d\varepsilon + \frac{1}{\mu_\varepsilon} dB \quad (14)$	$\frac{1}{\mu_\varepsilon} = \left(\frac{\partial H}{\partial B} \right)_\varepsilon \quad (16)$	$k^2 = \frac{4\pi \mu_\varepsilon h^2}{E_B} \quad (18)$						
IV B _c	$d\sigma = E_H d\varepsilon - \frac{e}{4\pi} dH \quad (19)$	$E_H = \left(\frac{\partial \sigma}{\partial \varepsilon} \right)_H \quad (21)$	$e = -4\pi \left(\frac{\partial \sigma}{\partial H} \right)_\varepsilon = \left(\frac{\partial B}{\partial \varepsilon} \right)_H \quad (23)$						
	$dB = e d\varepsilon + \mu_\varepsilon dH \quad (20)$	$\mu_\varepsilon = \left(\frac{\partial B}{\partial H} \right)_\varepsilon \quad (22)$	$k^2 = \frac{e^2}{e^2 + 4\pi \mu_\varepsilon E_H} \quad (24)$						
Związki między współczynnikami magnetostrykcyjnymi różnych układów									
$h = g E_B = \frac{e}{4\pi \mu} = \frac{d E_H}{4\pi \mu_\varepsilon} \quad (25)$		$k^2 = 1 - \frac{\mu_\varepsilon}{\mu_\sigma} = 1 - \frac{E_H}{E_B} = e g = \frac{d h}{1 + d h} \quad (26)$							
$d e = \mu_\sigma - \mu_\varepsilon \quad (27)$	$d g = \frac{1}{E_H} - \frac{1}{E_B} \quad (28)$	$e h = E_B - E_H \quad (29)$	$g h = \frac{1}{\mu_\varepsilon} - \frac{1}{\mu_\sigma} \quad (30)$						
Jednostki									
Wielkość	H	B	μ _ε , μ _σ	σ, E _B , E _H	ε, k	d	e	g	h
Jednostka	Oe	Gs	Gs/Oe	$\frac{\text{dyn}}{\text{cm}^2}$	—	$\frac{\text{Gs cm}^2}{\text{dyn}}$	Gs	1/Gs	$\frac{\text{dyn}}{\text{cm}^2 \text{Gs}}$

odpowiadają współczynnikom układu MKSA zr pomnożonym przez 10^9 , gdzie α wynika z doboru jednostek. Równania magnetostrykcyjne typu J_c , tzn. typu J w układzie CGSM, są rzadko stosowane, więc, aby nie zaćmiewać zagadnienia, nie będą tu przytaczane. Niektóre współczynniki równań tego typu spotyka się w literaturze japońskiej [26], [27].

Należy zaznaczyć, że podane w tablicach 2, 4 i 5 układy równań mogą być stosowane przy drganiach skrętnych przez podstawienie w miejsce σ — naprężeń ścinających τ , w miejsce ε — kątów odkształceń postaciowych γ

i w miejsce modułów sprężystości liniowych E_H i E_B lub (E_J) — modułów sprężystości postaciowej G_H i G_B (lub G_J).

W zasadzie najbardziej wierne z fizycznego punktu widzenia są układy równań typu J, które opisują materiał magnetostrykcyjny z wydzieleniem zjawisk zachodzących w powietrzu (lub w próżni). W praktyce zjawiska w materiale i w powietrzu występują łącznie i przy rozpatrywaniu z technicznego punktu widzenia nie ma potrzeby ich wydzielenia, gdyż wygodniej jest posługiwać się indukcją zamiast magnetyzacją, a co za tym idzie układem równań typu B.

Zdefiniowanie energii magnetycznej $\left(\frac{1}{2}BH\right)$ jako energii zmagazynowanej łącznie w materiale i próżni z praktycznego punktu widzenia jest poprawniejsze, gdyż o ile można mówić o energii zmagazynowanej w cewce $\left(\frac{1}{2}\mu_0 H^2\right)$ umieszczonej w próżni czy powietrzu, to nie można mówić o samej energii zmagazynowanej w materiale magnesowanym za pomocą solenoidu $\left(\frac{1}{2}JH\right)$ zaniedbując energię zmagazynowaną w samym solenoidzie powietrznym $\left(\frac{1}{2}\mu_0 H^2\right)$, a z kolei uwzględnianie tej ostatniej komplikuje przeliczenia. Dlatego wygodniej jest określać energię w zależności od indukcji B , chociaż wielkość jej charakteryzuje wtedy konkretny przetwornik, a nie materiał. Przyjęcie magnetyzacji J przy określaniu energii właściwej jest powodem rozbieżności w definicjach stosowanych w układzie równań typu B i typu J. I tak współczynnik sprzężenia magnetomechanicznego w pierwszym przypadku określony jest z zależności:

$$k^2(B) = 1 - \frac{\mu_e}{\mu_\sigma} = 1 - \frac{E_H}{E_B}, \quad (100)$$

a w drugim:

$$k^2(J) = 1 - \frac{\kappa_\epsilon}{\kappa_\sigma} = 1 - \frac{E_H}{E_J}. \quad (101)$$

Stąd:

$$k^2(B) = \frac{k^2(J)}{1 + \frac{\mu_0}{\kappa_\sigma}} \quad (102)$$

$$E_B = \frac{1 + \frac{\mu_0}{\kappa_\sigma}}{1 + \frac{\mu_0}{\kappa_\epsilon}} E_J. \quad (103)$$

Odpowiednio również różnią się niektóre współczynniki magnetomechaniczne (g i h). W tablicy 6 podano odpowiednie wzory przeliczeniowe.

Tablica 6

Związki między współczynnikami i zmiennymi równań typu B , J i B_c . Obok współczynników podano jednostki, w jakich należy podawać wartości danego współczynnika, aby zależności były słuszne zarówno wielkościowo, jak i ilościowo

MKSA _{zr}				CGSM		
B		J		B_c		
Wielkość	Jednostki	Wyrażenie równoważne	Jednostki	Wielkość odpowiadająca	Jednostki	Mnożnik przy przejściu do MKSA _{zr} *)
1	2	3	4	5	7	7
$d(B)$	Wb/N	$d(J)$	Wb/N	$d(B_c)$	$\frac{\text{Gs cm}^2}{\text{dyn}}$	10^{-3}
$e(B)$	Wb/m ²	$e(J)$	Wb/m ²	$e(B_c)$	Gs	10^{-4}
$g(B)$	m ² /Wb	$\frac{g(J)}{1 + \frac{\mu_0}{\mu_\sigma}}$	m ² /Wb	$g(B_c)$	1/Gs	10^4
$h(B)$	N/Wb	$\frac{h(J)}{1 + \frac{\mu_0}{\mu_\epsilon}}$	N/Wb	$h(B_c)$	$\frac{\text{dyn}}{\text{cm}^2 \text{ Gs}}$	10^3
E_B	N/m ²	$E_J \frac{1 + \frac{\mu_0}{\mu_\sigma}}{1 + \frac{\mu_0}{\mu_\epsilon}}$	N/m ²	E_B	$\frac{\text{dyn}}{\text{cm}^2}$	10^{-1}
E_H	N/m ²	E_H	N/m ²	E_H	$\frac{\text{dyn}}{\text{cm}^2}$	10^{-1}
μ_ϵ	H/m	$\mu_\epsilon + \mu_0$	H/m	$4\pi \mu_\epsilon$	Gs/Oe	10^{-7}
μ_σ	H/m	$\mu_\sigma + \mu_0$	H/m	$4\pi \mu_\sigma$	Gs/Oe	10^{-7}
$k(B)$	—	$\frac{k(J)}{1 + \frac{\mu_0}{\mu_\sigma}}$	—	k	—	1
B	Wb/m ²	$J + \mu_0 H$	Wb/m ²	B	Gs	10^{-4}
H	A/m	H	A/m	$\frac{1}{4\pi} H$	Oe	10^3
σ	N/m ²	σ	N/m ²	σ	$\frac{\text{dyn}}{\text{cm}^2}$	10^{-1}
ϵ	—	ϵ	—	ϵ	—	1

*) Np. gdy w CGSM $B = 15\,000 \text{ Gs}$, to w MKSA_{zr} $B = 15\,000 \cdot 10^{-4} \text{ Wb/m}^2 = 1,5 \text{ Wb/m}^2$, lub $d = 5 \frac{\text{Gs cm}^2}{\text{dyn}}$, to w MKSA_{zr} $d = 5 \cdot 10^{-6} 10^{-3} \text{ Wb/N} = 5 \cdot 10^{-9} \text{ Wb/N}$ lub $\mu_\sigma = 200 \text{ Gs/Oe}$, to w MKSA_{zr} $\mu_\sigma = 4\pi \cdot 200 \cdot 10^{-7} \text{ H/m} = 2,51 \cdot 10^{-4} \text{ H/m}$, a $\mu'_\sigma = \frac{\mu_\sigma}{\mu_0} = 200$.

Niezmienniczy pozostaje moduł sprężystości przy stałym natężeniu pola E_H , czułość naprężeń d i współczynnik e . Należy zaznaczyć, że dla każdego

układu równań magnetostrykcyjnych zależności wyprowadzone z tożsamości termodynamicznych potwierdzają wzór (102) na współczynnik sprzężenia magnetomechanicznego k . W tablicy 6 w piątej kolumnie podano odpowiednie wzory przeliczeniowe przy przechodzeniu z układu MKSA *zr* do układu jednostek CGSM niezracjonalizowanego. Obok definiowanych współczynników podano w kolumnie 2 ich jednostki (w MKSA *zr*), a w kolumnie 6 i 7 jednostki z układu CGSM i mnożnik, jaki należy wprowadzić, aby przy podstawianiu wartości liczbowe danej wielkości w danych jednostkach były równe wartościom liczbowym wielkości z układu MKSA *zr*. U dołu tablicy 6 podano kilka przykładów przeliczeniowych wyjaśniających sposób posługiwania się wielkościami podanymi w tablicy.

Wartości E_B i E_J w praktyce są identyczne. Wartości k , g i h dla materiałów o dużej przenikalności, np. ferrytów Ni-Zn o dużej zawartości cynku, również można wprost przenosić z układu typu B do układu typu J. W przypadkach ferrytów niklowych, dla których przenikalność względna spada poniżej 20, błąd przy porównywaniu k , g lub h przy zaniedbaniu współczynników przeliczeniowych jest już większy od 5%.

6. KILKA UWAG O PRAKTYCZNYM WYKORZYSTANIU OMAWIANYCH UKŁADÓW RÓWNAŃ MAGNETOSTRYKCYJNYCH

Przy wyprowadzaniu równań magnetostrykcyjnych, w celu uproszczenia rozważań, przyjęto pewne założenia. Często w praktyce założenia te nie są spełnione i wtedy zagadnienie nieco komplikuje się. Nie jest zamierzaniem autora wyczerpywać całkowicie to zagadnienie, gdyż brak miejsca nie pozwala na gruntowniejsze jego rozwinięcie.

Przy uwzględnianiu zjawiska samoodmagnesowania należy przy współczynnikach magnetostrykcyjnych wprowadzić odpowiednie mnożniki redukujące [4].

W przypadku niejednorodności materiału i rozpatrywaniu zjawisk zachodzących we wszystkich kierunkach należy wielkości podane w tablicach 2, 3, 4 i 5 zastąpić odpowiednio macierzami [25], [33], [39]. Na przykład dla rozważanego pręta o małym przekroju dla kierunku wzdłużnego słuszny jest następujący układ równań:

$$\varepsilon_3 = \frac{1}{E_{33B}} \sigma_3 + g_{33} B_3, \quad (104)$$

$$H_3 = -g_{33} \sigma_3 + \frac{1}{\mu_{33\sigma}} B_3. \quad (105)$$

Przy uwzględnianiu strat mechanicznych i magnetycznych należy uwzględnić kąt fazowy opóźnienia i wtedy omawiane równania przyjmą postać równań zespolonych [4], [25].

Mimo wymienionych zastrzeżeń ogólny charakter równań magnetostrykcyjnych pozostaje bez zmian, jedynie komplikują się przeliczenia matematyczne.

Poważne zmiany w postaci równań wprowadziłoby uwzględnienie zależności cieplnych. Wtedy należałoby uwzględnić wielką liczbę współczynników [30], [33], co często praktycznie jest niewykonalne.

W praktyce w wielu przypadkach można korzystać z podanych w części 5 równań, np. gdy badania wykonywane są na toroidach lub długich prętach, na materiale jednorodnym i przy małych stratach (np. w przypadku ferrytów).

Na zakończenie w tablicach 7, 8 i 9 podano zestawienia oznaczeń stosowanych przez różnych autorów oraz odpowiednie wzory przeliczeniowe dla jednostek układu MKSA *zr.* Autor uwzględnił większość publikacji polskich z dziedziny magnetostrykcji oraz najważniejsze (w Polsce dostępne) publikacje zagraniczne.

*Instytut Podstawowych Problemów Techniki PAN
Zakład Elektroniki*

WYKAZ LITERATURY

1. Bates L. F.: *Modern Magnetism*. — Cambridge University Press, III ed., 1951.
2. Biełow K. P.: *Упругие, тепловые и электрические явления в ферромагнитных материалах*. — GITTL, Moskwa 1951.
3. Bozorth R. M.: *Ferromagnetism*, III ed. — D. van Nostrand Company, New York 1955.
4. van der Burgt C. M.: *Dynamical Physical Parameters of the Magnetostrictive Excitation of Extensional and Torsional Vibrations in Ferrites*. — Philips Research Reports, vol. 8 (1953), 91—133.
5. van der Burgt C. M.: *Performance of Ceramic Ferrite Resonators as Transducers and Filter Elements*. — The Journal of the Acoustical Society of America, vol. 28, nr 6, 1020—32, Nov. 1956.
6. van der Burgt C. M.: *Ferroxcube Material for Piezomagnetic Vibrators*. — Philips Technical Review, vol. 18, 10, 285—98, 1956/7.
7. van der Burgt C. M.: *Controlled Crystal Anisotropy and Controlled Temperature Dependence of the Permeability and Elasticity of Various Cobalt-Substituted Ferrites*. — Philips Res. Repts, vol. 12, nr 2, 97—122, 1957.
8. van der Burgt C. M.: *Ferroxcube 7A1 and 7A2, New Ceramic Piezomagnetic Materials for Ultrasonic Power Transducers*. — Matronics, nr 15, 273—304, 1958.
9. van der Burgt C. M.: *Ferrites for Magnetic and Piezomagnetic Filter Elements with Temperature — Independent Permeability and Elasticity*. — The Proceedings of the Institution of Electrical Engineers, vol. 104, Part B, Supplement No 7, 550—7, 1957.
10. Clark C. A.: *The Dynamic Magnetostriction of Nickel-Cobalt Alloys*. — British Journal of Applied Physics, vol. 7, 10, 355—60, 1956.
11. Davis C. M. Jr., Ferebee S. F.: *Dynamic Magnetostrictive Properties of Alfenol*. — The Journal of the Acoustical Society of America, vol. 28, nr 2, 286—90, 1956.
12. Davis C. M. Jr., Ferebee S. F.: *Effect of Composition and Processing on the Activity of Some Magnetostrictive Materials*. — Journal of Applied Physics, suppl. to vol. 30, nr 4 (April), 1135—1155, 1959.

13. Davis C. M. Jr., Helms H. H., Ferebee S. F.: Dynamic Magnetostrictive Properties of Ni-Fe Alloys. — The Journal of the Acoustical Society of America, vol. 29, nr 4 (April), 431—4, 1957.
14. Ferebee S. F., Davis C. M. Jr.: Effect of Divalent Ion Substitutions on the Magnetomechanical Properties of Nickel Ferrite. — The Journal of the Acoustical Society of America, vol. 30, nr 8 (August), 747—50, 1958.
15. Filipczyński L.: Mechanism of Vibrations in the Magnetostrictive Transducer. — Proceeding of Vibration Problems, nr 1, 15—28, 1959.
16. Fryze S.: Racjonalizacja fizykalnych równań elektromagnetycznych i układów dymensyjnych. — Zeszyty Naukowe Politechniki Śląskiej, Elektryka nr 1, str. 7—26, 1954.
17. Golamina I. P.: Application of Ferrites to Electroacoustic Transducers. — Proceedings of the Conference on Electroacoustic Transducers, Krynica 1958, PWN, Warszawa 1961
18. Haskins J. F., Hickman J. S.: A Derivation and Tabulation of the Piezoelectric Equations of State. — The Journal of the Acoustical Society of America, vol. 22, nr 5, September, 584—8, 1950.
19. Henkel O.: Ein hochelastischer magnetostruktiver Werkstoff. — Nachrichtentechnik, 7 Jg., H. 8, 346—350, August, 1957.
20. Huber T.: Stereomechanika techniczna, cz. I. — Państwowe Zakłady Wydawnictw Szkolnych, Warszawa 1951.
21. Hueter T. F., Bolt R. H.: Sonics. — J. Wiley, New York 1955.
22. Kaczkowski Z.: Ferryty magnetostrykcyjne o dużej dobroci mechanicznej. — Rozprawy Elektrotechniczne, tom VI, z. 4, 1960, 451—474.
23. Kaczkowski Z., Smoliński A.: Ferryty magnetostrykcyjne. — Zeszyty Problemowe Nauki Polskiej, nr 20, 1960, 161—169.
24. Kaczkowski Z., Smoliński A.: Ferryty magnetostrykcyjne i ich zastosowania. — Rozpr. Elektrot., tom VI, z. 1/2, 1960, 77—105.
25. Katz H. W.: Solid State Magnetic and Dielectric Devices. — J. Wiley, New York 1959.
26. Kikuchi Y. i inni: Study on Ferrites for Use in Magnetostriction Vibrator, Part I, Ni-Zn Ferrite. — The Reports of the Research Institute of Electrical Communication Tohoku University, vol. 7, nr 1, 1—7 June, 1955.
27. Kikuchi Y. i inni: Study on Ferrites for Use in Magnetostriction Vibrator, Part II, Ni-Cu Ferrite. — The Reports of the Research Institute of Electrical Communication Tohoku University, vol. 7, nr 3, 171—8, Dec. 1955.
28. Konorski B.: Podstawy Elektrotechniki, tom I. — Trzaska, Evert i Michalski, Warszawa 1950.
29. Konorski B., Starczkow W., Wojciechowski S.: Równania i układy jednostek w elektrotechnice. — Państwowe Wydawnictwa Techniczne, Warszawa 1954.
30. Kwiek M.: Zjawisko magnetostrykcji i jego związki z termodynamiką oraz teorią elastyczności. — PAN, Instytut Fizyki, Sesja Naukowa Elektroniki Ciała Stałego, Czerwiec 1954.
31. Lee E. W.: Magnetostriction and Magnetomechanical Effect. — Reports on Progress in Physics, vol. 18, 184—229, 1955.
32. Mason W. P.: Electromechanical Transducers and Wave Filtres. — D. van Nostrand Co., New York 1949.
33. Mason W. P.: Physical Acoustics and the Properties of Solids. — D. van Nostrand Co., New York 1958.
34. Nodtvedt M.: Some Remarks on the Analysis of Magnetostrictive Transducers. — Acoustica, vol. 4, 432—8, 1954.

35. Popper P.: The Magnetostriction of Ferrites, Soft Magnetic Materials for Telecommunications. — C. E. Richards i A. C. Lynch, Pergamon Press Ltd., London 1953.
36. Root H. D., McDonald J.: Ceramic Ultrasonic Magnetostrictive Transducer. — Journal of American Ceramic Soc., vol. 30, nr 1, 1—5, 1957.
37. Round H. J.: Magnetostriction Transducer Measurements. — Wireless Engineer, vol. 29, nr 343 (April), 101—5, 1952.
38. Rozporządzenie Rady Ministrów z dnia 1 lipca 1953 r. w sprawie prawnie obowiązujących jednostek miar. — Dziennik Ustaw PRL, nr 35 z dnia 7.7.1953, poz. 148.
39. Schweizerhof S.: Über Ferrite für magnetostruktive Schwingen in Filterkreisen. — NTZ, Heft 4, 179—85, 1958.
40. Smoliński A.: Zjawisko magnetostrykcji i materiały magnetostrykcyjne. — Rozpr. Elektrot., tom V, nr 2, 211—36, 1959.
41. Smoliński A.: Magnetostriction and Magnetostrictive Materials. — Proceedings of the Conference on Electroacoustic Transducers, Krynica 1958. — PWN, 1961.
42. Sussman H., Erlich S. L.: Evaluation of the Magnetostrictive Properties of Hiperco. — The Journal of the Acoustical Society of America, vol. 22, nr 4, 499—506, 1950.
43. Suwalski R.: Wykorzystanie ferrytów magnetostrykcyjnych na przetworniki elektromechaniczne. — Zeszyty Problemowe Nauki Polskiej, nr 20, 1960, 247—257.
44. Suwalski R.: Obliczenie elektrycznego układu zastępczego obciążonego prętowego przetwornika magnetostrykcyjnego. — Archiwum Elektrotechniki, zeszyt 4, tom 7, 647—67, 1958.
45. Szczeniowski S.: Praktyczny układ elektromagnetycznych jednostek Giorgiego. — Postępy Fizyki, rok I, zeszyt 1-2, str. 1—20, 1949/50.
46. Thiele A. P.: Narrow Band Magnetostrictive Filtres Electronic. — Radio Engineer, 402—11, Nov. 1958.

3. КАЧКОВСКИ

МАГНИТОСТРИКЦИОННЫЕ УРАВНЕНИЯ И ИХ КОЭФФИЦИЕНТЫ

Резюме

Поданы основные магнитные и механические зависимости применяемые при рассматривании магнитомеханических явлений выступающих в магнитоstrictionонных материалах. Затем — выходя из положения, что эти процессы адиабатические (что подтверждается в большинстве случаев встречающихся в практике) поданы все возможные системы магнитоstrictionонных уравнений типа В и J. Уравнения типа В представлены в рационализированной системе единиц MKSA (таблица 2), а также в электромагнитной CGS несрационализированной (обозначенной CGSM) (таблица 5). В дальнейшей части работы разобраны зависимости между величинами входящими в склад уравнений типа В и J в системе MKSA рац (таблица 6). Поданы также связи между величинами типа В в системах MKSA рац и CGSM (таблица 6).

В заключению подан список наиболее часто применяемых магнитных (таблица 7), механических (таблица 8) и магнитомеханических (таблица 9) величин, а также их обозначение и их связи с величинами входящими в состав уравнений типа В в системе MKSA рац.

Z. KACZKOWSKI

MAGNETOSTRICTIVE EQUATIONS AND THEIR COEFFICIENTS

Summary

Basic magnetic and mechanic dependences applied at examining magneto-mechanic phenomena occurring in magnetostrictive materials have been given. Then, presuming, that conditions are adiabatic (which is right in majority of cases met with in practice) all possible systems of equations of the type B and J have been given. Equations of the type B have been presented in the rationalized system of units *MKSA* (meter, kilogramme, second, ampere) (Table 2) and in unrationalized electromagnetic system *CGS* (denoted as *CGSM*) — Table 5. In the further part of the article dependences between magnitudes of which equations of the type B consist, and magnitudes of equations of the type J in the system of units *MKSA r* — Table 6. Dependences between magnitudes of the type B in systems *MKSA r* and *CGSM* have been also given (Table 6).

At the end there is given a list of most often employed magnitudes: magnetic ones (Table 7), mechanic ones (Table 8) and magnetomechanic ones (Table 9), their denotations as well as their relationships to magnitudes of which consist equations of the type B in the system *MKSA r* (rationalized).

JAN KULIKOWSKI, MACIEJ NAŁĘCZ

Schematy zastępcze transduktorowych czujników magnetycznych o wyjściu impulsowym

Rękopis dostarczono 11. 7. 1960

W artykule niniejszym określony został schemat zastępczy transduktorowego czujnika magnetycznego-dwurdzeniowego. Podano również schematy zastępcze czujników typu mostkowego i dokonano porównania różnych konstrukcji czujników pod względem mocy na wyjściu.

1. WSTĘP

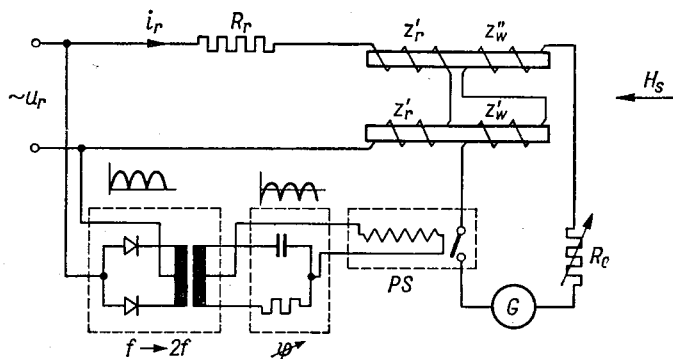
Transduktorowe czujniki magnetyczne o wyjściu impulsowym należą do najprostszych elementów mierzących stałe pola magnetyczne o natężeniu $10^{-3} - 100$ A/m. W artykule poprzednim [1] podano zależność czułości napięciowej od natężenia mierzonego pola magnetycznego oraz od rodzaju zasilania, bez uwzględnienia oporności obciążenia na wyjściu. Przeprowadzona poniżej analiza uwzględnia oporność obciążenia i prowadzi do określenia schematu zastępczego obwodu wyjściowego.

2. RÓWNANIA PODSTAWOWE

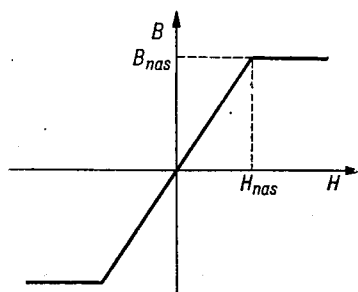
Dwurdzeniowy czujnik magnetyczny pokazany jest wraz z układem pomiarowym na rys. 1. Czujnik składa się z uzwojeń roboczych z'_r, z''_r oraz uzwojeń wyjściowych z'_w, z''_w połączonych przeciwsośnie w taki sposób, aby do obwodu wyjściowego nie transformowało się napięcie z obwodu roboczego, gdy składowa stała natężenia pola magnetycznego wzdłuż osi czujnika H_s wynosi zero. W obwodzie roboczym znajduje się duża oporność czynna, dzięki której przebieg prądu roboczego i_r jest sinusoidalnie zmienny w czasie:

$$i_r = I_m \sin \omega t. \quad (1)$$

Zakładamy, że uzwojenia nie mają indukcyjności rozproszeń oraz że charakterystyka magnesowania rdzeni jest taka jak na rys. 2. Pozostałe elementy schematu z rys. 1 to układ dyskryminatora fazy (prostownika fazowego), miernik magnetoelektryczny G oraz oporność obciążenia R w obwodzie wyjściowym. Jako prostownik fazowy użyty został przekaznik spolaryzowany PS , pobudzany przez podwajacz częstotliwości i przesuwnik fazowy, który umożliwia przepływ prądu w obwodzie wyjściowym tylko w określonych chwilach.



Rys. 1. Schemat układu pomiarowego z czujnikiem magnetycznym o wyjściu impulsowym



Rys. 2. Założona charakterystyka magnesowania rdzenia czujnika

Jak wiadomo [1], napięcie wyjściowe czujnika (a ściślej — napięcie indukowane w uzwojeniu wyjściowym czujnika) jest przebiegiem impulsowym (rys. 3d) o częstotliwości dwukrotnie większej niż częstotliwość źródła zasilania, przy czym impulsy nad osią czasu powstają w czasie nasycania rdzenia, a impulsy pod osią w czasie odsycania rdzeni. Można tak wyregulować prostownik fazowy, by jego styki zwarte były w chwilach pojawienia się impulsów nasycania, natomiast rozwarne w pozostałych momentach. Tylko wtedy, gdy styki są zamknięte, w obwodzie wyjściowym płynie prąd wyjściowy i_w . Obwód wyjściowy określony jest następującym układem równań:

$$z_w S \left(\frac{dB_1}{dt} - \frac{dB_2}{dt} \right) + (2r_w + R_o)i_w = 0, \quad (2)$$

$$B_1 = \mu_o \mu (h_r + H_s + h_w), \quad (3)$$

$$B_2 = \mu_o \mu (h_r - H_s - h_w), \quad (4)$$

gdzie:

B_1 — indukcja magnetyczna w pierwszym rdzeniu,

B_2 — indukcja magnetyczna w drugim rdzeniu,

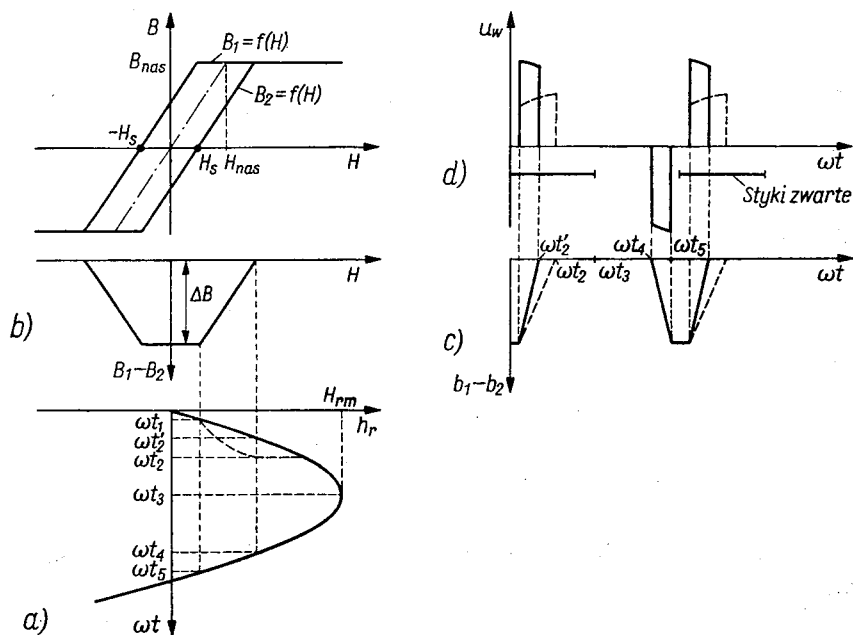
- z_w — liczba zwojów jednego z uzwojeń wyjściowych,
 S — przekrój rdzenia czujnika,
 r_w — oporność czynna obu uzwojeń wyjściowych (wraz z przewodami łączeniowymi),
 R_o — oporność obciążenia (wraz z opornością miernika),
 μ_0 — przenikalność magnetyczna próżni,
 μ — względna przenikalność magnetyczna obwodu,

oraz

$$h_r = \frac{i_r z_r}{l}, \quad (5)$$

$$h_w = \frac{i_w z_w}{l}, \quad (6)$$

gdzie: l — długość jednego z rdzeni czujnika. Na rys. 3 widać, że prąd



Rys. 3. Graficzny sposób określenia napięcia wyjściowego

roboczy i_r , który wytwarza pole magnetyczne h_r , powoduje okresowe nasycanie obu rdzeni.

Jak nadmieniono, rdzenie czujnika powinny być jednakowe, zatem prąd w obwodzie wyjściowym nie płynie, gdy oba rdzenie są nasycone. Ze względu na to, że czujnik umieszczony jest w stałym polu magnetycznym, jeden z rdzeni nasycony zostanie wcześniej — w momencie ωt_1

$$\omega t_1 = \arcsin \frac{H_{nas} - H_s}{H_{rm}}, \quad (7)$$

a wtedy $B_1 = B_{nas}$ i równanie (2) przekształca się do postaci

$$-z_w S \frac{dB_2}{dt} + (2r_w + R_0)i_w = 0, \quad (8)$$

a po podstawieniu (4) i uwzględnieniu (5), (6)

$$-\mu_o \mu \frac{z_w z_r S}{l} \frac{di_r}{dt} + \mu_o \mu \frac{z_w^2 S}{l} \frac{di_w}{dt} + (2r_w + R_0)i_w = 0. \quad (9)$$

Współczynniki występujące w równaniu (9) są odpowiednio: indukcyjnością wzajemną uzwojenia roboczego i wyjściowego

$$M = \mu_o \mu \frac{z_w^2 S}{l} = \frac{X_{rw}}{\omega} \quad (10)$$

oraz indukcyjnością własną uzwojenia wyjściowego

$$L_w = \mu_o \mu \frac{z_w^2 S}{l} = \frac{X_w}{\omega}. \quad (11)$$

Po wprowadzeniu tych oznaczeń i podstawieniu (1) otrzymuje się równanie

$$L_w \frac{di_w}{dt} + (2r_w + R_0)i_w = \omega M I_{rm} \cos \omega t, \quad (12)$$

którego rozwiązaniem jest

$$i_w = \frac{I_{rm} X_{rw}}{Z} \left[\cos(\omega t - \varphi) - \cos(\omega t_1 - \varphi) e^{\frac{\omega t - \omega t_1}{\text{tg } \varphi}} \right] \quad (13)$$

gdzie:

$$Z = \sqrt{(2r_w + R_0)^2 + X_w^2}, \quad (14)$$

$$\text{tg } \varphi = \frac{X_w}{2r_w + R_0}. \quad (15)$$

Po przedstawieniu do (4) i dokonaniu odpowiednich przekształceń otrzymuje się

$$\frac{B_2 + \mu_o \mu H_s}{H_{rm}} = \sin \omega t - \frac{X_{rw}}{Z} \left[\cos(\omega t - \varphi) - \cos(\omega t - \varphi) e^{\frac{\omega(t - t_1)}{\text{tg } \varphi}} \right] \quad (16)$$

Rdzeń drugi zostanie nasycony ($B_2 = B_{nas} = \mu_o \mu H_{nas}$) w momencie $\omega t_2 \leq \frac{\pi}{2}$.

Podstawiając $H_s = H_{nas}$ (wtedy $\omega t_1 = 0$), czyli największe natężenie mierzonego pola magnetycznego odpowiadające zakresowi teoretycznemu [1], otrzymuje się po przekształceniu

$$\frac{2H_{nas}}{H_{rm}} + \frac{X_{rw}}{Z} \left[\cos(\omega t_2 - \varphi) - \cos \varphi \cdot e^{-\frac{\omega t_1}{\text{tg } \varphi}} \right] = \sin \omega t_2. \quad (17)$$

Wyrażenie w nawiasie kwadratowym jest zawsze ułamkiem właściwym, który może mieć wartość bliską jedności dla małych wartości $(2r_w + R_o)$, gdyż $\varphi \rightarrow \frac{\pi}{2}$ i $\omega t_2 \rightarrow \frac{\pi}{2}$. Zachodzi wtedy przybliżony związek

$$\frac{X_{rw}}{Z} \approx \frac{X_{rw}}{X_w} = \frac{z_r}{z_w}, \quad (18)$$

a równanie (17) przyjmuje postać

$$\frac{2 H_{nas}}{H_{rm}} + \frac{z_r}{z_w} \approx \sin \omega t. \quad (19)$$

W skrajnym przypadku, gdy $z_r = z_w$, drugi składnik z lewej strony zależności (19) musi być znacznie mniejszy od jedności zatem:

$$H_{rm} \gg 2 H_{nas}. \quad (20)$$

W przypadku braku obciążenia warunek nasycenia drugiego rdzenia, a więc liniowej zależności między wartością średnią wyprostowanego napięcia wyjściowego a H_s , był mniej ostry [1]:

$$H_{rn} \gg 2 H_{nas}. \quad (20a)$$

Napięcie wyjściowe, określone wzorem

$$u_w = (2r_w + R_o)i_w = -z_w S \frac{d(B_1 - B_2)}{dt}, \quad (21)$$

posiada wartość średnią (po wyprostowaniu półokresowym)

$$U_w = -\frac{z_w S}{\pi} \int_{\omega t_1}^{\omega t_2} \frac{d(B_1 - B_2)}{dt} d\omega t = \frac{z_w S}{\pi} [B_2 - B_1]_{\omega t_1}^{\omega t_2} \quad (22)$$

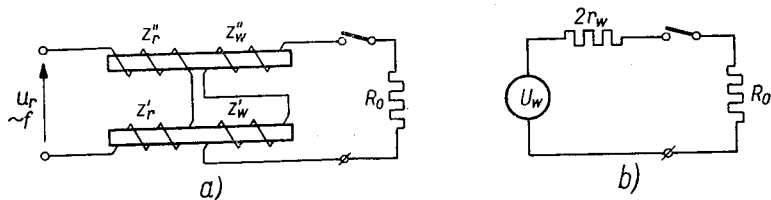
$$U_w = 2f z_w S \Delta B = 4f z_w S \mu_o \mu H_s. \quad (23)$$

Jak wykazano poprzednio [1], wartość średnia napięcia wyjściowego jest niezależna od rodzaju zasilania (sinusoidalny prąd roboczy czy też sinusoidalne napięcie robocze) oraz od niewielkich wahań napięcia roboczego.

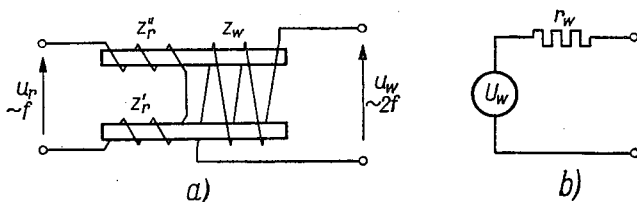
3. SCHEMAT ZASTĘPCZY OBWODU WYJŚCIOWEGO

Z zależności (21) wynika, że napięcie wyjściowe równa się sumie spadków napięcia na oporności obciążenia oraz oporności czynnej uzwojeń wyjściowych czujnika. Można więc podać prosty schemat zastępczy obwodu wyjściowego czujnika (rys. 4b), w którym występuje idealne źródło napięcia U_w , połączone w szereg z opornością czynną $2r_w$, która spełnia rolę oporności wewnętrznej oraz z opornością obciążenia R_o . Wartość średnia napięcia wyjściowego U_w jest wprost proporcjonalna do natężenia mierzonego pola magnetycznego H_s (dla $H_s \leq H_{nas}$) i nie zależy od oporności obciążenia R_o , gdy równanie (17) jest spełnione, co sprowadza się do warunku (20). Jak widać na rys. 3d, powierzchnia impulsu napięcia u_w nie zmienia się po załączeniu oporności obciążenia.

Należy podkreślić, że wartość średnia napięcia wyjściowego nie zależy od R_0 tylko w przypadku prostowania półokresowego jak na rys. 3d. Jak widać, prąd w obwodzie wyjściowym może płynąć tylko w chwilach nasycenia rdzeni, a prądu spowodowanego impulsami napięcia w chwilach odsycania nie ma. Poza tym zależności analityczne wyprowadzone zostały



Rys. 4. Czujnik z dwoma uzwojeniami wyjściowymi



Rys. 5. Czujnik z jednym uzwojeniem wyjściowym

przy założeniu, że indukcyjności rozproszeń nie występują. Założenie to jest dopuszczalne w zakresie niewielkich częstotliwości (50—500 Hz) w przypadku, gdy oporność obciążenia jest czynna (R_0).

Zgodnie z prostym schematem zastępczym (rys. 4) maksymalna moc czynna na oporności obciążenia $R_{opt} = 2r_w$ wyraża się wzorem

$$P_{omax} = \frac{U_w^2}{8r_w}. \quad (24)$$

Możliwe jest zastąpienie dwóch uzwojeń wyjściowych przez jedno, obejmujące oba rdzenie (rys. 5), nie zmienia to oczywiście zasady działania czujnika.

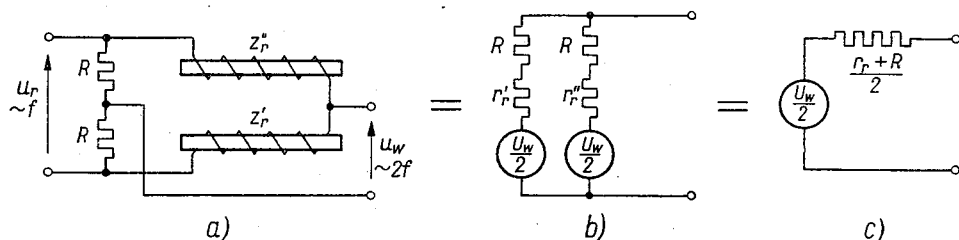
4. SCHEMATY ZASTĘPCZE CZUJNIKÓW MOSTKOWYCH

Oprócz opisanych wyżej czujników dwurdzeniowych z jednym lub dwoma uzwojeniami (rys. 4, 5), stosuje się często tzw. czujniki mostkowe (rys. 6, 7, 8, 9, 10), o których wzmianki znajdują się w publikacjach [2], [3], [4].

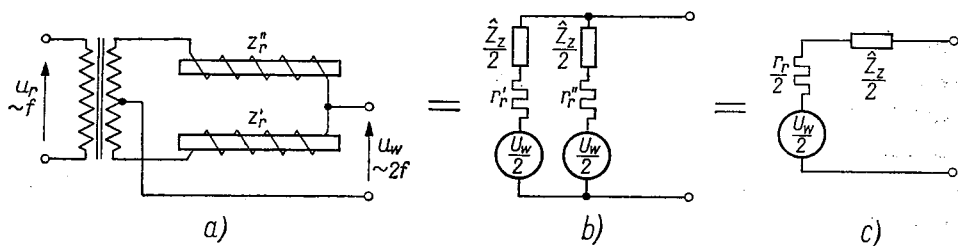
Pewną zaletą tych czujników jest brak uzwojenia wyjściowego, a więc gabaryt ulega zmniejszeniu, co jest ważne w niektórych konstrukcjach. Nazwa mostkowe czujniki magnetyczne wynika tylko z ich formalnego podobieństwa do mostków elektrycznych, gdyż w działaniu swym nie są podobne do tzw. mostków niezrównoważonych, co wynika z przedstawionych powyżej schematów zastępczych.

Na rys. 6 pokazany jest czujnik magnetyczny będący mostkiem, którego dwie gałęzie składają się z uzwojeń roboczych czujnika z_r' , z_r'' o opornościach czynnych r_r' , r_r'' , a pozostałe dwie stanowią oporniki o oporności czynnej R .

Ponieważ względem obwodu wyjściowego oba uzwojenia połączone są równolegle, więc na schemacie zastępczym źródła napięcia $\frac{U_w}{2}$ należy



Rys. 6. Czujnik mostkowy z opornikami dodatkowymi



Rys. 7. Czujnik mostkowy z transformatorem zasilającym

połączyć także równolegle (rys. 6b), co prowadzi ostatecznie do schematu zastępczego jak na rys. 6c. Widać, że oporności R powinny być jak najmniejsze, ale wtedy rośnie obciążenie źródła napięcia roboczego. W układzie tym uzwojenia robocze czujnika zasilane są napięciem sinusoidalnie zmiennym, gdyż przebieg taki wymuszony jest przez małą oporność R .

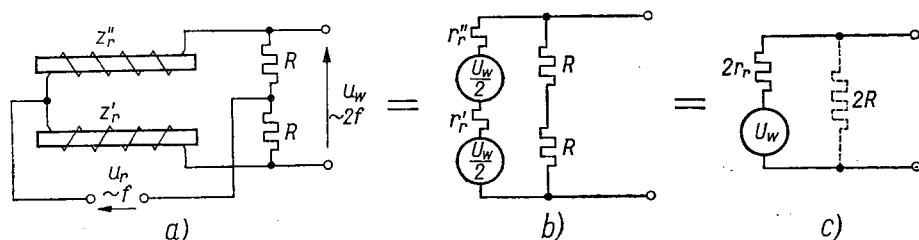
Odmianą układu z rys. 6 jest czujnik pokazany na rys. 7. Jako gałęzie mostka wykorzystuje się uzwojenie wtórne transformatora, składające się z dwóch symetrycznych połówek. Źródło napięcia roboczego nie jest w tym przypadku obciążone dodatkowymi opornikami. Od strony zacisków obwodu wyjściowego czujnika uzwojenie wtórne transformatora przedstawia tylko oporność zwarcia, co wynika z przeciwsobnego połączenia obu połówek uzwojenia wtórnego transformatora — przepływ prądu wyjściowego nie wytwarza strumienia w rdzeniu. Jest to bardzo korzystne, gdyż zapewniona jest mała oporność bez dodatkowego obciążenia źródła napięcia roboczego. Oporność zwarcia transformatora Z_z ma składową czynną (R_z) oraz indukcyjną (X_z):

$$Z_z = R_z + j X_z, \quad (25)$$

przy czym składowa indukcyjna może być czasem pominięta, zwłaszcza dla $f = 50$ Hz. Na każdą połówkę uzwojenia transformatora przypada oporność $Z_z/2$ (rys. 7b); ostateczny schemat zastępczy pokazany jest na rys. 7c.

W obu omówionych przypadkach, uzwojenia robocze z , z odgrywają rolę uzwojeń wyjściowych, dlatego też obliczając napięcie U_w należy podstawić do wzoru (23) $z_w = z_r$.

Innym rodzajem czujnika mostkowego jest układ na rys. 8. Powstał on jak gdyby przez zmianę zacisków źródła zasilającego i wyjściowych,



Rys. 8. Czujnik mostkowy z opornikami dodatkowymi na wyjściu

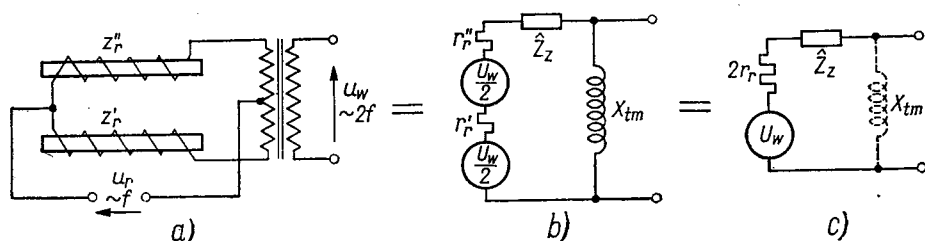
ale równocześnie zmienił się kierunek nawinięcia jednego z uzwojeń. W układzie tym istnieje obwód zamknięty dla parzystych harmonicznych nawet wtedy, gdy do zacisków wyjściowych nie jest załączony żaden odbiornik. Źródła napięcia wyjściowego pracują w tym układzie szeregowo (rys. 8b), a stąd wynika schemat zastępczy jak na rys. 8c. W tym przypadku oporności R powinny być jak największe, by spełnić warunek

$$R \gg X_r = \omega L_r, \quad (26)$$

gdzie:

$$L_r = \mu_0 \mu \frac{z_r^2 S}{l} \quad (27)$$

jest indukcyjnością własną uzwojenia roboczego. Warunek ten podyktowany jest ograniczeniem prądu wywołanego przez impulsy w momentach odsycania rdzeni. W tym przypadku, z uwagi na duże oporności R włą-



Rys. 9. Czujnik mostkowy z transformatorem na wyjściu

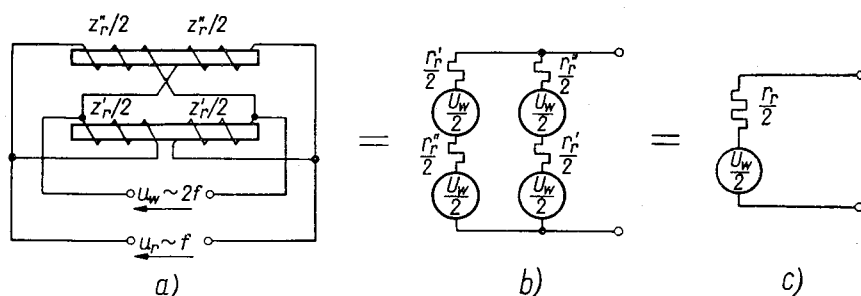
zione szeregowo z uzwojeniami roboczymi czujnika, wymuszony jest sinusoidalny przebieg prądu roboczego. Odmianą czujnika z rys. 8 jest układ, w którym zastępuje się oporności R transformatorem wyjściowym

(rys. 9). Do schematu zastępczego czujnika wchodzi wtedy schemat zastępczy transformatora np. typu I , jak na rys. 9b. Układ ten jest korzystniejszy od poprzedniego, gdyż opornością bocznikującą w tym przypadku jest oporność indukcyjna główna magnesowania transformatora X_{tm} (przy zaniedbaniu strat w żelazie), a oporność zwarcia Z_z połączona jest szeregowo z opornościami czynnymi uzwojeń. Oczywiście musi być spełniony warunek

$$X_{tm} \gg X_r. \quad (26a)$$

Od strony zacisków źródła zasilającego transformator przedstawia tylko oporność zwarcia, co jest korzystne, gdyż stanowi małe obciążenie dodatkowe dla źródła zasilania w porównaniu z opornikami R . W celu obliczenia napięcia wyjściowego do wzoru (20) podstawia się liczbę zwojów jednego z uzwojeń roboczych z_r .

Trzeci rodzaj układu mostkowego czujnika magnetycznego pokazany jest na rys. 10a. Wszystkie cztery gałęzie mostka stanowią połówki uzwo-



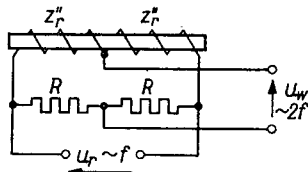
Rys. 10. Czujnik mostkowy czterouzwojeniowy

jeń roboczych czujnika. Na schemacie zastępczym uwzględnia się 4 źródła napięcia ($1/4 U_w$), przy czym pracują one w układzie szeregowo-równoległym (rys. 10b). Analizując otrzymany schemat zastępczy (rys. 10c) można by sądzić, że układ ten jest korzystniejszy od poprzednich, gdyż nie zawiera żadnych dodatkowych oporności. Tak też jest istotnie, jednak czujnik pokazany na rys. 10a jest trudny do wykonania (cztery jednakowe uzwojenia). Przy obliczaniu napięcia wyjściowego należy do wzoru (20) podstawić liczbę zwojów na jednym z rdzeni czujnika z_r .

Z przedstawionych powyżej schematów zastępczych wynika, że pod względem mocy czynnej na wyjściu wszystkie omówione wyżej czujniki są sobie mniej więcej równoważne, chociaż czujniki mostkowe z rys. 7 i 9 są mniej korzystne ze względu na występowanie oporności R w schematach zastępczych, co prowadzi do zmniejszenia mocy na wyjściu.

Czujniki pokazane na rys. 6—9 można zmodyfikować umieszczając oba uzwojenia robocze na jednym rdzeniu. Na rys. 11 pokazany jest przykładowo czujnik, odpowiadający układowi z rys. 6. Czujnik jednordzeniowy mierzy pole magnetyczne punktowo, w przeciwieństwie do dwurdzenio-

wych, które mierzą dwupunktowo. Przy założeniu rdzenia tych samych wymiarów oraz liczby zwojów obu uzwojeń $z'_r + z''_r$ tej samej co na każdym z rdzeni czujnika dwurdzeniowego, otrzymuje się teoretycznie napięcie

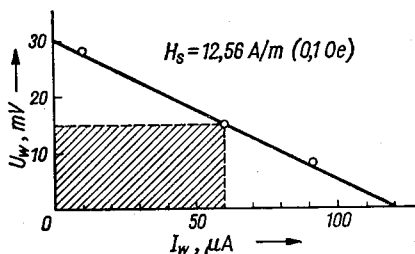


Rys. 11. Czujnik mostkowy jednordzeniowy

wyściowe 2 razy mniejsze, a moc wyjściową 4-krotnie mniejszą niż w układzie z rys. 6. Należy podkreślić, że w układzie z rys. 11 prąd roboczy wytwarza w obu połówkach rdzenia pole magnetyczne o przeciwnej biegunowości chwilowej, a zatem w środku istnieje strefa słabego przemagnesowania. Czujnik taki może namagnesować się w silnym polu magnetycznym, co równoważne jest przesunięciu zera, które fałszuje pomiary. Z tego względu czujniki te nie są stosowane.

5. WYNIKI DOŚWIADCZALNE

W celu sprawdzenia poprawności podanych schematów zastępczych zbadano różne czujniki stosując układ pomiarowy jak na rys. 1, przy częstotliwości zasilania $f = 50$ Hz. Przytoczone poniżej wyniki dotyczą



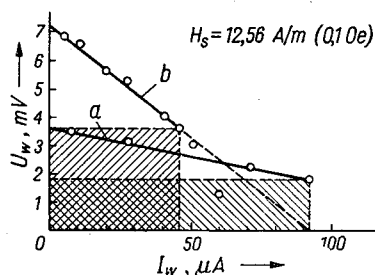
Rys. 12. Charakterystyka obciążenia czujnika z rys. 5

jednego z badanych czujników połączonego w układzie z jednym uzwojeniem wyjściowym (rys. 5) oraz w układach mostkowych (rys. 7, 9). Rdzenia czujnika wykonane z mumetalu posiadają wymiary $150 \times 3 \times 0,6$ mm, uzwojenia $z'_r = z''_r = 960$ zwojów, o oporności $r'_r = r''_r = 5,4 \Omega$, uzwojenie wyjściowe $z_w = 5000$ zwojów, $r_w = 250 \Omega$. W przypadku badania czujników mostkowych, użyty został transformator z rdzenia ferrytowego typu kubkowego o przekładni 1:1 i oporności zwarcia $Z_z = 60 \Omega$ ($R_z = 55 \Omega$).

Na rys. 12 pokazana jest charakterystyka obciążenia $U_0 = R_o I_w$ jako funkcja prądu wyjściowego I_w . Jak widać, jest to linia prosta, zatem pominięcie indukcyjności rozprożeń uzwojenia wyjściowego jest dopuszczal-

ne. Powierzchnia zakresłonego prostokąta proporcjonalna jest do mocy na optymalnej oporności obciążenia.

Na rys. 13 pokazane są charakterystyki obciążenia czujników mostkowych, które przebiegają w przybliżeniu liniowo dla oporności obciążenia $R_o \geq R_{opt}$, a wpływ składowej indukcyjnej oporności zwarcia zaznacza się



Rys. 13. Charakterystyki obciążenia:
a) czujnika z rys. 7, b) czujnika z rys. 9

dla oporności mniejszych, czyli dla części charakterystyki nie mającej praktycznego znaczenia. Widać, że maksymalna moc czynna na oporności obciążenia jest w obu przypadkach równa (równe powierzchnie prostokątów), co wskazuje równoważność obu konstrukcji.

*Instytut Podstawowych Problemów Techniki PAN
Zakład Elektrotechniki*

WYKAZ LITERATURY

1. Kulikowski J., Nałęcz M.: Podstawowe właściwości transduktorowych czujników magnetycznych o wyjściu impulsowym. — Rozprawy Elektrotechniczne, t. VI, z. 4, 1960.
2. Geyger W.: Magnetic Amplifier Circuits. — New York 1957, § 17.
3. Filipowicz B.: Sowriemiennye pribory dla izmierenija słabych magnitnyh polej. — Elektrichestwo, 1948, nr 5, str. 85—87.
4. Nowacki P. J., Nałęcz M., Kulikowski J.: Pomiar małych stałych pól magnetycznych czujnikiem magnetycznym. — Rozprawy Elektrotechniczne, t. IV, z. 3, 1958, str. 323—349.

Я. КУЛИКОВСКИ, М. НАЛЭЧ

СХЕМЫ ЗАМЕЩЕНИЯ МАГНИТОМОДУЛЯЦИОННЫХ ДАТЧИКОВ НАПРЯЖЕННОСТИ МАГНИТНОГО ПОЛЯ С ПИКОВЫМ ВЫХОДОМ

Резюме

В статье рассматриваются схемы замещения магнитомодуляционных датчиков магнитного поля с пиковым выходом построенных на двух сердечниках (рис. 4, 5) или с обмотками, соединёнными по мостовой схеме (рис. 6, 7, 8, 9, 10). Приведено сравнение этих типов по отношению к мощности на выходе.

J. KULIKOWSKI, M. NAŁĘCZ

EQUIVALENT CIRCUITS OF FLUX-GATE MAGNETOMETERS
WITH ALL-EVEN HARMONIC OUTPUT

Summary

Equivalent circuits of two-core flux-gate magnetometers with all-even harmonic pulse output is given (Figs. 4 and 5). There are shown equivalent circuits of magnetometers of bridge type too (Figs 6, 7, 8, 9 and 10) and these various types are compared with regard to their output power.

ANDRZEJ KOBUS

Powstawanie dyslokacji w kryształach

Rękopis dostarczono 11.9.1959

W artykule omówiono podstawowe mechanizmy powstawania dyslokacji w kryształach: 1) powstawanie dyslokacji w wyniku deformacji plastycznej, 2) powstawanie dyslokacji z zespołów luk, 3) powstawanie dyslokacji w wyniku nierównomiernego rozkładu domieszek (składników). W dalszym ciągu omówiono udział mechanizmu Franka-Reada w ostatecznym formowaniu układu dyslokacji w kryształach. Podano również warunki, w jakich poszczególne mechanizmy powstawania dyslokacji odgrywają decydującą rolę. W zakończeniu omówiono możliwości wykrywania dyslokacji powstałych w wyniku działania poszczególnych mechanizmów.

1. WSTĘP

Artykuł niniejszy omawia obecny stan wiedzy o podstawowych mechanizmach odpowiedzialnych za powstawanie dyslokacji w kryształach w czasie ich otrzymywania. Rozważania dotyczą w zasadzie kryształów półprzewodnikowych, jednak mechanizmy te biorą udział we wzroście wszelkich kryształów, chociaż znaczenie poszczególnych mechanizmów w konkretnych przypadkach może być różne.

Kryształy do celów półprzewodnikowych, pomimo że charakteryzują się strukturą najbliższą idealnej, zawierają wiele nieprawidłowości. Dyslokacje będące nieprawidłowościami o charakterze liniowym lub powierzchniowym zajmują drugie miejsce w hierarchii nieprawidłowości w kryształach półprzewodnikowych po domieszkach chemicznych. Ze względu na to, że dyslokacje z elektronicznego punktu widzenia są przede wszystkim centrami rekombinacji ograniczającymi czas życia nośników mniejszościowych oraz obszarami łatwej dyfuzji domieszek, stanowią ważny czynnik, który należy brać pod uwagę w technologii monokryształów półprzewodnikowych. Pomimo wielkich postępów osiągniętych przy otrzymywaniu kryształów o małej gęstości dyslokacji [1], [2], [3] zagadnienie to nie zostało jeszcze całkowicie rozwiązane, chociaż ilość prac poświęconych temu problemowi stale wzrasta.

Artykuł ten jest kontynuacją dwóch poprzednich artykułów [4], [5] poświęconych zagadnieniom dyslokacji w kryształach półprzewodnikowych. Dlatego też nie będą tu omawiane podstawowe problemy związane z dyslokacjami, jak pojęcie obiegu i wektora Burgersa. Geometryczne

modele dyslokacji zostaną podane w skrócie. Szersze wiadomości o podstawach teorii dyslokacji oraz o ujawnianiu dyslokacji Czytelnik znajdzie w podanych wyżej publikacjach. Natomiast podstawowymi źródłami wiadomości o dyslokacjach są książki: Cottrella [6] i Reada [31].

W dalszych rozdziałach zostaną omówione trzy podstawowe mechanizmy powstawania dyslokacji:

- 1) Powstawanie dyslokacji w wyniku deformacji plastycznej.
- 2) Powstawanie dyslokacji z zespołów luk.
- 3) Powstawanie dyslokacji w wyniku segregacji domieszek.

2. POWSTAWANIE DYSLOKACJI W WYNIKU DEFORMACJI PLASTYCZNEJ

Właściwości plastyczne kryształów są obecnie tłumaczone istnieniem i ruchem dyslokacji. Wiadomo, że monokryształy otrzymywane w praktyce posiadają znacznie mniejszą wytrzymałość mechaniczną niż wynikałoby to z teoretycznie przewidywanych właściwości absolutnie idealnego kryształu. Za rozbieżność tę są odpowiedzialne dyslokacje istniejące w kryształach, które są przyczyną łatwego poślizgu pewnych obszarów materiału.

Poślizg kryształów jest zjawiskiem anizotropowym i odbywa się wzdłuż pewnych płaszczyzn. Uprzywilejowane są płaszczyzny i kierunki o najgęstszym upakowaniu atomów w sieci krystalicznej. Na przykład w sieci diamentu odpowiada to czterem płaszczyznom $\{111\}$ i sześciu kierunkom $\langle 110 \rangle$. Poślizg zachodzi wtedy, gdy naprężenia działające w płaszczyźnie i kierunku poślizgu przekroczą wartość krytyczną. Płaszczyzny i kierunki łatwego poślizgu mają oczywiście najmniejszą wartość naprężenia krytycznego dla danej sieci krystalicznej.

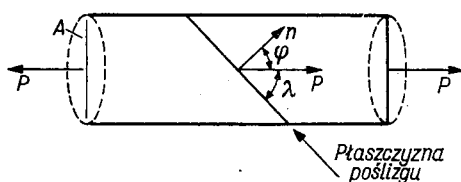
Rozpatrzmy teraz jako przykład przypadek rozciągania walca, w którym płaszczyzny łatwego poślizgu są nachylone do osi walca pod kątem λ , jak na rys. 1. Siła rozciągająca P przyłożona jest osiowo. Składowa siły działająca w płaszczyźnie poślizgu będzie $P \cos \lambda$. Siła ta rozłożona będzie na powierzchnię poślizgu równą $\frac{A}{\cos \varphi}$, wobec tego naprężenie działające w płaszczyźnie poślizgu będzie wynosiło [6]:

$$\sigma = \frac{P}{A} \cos \varphi \cos \lambda. \quad (1)$$

Jak widać, poślizg będzie zależał nie tylko od wielkości siły przyłożonej P , lecz i od kierunku jej działania względem płaszczyzny poślizgu. W rozpatrywanym przypadku naprężenia działające w płaszczyźnie poślizgu osiągną maksimum dla $\lambda = 45^\circ$ i będą wtedy wynosiły $\frac{1}{2} \frac{P}{A}$, gdy dla kątów $\lambda = 0^\circ$ i 90° będą równe zeru. W innych konkretnych przypadkach należy również przeprowadzić podobną analizę, aby określić, w jaki sposób naprężenie powodujące poślizg będzie zależało od sposobu przyło-

żenia siły względem kierunków płaszczyzn poślizgu. Przyłożenie jednej siły do kryształu może spowodować poślizg wielokrotny, o ile składowe siły przyłożonej w kilku odpowiednich kierunkach wywołają naprężenia większe od krytycznego.

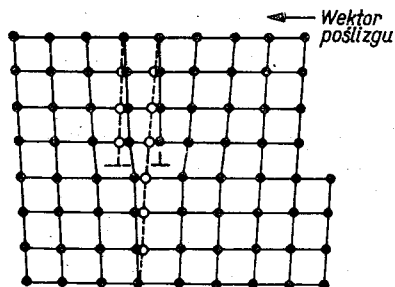
Krytyczne naprężenie poślizgu jest strukturalnie czułe. W praktycznych przypadkach poślizg zaczyna się od wszelkich nieprawidłowości sieci



Rys. 1. Przykład rozciągania kryształu w kształcie walca

i w tych przypadkach naprężenie krytyczne może być znacznie mniejsze. Naprężenie krytyczne maleje również z temperaturą. W przypadku germanu w temperaturach większych od 800°C naprężenia te są tak małe we wszystkich kierunkach, że deformacja plastyczna jest jednakowo prawdopodobna we wszystkich kierunkach [8].

Poślizg nie odbywa się przez jednoczesny ruch wszystkich atomów po obu stronach płaszczyzny poślizgu, ale stopniowo, w pewien uporządkowany sposób. Odpowiada to wprowadzeniu i ruchowi dyslokacji na płaszczyźnie poślizgu. Dyslokacja taka przesuwa się wzdłuż linii poślizgu

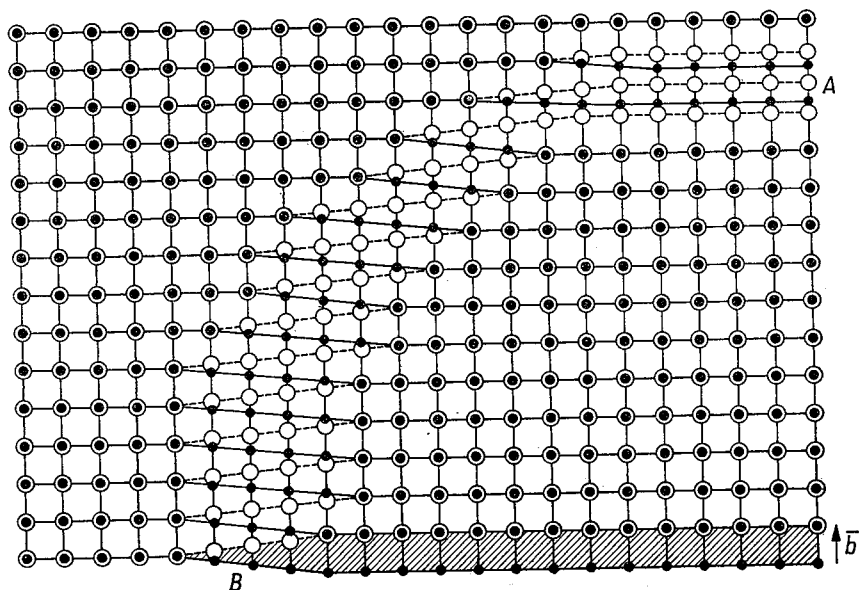


Rys. 2. Poślizg kryształu przedstawiony jako ruch dyslokacji krawędziowej

w postaci pewnego zaburzenia sieci w sposób podobny do ruchu falowego. Obszar zaburzenia jest niewielki, średnicę jego szacuje się na 3 stałe sieci (tj. ok. kilkunastu $\text{\AA}$). Jednak zaburzenie może poruszać się na duże odległości rzędu centymetrów. Na rys. 2 pokazano fragment poślizgu jako ruch dyslokacji krawędziowej. W ten sposób mikroskopowy ruch elementów sieci krystalicznej powoduje makroskopowy ruch pewnych obszarów kryształu.

W wyniku działania poślizgu mogą w kryształach powstawać różne rodzaje dyslokacji. Dyslokację opisuje się przy pomocy określenia jej orientacji, tj. określenia kąta zawartego między kierunkiem linii dyslokacyjnej

a kierunkiem poślizgu (tj. kierunkiem wektora Burgersa) i wielkości wektora Burgersa względem translacyjnego wektora sieci. Dla dyslokacji całkowitych, tj. gdy wektor Burgersa równy jest wektorowi translacyjnemu sieci, jeżeli kąt określający orientację równy jest zeru, mamy do czynienia z dyslokacją o orientacji śrubowej, gdy kąt ten wynosi $\frac{\pi}{2}$ jest to dyslokacja krawędziowa. Pozostałe kąty charakteryzują dyslokacje



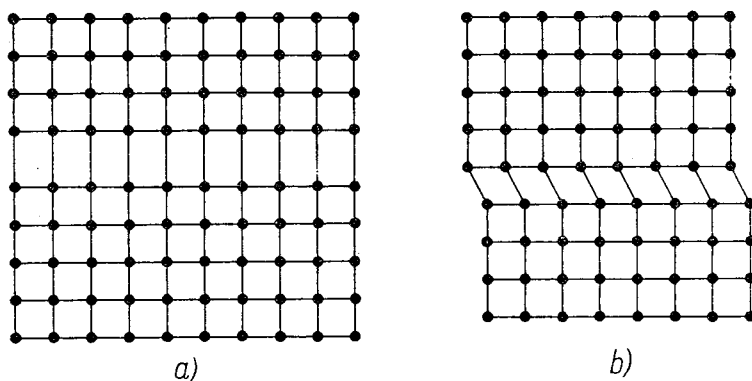
Rys. 3. Model krzywoliniowej dyslokacji w prostej sieci kubicznej

o orientacji pośredniej. Na rys. 3 pokazano schematycznie układ atomów w prostej sieci kubicznej nad-(kropki) i pod-(kółka) płaszczyznę poślizgu w dyslokacji krzywoliniowej. Jak widać, w punkcie A dyslokacja posiada orientację krawędziową, w punkcie B — śrubową, w punktach pośrednich przebiegu krzywoliniowego orientacja jest mieszana.

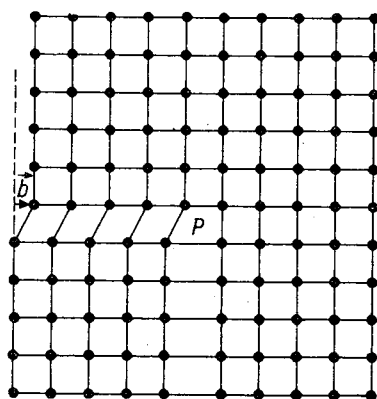
W wyniku poślizgu mogą powstawać również tzw. dyslokacje częściowe. Charakteryzują się one tym, że ich wektor Burgersa jest mniejszy od translacyjnych wektorów sieci. Istnienie dyslokacji częściowych związane jest z obecnością w kryształach tzw. naruszeń, przy pomocy których Heidenreich i Shockley w 1948 r. [7] wyjaśniali pewne własności kryształów metali i które następnie były badane przez Barreta w 1951 r. [9].

Naruszeniami nazywa się niskoenergetyczne nieprawidłowości powierzchniowe. Przykłady naruszeń dla prostej sieci kubicznej podano schematycznie na rys. 4. Jeżeli naruszenie kończy się w kryształach, granicą naruszenia będzie dyslokacja częściowa.

W przypadku niejednorodnego ścinania w kryształach może powstać układ atomów jak na rys. 5. W płaszczyźnie poślizgu powstanie naruszenie



Rys. 4. Przykłady naruszeń w prostej sieci kubicznej



Rys. 5. Model częściowej dyslokacji Shockleya w prostej sieci kubicznej

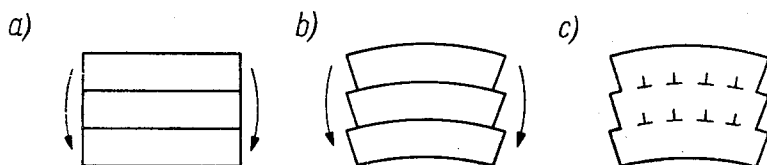
spowodowane przesunięciem dwóch obszarów kryształu względem siebie o wielkość równą wektorowi poślizgu, który jednak jest mniejszy niż wektor translacyjny sieci. Granicą powierzchni poślizgu zgodnie z definicją podaną powyżej będzie dyslokacja częściowa. W tym przypadku wektor Burgersa (tj. wektor poślizgu) leży w płaszczyźnie naruszenia; dyslokacja taka nazywa się dyslokacją częściową Shockleya [7]. Inne rodzaje dyslokacji częściowych będą omówione w następnym rozdziale.

Poślizg w kryształach można wywołać w następujący sposób:

- 1) przez przyłożenie sił zewnętrznych,
- 2) przez wywołanie naprężeń termicznych wskutek nierównomiernego chłodzenia ogrzanego kryształu.

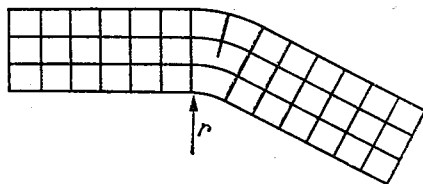
ad 1. Przez przyłożenie sił zewnętrznych, wywołujących naprężenia w płaszczyźnie poślizgu większe od krytycznych, można spowodować poślizg i zatem wprowadzenie pewnej ilości dyslokacji. W zasadzie dowolne przyłożenie odpowiednio dużych sił do kryształu może być przyczyną wprowadzenia dyslokacji. W dalszym ciągu zostanie omówiony tylko jeden szczególny przypadek przyłożenia sił zewnętrznych — zginanie, który ma jednak duże znaczenie praktyczne w razie potrzeby badania

własności tylko dyslokacji krawędziowych [6]. Jeżeli sztabka kryształu jak na rys. 6a zostanie poddana zginaniu plastycznemu, z mikroskopowego punktu widzenia można to porównać ze zgięciem pliku arkuszy papieru (rys. 6b). W kryształe powstanie szereg warstw oddzielonych płaszczyznami poślizgu. Ponieważ sieć krystaliczną trudno jest zdeformować w znacz-



Rys. 6. Wprowadzanie dyslokacji krawędziowych przez zginanie kryształu

nym stopniu, stabilnym układem energetycznym zgiętego kryształu będzie wprowadzenie dyslokacji krawędziowych (rys. 6c). Przypadek zgięcia w układzie prostej sieci kubicznej pokazano schematycznie na rys. 7. Jak z niego wynika, plastyczne zgięcie kryształu w odpowiednim kierunku spowoduje wprowadzenie dyslokacji krawędziowych. Gęstość ich można wyznaczyć w następujący sposób. Warstwa kryształu o grubości d zgięta



Rys. 7. Dyslokacyjny model zgiętego kryształu; promień krzywizny zginania wynosi r została tak, że promień krzywizny wynosi r . Ogólna wielkość przemieszczenia sieci pomiędzy warstwami na jednostkę długości będzie $\frac{d}{r}$. Przemieszczenie to jest równoważne wprowadzeniu $\frac{d}{rb}$ dyslokacji krawędziowych o nateżeniu b . Jeżeli na jednostkę grubości przypada $\frac{1}{d}$ część warstwy, to gęstość dyslokacji będzie:

$$N_D = \frac{1}{r \cdot b}. \quad (2)$$

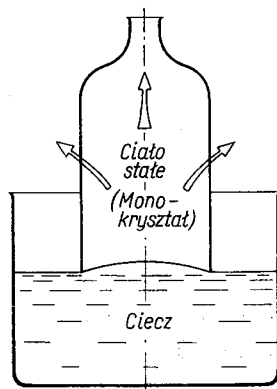
Opisane doświadczenie posiada duże znaczenie praktyczne, gdyż przy doborze odpowiednich osi zginania wprowadza tylko dyslokacje krawędziowe. Z elektrycznego punktu widzenia jest to wielka zaleta, gdyż pozwala na określenie stanów akceptorowych w przypadku badania przewodnictwa lub rekombinacji na dyslokacjach.

Pearson, Read i Morin [10] badali próbki germanu poddawane deformacji plastycznej przez zginanie w temperaturze 650°C . Dla promienia

$r = 5$ cm otrzymali oni gęstość dyslokacji określaną przez trawienie chemiczne powierzchni $\{111\}$ jako $3 \cdot 10^6 \text{ cm}^{-2}$. Jest to wartość nieco mniejsza od wynikającej ze wzoru (2), co prawdopodobnie należy przypisać trudności odróżnienia gęsto ułożonych wgłębień dyslokacyjnych. Próbki kryształów o znanej gęstości dyslokacji służyły następnie do badania wpływu dyslokacji na ruchliwość, przewodność i gęstość nośników w germanie.

ad 2. Badaniem wpływu naprężeń termicznych na gęstość dyslokacji w kryształach półprzewodnikowych zajmowało się wielu badaczy [1—3], [11—17]. Doświadczenia były prowadzone w dwóch kierunkach: 1) badano rozkład dyslokacji w zależności od warunków otrzymywania kryształów metodą Czochralskiego lub krystalizacji postępującej [1—3], [11], [12], [14—17], 2) gotowe kryształy poddawano specjalnej obróbce termicznej [13].

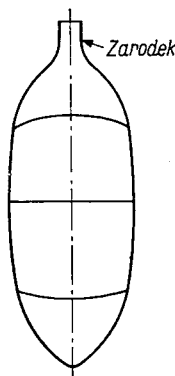
Na wstępie zostaną omówione warunki otrzymywania monokryształów metodą Czochralskiego. Schemat tej metody pokazano na rys. 8. Materiał



Rys. 8. Schemat otrzymywania monokryształów metodą Czochralskiego

w fazie ciekłej znajduje się w tyglu. Zarodek o odpowiedniej orientacji zanurza się w cieczy i następnie wolnym ruchem przesuwają się do góry. Na zarodku krystalizuje ciecz tworząc kryształ o orientacji narzuconej przez zarodek. Zestalona część materiału otrzymuje od cieczy pewną ilość ciepła, które jest odprowadzane przez przewodnictwo do zarodka i przez promieniowanie powierzchni bocznych, gdy kryształ jest wyciągany w próżni. Jeżeli proces ten odbywa się w atmosferze gazu obojętnego, należy jeszcze uwzględnić odprowadzanie ciepła przez konwekcję. Odprowadzanie ciepła w różnorodny sposób pociąga za sobą nieuniknione nierównomierności chłodzenia, a więc i temperatury. Na podstawie powyższych rozważań należałoby się spodziewać zarówno wypukłej powierzchni przejściowej ciecz — ciało stałe, jak i pozostałych izoterm w kryształach. Doświadczenie potwierdza ten fakt w szeregu przypadków. Jednak w praktyce można otrzymać płaskie, a nawet wklęsłe powierzchnie przejściowe. Można to spowodować albo dodatkowymi grzejnikami, albo przez dogrze-

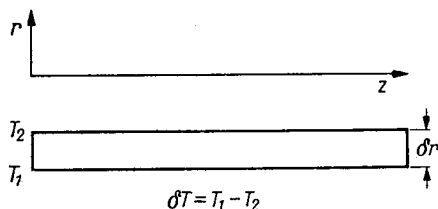
wające działanie ścian tygla, gdy opadnie poziom cieczy po wyciągnięciu pewnej ilości kryształu z tygla. Zjawisko to zostało potwierdzone doświadczalnie. Billig [11] wykonywał złącza p - n w kryształach krzemu przez domieszkowanie cieczy w czasie wyciągania. Ujawnione złącze jest wskaz-



Rys. 9. Typowy kształt izoterm w różnych punktach monokryształu otrzymanego metodą Czochralskiego w wysokim tyglu

nikiem kształtu powierzchni przejściowej. Typowe kształty powierzchni przejściowych pokazano na rys. 9. Podobne rezultaty otrzymano przez badanie przy pomocy zjawiska fotowoltaicznego świadomie wykonanych lub przypadkowych gradientów koncentracji domieszek w kryształach [18], [19].

Gęstość dyslokacji wytworzonych pewnym gradientem temperaturowym występującym w kryształach Billig oszacował w następujący sposób.



Rys. 10. Model służący do powiązania gradientu temperaturowego z gęstością dyslokacji powstałych w wyniku jego istnienia w kryształach

Rozpatruje się cienką płaską płytkę o długości z i grubości δr jak na rys. 10. Między powierzchniami dolną i górną płytki istnieje różnica temperatur δT . Wytworzony gradient temperaturowy $\frac{\delta T}{\delta r}$ spowoduje „spaczenie się” płytki wskutek niejednakowego wydłużenia, które po stronie cieplejszej będzie:

$$\delta z = \alpha z \delta T, \quad (3)$$

gdzie α — współczynnik rozszerzalności liniowej.

Wobec tego krzywizna spaczenia jest:

$$\frac{1}{R} = \frac{1}{Z} \frac{\delta z}{\delta r} = a \frac{\delta T}{\delta r} \quad (4)$$

Jeżeli teraz będzie rozpatrywana płytka, która jest ściśnięta tak, że nie może ulec spaceniu, cieplejsza warstwa będzie podlegała naprężeniom ściskającym, a zimniejsza — rozciągającym. Jeżeli kryształ znajduje się w zakresie temperatur plastycznych, wystąpi odkształcenie odpowiadające wprowadzeniu pewnej ilości dyslokacji krawędziowych. Korzystając ze wzorów (2) i (4) otrzymamy wyrażenie wiążące gęstość dyslokacji z gradientem temperaturowym:

$$N_D = \frac{a}{b} \frac{\delta T}{\delta r} \quad (5)$$

Równanie to będzie określało w pierwszym przybliżeniu gęstość dyslokacji w monokryształe wyciąganym z fazy ciekłej. Jak wykazał Wagner [15], decydującym czynnikiem o powstawaniu dyslokacji przy wyciąganiu metodą Czochralskiego jest promieniowy rozkład temperatury. Wyciągał on monokryształy z dużą i małą prędkością, a następnie chłodził szybko lub powoli. Kryształy chłodzone powoli (dzięki dogrzewaniu przez dodatkowe grzejniki), wyciągane zarówno z dużą, jak i małą prędkością miały znacznie mniejszą gęstość dyslokacji, gdy kryształy szybko ochłodzone po wyciągnięciu wykazywały znaczny wzrost gęstości dyslokacji z $5 \cdot 10^3 \text{ cm}^{-2}$ do 10^5 cm^{-2} .

Powiązanie gradientu temperaturowego z gęstością powstałych dyslokacji w kryształach można dokładniej przeprowadzić przy pomocy rachunku tensorowego [20]—[24]. Niestety, przy analizie konkretnych sieci kryształów rachunek ten bardzo się komplikuje i dla najbardziej nas interesującej sieci diamentu brak jeszcze szczegółowych wyników.

Tablica 1

Zakresy plastyczności w kryształach półprzewodnikowych

Materiał	T_{top}	Zakres deformacji plast. wzdłuż płaszczyzn poślizgu	Zakres bezwładności deformacji plastycz.
Krzem	1415 °C	500—1000 °C [25]	1000—1415 °C
German	960 °C	450—800 °C	800—660 °C
Antymonek indu	525 °C	80—400 °C [25]	400—425 °C
Antymonek galu	725 °C	100—600 °C [25]	600—725 °C

Deformacja plastyczna kryształu może zachodzić tylko w zakresie temperatur, w których jest on plastyczny. W tablicy 1 podano zakresy temperatur plastyczności różnych kryształów półprzewodnikowych. W drugiej kolumnie podano zakres temperatur, w którym deformacja plastyczna odbywa się wzdłuż płaszczyzn poślizgu, tzn. naprężenie krytyczne wzdłuż

plaszczyn o gęstym upakowaniu atomów jest znacznie mniejsze niż na pozostałych. W trzeciej kolumnie podano zakres temperatur, w którym naprężenia krytyczne we wszystkich kierunkach są małe i mają podobne wartości, wskutek czego deformacja zachodzi w sposób bezładny.

3. POWSTAWANIE DYSLOKACJI Z UKŁADÓW LUK

Kryształy otrzymywane przez wyciąganie z fazy ciekłej mają pewną koncentrację luk jak i atomów w pozycjach międzywęzłowych (defektów Frenkla i Schottky'ego), odpowiadającą równowadze termodynamicznej kryształu w danej temperaturze. Koncentracja defektów w sieci będzie największa w temperaturze topnienia. W praktycznych przypadkach w literaturze podaje się, że dla temperatury topnienia kryształ zawiera 10^{-2} — 10^{-30} % luk w stosunku do ilości węzłów sieci [26]. Dla temperatur niższych koncentracja defektów odpowiadająca innym warunkom równowagi termodynamicznej jest niższa. Należy teraz rozpatrzyć, co się dzieje z defektami, gdy temperatura kryształu spada od temperatury topnienia do niższych.

Istnieją cztery mechanizmy ustalania się równowagi defektów przy obniżaniu temperatury kryształu.

1. Zajęcie pozycji lukowej przez atom międzywęzłowy (w przypadku defektu Frenkla) lub dyfuzja atomu z powierzchni do pozycji lukowej (w przypadku defektu Schottky'ego).
2. Dyfuzja luki do powierzchni kryształu (w przypadku defektu Frenkla).
3. Dyfuzja luki w kierunku dyslokacji krawędziowej.
4. Kondensacja pewnej ilości luk w układ (np. płytkę lub dysk).

Pokrótce zostaną omówione trzy pierwsze mechanizmy, natomiast czwarty będzie tematem tego rozdziału, gdyż jest odpowiedzialny za powstawanie dyslokacji.

ad 1. Wydaje się, że unicestwienie defektu sieci odgrywa pewną rolę tylko w kryształach o dużych wymiarach, małej gęstości dyslokacji i przy bardzo wolnym chłodzeniu. W takich idealnych warunkach udział innych mechanizmów ustalania się równowagi termodynamicznej defektów mógłby być wyeliminowany. W praktycznych przypadkach: istnienia dyslokacji i dużych prędkości chłodzenia mechanizm ten ma prawdopodobnie mały udział w zanikaniu defektów.

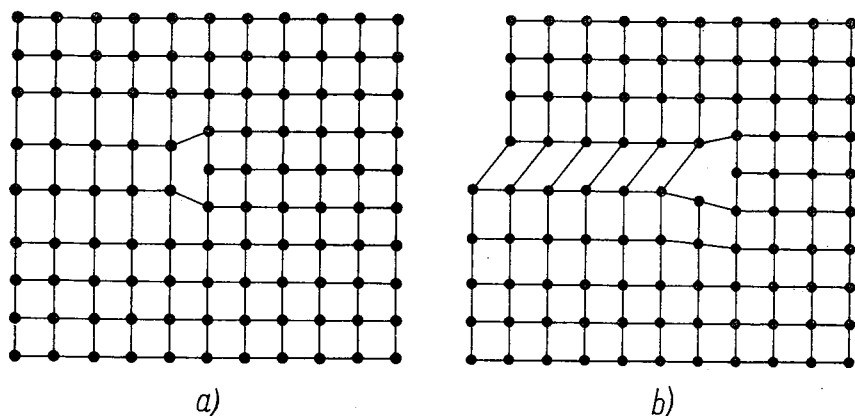
ad 2. Możliwość zanikania luk przez dyfuzję do powierzchni podał Nabarro [27]. Uważa się, że ze względu na ograniczoną ruchliwość luk mechanizm ten może mieć praktyczne znaczenie tylko w kryształach o bardzo małych wymiarach.

ad 3. Występowanie sił oddziałujących między dyslokacją a punktowymi źródłami naprężeń w kryształach wykazali Cottrell i Bilby w 1949 r. [28]. Takim źródłem naprężeń są również luki, na które będą działać siły se-

gregacji w kierunku linii dyslokacyjnej. Mechanizm ten posiada duży wpływ na rozkład luk w rzeczywistym kryształe.

ad 4. Możliwość powstawania zespołów luk w przypadku kryształów alkalihaloidków podał Seitz [29]. Crussard [30] wykazał, że najbardziej stabilną postacią zespołu luk jest płaski dysk. Brooks [wg 31] stwierdził, że do powstania zespołu luk konieczne jest istnienie nadmiaru ich koncentracji, tj. „przechłodzenie” kryształu. Kondensacja luk będzie postępowała dopóki nie zmaleje temperatura kryształu, tzn. dopóki ruchliwość luk jest dostatecznie duża, aby powstało prawdopodobieństwo utworzenia zespołu. Według Reada [31] dyski mają tendencje do powstawania na płaszczyznach o gęstym upakowaniu atomów. Dla sieci diamentu, w której krystalizuje duża część kryształów półprzewodnikowych płaszczyznami tymi są płaszczyzny $\{111\}$.

W wyniku odpowiednich warunków, tzn. istnienia odpowiedniej temperatury w dostatecznie długim czasie, będzie odbywał się wzrost dysków

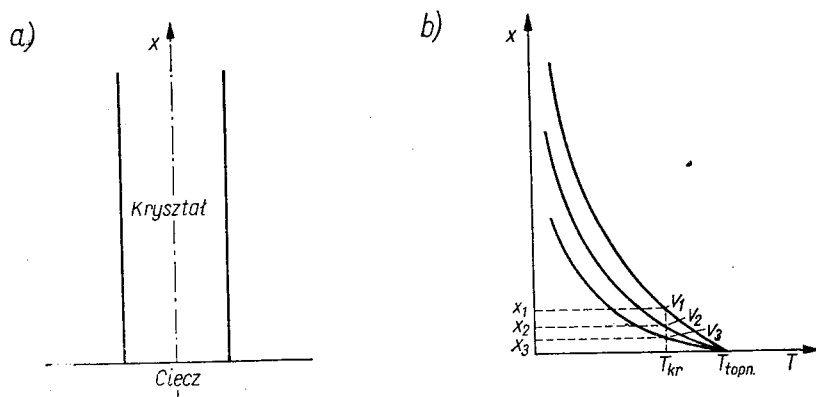


Rys. 11. Powstawanie dyslokacji przez zapadnięcie się dysku lukowego: a) pełna dyslokacja krawędziowa, b) częściowa dyslokacja Franka

lukowych aż do osiągnięcia pewnej wielkości krytycznej, przy której zachodzi zapadnięcie się sieci na obszarze dysku. Zapadnięcie może odbyć się w dwojaki sposób: 1) sąsiednie płaszczyzny mogą zbliżać się po normalnej bez przesunięcia prostopadłego tworząc w przypadku prostej sieci kubicznej dyslokację krawędziową (rys. 11a) i 2) sąsiednie płaszczyzny atomowe mogą zbliżyć się do siebie i przesunąć w kierunku stycznym tworząc tzw. dyslokację częściową Franka (rys. 11b).

Dyslokacjami częściowymi Franka nazywa się dyslokacje posiadające wektor Burgersa nierównoległy do płaszczyzny naruszenia. Analogiczna dyslokacja, ale o przeciwnym wektorze Burgersa, powstanie w przypadku obecności części dodatkowej płaszczyzny atomów. Układ atomów będzie się różnił przeciwnymi kierunkami dodatkowych półpłaszczyzn atomowych.

W przypadku otrzymywania kryształów z fazy ciekłej, np. metodą Czochralskiego, schematycznie pokazanej na rys. 12a, proces powstawania dysków lukowych i następnie pętli dyslokacyjnych będzie ograniczony. Ograniczenie to wynika z określonej prędkości przesuwania powierzchni przejściowej ciecz — ciało stałe. W wyniku odprowadzania ciepła przez



Rys. 12. Warunki termiczne powstawania zespołów luk i dyslokacji w kryształach otrzymywanych metodą Czochralskiego: a) schemat kryształu, b) osiowy rozkład temperatury w kryształach przy różnych prędkościach wyciągania

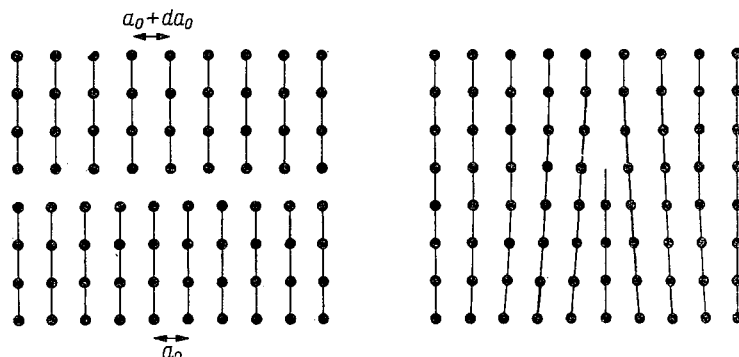
przewodnictwo do zarodka i promieniowanie z powierzchni zestalonej części kryształu powstaje temperaturowy gradient osiowy. Rozkład temperatury w kryształach schematycznie pokazano na rys. 12b, gdzie przez T_{min} oznaczono temperaturę, w której ruchliwość luk jest tak mała, że nie mogą one tworzyć dysków. Ze wzrostem prędkości wyciągania kryształu rozkład temperatury zmienia się tak, że czas, w którym istnieją warunki do powstawania dysków, maleje. Wobec tego ze wzrostem prędkości kryształizacji będzie najpierw malała wielkość „zamrożonych” dysków lukowych, a zatem i pętli dyslokacyjnych, a następnie i ich ilość [32]. Jak więc wynika z wyżej przeprowadzonych rozważań i przyjmując powyższy mechanizm, w celu uzyskania małych gęstości dyslokacji należy stosować duże prędkości wyciągania monokryształów.

Ten jakościowo przedstawiony mechanizm znajduje potwierdzenie w doświadczeniu. Mianowicie autor [33] badał rozkład „małych” wgłębień dyslokacyjnych, odpowiadających prawdopodobnie dyslokacjom powstałym z zespołów luk [5], [41], [43] w monokryształach germanu otrzymanym metodą Czochralskiego w wysokim tyglu. W początkowej chwili wyciągania monokryształu był nagrzewany praktycznie tylko przez przewodnictwo ciepła od cieczy. Natomiast w miarę opadania poziomu cieczy wyciągany monokryształ był nagrzewany również przez promieniowanie ścian tygla. Odpowiada to zmniejszaniu osiowego gradientu temperatury w kryształach. Stwierdzono, że w takim monokryształach gęstość małych

wgłębień dyslokacyjnych wzrastała z odległością od początku (zarodka) kryształu.

4. POWSTAWANIE DYSLOKACJI W WYNIKU ZMIAN KONCENTRACJI DOMIESZEK

Obecność w kryształach domieszek o promieniu atomowym różnym od atomów ośrodka powoduje zmiany stałej sieci materiału. W przypadku kryształów z czystych półprzewodników elementarnych zmiany stałej sieci są do pominięcia, jednak w przypadku związków lub stopów zmiany stałej sieci mogą być znaczne. Rozważymy granicę dwóch obszarów kryształu, w których średnia stała sieci znacznie się różni. Jeżeli granica między tymi



Rys. 13. Model powstawania dyslokacji krawędziowej na granicy obszaru o różnych stałych sieci

obszarami kryształu będzie dostatecznie ostra, stabilny układ sieci musi zawierać dyslokacje krawędziowe (rys. 13). Dyslokacje takie powstaną ze względu na to, że redukują one energię odkształceń elastycznych pomiędzy sąsiednimi warstwami różniącymi się stałą sieci.

Maksymalna liniowa gęstość dyslokacji, jaka powstanie przy zmianie stałej sieci w wyniku zmiany koncentracji domieszek (ewentualnie składnika stopu lub związku) na powierzchni przejściowej kryształu, może być wyznaczona w następujący sposób, przy założeniach, że: 1) wszystkie makroskopowe odkształcenia elastyczne są uwolnione przez powstanie dyslokacji krawędziowych; 2) dyslokacja krawędziowa uwalnia odkształcenia elastyczne równe modułowi wektora Burgersa — b . W ogólnym przypadku maksymalna gęstość dyslokacji krawędziowych na granicy obszarów o różnych stałych sieci będzie wyrażała się wzorem [34]:
gdzie:

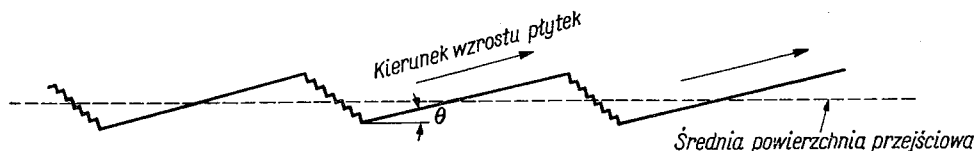
$$N_D = \Delta C \left(\frac{d a_0}{d C} \right) \frac{1}{b a_0}, \quad (6)$$

$\frac{d a_0}{d C}$ — zmiana stałej sieci na jednostkową zmianę koncentracji domieszek,

ΔC — przyrost koncentracji domieszek,
 a_0 — średnia stała sieci.

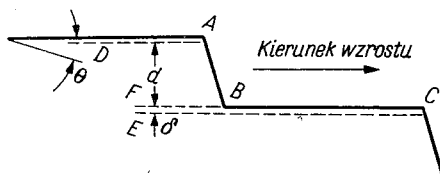
Stwierdzono doświadczalnie [34], że mechanizm ten powoduje powstawanie dyslokacji dopiero dla pewnej minimalnej zmiany koncentracji domieszek. Gdy zmiana koncentracji jest zbyt mała, może powstać tylko odkształcenie elastyczne sieci a dyslokacje nie powstaną.

Tiller [35] rozwinął to zagadnienie i podał teorię powstawania dyslokacji w wyniku segregacji domieszek w czasie płytkowego i komórkowego



Rys. 14. Schemat powierzchni przejściowej przy płytkowym wzroście kryształów

wzrostu kryształów. Następnie zostaną omówione oba przypadki wzrostu kryształów. Wzrost płytkowy jest ogólną postacią wzrostu kryształów. Ten sposób wzrostu został stwierdzony doświadczalnie w szeregu metali i NaCl. Na rys. 14 pokazano schematycznie wycinek powierzchni przejściowej kryształu, którego zestalanie odbywa się przez wzrost płytkowy. Mechanizm ten polega na tym, że kryształ narasta nie prostopadle do powierzchni przejściowej, a przez wzrost krawędzi płytek w kierunku równoległym do ich powierzchni, jak pokazano strzałką na rys. 14. Wobec tego faktyczna powierzchnia wzrostu płytek przecina się pod ką-

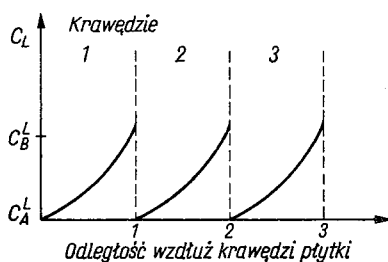


Rys. 15. Schemat powstawania dyslokacji krawędziowych przy płytkowym wzroście kryształów

tem θ ze średnią powierzchnią przejściową. Płaszczyzny płytek mają zazwyczaj orientację o niskim indeksie i są to najczęściej powierzchnie $\{100\}$ lub $\{111\}$. Grubość płytek wynosi w zależności od warunków wzrostu od 0,1 do 1 μ .

W dalszych rozważaniach będzie mowa o wzroście kryształów zawierających domieszki, których współczynnik segregacji k jest mniejszy od jedności, tzn. domieszki są jakby odrzucane od krawędzi zestalonego kryształu i dyfundują do wnętrza cieczy. Na rys. 15 pokazano schematycznie powiększony odcinek powierzchni przejściowej. Domieszki znacznie łatwiej będą przechodziły do cieczy z punktu A niż z punktu B, który jest bardziej „osłonięty”. W wyniku tego zjawiska koncentracja domieszek

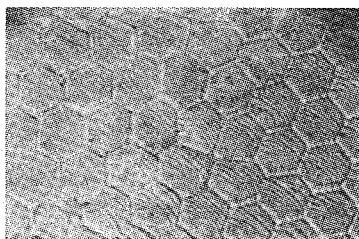
w cieczy w punkcie B będzie większa niż w punkcie A . Na grubości płytki d wystąpi gradient koncentracji domieszek. Ponieważ zakładano periodyczną strukturę powierzchni przejściowej, koncentracja domieszek w punkcie C będzie taka jak w punkcie A . Wobec tego największa zmiana koncentracji domieszek wystąpi na bardzo cienkiej warstwie δ oddzielającej poszczególne płytki. Rozkład koncentracji na krawędziach płytek będzie taki, jak pokazano na rys. 16. Jeżeli na małej odległości występuje duża zmiana koncentracji domieszek, a zatem i średniej stałej sieci, fakt



Rys. 16. Zmiany koncentracji domieszek na krawędziach płytki przy płytkowym wzroście kryształów

ten będzie odpowiadał wbudowaniu pewnej ilości dyslokacji krawędziowych, aby usunąć powstałe naprężenia elastyczne w sieci krystalicznej. Jak obliczył Tiller [35], gęstości dyslokacji powstałe w wyniku działania omawianego mechanizmu mogą osiągać wartości nawet 10^6 cm^{-2} .

W pewnych warunkach wzrostu kryształu (tzw. „przechłodzenia strukturalnego”) powstaje tzw. komórkowa struktura kryształu. Powierzchnię przejściową, która utworzyła taką strukturę, pokazano na rys. 17. Struk-



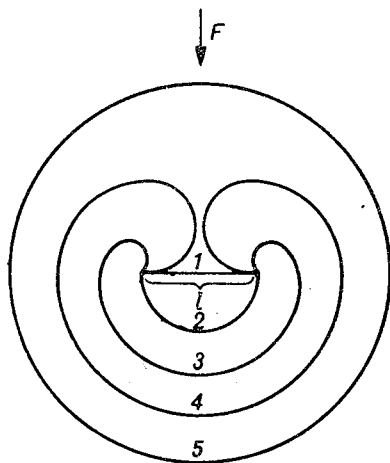
Rys. 17. Obraz metalograficzny powierzchni kryształu, który krystalizował w postaci komórkowej [35]

tura komórkowa charakteryzuje się segregacją domieszek podobną do struktury płytkowej, ale różniącą się geometrycznym układem powierzchni bogatych w domieszki. W tym przypadku powierzchniami tymi są ściany kolumn sześciokątnych równoległych do osi wzrostu kryształów. Analogicznie jak we wzroście płytkowym pewna ilość dyslokacji musi być wprowadzona do obszaru granic poszczególnych komórek, aby usunąć powstałe naprężenia elastyczne. Gęstości dyslokacji towarzyszące powsta-

niu struktury komórkowej również wynoszą 10^6 — 10^8cm^{-2} . Również ważnym faktem jest, że struktura komórkowa powstaje w przypadku obecności domieszek w koncentracji większej od pewnej wartości krytycznej. Wartość krytyczna koncentracji domieszki jest funkcją rodzaju domieszki, warunków wzrostu i waha się w granicach 10^{-2} — $10^{-4} \%$ w stosunku atomowym.

5. MECHANIZM FRANKA-READA

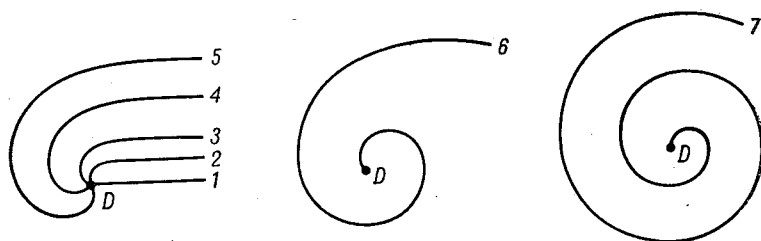
W poprzednich rozdziałach omawiano powstawanie dyslokacji w wyniku różnych zjawisk fizycznych zachodzących w kryształach. Jednak wytworzenie dyslokacji w jeden z wyżej omówionych sposobów nie oznacza zakończenia kształtowania ostatecznej struktury wewnętrznej kryształu. Dopóki istnieje temperatura taka, że siły międzyatomowe są stosunkowo niewielkie i kryształ jest plastyczny, może odbywać się proces zarówno powstawania nowych dyslokacji, jak i ruch uprzednio powstałych. Frank i Read [36] zaproponowali pewien mechanizm, który wyjaśnił



Rys. 18. Schemat działania źródła Franka-Reada w przypadku dyslokacji zakotwiczonej w dwóch punktach

możliwość powstawania nowych dyslokacji w wyniku ruchu uprzednio powstałych. Mechanizm ten jest nazywany mechanizmem Franka-Reada i polega na tym, że dyslokacja poruszając się pod wpływem siły zewnętrznej może zostać zakotwiczona w jednym lub więcej punktach. Przez zakotwiczenie rozumiemy unieruchomienie pewnych punktów dyslokacji, gdy pozostałe części dyslokacji poruszają się tworząc początkowo odcinki w kształcie łuku. Punktami zakotwiczenia mogą być węzły dyslokacyjne, atomy domieszkowe, powierzchnia kryształu lub wewnętrzne pęknięcia. Jeżeli pewien odcinek dyslokacji zostanie zakotwiczony w dwóch punktach, a naprężenie powodujące ruch dyslokacji działa w dalszym ciągu, unieruchomiony odcinek dyslokacji może być źródłem szeregu pętli dyslo-

kacyjnych. Poszczególne fazy powstawania pętli dyslokacyjnej pokazane są na rys. 18. Ostatnią fazą jednego cyklu działania źródła Franka-Reada jest obecność początkowego odcinka zakotwiczonego dyslokacji i pętli dyslokacyjnej, która została „wyemitowana” do kryształu. Cykl taki może



Rys. 19. Schemat działania źródła Franka-Reada w przypadku dyslokacji zakotwiczonej w dwóch punktach

powtarzać się wielokrotnie. Mechanizm Franka-Reada działa również w przypadku zakotwiczenia w jednym punkcie, jednak wynikiem działania będzie dyslokacja spiralna jak na rys. 19.

Aby mechanizm Franka-Reada zaczął działać, musi powstać naprężenie ścinające, którego składowa w kierunku prostopadłym do kierunku dyslokacji zakotwiczonej na węzłach dyslokacyjnych jest większa od [36]:
gdzie:

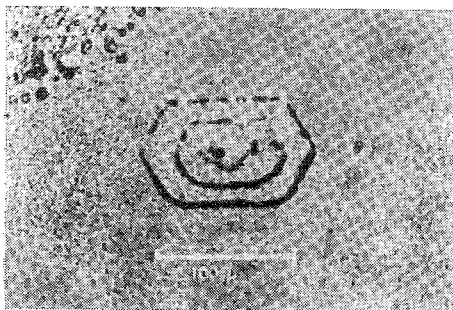
$$\tau = \frac{aGb}{l}, \quad (7)$$

- a — współczynnik liczbowy zależny od stałej elastyczności materiału i częściowo od rodzaju dyslokacji,
- rzędu jedności,
- G — moduł ścinania,
- b — moduł wektora Burgersa,
- l — długość zakotwiczonego odcinka dyslokacji.

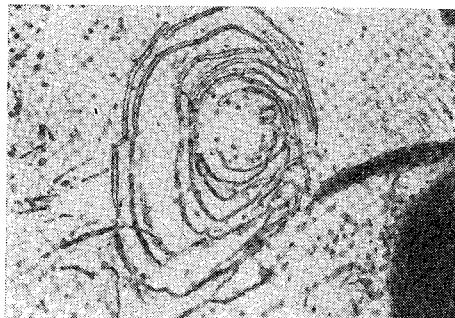
W przypadku zakotwiczenia dyslokacji na powierzchni kryształu lub na atomach domieszkowych krytyczne naprężenie ścinające potrzebne do „uruchomienia” mechanizmu Franka-Reada będzie mniejsze, niż wynikałoby ze wzoru (7) [37 i wg 38].

Mechanizm Franka-Reada został wykryty doświadczalnie w kryształach germanu i krzemu. Na kryształach germanu metodą wykrywania dyslokacji położonych równolegle do powierzchni kryształu opracowali Tyler i Dash [8]. Stosowali oni dyfuzję litu do zdeformowanych kryształów i następnie chemicznie trawienie powierzchni. Otrzymane obrazy działania mechanizmu Franka-Reada pokazano na rys. 20 i 21. Pierwszy obraz uzyskano na kryształ zdeformowanym w temperaturze 550°C, w wyniku czego poślizg zachodził po płaszczyznach o gęstym upakowaniu atomów. W rezultacie otrzymano regularny kształt linii dyslokacyjnych.

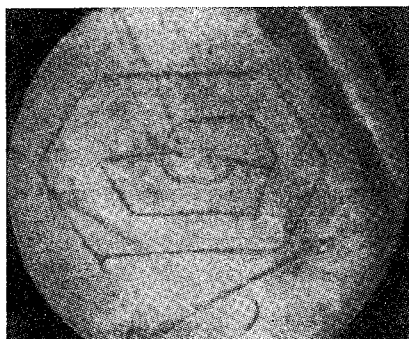
Na rys. 21 pokazano obraz uzyskany na kryształach zdeformowanych w 800°C . Z przyczyn omawianych w rozdziale 2 deformacja przebiegała bezładnie i wobec tego brak jest wyróżnionych kierunków poślizgu. Źródła Franka-Reada na krzemie zostały wykryte przez Dasha [39] przez prześwietlanie płytek krzemowych podczerwienią oraz badanie obrazu przy pomocy



Rys. 20. Przykład ujawnionego źródła Franka-Reada w germanie; deformację przeprowadzono w temperaturze 350°C [8]



Rys. 21. Przykład ujawnionego źródła Franka-Reada w germanie; deformację przeprowadzono w temperaturze 800°C [8]



Rys. 22. Przykład ujawnionego źródła Franka-Reada w krzemie metodą mikroskopii podczerwonej [8]

przetwornika podczerwieni i mikroskopu optycznego. Obraz uzyskany w ten sposób pokazano na rys. 22.

6. DYSKUSJA

Istnieją dwie zasadnicze metody badania dyslokacji w kryształach: 1) trawienie chemiczne niskoindeksowych powierzchni kryształu, w wyniku czego powstają wgłębienia w punktach przecięcia się linii dyslokacyjnej z powierzchnią kryształu [40] i 2) badanie rozpraszania odbitego od powierzchni kryształu wiązki promieni rentgenowskich. Poza tym istnieje grupa metod wykorzystujących powstawanie atmosfery domieszek wokół dyslokacji. Polega ona na dyfuzji atomów domieszkowych (np. litu

lub miedzi), których atomy umieszczają się wokół dyslokacji. Następnie w przypadku germanu trawienie chemiczne ujawnia linie dyslokacyjne biegnące równolegle do badanej powierzchni, gdyż obszary nasycone domieszkami wytrawiają się szybciej [8], natomiast w przypadku krzemu prześwietlanie kryształu promieniami podczerwonymi i oglądanie przy pomocy przetwornika podczerwieni i mikroskopu optycznego pozwala na dokładne śledzenie przebiegu dyslokacji w kryształach [39]. Tę grupę metod badania dyslokacji można w zasadzie traktować jako rozszerzenie metod trawienia chemicznego.

Metody trawienia chemicznego mają tę zaletę, że pozwalają określić zarówno gęstość, jak i rozkład dyslokacji w kryształach [40], [41], [5]. Poza tym niektóre rodzaje trawień pozwalają na określenie różnych typów dyslokacji [41], [5]. Wydaje się, że „duże” wgłębienia otrzymywane przy trawieniu większością mieszanek, jak CP4 , $\text{K}_3\text{Fe}(\text{CN})_6$, $\text{HF}-\text{HNO}_3-\text{H}_2\text{O}$ [5] odpowiadają dyslokacjom liniowym o charakterze krawędziowym lub złożonym, powstałym w wyniku poślizgu plastycznego. Za tym twierdzeniem przemawia fakt korelacji wgłębień otrzymanych w wyniku trawienia chemicznego i dyslokacji badanych przy pomocy prześwietlania promieniami podczerwonymi [2], [3] oraz doświadczenie Vogela i in. [40]. Wydaje się, że ujawnienie dyslokacji o charakterze śrubowym lub przebiegających prawie równolegle do badanej powierzchni jest niemożliwe do uzyskania bez specjalnych zabiegów [8]. W sposób opisany powyżej ujawniają się również dyslokacje powstałe w wyniku nierównomiernego rozkładu domieszek [34]. Natomiast niektóre mieszanki (CP4 , wodny roztwór $\text{HF}-\text{H}_2\text{O}_2-\text{H}_2\text{O}$) wykazują również istnienie małych wgłębień [41], [42], [5], których gęstość wynosi 10^5-10^7cm^{-2} dla przeciętnych kryształów germanu i krzemu. Wydaje się, że wgłębienia te można przypisać obecności dyslokacji powstałych w wyniku tworzenia zespołów luk. Dowód tej hipotezy przeprowadzili dla kryształów glinu Doherty i Davis [43]. Badania autora na kryształach germanu również potwierdzają tę hipotezę [33].

Metody rentgenowskie pozwalają na określenie gęstości dyslokacji odpowiadających gęstości „małych” wgłębień [42], [44]. Jednak interpretacja wyników otrzymanych z metod rentgenowskich przedstawia jeszcze znaczne trudności [44], [45]. Natomiast prace prowadzone przez J. Auleytnera [46] w Instytucie Fizyki PAN pozwalają przypuszczać, że metody rentgenowskie będą mogły służyć również do określania gęstości dyslokacji odpowiadających „dużym” wgłębieniom.

Udział poszczególnych mechanizmów powstawania dyslokacji w różnych materiałach jest różny. Deformacja plastyczna jako przyczyna powstawania dyslokacji będzie odgrywała większą rolę ze wzrostem temperatury obróbki cieplnej, a więc i temperaturą topnienia, ponieważ będą narastały trudności w uzyskaniu jednorodnego rozkładu temperatury w obrabianym kryształach. Powstawanie dyslokacji w wyniku deformacji plastycznej będzie wobec tego uprzywilejowane w przypadku materiałów

o wysokiej temperaturze topnienia. Poza tym deformacja plastyczna będzie również uprzywilejowana w materiałach o dużym współczynniku rozszerzalności cieplnej oraz małej przewodności cieplnej. Jednak należy również stwierdzić, że ten mechanizm powstawania dyslokacji może być w praktyce wyeliminowany. Ostatnie prace Dasha [1], [2], [3] omawiają otrzymywanie monokryształów krzemu wolnych od dyslokacji powstałych w wyniku deformacji plastycznej.

Drugi mechanizm powstawania dyslokacji jest chyba najmniej znany. Wynika to z trudności jego badania i fizycznej strony jego istnienia. Obecności wysokiej koncentracji luk w temperaturze topnienia uniknąć nie można. Natomiast zwiększenie prędkości wzrostu kryształów w celu uniknięcia zespołów luk jest ograniczone określoną prędkością krystalizacji dla danego materiału. Wydaje się, że ten mechanizm będzie stanowił najtrudniejszy problem przy otrzymywaniu kryształów o małej gęstości dyslokacji.

Trzeci mechanizm — powstawanie dyslokacji w wyniku zmiany koncentracji domieszek lub składników — odgrywa pierwszorzędą rolę w półprzewodnikach dwu- i więcej składnikowych, jak np. związki typu $A_{III}-B_V$. W wyniku niedokładnego mieszania fazy ciekłej i różnych współczynników segregacji poszczególnych składników mogą powstać znaczne gęstości dyslokacji będące również przyczyną trudności zastosowania tych materiałów do techniki tranzystorowej. W półprzewodnikach jednorodnych, jak krzem czy german, ten mechanizm powstawania dyslokacji może być w większości przypadków pomijany.

7. ZAKOŃCZENIE

Badanie przyczyn powstawania dyslokacji jest dopiero w początkowej fazie, pomimo dużych osiągnięć praktycznych w tej dziedzinie. Opracowanie ogólnej teorii tego zagadnienia przedstawia duże trudności już na wstępie, gdyż dokładne rozwiązanie równania przewodnictwa cieplnego w celu poznania rozkładu temperatury w kryształach nastrocza wielkie trudności matematyczne. Istniejące przybliżone rozwiązania nie wystarczają do interpretacji osiągniętych wyników doświadczalnych. Poza tym istnieje szereg dalszych trudności zarówno matematycznych, jak i wynikających z braku doświadczeń ze względu na złożoność procesów powstawania różnych rodzajów dyslokacji. Należy się jednak spodziewać, że już najbliższa przyszłość przyniesie znaczne postępy w rozwiązywaniu tego zagadnienia.

W zakończeniu pragnąłbym podziękować p. drowi S. Sikorskiemu z Zakładu Elektroniki IPPT PAN za liczne dyskusje nad zagadnieniami związanymi z powyższym tematem oraz p. drowi J. Auleytnerowi z Instytutu Fizyki Uniw. Warszawskiego za życzliwą krytykę niniejszej pracy.

WYKAZ LITERATURY

1. Dash W. C.: Silicon Crystals free of dislocations. — J. Appl. Phys. 29 (1958), str. 736.
2. Dash W. C.: The Growth of Silicon Crystals Free from Dislocations w „Growth and Perfection of Crystals”, ed. R. M. Doremus, J. Wiley and Sons, 1958.
3. Dash W. C.: Growth of Silicon Crystals Free form Dislocations. — J. Appl. Phys. 30 (1959), nr 4, str. 459.
4. Kobus A.: Podstawy geometrycznej teorii dyslokacji w kryształach. — Rozpr. Elektrot., t. VI, 1960, z. 3, str. 213.
5. Kobus A.: Chemiczne metody ujawniania dyslokacji w kryształach germanu. — Rozpr. Elektrot., t. VI, 1960, z. 3, str. 239.
6. Cottrell A. H.: Dislocations and Plastic Flow in Crystals. Oxford, At the Clarendon Press, 1953.
7. Heidenreich R. D., Shockley W.: Study of Slip in Aluminium Crystals by Electron Microscope and Electron Diffraction Methods. Repts. of Conf. on Strength of Solids. — Phys. Soc. 1948, str. 57.
8. Tyler W. W., Dash W. C.: Dislocation Arrays in Germanium. — J. Appl. Phys. 28 (1957), nr 11, str. 1221.
9. Barret C. S.: Imperfection from Transformation and Deformation w “Imperfections in Nearly Perfect Crystals”, ed. W. Shockley i in., J. Wiley and Sons, New York 1952.
10. Pearson G. L., Read W. T. jr., Morin F. J.: Dislocations in Plastically Deformed Germanium. — Phys. Rev. 93 (1954), nr 4, str. 666.
11. Billig E.: Some Defects in Crystals Grown from the Melt. I. Defects Caused by Thermal Stresses. — Proc. Roy. Soc. A235 (1956), nr 1200, str. 37.
12. Billig E.: Defects in Germanium Crystals Grown from the Melt. — Brit. J. Appl. Phys. 7 (1956), str. 375.
13. Penning P.: Generation of Imperfections in Germanium Crystals by Thermal Strain. — Philips Res. Rep. 13 (1958), nr 1, str. 79.
14. Bennet D. C., Sawyer B.: Single Crystals of Exceptional Perfection and Uniformity by Zone Melting. — Bell Syst. Techn. J. 35 (1956), nr 3, str. 637.
15. Wagner R. S.: Production of Dislocations in Germanium by Thermal Shock. — J. Appl. Phys. 29 (1958), nr 12, str. 1679.
16. Brochocki A., Kobus A.: Wpływ warunków wyciągania monokryształów germanu na rozkład dyslokacji krawędziowych. — Przegląd Elektroniki, t. I, 1960, z. 2, str. 144.
17. Kobus A., Brochocki A.: Powstawanie dyslokacji w germanie w wyniku zmian prędkości wyciągania monokryształu. — Archiwum Elektrot. 9 (1960), z. 4, str. 717.
18. Brochocki A., Świdorski J.: Wyznaczanie kształtu powierzchni przejściowej krystalizacji germanu metodą wykorzystującą objętościowe zjawisko fotowoltaiczne. — Przegl. Elektroniki 1 (1960), nr 4, str. 380.
19. Brochocki A., Świdorski J.: Photoelectric Method for Determining Interface Trace Shape in Germanium Monocrystals. — Bull. Acad. Pol. Sci. Ser. techn. 8 (1960), nr 9, str. 537.
20. Nye J. F.: Some Geometrical Relations in Dislocated Crystals. — Acta Metallurg. 1 (1953), str. 153.
21. Kröner E.: Der fundamentale Zusammenhang zwischen Versetzungsdichte und Spannungsfunktionen. — Zf. Phys. 142 (1955), nr 4, str. 463.
22. Kröner E., Rieder G.: Kontinuumstheorie der Versetzungen — Zf. Phys. 145 (1956), nr 4, str. 424.

23. Indenbom W. L.: Makroskopическая теория образования дислокации при росте кристаллов. — Кристаллография 2 (1957), nr 5, str. 594.
24. Indenbom W. L.: Dyslokacje w kryształach. — Post. Fiz., 10 (1959), nr 6, str. 637.
25. Churchman A. T., Geach G. A., Winton J.: Deformation Twinning in Materials of the A4 (Diamond) Crystal Structure. — Proc. Roy. Soc. A238 (1956), str. 194.
26. Seitz F.: Imperfections in Nearly Perfect Crystals. A synthesis w "Imperfections in Nearly Perfect Crystals", ed. W. Shockley i in., J. Wiley and Sons, 1952, str. 3.
27. Nabarro F. R. N.: Deformation of Crystals by the Motion of Single Ions. — Repts. of a Conf. on Strength of Solids. — Phys. Soc., London 1948, str. 75.
28. Cottrell A. H., Bilby B. A.: Dislocation Theory of Yielding and Strain Ageing of Iron. — Proc. Phys. Soc. A62 (1949), st. 49.
29. Seitz F.: On the Formation of Dislocation from Vacancies. — Phys. Rev. 79 (1950), str. 890.
30. Crusard P.: Diffusion et interaction des impuretés et des défauts de structure dans les métaux. — Métaux et Corrosion, 1950.
31. Read W. T. jr.: Dislocations in Crystals. J. Wiley and Sons, New York 1953.
32. Teghtsoonian E., Chalmers B.: The Macromosaic Structure of Tin Single Crystals. — Can. J. Phys. 29 (1951), nr 5, str. 370.
33. Kobus A.: dane nie opublikowane.
34. Goss A. J., Benson K. K., Pfann W. G.: Dislocation at Compositional Fluctuations in Germanium—Silicon Alloys. — Acta Metallurg. 4 (1956), nr 3, str. 332.
35. Tillier W. A.: Production of Dislocations During Growth from the Melt. — J. Appl. Phys. 29 (1958), nr 4, str. 611.
36. Frank F. C., Read W. T.: Multiplication Processes for Slow Moving Dislocations. — Phys. Rev. 79 (1950), str. 722.
37. Hart E. W.: The Slip Process During Yield—Point Deformation. — Acta Metallurg. 2 (1954), str. 416.
38. van Bueren H. G.: Influence Lattice Defects on the Electrical Properties of Cold—Worked Metals. — Philips Res. Rep. 12 (1957), str. 1 i 120.
39. Dash W. C.: The Observation of Dislocations in Silicon w "Dislocations and Mechanical Properties of Crystals", ed. W. Fisher i in., J. Wiley and Sons, 1957.
40. Vogel F. L., Pfann W. G., Corey H. E., Thomas E. E.: Observations of Dislocations in Lineage Boundaries in Germanium. — Phys. Rev. 90 (1953), str. 489.
41. Ellis S. G.: Dislocations in Germanium. — J. Appl. Phys. 26 (1956), str. 1140.
42. Kurtz A. D., Kulin S. A., Averbach B. L.: Effect of Dislocations on the Minority Carrier Lifetime in Semiconductors. — Phys. Rev. 101 (1956), nr 4, str. 1285.
43. Doherty P. E., Davis R. S.: The Formation of Surface Pits by the Condensation of Vacancies. — Acta Metallurg. 7 (1959), nr 2, str. 118.
44. Gay P., Hirsch P. B., Kelly A.: The Estimation of Dislocation Densities in Metals from X—Ray Data. — Acta Metallurg. 1 (1953), str. 316.
45. Batterman B. W.: X—Ray Integrated Density of Germanium Effect of Dislocations and Chemical Impurities. — J. Appl. Phys. 30 (1959), str. 508.
46. Auleytner J.: Informacja prywatna, 1959.

А. КОВУС

ВОЗНИКНОВЕНИЕ ДИСЛОКАЦИИ В КРИСТАЛЛАХ

Резюме

В статье рассматриваются причины, вызывающие дислокацию в кристаллах. Существуют три основных механизма, вызывающие дислокации: 1) пластическая деформация, 2) возникновение комплекса колонии вакансий и его провал, 3) неравномерность разложения примеси (или элементов).

Пластическая деформация кристалла может возникнуть при получении кристалла во время его охлаждения до температуры среды. Вследствие неравномерной отдачи тепла массой кристалла возникают разницы температур между поверхностью и центром, которые способствуют возникновению термических напряжений. Если эти напряжения превысят сопротивление кристалла, — наступает скольжение некоторых участков кристалла, вследствие чего вводятся дислокации. Это явление выступает преимущественно в материалах, отличающихся высокой температурой термической обработки, высоким коэффициентом линейного расширения и низкой теплопроводностью. При соответствующих условиях, в результате термических колебаний сетки, в кристалле всегда наблюдается правдоподобие перехода части атомов из узлового положения на поверхность кристалла или в междузловое положение. Возникшие, таким образом, вакансии могут создавать комплексы в виде дисков, которые вследствие провала в результате превышения критической величины — создают дислокационные петли. Этот механизм будет привилегирован только в таких условиях, когда вакансий будут обладать достаточной подвижностью в течение достаточно продолжительного времени для создания дисков. Такое явление происходит в тех кристаллах, которые долго находились в высоких температурах.

Третье явление, вызывающее дислокацию, выступает при наличии неравномерности распада примесей в элементарном полупроводнике (германий, кремний), или в случае отступления от стехиометрического состояния полупроводников, состоящих из двух или более элементов. На границе сфер с разным содержанием примесей или элементов сталкиваются сферы кристаллов с разной постоянной сетки. Если эта разница достаточно велика, — возникает ряд краевых дислокаций, которые ослабляют напряжения, появившиеся в сетке в результате разницы постоянных. Эти дислокации будут возникать главным образом в полупроводниках, состоящих из двух или более элементов.

Кроме того, за окончательное разложение и плотность дислокации отвечает механизм увеличения многократности дислокации Франка-Рида. Он состоит в „эмитировании” дислокационных петель с участка „закрепленной”, уже существующей дислокации в результате действия разных напряжений в кристалле, которые могут быть меньше напряжений, вызывающих пластическую деформацию идеального кристалла.

А. KOBUS

ORIGIN OF DISLOCATIONS IN CRYSTALS

Summary

Basic ways of origin of dislocations in crystals have been discussed in the article: 1) origin of dislocations as a result of plastic deformation, 2) origin of dislocations from vacancy clusters, 3) origin of dislocations due to disproportional distribution of impurities (of components). The part that the Frank-Read mechanism

takes in the final forming of arrangement of dislocations in crystals has been discussed then. The conditions in which respective mechanisms of origin of dislocations play a decisive role have been given. The possibilities of detecting dislocations which arose as a result of operating of respective mechanisms have been discussed at the end.

A. KOBUS

FORMATION DES DISLOCATIONS DANS LES CRISTAUX

Résumé

On discute les principaux mécanismes formants des dislocations dans les cristaux: 1) formation de dislocation à cause de déformation plastique, 2) formation de dislocation du groupe des lacunes, 3) formation des dislocations à cause d'inégal placement des impuretés. Ensuite on discute la participation du mécanisme de Frank-Read en formation finale du système des dislocations dans les cristaux. On présente aussi les conditions, sous quelles les mécanismes particuliers formants les dislocations, jouent un rôle décisif.

Pour terminer on discute les possibilités de la découverte des dislocations formées à cause de l'action des mécanismes particuliers.

A. KOBUS

DAS ENTSTEHEN DER VERSETZUNGEN IN DEN KRISTALLEN

Zusammenfassung

Die Hauptmechanismen des Entstehens der Versetzungen in den Kristallen wurden im Artikel besprochen: 1) das Entstehen der Versetzung infolge der plastischen Verformung, 2) das Entstehen der Versetzung aus den Leerstellengruppen, 3) das Entstehen der Versetzungen infolge des ungleichmässigen Verlegung der Verunreinigungen (Bestandteile). Zunächst wurde die Beteiligung des Frank-Read Mechanismus an der endgültigen Gestaltung des Versetzungssystems in den Kristallen besprochen. Es wurden ebenfalls Bedingungen angegeben, in welchen einzelne Versetzungsmechanismen eine entscheidende Rolle spielen.

Zum Schluss wurden die Möglichkeiten des Entdeckens der infolge der Wirkung einzelner Mechanismen entstandenen Versetzungen besprochen.

POLSKA AKADEMIA NAUK
INSTYTUT PODSTAWOWYCH PROBLEMÓW TECHNIKI

ROZPRAWY ELEKTROTECHNICZNE

TOM VII • ZESZYT 3

KWARTALNIK

PAŃSTWOWE WYDAWNICTWO NAUKOWE
WARSZAWA 1961

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*

PAŃSTWOWE WYDAWNICTWO NAUKOWE
Warszawa, Miodowa 10

Nakład 640 (503 + 137)	Oddano do składania 15.VII.1961
Ark. wyd. 8, ark. druk. 7,5	Podpisano do druku w listopadzie 1961
Papier druk. sat. kl. III, 80 g, 70 × 100	Druk ukończono w listopadzie 1961
Zamówienie nr 1175/61	Cena zł 25.— S-24

JULIUSZ KULIKOWSKI

O pewnych zagadnieniach związanych z wyliczeniem elementów odwróconej macierzy kowariancji pasywnych zakłóceń radiolokacyjnych

Rękopis dostarczono 12.9.1960

W artykule podano metodę wyliczenia elementów odwróconej macierzy kowariancji zakłóceń, opartą o tzw. „dwustronną transformację typu Z”. Metoda ta ma zastosowanie w teorii odbioru sygnałów radiolokacyjnych na tle zakłóceń pasywnych. Metodę rozszerzono na wypadki odbioru wielokanałowego oraz odbioru sygnałów ze zmienną częstotliwością powtarzania impulsów.

1. WSTĘP

W teorii odbioru sygnałów radiolokacyjnych na tle zakłóceń pasywnych zachodzi konieczność określenia N -wymiarowej funkcji gęstości prawdopodobieństwa zakłóceń. Można wykazać, że o ile w kanale odbiorczym nie zachodzą procesy nieliniowe (z wyłączeniem takich procesów, jak przemiana częstotliwości lub dwupółówkowa detekcja koherentna sygnału, które powodują tylko przesunięcie widma sygnału na skali częstotliwości), to wspomniana funkcja gęstości ma postać funkcji Gaussa. Ze względu na możliwość eliminacji zakłóceń pasywnych szczególnie ważna jest lokalna funkcja gęstości zakłócenia, wzięta dla danego punktu analizowanej przestrzeni, rozważanego w przedziale czasu $(0, T)$. Funkcja ta ma postać:

$$w(\mathbf{z}) = \frac{1}{(2\pi)^N \sqrt{|\tilde{Q}|}} \exp \left[-\frac{1}{2} \mathbf{Q}^{-1}(\mathbf{z}) \right], \quad |\tilde{Q}| \neq 0, \quad (1)$$

gdzie:

$\mathbf{z} = \mathbf{x} + j\mathbf{y} = \{x_1 + jy_1, x_2 + jy_2, \dots, x_N + jy_N\}$ — wektor zakłóceń.

Dla określenia wielkości $|\tilde{Q}|$ oraz $\mathbf{Q}^{-1}(\mathbf{z})$ wprowadźmy wpierw pojęcie macierzy współczynników kowariancji $\mathbf{Q} = \|\mathbf{R}_{\mu\nu}\|$, gdzie

$$\mathbf{R}_{\mu\nu} = \overline{(z_\mu - \bar{z}_\mu)(z_\nu^* - \bar{z}_\nu^*)}. \quad (1a)$$

Kreską poziomą oznaczono tu operację statystycznego uśrednienia, a gwiazdka — wielkość zespolono-sprzężoną. Zaznaczmy, że ponieważ proces przypadkowy $z(t)$ traktujemy jako proces zespolony (tj. uwzględniamy sinusoidalną i kosinusoidalną składową zakłóceń), to współczynniki kowariancji $R_{\mu\nu}$ są również liczbami zespolonymi, spełniającymi zależność:

$$R_{\mu\nu} \equiv R_{\nu\mu}^* = r_{\mu\nu} e^{j(\varphi_\mu - \varphi_\nu)}. \quad (1b)$$

W zapisie macierzowym wygodniej jest niekiedy przedstawić obie składowe wektorów oddzielnie, gdyż charakteryzują one składowe sygnały występujące oddzielnie w dwóch kwadraturowych kanałach odbiornika koherentnego. Zgodnie z tzw. izomorfizmem liczb zespolonych i ich reprezentacji macierzowych zachodzi równoważność następujących zapisów:

$$R_{\mu\nu} \Leftrightarrow \begin{vmatrix} r_{\mu\nu} \cos(\varphi_\mu - \varphi_\nu) & r_{\mu\nu} \sin(\varphi_\mu - \varphi_\nu) \\ -r_{\mu\nu} \sin(\varphi_\mu - \varphi_\nu) & r_{\mu\nu} \cos(\varphi_\mu - \varphi_\nu) \end{vmatrix}, \quad (1c)$$

$$R_{\nu\mu} \Leftrightarrow \begin{vmatrix} r_{\nu\mu} \cos(\varphi_\nu - \varphi_\mu) & r_{\nu\mu} \sin(\varphi_\nu - \varphi_\mu) \\ -r_{\nu\mu} \sin(\varphi_\nu - \varphi_\mu) & r_{\nu\mu} \cos(\varphi_\nu - \varphi_\mu) \end{vmatrix}, \quad (1d)$$

przy czym:

$$r_{\mu\nu} \equiv r_{\nu\mu}. \quad (1e)$$

Macierz kowariancji $\tilde{Q}$ można teraz zapisać w postaci:

$$\tilde{Q} = \tilde{T} \cdot \tilde{R} \cdot \tilde{T}^{-1} = \begin{vmatrix} e^{j\varphi_1} & 0 \\ & e^{j\varphi_2} \\ & & \ddots \\ 0 & & & e^{j\varphi_N} \end{vmatrix} \cdot \begin{vmatrix} r_{11} & r_{12} & \dots & r_{1N} \\ r_{21} & r_{22} & \dots & r_{2N} \\ \vdots & \vdots & & \vdots \\ r_{N1} & r_{N2} & \dots & r_{NN} \end{vmatrix} \cdot \begin{vmatrix} e^{-j\varphi_1} & & & 0 \\ & e^{-j\varphi_2} & & 0 \\ & & \ddots & \\ 0 & & & e^{-j\varphi_N} \end{vmatrix} \quad (1f)$$

Wyrażenie $Q^{-1}(\mathbf{z})$ występujące we wzorze (1) jest tzw. formą kwadratową Hermite'a zmiennych $\mathbf{z}$:

$$Q^{-1}(\mathbf{z}) = \sum_{\mu=1}^N \sum_{\nu=1}^N z_\mu \xi_{\mu\nu} z_\nu^*, \quad (1g)$$

gdzie współczynniki $\xi_{\mu\nu}$ są elementami macierzy $\tilde{P} = \|\xi_{\mu\nu}\|$ związanej z macierzą kowariancji $\tilde{Q}$ zależnością:

$$\tilde{P} = \tilde{Q}^{-1} \quad \text{lub} \quad \tilde{P} \cdot \tilde{Q} = E, \quad (1h)$$

przy czym iloczyn rozumie się w sensie macierzowym, a E jest macierzą jednostkową. Macierz $\tilde{P}$ nazywać będziemy odwróconą macierzą kowariancji. Wielkość $|\tilde{Q}|$ oznacza wyznacznik macierzy kowariancji $\tilde{Q}$.

W praktycznych zastosowaniach do teorii syntezy optymalnych układów odbiorczych funkcje gęstości prawdopodobieństw występują jako składniki pewnych różnic lub ilorazów, w związku z czym mogą być wyznaczone z dokładnością do stałych współczynników (równych dla obu składników). Dlatego też wyznaczenie wielkości wyznacznika $|\tilde{Q}|$ nie odgrywa w tych zastosowaniach zasadniczej roli.

Ze względu na dalsze operacje na macierzach zakładamy, że $|\tilde{Q}| \neq 0$.

W celu obliczenia macierzy $\tilde{Q}^{-1}$ zauważmy po pierwsze, że postać wzoru (1 f) narzuca możliwość wykorzystania następującej zależności macierzowej:

$$\tilde{Q}^{-1} = (\tilde{T} \cdot \tilde{R} \cdot \tilde{T}^{-1})^{-1} = \tilde{T} \cdot \tilde{R}^{-1} \cdot \tilde{T}^{-1}, \quad (2)$$

tj. nasze rozważanie wystarczy w zasadzie ograniczyć do sposobu obliczenia macierzy $\tilde{R}^{-1}$. Z zależności $\tilde{R}^{-1} \cdot \tilde{R} = E$ wynika, że elementy macierzy $\tilde{R}^{-1} = \|\zeta_{mn}\|$ powinny spełniać układ następujących zależności algebraicznych:

$$\sum_{k=1}^N \zeta_{mk} r_{kn} = \delta_{mn}, \quad (m, n = 1, 2, \dots, N). \quad (3)$$

Rozwiązanie podobnego układu N^2 równań liniowych byłoby niezmiernie pracochłonne (zauważmy, że stopień macierzy N może sięgać wielu tysięcy!). Dlatego też niektórzy autorzy (np. [1]) proponują zastąpienie układu (3) odpowiednim przybliżającym układem równań całkowych, które w pewnych wypadkach można rozwiązać stosunkowo prosto. Metoda taka nasuwa jednak pewne trudności formalne, związane z aproksymacją ciągów wielkości dyskretnych przez funkcje ciągłe, oraz funkcje uogólnione typu δ Diraca. Najbardziej właściwym aparatem dla rozwiązania sformułowanego zagadnienia wydaje się aparat tzw. „transformacji typu Z” rozpracowany w zastosowaniu do analizy obwodów impulsowych [2], który jak pokażemy dalej, pozwala rozwiązać szereg praktycznie ważnych przykładów.

2. WYPADEK DWUSTRONNIE NIESKOŃCZONEJ MACIERZY STACJONARNEJ

Jest to wypadek rachunkowo najprostszy, zachodzący wówczas, gdy czas obserwacji sygnału jest obustronnie nieograniczony (praktycznie oznacza to, że jest on znacznie większy od czasu korelacji zakłóceń), a sam ciąg stochastyczny $\{z_n\}$ jest stacjonarny. Ostatni warunek jest spełniony, jeżeli częstotliwość repetycji radiolokatora $f_p = \frac{1}{T_p}$ jest stała, a zakłócenie pasywne jest stacjonarne w szerokim sensie.

Wówczas:

$$R_{\mu\nu} \equiv R_{|\mu-\nu|} \quad (4)$$

i równanie (3) przybiera postać:

$$\sum_{k=-\infty}^{\infty} r(m T_p - k T_p) \cdot \zeta_{kn} = \delta_{mn}, \quad (m, n = \dots, -2, -1, 0, 1, 2, \dots). \quad (5)$$

Dokonajmy następującej transformacji funkcjonalnej obu części równa-

nia (5): względem indeksu m :

$$\sum_{m=-\infty}^{\infty} Z^m \cdot \sum_{k=-\infty}^{\infty} r(m T_p - k T_p) \cdot \zeta_{kn} = \sum_{m=-\infty}^{\infty} Z^m \delta_{mn}, \quad (n = \dots - 2, -1, 0, 1, 2, \dots). \quad (6)$$

Operacja jest dopuszczalna, jeżeli założymy, że wchodzące do wyrażenia (5) ciągi liczbowe zależne od indeksu m zbieżne są nie wolniej od pewnego ciągu geometrycznego $\{Z_0^m\}$ o promieniu zbieżności R_0 równym jedności. W praktyce warunek ten będzie zawsze spełniony, gdyż funkcja korelacji zakłóceń pasywnych tłumiona jest co najmniej jak funkcja wykładnicza. Jest to zatem jedna z możliwych postaci warunku dopuszczalności operacji algebraicznych wykonywanych na macierzach nieskończonych, jakie zostały omówione szczegółowo w pracy [3]. Zatem z wyrażenia (6) otrzymamy:

$$R(Z) \cdot X_n(Z) = Z^n, \quad (7)$$

gdzie:

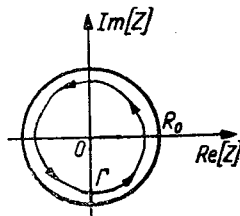
$$R(Z) = \sum_{m=-\infty}^{\infty} Z^m \cdot r(m T_p), \quad (7a)$$

$$X_n(Z) = \sum_{m=-\infty}^{\infty} Z^m \cdot \zeta_{mn}. \quad (7b)$$

Stąd:

$$X_n(Z) = Z^n \cdot [R(Z)]^{-1}. \quad (8)$$

Jak wynika z własności „transformacji typu Z ”, gdzie Z jest operatorem opóźnienia procesu transformowanego o czas T_p , różnica między ciągami $\{\zeta_{ko}\}$ i $\{\zeta_{kn}\}$ sprowadza się wyłącznie do przesunięcia o n pozycji wy-



Rys. 1

razów tego ostatniego, a więc macierz $\|\zeta_{mn}\|$ jest również (w rozważanym przez nas przypadku) macierzą o własności $\zeta_{mn} \equiv \zeta_{|m-n|}$. Wyrazy ciągu $\{\zeta_{kn}\}$ można znaleźć stosując znany wzór całkowy Cauchy'ego:

$$\zeta_{kn} = \frac{1}{2\pi j} \int_{\Gamma} \frac{X_n(Z)}{Z^{k+1}} dZ, \quad (9)$$

gdzie droga całkowania Γ okrąży początek współrzędnych odwrotnie do biegu wskazówek zegara i znajduje się wewnątrz koła o promieniu R_0 (patrz rys. 1).

W zastosowaniu do syntezy optymalnych przeddetekcyjnych filtrów impulsowych, których funkcja przenoszenia, jak wykazano w pracy [4], identyczna jest z funkcją $X(Z)$, wygodniej jest poprzestać na wyniku użytym w formie (8).

W celu określenia funkcji $R(Z)$ według wzoru (7a), zauważmy, że ze względu na symetrię funkcji korelacyjnej $r(\tau)$ funkcję $R(Z)$ można przedstawić w postaci:

$$R(Z) = R_+(Z) + R_-(Z) - r(0), \quad (10a)$$

gdzie:

$$R_+(Z) = \sum_{m=0}^{\infty} Z^m r(m T_p); \quad R_-(Z) = \sum_{m=0}^{\infty} Z^m r(m T_p), \quad (10b)$$

$$R_+(Z) \equiv R_-(Z^{-1}). \quad (10c)$$

Dla określenia funkcji $R_+(Z)$ można z kolei posłużyć się metodą tzw. „transformacji typu D” określającej związek między transformacją Laplace’a $\tilde{R}_+(p) = \mathcal{L}[r(\bar{t})]$ funkcji ciągłej $r_+(\bar{t})$ zadanej na półosi dodatniej $\bar{t}$ (gdzie $\bar{t}$ jest unormowanym czasem ciągłym, określonym wzorem $\bar{t} = \frac{t}{T_p}$) a transformacją dyskretną tej funkcji $\tilde{R}_+^*(q)$, wziętą z okresem próbkowania $T = 1$ [2]:

$$\tilde{R}_+^*(q) = D_T[\tilde{R}_+(p)] \frac{1}{2} r(0) + \sum_{\nu=-\infty}^{\infty} \tilde{R}_+(q + 2\pi j \nu). \quad (11)$$

Przejście od funkcji $\tilde{R}_+^*(q)$ do szukanej funkcji $R(Z)$ odbywa się następnie przez podstawienie nowej zmiennej $Z = e^{-q}$ do wyrażenia (11).

W wypadku gdy funkcja korelacji zadana jest w postaci naogólnionego szeregu wykładniczego:

$$r(\tau) = \sum_k D_k e^{\vartheta_k |\tau|}, \quad (12)$$

gdzie D_k, ϑ_k — współczynniki zespolone, bardziej wygodna jest metoda bezpośrednia polegająca na sumowaniu szeregów geometrycznych, jak to zostanie pokazane na przykładzie.

Przykład 1.

Niech funkcja korelacyjna ma postać $r(\tau) = e^{-\frac{|\tau|}{\tau_0}}$. Transformacja Laplace’a na dodatniej półosi τ tej funkcji ma postać $R(p) = \frac{\tau_0}{1 + p\tau_0}$. Ponieważ okres próbkowania w naszym wypadku wynosi T_p , więc przetransformujemy funkcję $R(p)$ tak, aby odpowiadała ona funkcji $r_1(\tau) = r\left(\frac{\tau}{T_p}\right)$, której okres próbkowania równy jest jedności. Dla funkcji $r_1(\tau)$ mamy na-

stępujące wyrażenie transformacji Laplace'a.

$$R_1(p) = \frac{1}{T_p} R\left(\frac{p}{T_p}\right) = \frac{\frac{\tau_0}{T_p}}{1 + p \frac{\tau_0}{T_p}}.$$

Stosując wzór (11) otrzymamy:

$$\tilde{R}_+(q) = \frac{1}{2} + \sum_{\nu=-\infty}^{\infty} \frac{\frac{\tau_0}{T_p}}{1 + (q + 2\pi j\nu) \frac{\tau_0}{T_p}} = \frac{1}{2} \left[\operatorname{cth} \frac{1}{2} \left(\frac{T_p}{\tau_0} + q \right) + 1 \right],$$

gdzie ostatnie zależności otrzymano korzystając ze znanych rozkładów w szeregi nieskończone (patrz [5] wzór 1.421.4, str. 50). Podstawiając $Z = e^{-q}$ oraz $Z_0 = e^{-\frac{T_p}{\tau_0}}$ otrzymamy:

$$R_+(Z) = \frac{1}{1 - Z_0 Z},$$

a stosując (10c) uzyskamy stąd:

$$R_-(Z) = \frac{1}{1 - Z_0 Z^{-1}}.$$

Sumując następnie zgodnie ze wzorem (10a) otrzymamy:

$$R(Z) = R_+(Z) + R_-(Z) - r(0) = \frac{1 - Z_0^2}{1 + Z_0^2 - Z^{-1} Z_0 - Z Z_0}.$$

Stosując wreszcie wzór (8) otrzymamy wyrażenie dla widma odwróconej macierzy kowariancji:

$$X_n(Z) = \frac{Z^n}{1 - Z_0^2} (1 + Z_0^2 - Z^{-1} Z_0 - Z Z_0).$$

Wyrażenie dla elementów odwróconej macierzy kowariancji ξ_{kn} można uzyskać w tym wypadku bezpośrednio, tj. nie stosując odwrotnego przekształcenia całkowego (9). W tym celu wystarczy przypomnieć, że wyraz Z jest operatorem przesunięcia ciągu wartości dyskretnych o jedno miejsce w prawo. Zatem współczynniki ξ_{kn} mają postać następującą:

$$\xi_{kn} = \begin{cases} \frac{1 + Z_0^2}{1 - Z_0^2} & \text{dla } k = n, \\ -\frac{Z_0}{1 - Z_0^2} & \text{dla } k = n \pm 1, \\ 0 & \text{dla pozostałych } k. \end{cases}$$

Wynik powyższy jest zgodny ze znanymi w literaturze wyrażeniami dla

elementów odwróconej macierzy kowariancji przy wykładniczej funkcji korelacji [6].

Przykład 2.

Rozważony przykład można łatwo rozszerzyć na wypadek sumy zakłóceń pasywnych o wykładniczej funkcji korelacji oraz szumów białych. Wówczas z dokładnością do stałej:

$$R_1(Z) = R(Z) + N_0,$$

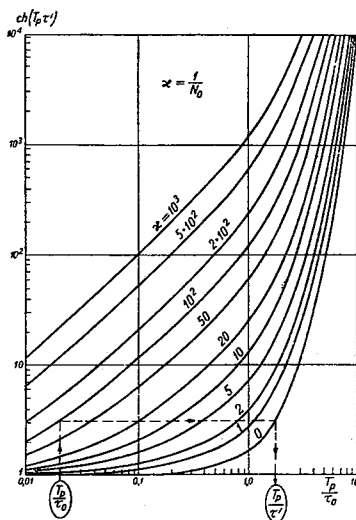
gdzie $R(Z)$ określono w przykładzie poprzednim, a N_0 oznacza stosunek mocy szumów do mocy zakłóceń korelowanych. Po kilku przekształceniach uzyskamy widmo odwróconej macierzy kowariancji w postaci:

$$X_n(Z) = \frac{Z^n}{N_0} \left[1 - \frac{2 \operatorname{sh} \frac{T_p}{\tau_0}}{N_0(Z_1^{-2} - 1)} \cdot \frac{1}{Z - Z_1} - \frac{2 \operatorname{sh} \frac{T_p}{\tau_0}}{N_0(Z_1^{-1} - Z_1)} \cdot \frac{1}{1 - Z_1 Z} \right],$$

gdzie stałe Z_1 określone są przez związki:

$$\operatorname{ch} \frac{T_p}{\tau'} = \frac{1}{2} (Z_1 + Z_1^{-1}) = \operatorname{ch} \frac{T_p}{\tau_0} + \frac{1}{N_0} \operatorname{sh} \frac{T_p}{\tau_0}, \quad Z_1 = e^{\frac{T_p}{\tau'}}.$$

Znalezienie pierwiastków powyższego równania może być ułatwione przez zastosowanie nomogramu przedstawionego na rys. 2.



Rys. 2.

Przykład 3.

W przypadku funkcji korelacyjnych aproksymowanych szeregami wykładniczymi typu (12) wygodniej jest zastosować metodę bezpośrednią obliczania widma $R(Z)$. Rozważmy wypadek gaussowskiej funkcji kore-

lacji $r(\tau) = e^{-\frac{(\tau T_p)^2}{2\tau_0^2}}$. Funkcja ta może być aproksymowana wyrażeniem [4]:

$$\begin{aligned} r(\tau) &= 1,637 e^{-1,098 \frac{|\tau| T_p}{\tau_0}} \cdot \cos \left[0,8087 \frac{|\tau| T_p}{\tau_0} - 0,91 \right] = \\ &= 0,8185 [\exp(\vartheta \cdot |\tau|) + \exp(\vartheta^* \cdot |\tau|)], \end{aligned}$$

gdzie:

$$\vartheta = -1,098 \frac{T_p}{\tau_0} + j \left(0,8087 \frac{T_p}{\tau_0} - 0,9 \right).$$

Jej widmo dyskretne można wyliczyć stosując znane wyrażenia na sumę szeregu geometrycznego:

$$\begin{aligned} R_+(Z) &= \sum_{m=0}^{\infty} Z^m r(m\tau) = 0,8185 \left[\sum_{m=0}^{\infty} Z^m e^{m\vartheta\tau} + \sum_{m=0}^{\infty} Z^m e^{m\vartheta^*\tau} \right] = \\ &= 0,8185 \left[\frac{1}{1 - Z e^{\vartheta\tau}} + \frac{1}{1 - Z e^{\vartheta^*\tau}} \right]. \end{aligned}$$

Odpowiednio też:

$$R_-(Z) = 0,8185 \left[\frac{1}{1 - Z^{-1} e^{\vartheta\tau}} + \frac{1}{1 - Z^{-1} e^{\vartheta^*\tau}} \right].$$

Stąd według (10a) i po przekształceniach otrzymamy:

$$R_1(Z) = R(Z) + N_0 = \frac{aq^2 + bq + c}{dq^2 + eq + f},$$

gdzie:

$$q = \frac{Z + Z^{-1}}{2},$$

$$a = (8 - 4K) Z_1 Z_1^*,$$

$$b = (4K - 12) \frac{Z_1 + Z_1^*}{2} + (4K - 4) Z_1 Z_1^* \frac{Z_1 + Z_1^*}{2},$$

$$c = 4 - K + (4 - 2K) \frac{Z_1^2 + Z_1^{*2}}{2} - K Z_1^2 Z_1^{*2},$$

$$d = 4 Z_1 Z_1^*,$$

$$e = -2(Z_1 + Z_1^*)(1 + Z_1 Z_1^*),$$

$$f = 1 + Z_1^2 Z_1^{*2} + Z_1^2 + Z_1^{*2},$$

$$K = N_0 - 1, \quad Z_1 = e^{\vartheta\tau}.$$

Dalej sposób liczenia macierzy odwrotnej nie różni się zasadniczo od poprzednio opisanego.

3. WYPADEK DWUSTRONNIE NIESKOŃCZONYCH MACIERZY KOWARIANCJI ZAKŁÓCEŃ W UKŁADZIE WIELOKANALOWYM

Wzory ogólne przytoczone we wstępie zachowują swą moc również w wypadku pracy wielokanałowej, tj. przy pracy w systemie tzw. „diversity częstotliwości”, „diversity przestrzennego” lub „diversity polaryzacji”. W każdym z tych wypadków ciągi próbek sygnałowych odbierane w poszczególnych kanałach charakteryzowane są przez macierze swych współczynników korelacji własnej oraz współczynników korelacji wzajemnych, międzykanałowych. Wartości współczynników korelacji są podobnie jak w wypadku jednokanałowym dyskretnymi wartościami odpowiednich funkcji auto- i wzajemnokorelacyjnych. Odpowiednią macierz kowariancji można ponownie przedstawić w formie (1f), gdzie jednak między wykładnikami elementów macierzy $\tilde{T}$ nie zachodzą już relacje zwykłej wielokrotności: $\varphi_n = n \cdot \varphi_1$, o ile odpowiednie fazy nie odnoszą się do sygnałów obserwowanych w tym samym kanale. W ogólnym wypadku mamy:

$$\varphi_{k;n} = \varphi_{k;0} + t_{k;n} \cdot \omega_k, \quad (k = 1, 2, \dots, K), \quad (13)$$

gdzie:

- $\varphi_{k;n}$ — średnia faza zakłócenia obserwowana w k -tym kanale w momencie czasu $t_{k;n}$,
- $\varphi_{k;0}$ — pewna stała faza początkowa zakłócenia obserwowanego w k -tym kanale;
- ω_k — średnia częstotliwość dopplerowska chmury obserwowanej w k -tym kanale.

W wypadku „diversity polaryzacji” (tj. np. odbioru dwóch ciągów sygnałowych spolaryzowanych liniowo w różnych płaszczyznach) można przyjąć, że częstotliwości ω_k są w obu kanałach jednakowe. W wypadku „diversity częstotliwości” stosunki tych częstotliwości odpowiadają stosunkom użytych częstotliwości nośnych. Wreszcie w wypadku „diversity przestrzennego” ich wartości są równe kombinacjom liniowym trzech ortogonalnych składowych wektora szybkości celu, przy czym współczynniki kombinacji zależą od geometrycznego usytuowania linii obserwacji poszczególnych anten.

Odnosnie faz początkowych $\varphi_{k;0}$ będziemy zakładali, że są one stałe, lecz nieznanne a priori, tj., że posiadają niezależne rozkłady równomierne w przedziałach $(-\pi, \pi)$.

Uogólniona na wypadek K -kanałowy macierz współczynników korelacji (f) przybiera postać

$$\tilde{Q}_{(K)} = \tilde{T}_{(K)} \cdot \tilde{R}_{(K)} \cdot \tilde{T}_{(K)}^{-1} = \begin{vmatrix} \tilde{T}_1 & 0 \\ & \tilde{T}_2 \\ & & \ddots \\ 0 & & & \tilde{T}_K \end{vmatrix} \cdot \begin{vmatrix} \tilde{R}_{11} & \tilde{R}_{12} & \dots & \tilde{R}_{1K} \\ \tilde{R}_{21} & \tilde{R}_{22} & \dots & \tilde{R}_{2K} \\ \vdots & \vdots & \ddots & \vdots \\ \tilde{R}_{K1} & \tilde{R}_{K2} & \dots & \tilde{R}_{KK} \end{vmatrix} \cdot \begin{vmatrix} \tilde{T}_1^{-1} & 0 \\ & \tilde{T}_2^{-1} \\ & & \ddots \\ 0 & & & \tilde{T}_K^{-1} \end{vmatrix} \quad (14)$$

gdzie:

$\tilde{T}_K$ — podmacierz diagonalna postaci analogicznej jak w (1f), o elementach postaci $e^{j\phi k; n}$

$\tilde{R}_{ij}$ — podmacierz współczynników korelacji wzajemnej zakłóceń w i -tym i j -tym kanale.

Dzięki zastosowaniu wzorów (2) zagadnienie wyliczenia macierzy kowariancji sprowadza się znów do wyliczenia macierzy $\tilde{P}_{(K)}$ spełniającej równanie macierzowe:

$$\tilde{P}_{(K)} \cdot \tilde{Q}_{(K)} = E, \quad (15)$$

gdzie E oznacza znów macierz jednostkową, zaś macierz $\tilde{P}_{(K)}$ posiada strukturę „klatkową”, tj. złożona jest z podmacierzy $\tilde{P}_{kl}$:

$$\tilde{P}_{(K)} = \begin{bmatrix} \tilde{P}_{11} & \tilde{P}_{12} & \dots & \tilde{P}_{1K} \\ \tilde{P}_{21} & \tilde{P}_{22} & \dots & \tilde{P}_{2K} \\ \vdots & \vdots & \ddots & \vdots \\ \tilde{P}_{K1} & \tilde{P}_{K2} & \dots & \tilde{P}_{KK} \end{bmatrix}. \quad (15a)$$

Dla niewielkiej liczby kanałów K podmacierze $\tilde{P}_{kl}$ najwygodniej jest liczyć wprost, uwzględniając jednak niekomutatywność mnożenia macierzowego. Zobrazujemy to na przykładzie $K = 2$.

Przykład 4.

Niech macierz współczynników korelacji $\tilde{Q}_{(2)}$ ma postać:

$$\tilde{Q}_{(2)} = \begin{bmatrix} \tilde{Q}_{11} & \tilde{Q}_{12} \\ \tilde{Q}_{21} & \tilde{Q}_{22} \end{bmatrix}, \quad (16)$$

gdzie $\tilde{Q}_{11}$, $\tilde{Q}_{22}$ — podmacierze współczynników korelacji własnej sygnałów w pierwszym i drugim kanale, $\tilde{Q}_{12} = [-\tilde{Q}_{21}]^{(T)}$ — macierze współczynników korelacji wzajemnej sygnałów w pierwszym i drugim kanale. Odwrócona macierz kowariancji $\tilde{P}_{(2)}$ ma postać:

$$\tilde{P}_{(2)} = \begin{bmatrix} \tilde{P}_{11} & \tilde{P}_{12} \\ \tilde{P}_{21} & \tilde{P}_{22} \end{bmatrix}, \quad (17)$$

gdzie podmacierze $\tilde{P}_{ik}$ określone są wzorami:

$$\tilde{P}_{11} = \tilde{Q}_{11}^{-1} + \tilde{Q}_{11}^{-1} \tilde{Q}_{12} (\tilde{Q}_{22} - \tilde{Q}_{21} \tilde{Q}_{11}^{-1} \tilde{Q}_{12})^{-1} \tilde{Q}_{21} \tilde{Q}_{11}^{-1}, \quad (17a)$$

$$\tilde{P}_{12} = (\tilde{Q}_{12} \tilde{Q}_{22}^{-1} \tilde{Q}_{21} - \tilde{Q}_{11})^{-1} \tilde{Q}_{12} \tilde{Q}_{22}^{-1}, \quad (17b)$$

$$\tilde{P}_{21} = (\tilde{Q}_{21} \tilde{Q}_{11}^{-1} \tilde{Q}_{12} - \tilde{Q}_{22})^{-1} \tilde{Q}_{21} \tilde{Q}_{11}^{-1}, \quad (17c)$$

$$\tilde{P}_{22} = \tilde{Q}_{22}^{-1} + \tilde{Q}_{22}^{-1} \tilde{Q}_{21} (\tilde{Q}_{11} - \tilde{Q}_{12} \tilde{Q}_{22}^{-1} \tilde{Q}_{21})^{-1} \tilde{Q}_{12} \tilde{Q}_{22}^{-1}. \quad (17d)$$

Prawdziwość powyższych wzorów można sprawdzić wprost przez przemnożenie macierzy $\tilde{P}_2$ i $\tilde{Q}_2$, z zachowaniem reguł mnożenia macierzowego. Jak widać, określenie współczynników macierzy kowariancji (a więc

i określenie optymalnej reakcji impulsowej filtru przeddetekcyjnego) już w przypadku pracy dwukanałowej staje się dość skomplikowane. Celowe jest zatem wprowadzenie pewnych założeń upraszczających. W szczególności, jeśli przyjąć, że ciąg obserwacji jest obustronnie nieskończony, zakłócenia obserwowane w poszczególnych kanałach są procesami stacjonarnymi i stacjonarnie związanymi, a funkcje korelacji własnej i międzykanałowej można aproksymować w postaci (12), to — podobnie jak w wypadku jednokanałowym — działania algebraiczne na podmacierzach stają się równoważne zwykłym działaniom algebraicznym na widmach tych podmacierzy. Wynika stąd, że układowi K^2 równań macierzowych (15) dla podmacierzy $\tilde{P}_{mn}$ odpowiada układ K^2 równań algebraicznych liniowych dla widm tych podmacierzy. Oznaczając przez $R_{ij}(Z)$ i $X_{mn}(Z)$ odpowiednio widma podmacierzy $\tilde{Q}_{ij}$ i $\tilde{P}_{mn}$, wyrazimy widmo $X_{mn}(Z)$ na zasadzie znanych zależności algebraicznych jak następuje:

$$X_{mn}(Z) = \frac{\text{Ad}[R_{ij}(Z)]}{\text{Det} \parallel R_{ij}(Z) \parallel}, \quad (18)$$

gdzie $\text{Ad}[R_{ij}(Z)]$ — dopełnienie algebraiczne elementu R_{ij} w macierzy widm $\parallel R_{ij}(Z) \parallel$; $\text{Det} \parallel R_{ij}(Z) \parallel$ — wyznacznik macierzy $\parallel R_{ij}(Z) \parallel$ (z założenia różny od zera).

W tym przypadku więc synteza filtrów impulsowych nie nastęrcza istotnych trudności rachunkowych.

4. PRZYPADEK ZMIENNEJ CZĘSTOŚCI POWTARZANIA IMPULSÓW SONDUJĄCYCH

W dotychczasowych rozważaniach zakładaliśmy, że częstotliwość powtarzania impulsów jest stała, tj., że ciąg próbek sygnałowych jest (przy założeniu stacjonarności zakłócenia) stacjonarny. Znalazło to wyraz w wyrażeniu (4). W przypadku zmiennej częstotliwości powtarzania współczynniki korelacji $R_{\mu\nu}$ można wyrazić następująco:

$$R_{\mu\nu} = R(t_\mu - t_\nu), \quad (19)$$

gdzie ciąg liczb $\{t_\nu\}$ związany jest z przyjętym systemem zmian częstości powtarzania impulsów. Zależności (5) przybiorą zatem postać ogólniejszą:

$$\sum_{k=-\infty}^{\infty} r(t_m - t_k) \zeta_{kn} = \delta_{mn}, \quad (m, n = \dots -2, -1, 0, 1, 2, \dots). \quad (20)$$

Aby określić elementy macierzy odwrotnej $\tilde{P}$, wyrazimy macierz kowariancji $\tilde{Q}$ w formie następującej:

$$\tilde{Q} = \tilde{A} \cdot \tilde{Q}' \cdot \tilde{B}, \quad (21)$$

gdzie: $\tilde{Q}$ — macierz współczynników korelacji dla okresowego (o średnim okresie T) próbkowania sygnału,

$\tilde{A}$, $\tilde{B}$ — pewne macierze — przekształceń, związane ze sposobem zmian częstości powtarzania impulsów sondujących.

Można wykazać, że przy założeniu istnienia wszystkich pochodnych funkcji korelacyjnej (warunki te spełnia np. funkcja Gaussa), a więc istnienia wszystkich pochodnych procesu przypadkowego, opisanego tymi funkcjami, macierze $\tilde{A}$ i $\tilde{B}$ można przedstawić w postaci następujących dwóch funkcji macierzowych:

$$\left. \begin{aligned} \tilde{A} &= e^{M\Delta}, \\ \tilde{B} &= e^{-\Delta M}, \end{aligned} \right\} \quad (22)$$

gdzie:

M — diagonalna macierz „modulacyjna”:

$$M = \begin{pmatrix} & & & 0 \\ & m_{-1} & & \\ & & m_0 & \\ 0 & & & m_1 \end{pmatrix}, \quad m_i = \frac{\Delta t_i}{T_p} \text{ — względne przesunięcie } i\text{-go impulsu} \\ \text{sondującego w stosunku do położenia} \quad \text{środkowego;} \quad (23)$$

Δ — macierz — operator wyliczenia różnicy pierwszego rzędu:

$$\Delta = \frac{1}{2} \begin{pmatrix} & & & & \\ & 0 & -1 & & 0 \\ & 1 & 0 & -1 & \\ & & 1 & 0 & -1 \\ 0 & & & 1 & 0 \end{pmatrix} \quad (24)$$

Pomnożenie dowolnej macierzy $\tilde{Q}'$ z lewa przez iloczyn macierzy $M \cdot \Delta$ daje w wyniku macierz, której elementami są różnice rzędu pierwszego odpowiednich elementów macierzy $\tilde{Q}'$, wzięte w kierunku pionowym i pomnożone przez odpowiednie wartości m_i . Analogicznie, pomnożenie dowolnej macierzy $\tilde{Q}'$ z prawa przez iloczyn macierzy $\Delta \cdot M$ daje w wyniku macierz, której elementami są różnice rzędu pierwszego odpowiednich elementów macierzy $\tilde{Q}'$, wzięte w kierunku poziomym i pomnożone przez odpowiednie wartości m_i . Rozkładając wyrażenia $\tilde{A}$ i $\tilde{B}$ w szeregi potęgowe względem $M\Delta$ i ΔM a następnie stosując wielokrotnie i superponując operację obliczania różnic skończonych otrzymamy formalne uzasadnienie wyrażenia wielkości $\tilde{A}$, $\tilde{B}$ w formie (22).

Jak łatwo sprawdzić, macierz odwrotna względem $\tilde{Q}$ [patrz (21)] ma postać:

$$\tilde{P} = \tilde{B}^{-1} \cdot \tilde{Q}'^{-1} \cdot \tilde{A}^{-1}. \quad (25)$$

Zaletą zapisu (21) polega na tym, że wyliczenie każdej z osobna macierzy odwrotnej w wyrażeniu (25) nie przedstawia zasadniczych trudności.

Istotnie, obliczenie macierzy $\tilde{Q}'^{-1}$ może się odbywać analogicznie, jak to zostało przedstawione w rozważaniach poprzednich.

Macierze odwrotne względem A i B mają postać:

$$\tilde{A}^{-1} = e^{-M\Delta}, \quad (26)$$

$$\tilde{B}^{-1} = e^{\Delta M}. \quad (27)$$

Wstawiając wyrażenia (26), (27) do (25) otrzymamy:

$$\tilde{P} = e^{\Delta M} \cdot \tilde{Q}'^{-1} \cdot e^{-M\Delta}. \quad (28)$$

Budowa macierzy $\tilde{P}$ jest analogiczna do budowy macierzy $\tilde{Q}$; różnica sprowadza się wyłącznie do mnożnika modulującego $(E + M)$. Jeżeli założyć, że $m_i = \frac{\Delta t_i}{T_p} \ll 1$ ($i = \dots -2, -1, 0, 1, 2, \dots$), to można przyjąć, że $(E + M\Delta)^{-1} \approx E - M\Delta$. Wówczas macierz $\tilde{P}$ w przybliżeniu może być otrzymana jako macierz odwrotna $\tilde{Q}'^{-1}$ o elementach (tj. wierszach i kolumnach) przesuniętych odwrotnie niż elementy macierzy korelacyjnej $\tilde{Q}$.

W przypadku gdy funkcja korelacyjna jest różniczkowalna N -krotnie, stosowanie operatora przesunięcia postaci e^{Δ} napotyka na trudności natury formalnej. Wówczas jednak operator ten można zastąpić pierwszymi członami jego rozkładu w szereg Taylora:

$$e^{\Delta} = \sum_{k=0}^{\infty} \frac{1}{k!} \cdot \Delta^k \approx \sum_{k=0}^N \frac{1}{k!} \Delta^k, \quad (29)$$

gdzie Δ^k jest macierzowym operatorem wyliczenia k -tej różnicy wierszy (lub kolumn) macierzy $\tilde{Q}$. Można łatwo sprawdzić, że macierz $\sum_{k=0}^N \frac{1}{k!} \Delta^k$

jest macierzą, której wszystkie elementy skupione są na pozycjach oddległych nie więcej niż o N miejsc po obu stronach głównej przekątnej, a także są identyczne na liniach równoległych do głównej przekątnej. Wyliczenie macierzy odwrotnej dla takiego wypadku przy pomocy zwykłych metod algebraicznych nie nastęrcza zwykle trudności, zwłaszcza że przy małym stosunku czasu repetycji do czasu korelacji zakłócenia operator e^{Δ} we wzorze (29) wystarczy aproksymować jednym lub dwoma wyrazami.

Dla ilustracji rozważań przytoczymy następujący przykład.

Przykład 5.

Założmy, że funkcja korelacyjna zakłóceń ma postać:

$$r(\tau) = A_1 \left(e^{-\frac{|\tau|}{\tau_0}} + A_2 e^{-\frac{|\tau|}{\theta}} \right),$$

gdzie stałe A_1 i A_2 dobieramy tak, aby były spełnione warunki

$$r(0) = 1, \quad r'(0) = 0.$$

Ostatni warunek związany jest z żądaniem dwukrotnej różniczkowalności funkcji korelacji, tj. z żądaniem jednokrotnej różniczkowalności procesu przypadkowego. Można wykazać, że stałe A_1 i A_2 powinny spełniać równania:

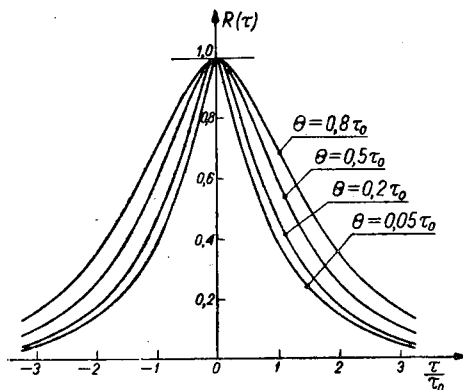
$$A_1 = \frac{\tau_0}{\tau_0 - \theta},$$

$$A_2 = -\frac{\theta}{\tau_0},$$

tj.

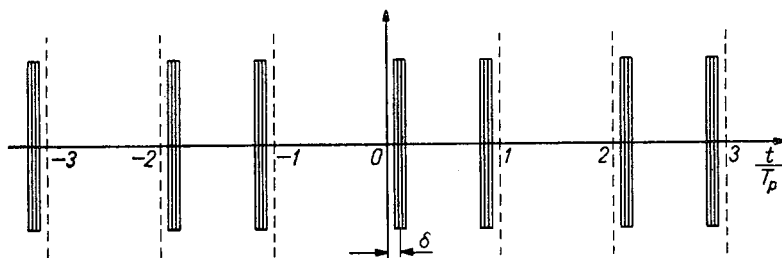
$$r(\tau) = \frac{\tau_0}{\tau_0 - \theta} e^{-\frac{|\tau|}{\tau_0}} - \frac{\theta}{\tau_0 - \theta} e^{-\frac{|\tau|}{\theta}}.$$

Typowe postaci funkcji $r(\tau)$ dla kilku stosunków τ/θ pokazano na rys. 3.



Rys. 3.

Niech średni okres powtarzania stacji radiolokacyjnej wynosi T_p . Impulsy sondujące nie są jednak wysyłane okresowo, lecz z odchyleniami od położenia średniego kT_p ($k = \dots, -2, -1, 0, 1, 2, \dots$), wynoszącymi na przemian $+\delta$ i $-\delta$ (rys. 4). Określimy postać odwróconej macierzy ko-



Rys. 4.

wariancji przy założeniu dwustronnie nieograniczonego ciągu impulsów odebranych.

W przybliżeniu liniowym macierz współczynników korelacji może być przedstawiona w postaci:

$$\tilde{Q} \approx \tilde{A} \cdot \tilde{Q}' \cdot \tilde{B},$$

gdzie:

$$\tilde{Q}' = \|r_{ij}\| = \|\bar{r}[(i-j)T_p]\|,$$

$$\tilde{A} = E + M\Delta + O[(M\Delta)^2],$$

$$\tilde{B} = E - \Delta M + O[(\Delta M)^2],$$

$$M = \begin{vmatrix} & & & \\ & -m & & 0 \\ & +m & & \\ 0 & & -m & \\ & & +m & \end{vmatrix}, \quad m = \frac{\delta}{T_p}, \quad m_{00} = +m,$$

a macierz Δ określona jest wzorem (24).

Wyliczmy widmo macierzy współczynników korelacji zadanych przez funkcję $\bar{r}(\tau)$. Posługując się jedną z metod opisanych w rozdz. 1 otrzymamy wynik, który można sprowadzić do postaci:

$$R(Z) = \frac{1}{\tau_0 - \theta} \cdot \frac{A + Bq}{C + Dq + Eq^2},$$

gdzie:

$$q = \frac{1}{2}(Z + Z^{-1}),$$

$$A = \tau_0(1 - Z_0^2)(1 + Z_1^2) - \theta(1 - Z_1^2)(1 + Z_0^2),$$

$$B = 2[Z_0\theta(1 - Z_1^2) - Z_1\tau_0(1 - Z_0^2)],$$

$$C = 1 + Z_1^2 + Z_0^2 + Z_0^2 Z_1^2,$$

$$D = -2(Z_0 + Z_1 + Z_0^2 Z_1 + Z_1^2 Z_0),$$

$$E = 4Z_0 Z_1,$$

$$Z_0 = e^{-T_p/\tau_0},$$

$$Z_1 = e^{-T_p/\theta}.$$

Widmo macierzy odwrotnej ma postać:

$$\chi(Z) = (\tau_0 - \theta) \frac{C + Dq + Eq^2}{A + Bq} = \alpha q + \beta + \frac{\gamma}{q + \delta},$$

gdzie:

$$\alpha = (\tau_0 - \theta) \frac{E}{B},$$

$$\beta = (\tau_0 - \theta) \frac{E}{B} \left(\frac{D}{E} - \frac{A}{B} \right),$$

$$\gamma = (\tau_0 - \theta) \frac{E}{B} \left[\frac{C}{E} - \frac{A}{B} \left(\frac{D}{E} - \frac{A}{B} \right) \right],$$

$$\delta = \frac{A}{B}.$$

Jak można wykazać, współczynniki macierzy odwrotnej $(\tilde{Q}')^{-1}$ wynoszą:

$$\xi'_{k,k} = \beta + \frac{\gamma}{\text{sh } \lambda},$$

$$\xi'_{k,k+1} = \xi'_{k,k-1} = \frac{1}{2} \alpha$$

$$\xi'_{k,l} = \xi'_{k,-l} = \frac{\gamma}{\text{sh } \lambda} e^{-|l-k|\lambda} (|l-k| = 2, 3, 4, \dots),$$

gdzie:

$$\text{ch } \lambda = -\delta.$$

Dla wyliczenia macierzy odwrotnej posłużymy się liniowym przybliżeniem wykładniczej funkcji macierzowej:

$$e^{\Delta M} = E + \Delta M + O[(\Delta M)^2], \quad e^{-\Delta M} = E - \Delta M + O[(\Delta M)^2],$$

gdzie wyrazy rzędu drugiego są pomijalne ze względu na to, iż $m = \frac{\delta}{T_p} \ll 1$.

Przy pomocy bezpośredniego rachunku można sprawdzić, że w wyniku podobnego przybliżenia funkcji wykładniczych otrzymamy:

$$(E + \Delta M)(E - \Delta M) = \mathcal{C},$$

gdzie $\mathcal{C}$ jest macierzą o wyrazach: ... 0, 0, $-m^2$, 0, 1, 0, $-m^2$, 0, 0, ... zamiast ... 0, 0, 1, 0, 0, ... jak powinno być w przypadku macierzy jednostkowej E .

Zgodnie z wyrażeniem (28) macierz odwrotna powinna mieć postać:

$$\tilde{P} = e^{\Delta M} \cdot \tilde{Q}'^{-1} \cdot e^{-\Delta M} \approx (E + \Delta M) \cdot \tilde{Q}'^{-1} \cdot (E - \Delta M).$$

Ponieważ jednak macierz M jest macierzą diagonalną, a więc spełniającą warunek komutatywności względem mnożenia, zatem ostatnie wyrażenie można też zapisać jako:

$$\tilde{P} \approx (E + \Delta M) \cdot \tilde{Q}'^{-1} \cdot (E - \Delta M).$$

Widać stąd, że zasada tworzenia współczynników macierzy $\tilde{P}$ ze współczynników macierzy $\tilde{Q}'^{-1}$ jest identyczna jak zasada tworzenia współczynników macierzy kowariancji $\tilde{Q}'$ ze współczynników „stacjonarnej” macierzy kowariancji $\tilde{Q}$. Inaczej mówiąc, z uzyskanych poprzednio wyra-

żeń na $\xi'_{k,l}$ oraz z założonej budowy macierzy modulacyjnej M otrzymamy:

$$\xi_{kl} = \xi'_{kl} + (-1)^l \cdot \frac{m}{2} (\xi'_{k, l+1} - \xi'_{k, l-1}) + (-1)^k \frac{m}{2} (\xi'_{k+1, l} - \xi'_{k-1, l}) + 0(m^2).$$

Wyrażenie powyższe określa w przybliżeniu właściwości optymalnego filtru impulsowego o parametrach zmiennych, dostrojonego do zmian okresu powtarzania impulsów sondujących.

Jak widać z powyższych rozważań, zastosowanie stosunkowo prostych środków matematycznych pozwala na określenie odwróconej macierzy kowariancji zakłóceń dla dostatecznie szerokiej klasy reżymów pracy urządzenia radiolokacyjnego, w związku z czym zagadnienie syntezy optymalnych układów odbiorczych w dużym stopniu można uważać za ułatwione wszędzie tam, gdzie filtracja przeddetekcyjna stanowi istotny składnik optymalnej obróbki sygnału odbieranego.

*Politechnika Warszawska
Katedra Techniki Fal Ultrakrótkich*

WYKAZ LITERATURY

1. Helstrom Carl W.: Statistical Theory of Signal Detection. Pergamon Press 1960.
2. Цыпкин J. Z.: Теория импульсных систем. Москва 1959.
3. Cooke Richard G.: Infinite Matrices and Sequence Spaces. London 1950.
4. Kulikowski J. E.: Woprosy optimalnego obnarużenia i selekcji cielej. Dysertacja kandydacka. Moskwa 1959.
5. Ryżik I. M., Gradštejn I. S.: Tablicy integralow, summ, riadow i proizwiedienij. Moskwa 1951.
6. Reich Edgar, Swerling Peter: The Detection of a Sine Wave in Gaussian Noise. — Journ. Appl. Phys. vol. 24, Nr. 3, 1953.

Ю. КУЛИКОВСКИ

О НЕКОТОРЫХ ВОПРОСАХ СВЯЗАННЫХ С ВЫЧИСЛЕНИЕМ ОБРАТНОЙ МАТРИЦЫ КОВАРИАНЦИИ ПАССИВНЫХ РАДИОЛОКАЦИОННЫХ ПОМЕХ

Резюме

В статье рассматриваются методы практического определения коэффициентов матрицы, определяющей квадратичную эрмитову форму, выступающую в показателе многомерного гауссового распределения вероятностей. Указанные коэффициенты определяют непосредственно оптимальную переходную функцию импульсного переддетекторного фильтра в когерентном приёмнике. Метод расчёта основан на использовании свойств „двустороннего Z — преобразования”. Рас-

смотрены случаи многоканального приёма и приёма с меняющейся частотой повторения зондирующих импульсов радиолокатора. Выводы иллюстрируются на ряде практических примеров.

J. KULIKOWSKI

ON SOME PROBLEMS CONCERNING THE DETERMINATION OF INVERSE COVARIANCE-MATRIX COEFFICIENTS OF RADAR SIGNALS IN CLUTTER

Summary

Some practical methods of determining the coefficients of a hermitian quadratic-form are considered. The coefficients are necessary for giving an exact expression for a multi-dimensional normal probability density-function of a radar signal. They make also possible obtaining of an optimal transfer-function of a pre-detection pulse-filter for a coherent receiver. The method is based on the so called „two-side Z-transform”. The case of a multi-channel radar operating-mode as well as that of a variable repetition-frequency are considered. The theoretical assumptions are illustrated by some practical examples.

JERZY SAWICKI

Skompensowany mostek Thomsona z podstawieniem

Rękopis dostarczono 21. 12. 1960

Wysoką dokładność pomiaru w zrównoważonym mostku Thomsona osiąga się przez zastosowanie podstawienia (substytucji) przy użyciu regulowanego bocznika. Metoda ta daje się zrealizować dla oporów nie mniejszych niż $0,01 \Omega$. W pozostałych przypadkach trzeba uciec się do innych modyfikacji mostka, wymagających sprzętu specjalnego (jak oporniki w układzie prof. Krukowskiego), bądź też do metody kompensatora wzorcowego.

Podstawienie w mostku skompensowanym daje możliwość precyzyjnego porównywania oporników przy użyciu wyłącznie sprzętu typowego. Zakres pomiaru w tej metodzie rozciąga się na cały obszar stosowności układu Thomsona. Proponowany sposób pomiaru wykazuje podobne zalety, jak substytucja w skompensowanym mostku Wheatstone'a ([3]) przy sprawdzaniu oporów średnich i dużych.

Artykuł obejmuje wyprowadzenie wzorów na wielkość mierzoną oraz jej uchyby. Dalsze rozważania wskazują warunki, których spełnienie pozwoli osiągnąć założoną dokładność pomiaru. Przytoczono również przykład porównania dwu oporników normalnych o wartości $0,1 \Omega$ przy zastosowaniu podstawienia w mostku skompensowanym i zrównoważonym.

1. WSTĘP

Skompensowany mostek Thomsona różni się od zrównoważonego tym, że warunek równowagi nie jest ściśle spełniony. Galwanometr zostaje jednak doprowadzony do wskazania zerowego dzięki włączonej w szereg z nim *SEM* kompensującej. Wartość tej siły elektromotorycznej musi być znana i regulowana. Zadanie to spełnia bardzo dobrze regulowany dzielnik napięcia zasilany z odrębnego akumulatora.

Pomiar metodą podstawienia w skompensowanym mostku Thomsona składa się z dwu zasadniczych części.

W pierwszym etapie umieszcza się opornik wzorcowy w tym miejscu układu, gdzie zwykle włącza się obiekt mierzony. Mostek należy dokładnie zrównoważyć przez regulację odpowiednich oporników przy wyłącznej *SEM* kompensującej. Prawidłowość zrównoważenia kontroluje się przy przerwaniu połączenia pomiędzy małymi oporami, a w razie potrzeby przeprowadza się regulację uzupełniającą. Układ powinien wyka-

zywać równowagę zarówno przy zamkniętym, jak i otwartym połączeniu między małymi opornikami.

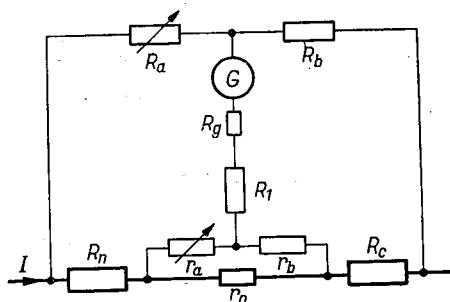
Przed przystąpieniem do drugiego etapu usuwa się wzorzec i w zwykłe miejsce włącza obiekt badany. Obecnie wykonuje się kompensację układu przy pomocy dodatkowej SEM, nie zmieniając już jednak żadnego z oporów mostka.

Różnicę między oporem badanym a wzorcowym oblicza się z wartości SEM kompensującej, prądu zasilającego i stosunków między znanymi oporami mostka. Wartości te nie muszą być znane bardzo dokładnie, gdyż dopuszczalne uchyby leżą w okolicy 10%.

Przedstawioną metodę można stosować do sprawdzania oporów precyzyjnych o małych wartościach względnie wykonanych jako czterozaciskowe. Używa się przy tym wyłącznie typowego sprzętu laboratoryjnego. Oczywiście źródła prądu zasilającego jak i SEM kompensującej powinny wykazywać wystarczającą stałość napięcia. Akumulatory ołowiowe odpowiednio dobranej pojemności spełniają to wymaganie zupełnie dobrze.

2. WYPROWADZENIE WZORU OBLICZENIOWEGO

Układ zwykłego mostka Thomsona, który występuje w pierwszym etapie pomiaru, pokazano na rys. 1. Źródło SEM kompensującej jest wyłączone, wobec czego nie zostało umieszczone na schemacie. R_n oznacza



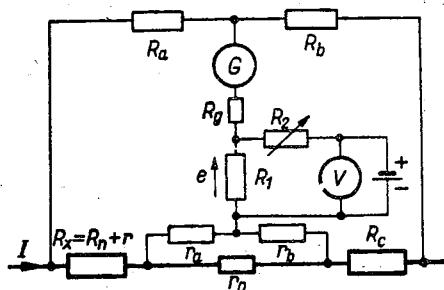
Rys. 1. Pierwszy etap pomiaru

opornik wzorcowy, z którym porównujemy badany R_x o tej samej wartości nominalnej. Opory R_a , R_b , R_c , r_a , r_b są znane. Oporność połączenia pomiędzy prawym zaciskiem napięciowym oporu wzorcowego a lewym napięciowym oporu pomocniczego R_c została oznaczona przez r_o .

Metodą kolejnego wyrównywania w układzie Thomsona i Wheatstone'a (tj. przy przerwaniu połączenia r_o) sprowadza się mostek do ścisłej równowagi, której warunkiem jest spełnienie zależności

$$\frac{R_n}{R_c} = \frac{R_a}{R_b} = \frac{r_a}{r_b}. \quad (1)$$

Układ połączeń w drugim etapie przedstawia rys. 2. Oporność połączeniowa r_o może tu mieć inną wartość niż poprzednio ze względu na włączenie R_x zamiast R_n . Oporność między zaciskiem prądowym i przyna-



Rys. 2. Drugi etap pomiaru

leżnym napięciowym zależy bowiem od wartości nominalnej opornika i jego wykonania (producent).

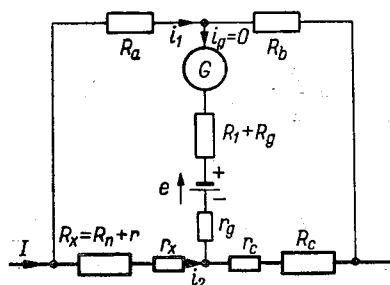
Trójkąt oporów r_a , r_b , r_o przekształcamy na równoważną gwiazdę o opornościach

$$r_x = \frac{r_o \cdot r_a}{r_a + r_b + r_o}, \quad r_c = \frac{r_o \cdot r_b}{r_a + r_b + r_o}, \quad r_g = \frac{r_a \cdot r_b}{r_a + r_b + r_o} \quad (2)$$

i otrzymujemy skompensowany mostek Wheatstone'a, przedstawiony na rys. 3. Wprowadzono tu oznaczenie

$$R_x = R_n + r, \quad (3)$$

gdzie r oznacza mierzoną różnicę między oporem badanym a wzorcowym. Zwrot SEM kompensującej e na rys. 3 odpowiada przypadkowi $r > 0$.



Rys. 3. Zastępczy skompensowany mostek Wheatstone'a

W stanie równowagi, tj. dla $i_g = 0$, spełnione są zależności

$$\left. \begin{aligned} I &= i_1 + i_2 \\ i_2(R_n + r + r_x + r_c + R_c) &= i_1 \cdot (R_a + R_b) \\ i_2 \cdot (R_n + r + r_x) &= i_1 \cdot R_a + e \end{aligned} \right\} \quad (4)$$

Po wyeliminowaniu prądów i_1 oraz i_2 w wyniku przekształceń otrzymujemy poniższe wyrażenie:

$$r \cdot \left(R_b - \frac{e}{I} \right) = \frac{e}{I} \cdot (R_a + R_b + R_c + R_n + r_c + r_x) + R_a(R_c + r_c) - R_b(R_n + r_x). \quad (5)$$

Dwa wyrazy prawej strony wzoru (5) przekształcimy w oparciu o (2) w sposób następujący:

$$\begin{aligned} R_a(R_c + r_c) - R_b(R_n + r_x) &= R_a R_c + R_a r_b \frac{r_o}{r_a + r_b + r_o} - R_b R_n + \\ &- R_b r_a \frac{r_o}{r_a + r_b + r_o} = R_b R_c \left(\frac{R_a}{R_b} - \frac{R_n}{R_c} \right) + \frac{R_b r_b r_o}{r_a + r_b + r_o} \left(\frac{R_a}{R_b} - \frac{r_a}{r_b} \right) = 0 \end{aligned}$$

na podstawie zależności (1). Uwzględniając, że

$$r_x + r_c = r_o \frac{r_a + r_b}{r_a + r_b + r_o} \approx r_o \quad (\text{gdyż } r_o \ll r_a + r_b),$$

wzór (5) możemy sprowadzić do postaci

$$r = \frac{e}{I \cdot R_b - e} (R_a + R_b + R_c + R_n + r_o). \quad (6)$$

Wprowadzimy teraz następujące oznaczenia:

$$m = \frac{R_c}{R_b}, \quad n = \frac{R_a}{R_b}, \quad p = \frac{r_b}{R_b}, \quad q = \frac{R_g + R_1}{R_b}, \quad s = \frac{r_o}{R_c}. \quad (7)$$

Wzór (6) upraszczamy do postaci przybliżonej, przyjmując że

$$I R_b - e \approx I R_b \quad \text{oraz} \quad R_a + R_b + R_c + R_n + r_o \approx R_a + R_b + R_c + R_n.$$

Po zastosowaniu tego przybliżenia oraz oznaczeń (7) przyjmie on postać

$$r' = \frac{e}{I} (1 + m)(1 + n). \quad (8)$$

Względny uchyb wprowadzonego przybliżenia można wyrazić jako

$$\delta' = - \frac{I R_b}{I R_b - e} \cdot \frac{r' + r_o}{R_a + R_b + R_c + R_n} \approx \frac{r' + r_o}{R_a + R_b} \approx - \frac{r + r_o}{(1 + n) R_b}. \quad (9)$$

Wartości uchybu przybliżenia obliczymy przy założeniu najbardziej niekorzystnych warunków. Przyjmujemy, że $R_b = 10\Omega$, wobec czego $R_a + R_b = 10(1 + n)\Omega$, natomiast $r' = 5 \cdot 10^{-3} R_x \approx 5 \cdot 10^{-3} R_n$. Wartość oporu połączeniowego r_o przyjęto jako równą oporności doprowadzeń w czterozaciskowym oporniku normalnym. Przez „oporność doprowadzeń” rozumiemy tutaj różnicę pomiędzy opornością mierzoną na zaciskach prądowych i napięciowych. Autor zbadał metodą kompensacyjną

25 oporników o wartościach od 10° do $10^{-4}\Omega$ różnych firm (Wolff, Siemens, Schlotheim, Etalon, RFT, Zakł. Optyki i Mechaniki Prec. Polit. Śląskiej) i otrzymał wartości od $0,4 \text{ m}\Omega$ do $3 \text{ m}\Omega$.

Przy tych założeniach wartości uchybu przybliżenia kształtują się następująco:

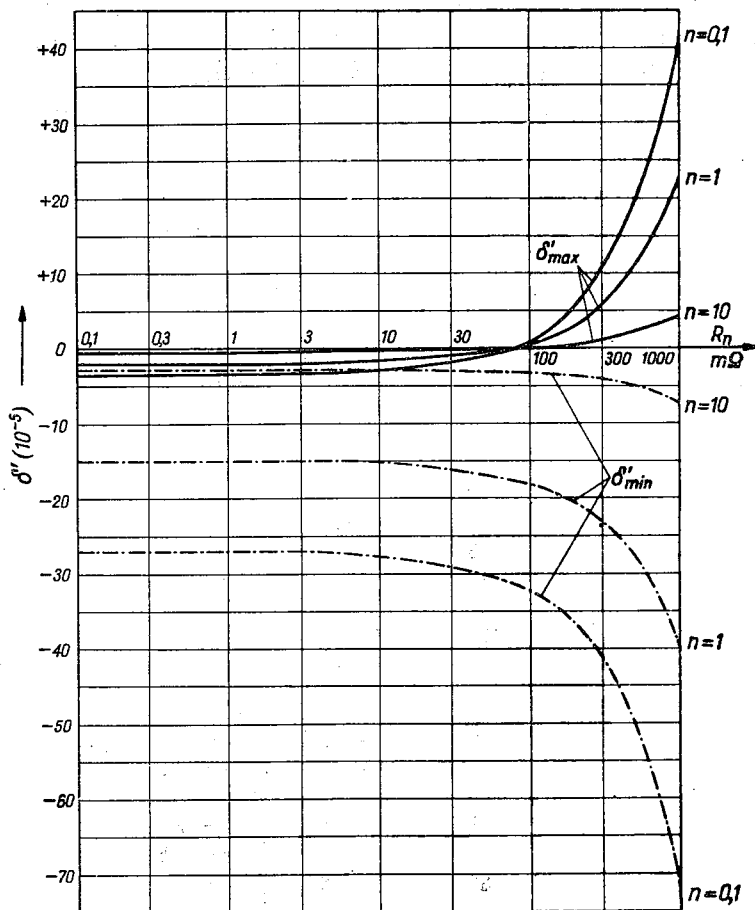
Tablica 1

$$\delta' = \frac{r' - r}{r'} = f(R_n; n)$$

n	R_n	$(\delta')_{\min}$	$(\delta')_{\max}$
	$\text{m}\Omega$	10^{-5}	10^{-5}
0,1	0,1	— 27	— 3,6
	0,3	— 27	— 3,6
	1	— 27	— 3,6
	3	— 27	— 3,6
	10	— 28	— 2,7
	30	— 29	— 1,8
	100	— 32	+ 0,9
	300	— 41	+ 10
	1000	— 73	+ 42
1	0,1	— 15	— 2,0
	0,3	— 15	— 2,0
	1	— 15	— 2,0
	3	— 15	— 2,0
	10	— 16	— 1,5
	30	— 16	— 1,0
	100	— 18	+ 0,5
	300	— 23	+ 5,5
	1000	— 40	+ 23
10	0,1	— 2,7	— 0,4
	0,3	— 2,7	— 0,4
	1	— 2,7	— 0,4
	3	— 2,7	— 0,4
	10	— 2,8	— 0,3
	30	— 2,9	— 0,2
	100	— 3,2	+ 0,1
	300	— 4,1	+ 1,0
	1000	— 7,3	+ 4,2

Przebieg powyższej zależności ilustrują krzywe na rys. 4. W przypadku zastosowania większej wartości R_b (zwykle 100 omów) przytoczone uchyby przybliżenia zmniejszają się (10-krotnie).

Z dalszych rozważań wynika, że średni uchyb systematyczny poprawki r w najkorzystniejszych warunkach jest rzędu $100 \cdot 10^{-5}$. Wpływa stąd



Rys. 4. Względny uchyb przybliżenia

$$\delta' = f(R_n \cdot n); R_b = 10 \, \Omega, \quad |r| = 5 \cdot 10^{-3} R_x; \quad 4 \cdot 10^{-4} \leq r_o \leq 3 \cdot 10^{-3} \, \Omega$$

wniosek, że przybliżenie można stosować zawsze, o ile $R_b = 100$ omów. W przypadku $R_b = 10$ omów należy unikać wartości $n = 0,1$ (która również nie jest wskazana ze względu na duży wpływ oporności przewodów łączących opór $R_a = 10 \cdot 0,1 = 1 \, \Omega$ z resztą układu).

3. UCHYB SYSTEMATYCZNY POPRAWKI r

Względny średni uchyb systematyczny obliczamy na podstawie zależności $y = f(a_1; a_2; \dots a_q)$ przy pomocy wzoru

$$[\delta_s(r)]^2 = \sum_{i=1}^{i=q} \left[\frac{\partial(f)}{\partial(a_i)} \cdot \frac{a_i}{f} \delta_s(a_i) \right]^2.$$

Wychodząc ze wzoru (8) otrzymujemy

$$[\delta_{s(r)}]^2 = [\delta_{s(e)}]^2 + [\delta_{s(I)}]^2 + \left[\frac{m}{1+m} \delta_{s(m)} \right]^2 + \left[\frac{n}{1+n} \delta_{s(n)} \right]^2 \quad (10)$$

Stosownie do układu pokazanego na rys. 2 wartość e oblicza się jako

$$e = U \frac{R_1}{R_1 + R_2}, \quad (11)$$

wobec czego przy uwzględnieniu, że $R_2 \gg R_1$, mamy

$$[\delta_{s(e)}]^2 = [\delta_{s(U)}]^2 + [\delta_{s(R_1)}]^2 + [\delta_{s(R_2)}]^2. \quad (12)$$

Traktując klasę dokładności miernika wskazówkowego jako miarę uchybu wiarygodnego otrzymujemy wzory

$$|\delta_{s(U)}| = \frac{k_V}{300} \cdot \frac{U_n}{U}; \quad |\delta_{s(I)}| = \frac{k_A}{300} \cdot \frac{I_n}{I}. \quad (13)$$

Zgodnie z definicjami oznaczeń $m, n \dots$ (7) jest

$$[\delta_{s(m)}]^2 = [\delta_{s(R_b)}]^2 + [\delta_{s(R_c)}]^2, \quad (14)$$

$$[\delta_{s(n)}]^2 = [\delta_{s(R_a)}]^2 + [\delta_{s(R_b)}]^2. \quad (15)$$

Po uwzględnieniu zależności (12), (14) i (15) wzór (10) przyjmuje postać

$$[\delta_{s(r)}]^2 = [\delta_{s(I)}]^2 + [\delta_{s(U)}]^2 + [\delta_{s(R_1)}]^2 + [\delta_{s(R_2)}]^2 + \\ + \left(\frac{n}{1+n} \right)^2 \{ [\delta_{s(R_a)}]^2 + [\delta_{s(R_b)}]^2 \} + \left(\frac{m}{1+m} \right)^2 \{ [\delta_{s(R_b)}]^2 + [\delta_{s(R_c)}]^2 \}. \quad (16)$$

W powyższym wzorze ostatnie wyrażenie w nawiasach $\{ \}$ można pominąć. Świadczy o tym poniższy rachunek. Początkowe człony wzoru (16) przybierają minimalne wartości przy założeniu, że

$$I_n/I = U_n/U = 1, \quad k_A = k_V = 0,2, \quad k_{R_1} = k_{R_2} = k_{R_a} = k_{R_b} = 0,01, \quad n = 0,1.$$

Suma początkowych wyrazów wynosi wówczas $90 \cdot 10^{-8}$.

Maksymalna wartość ostatniego członu występuje, gdy

$$m_{\max} = R_{c \max} : R_{b \min} = 1 \Omega : 10 \Omega = 0,1$$

$k_{R_c} = 0,5$; przyjmujemy również $k_{R_b} = 0,5$ — chociaż poprzednio założyliśmy odmiennie. W takim przypadku końcowy człon osiąga wartość poniżej $14 \cdot 10^{-8}$.

Uchyb systematyczny, wyliczony z pełnego wzoru (16), wynosi zatem $\pm 10,2 \cdot 10^{-4}$, a z początkowych członów $\pm 9,5 \cdot 10^{-4}$, co stanowi różnicę całkowicie pomijalną w rachunku uchybów. W warunkach rzeczywistych różnica ta będzie jeszcze mniejsza.

Z wystarczającą dokładnością można więc napisać:

$$[\delta_{s(r)}]^2 = [\delta_{s(I)}]^2 + [\delta_{s(U)}]^2 + [\delta_{s(R_1)}]^2 + [\delta_{s(R_2)}]^2 + \left(\frac{n}{1+n} \right)^2 \{ [\delta_{s(R_a)}]^2 + [\delta_{s(R_b)}]^2 \}. \quad (17)$$

4. UCHYB OD NIECZUŁOŚCI

Sredni względny uchyb od nieczułości oblicza się na podstawie zależności $y = f(a_1; a_2; \dots a_q; a)$ przy pomocy wzoru

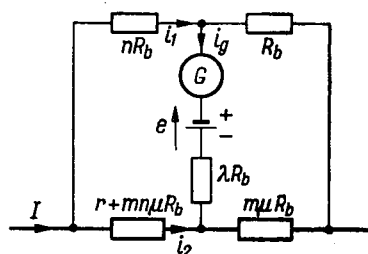
$$|\delta_n(y)| = \left[\frac{\partial(f)}{\partial(a)} \right]_{a=0} \cdot \frac{\Delta_a}{[f]_{a=0}}.$$

Wyprowadzenie potrzebnego związku pomiędzy mierzoną wartością r a wychyleniem galwanometru α oprzemy na układzie zastępczego mostka Wheatstone'a. Po wprowadzeniu oznaczeń (7) oraz wyrażeń pomocniczych

$$\lambda = q + \frac{np^2}{p(1+n) + sm}, \quad \mu = 1 + \frac{sp}{p(1+n) + sm} \quad (8)$$

schemat z rys. 3 przyjmie postać przedstawioną na rys. 5. Obowiązują tutaj następujące równania wyjściowe:

$$\left. \begin{aligned} I &= i_1 + i_2, \\ i_1(nR_b + R_b) - i_g R_b &= i_2(r + mn\mu R_b + m\mu R_b) + i_g m\mu R_b \\ (i_2 + i_g)m\mu R_b + i_g \lambda R_b + e - (i_1 - i_g)R_b &= 0 \end{aligned} \right\}.$$



Rys. 5. Układ do obliczenia uchybu od nieczułości

Układ ten można zapisać w formie

$$\begin{aligned} i_1 \cdot 1 &+ i_2 \cdot 1 &+ i_g \cdot 0 &= I \\ i_1(1+n) + i_2 \left[-\frac{r}{R_b} - m\mu(1+n) \right] + i_g [-(1+m\mu)] &= 0 \\ i_1(-1) + i_2 m\mu &+ i_g(\lambda + 1 + m\mu) &= -\frac{e}{R_b}. \end{aligned}$$

Wyznacznik główny D przedstawionego układu równań ma postać

$$D = -\frac{r}{R_b}(\lambda + 1 + m\mu) - (1 + m\mu)[\lambda(1+n) + n(1+m\mu)].$$

Wyznacznik D_{i_g} niewiadomej i_g wyraża się następująco:

$$D_{i_g} = -\frac{r}{R_b} \left(I - \frac{e}{R_b} \right) + \frac{e}{R_b} (1+n)(1+m\mu).$$

Wobec tego prąd w galwanometrze wynosi

$$i_g = \frac{\frac{r}{R_b} \left(I - \frac{e}{R_b} \right) - \frac{e}{R_b} (1+n)(1+m\mu)}{\frac{r}{R_b} (\lambda + 1 + m\mu) + (1+m\mu)[\lambda(1+n) + n(1+m\mu)]} = C_i \alpha_i$$

Przekształcając dalej powyższy wzór otrzymujemy wyrażenie

$$r = R_b \cdot (1+m\mu) \frac{\frac{e}{R_b} (1+n) + C_i \alpha [\lambda(1+n) + n(1+m\mu)]}{I - \frac{e}{R_b} - C_i \alpha (\lambda + 1 + m\mu)}. \quad (19)$$

Różniczkujemy wzór (19) względem odchylenia galwanometru, a następnie przyjmujemy $\alpha = 0$:

$$\left[\frac{\partial r}{\partial \alpha} \right]_{\alpha=0} = \frac{R_b (1+m\mu)}{\left[I - \frac{e}{R_b} \right]^2} C_i \left\{ \left(I - \frac{e}{R_b} \right) [\lambda(1+n) + n(1+m\mu)] + (\lambda + 1 + m\mu) \frac{e}{R_b} (1+n) \right\}. \quad (20)$$

Wartość r dla zerowego odchylenia galwanometru wyraża się jako

$$[r]_{\alpha=0} = \frac{R_b (1+m\mu)}{I - \frac{e}{R_b}} \cdot \frac{e}{R_b} \cdot (1+n). \quad (21)$$

Po podzieleniu (20) przez (21) i wykonaniu redukcji otrzymujemy

$$\begin{aligned} \left[\frac{\partial r}{\partial \alpha} \right]_{\alpha=0} \cdot \frac{1}{[r]_{\alpha=0}} &= C_i \frac{I\lambda(1+n) + (1+m\mu) \left(In + \frac{e}{R_b} \right)}{\left(I - \frac{e}{R_b} \right) \frac{e}{R_b} (1+n)} \approx \\ &\approx C_i \frac{I\lambda(1+n) + (1+m\mu) In}{I \frac{e}{R_b} (1+n)} = \frac{C_i R_b}{e} \left[\lambda + \frac{n}{1+n} \cdot (1+m\mu) \right]. \end{aligned} \quad (22)$$

Do wzoru (22) wstawiamy teraz oznaczenia pomocnicze (18). Po zgrupowaniu odpowiednich symboli możemy napisać

$$|\delta_{n(r)}| = \frac{C_i A_\alpha R_b}{e} \cdot \left[q + \frac{n}{1+n} + p \left(1 - \frac{1}{1+n + \frac{sm}{p}} \right) \right].$$

Powyższy wzór można uprościć przez pominięcie wyrazu $\frac{sm}{p}$. Wynika to

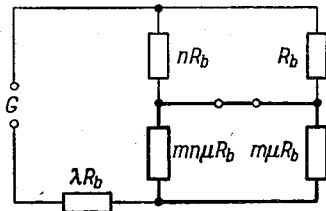
stąd, że $\left(\frac{s \cdot m}{p} \right)_{\max} = (r_o)_{\max} : (r_b)_{\min} = 3 \cdot 10^{-3} : 10 = 3 \cdot 10^{-4}$, co jest pomijalnie małe wobec 1.

Wzór na uchyb od nieczułości przyjmuje obecnie postać

$$|\delta_{n(r)}| = \frac{C_i \Delta_a R_b}{e} \left(q + n \frac{1+p}{1+n} \right) \quad (23)$$

5. OPORNOŚĆ ZASTĘPCZA MOSTKA WIDZIANA Z ZACISKÓW GALWANOMETRU

Schemat zastępczego mostka Wheatstone'a z rys. 5 przekształcamy na układ przedstawiony na rys. 6. Oporność wewnętrzna galwanometru w tym układzie została włączona do oporu λR_b . Przez stopniowe upraszczanie schematu łatwo dochodzimy do wyrażenia dla oporności zastępczej



Rys. 6. Układ do obliczenia oporności zastępczej

$$R_M = R_b \left[q + n \cdot \frac{1+m}{1+n} + n p \frac{1 + \frac{sm}{p(1+n)}}{1 + n + \frac{sm}{p}} \right]$$

Powyżej wykazaliśmy, że $\frac{sm}{p} \ll 1$, wobec czego wyrażenie na oporność zastępczą upraszcza się do postaci

$$R_M = R_b \left(q + n \frac{1+m+p}{1+n} \right) \approx R_b \cdot \left(q + n \frac{1+p}{1+n} \right) \quad (24)$$

Dokonane przybliżenie jest zwykle wystarczająco dokładne, gdyż rzadko spotykamy się z wartością $m_{\max} = 0,1$. W znacznej większości przypadków spotykanych w praktyce jest $m \leq 0,01$.

Korzystając z wyprowadzonego wyrażenia dla oporności zastępczej mostka widzianej z zacisków galwanometru idealnego, możemy przepisać wzór na uchyb od nieczułości w nowej postaci:

$$|\delta_{n(r)}| = \frac{C_i \cdot \Delta_a R_M}{e} \quad (25)$$

6. DYSKUSJA UCHYBU SYSTEMATYCZNEGO

Dyskusja uchybu systematycznego ma na celu danie wytycznych dla takiego doboru dokładności elementów układu, by uzyskać możliwie dużą dokładność przy pomocy sprzętu typowego.

Dla uproszczenia rozważań przyjmujemy, że amperomierz i woltomierz mają jednakową klasę dokładności k_m , a zakresy ich są tak dobrane, iż oba mierniki pracują ze wskazaniem względnym $a/a_{\max} = \frac{2}{3}$. Zakładamy również, że oporniki R_1, R_2, R_a, R_b są jednakowej klasy dokładności k_r .

Wzór (17) na uchyb systematyczny przyjmie wówczas postać

$$[\delta_{s(r)}]^2 = 2 \cdot \left(\frac{k_m}{300} \cdot \frac{3}{2} \right)^2 + 2 \left(\frac{k_r}{\sqrt{3} \cdot 100} \right)^2 + \left(\frac{n}{1+n} \right)^2 \left[2 \cdot \left(\frac{k_r}{\sqrt{3} \cdot 100} \right)^2 \right] = \\ = \frac{2}{4} \left(\frac{k_m}{100} \right)^2 + \frac{2}{3} \left(\frac{k_r}{100} \right)^2 \left[1 + \left(\frac{n}{1+n} \right)^2 \right].$$

Zbadamy teraz, kiedy drugi człon powyższej sumy jest pomijalnie mały wobec pierwszego. Przyjmujemy, że w obliczeniu uchybu można dopuścić błąd o wartości $\pm 5\%$. Błąd taki powstanie, jeżeli odrzucimy człon o wartości 10-krotnie mniejszej od członu zachowanego. Wynika stąd wymaganie

$$\frac{2}{3} \left(\frac{k_r}{100} \right)^2 \left[1 + \left(\frac{n}{1+n} \right)^2 \right] \leq \frac{1}{10} \cdot \frac{2}{4} \left(\frac{k_m}{100} \right)^2,$$

które zapiszemy w postaci

$$\frac{k_r}{k_m} \leq \sqrt{\frac{3}{40} \cdot \frac{(1+n)^2}{(1+n)^2 + n^2}}. \quad (26)$$

Dla różnych współczynników n wartości wymagania (26) kształtują się następująco:

$$\begin{array}{ccc} n = 0,1 & n = 1 & n = 10 \\ k_r/k_m \leq 0,27 & k_r/k_m \leq 0,25 & k_r/k_m \leq 0,20 \end{array}$$

Widać zatem, że niezależnie od współczynnika n spełnienie zależności $k_r/k_m \leq 1/5$ pozwala obliczać uchyb systematyczny ze wzoru uproszczonego

$$[\delta_{s(r)}]^2 = \left(\frac{k_m}{300} \cdot \frac{I_n}{I} \right)^2 + \left(\frac{k_m}{300} \cdot \frac{U_n}{U} \right)^2 \quad (27)$$

pod warunkiem, że wskazania obu mierników nie są mniejsze niż $2/3$ użytego zakresu pomiaru.

7. DYSKUSJA UCHYBU OD NIECZUŁOŚCI

Dyskusja uchybu od nieczułości ma za zadanie wskazać taki dobór sprzętu i warunków pomiaru, by możliwie łatwo uzyskać czułość potrzebną dla zachowania założonej dokładności.

Ze względu na szybkość i wygodę pomiaru wskazane jest zastosowanie mostka i galwanometru tak dopasowanych, by warunki tłumienia były krytyczne (graniczne). Wynika stąd żądanie, by oporność krytyczna całkowita R_{kc} galwanometru była równa oporności zastępczej mostka, widzianej z zacisków wskaźnika zerowego R_M . Zapiszemy to jako

$$R_{kc} = R_b \left(q + n \frac{1+p}{1+n} \right). \quad (28)$$

Stała napięciowa galwanometru w warunkach tłumienia krytycznego C_{uk} wyraża się jako $C_{uk} = C_i R_{kc}$, wobec czego wzór (25) na uchyb od nieczu-

łości przy spełnieniu wymagania (28) przybiera postać

$$|\delta_{n(r)}| = \frac{C_{uk} \Delta_a}{e}. \quad (29)$$

Zbadamy obecnie, kiedy przy obliczaniu uchybu całkowitego uchyb od nieczułości jest pomijalnie mały wobec uchybu systematycznego. Uchyb systematyczny wg wzoru (27) przy założeniu wskazań względnych obu mierników w wysokości $\frac{2}{3}$ wynosi:

$$[\delta_s(r)]^2 = 2 \left[\frac{k_m}{300} \cdot \frac{3}{2} \right]^2 = \frac{1}{2} \left[\frac{k_m}{100} \right]^2.$$

Zgodnie z uwagą w rozdz. 6 przy rozpatrywaniu wzoru (17) należy spełnić wymaganie

$$\left(\frac{C_{uk} \Delta_a}{e} \right)^2 \leq \frac{1}{10} \cdot \frac{1}{2} \left(\frac{k_m}{100} \right)^2.$$

Przyjmując ostrożnie $\Delta_a = 1/4$ działki, powyższą zależność przedstawimy w postaci uproszczonej

$$C_{uk} \leq \frac{e k_m}{100}. \quad (30)$$

Wartość e w poprzednim wyrażeniu jest zależna od mierzonej wartości r oraz innych czynników wg wzoru (8). Występujące tam natężenie prądu zasilającego układ I jest uwarunkowane z kolei obciążalnością elementów mostka. Zakładamy, że jest ona równa dla obu oporów R_x oraz R_c i wynosi P_d [W].

W przypadku, gdy $n \leq 1$ jest $R_c \geq R_x$, wobec czego $i_{2d} = \sqrt{P_d/R_c} = \sqrt{P_d/R_x} \cdot \sqrt{n}$, zaś dopuszczalny prąd zasilający wynosi wówczas w przybliżeniu $I_d \approx i_{2d} (1+m) = \sqrt{P_d/R_x} \sqrt{n} \cdot (1+m)$.

Analogicznie dla przypadku $n \geq 1$ mamy $R_x \geq R_c$, czyli $i_{2d} = \sqrt{P_d/R_x}$ oraz $I_d \approx \sqrt{P_d/R_x} (1+m)$.

Wyrażając wartość e z równania (8) i podstawiając powyżej obliczone dopuszczalne natężenia prądu zasilającego I_d otrzymujemy:

$$\begin{aligned} \text{dla } n \leq 1: \quad e &= \left(\frac{r}{R_x} \sqrt{P_d R_x} \right) \frac{\sqrt{n}}{1+n}, \\ \text{dla } n \geq 1: \quad e &= \left(\frac{r}{R_x} \sqrt{P_d R_x} \right) \frac{1}{1+n}. \end{aligned} \quad (31)$$

Porównując wzory (30) i (31) widzimy, że warunkiem łatwego uzyskania potrzebnej czułości układu jest możliwie duża wartość mnożnika stojącego za nawiasem po prawej stronie równości częściowych (31).

Dla $n = 0,1$ mnożnik ten wynosi 0,29, dla $n = 1$ — 0,50, zaś dla $n = 10$ — tylko 0,09. Wynika stąd wniosek, że należy stosować $n = 1$. Po uwzględ-

nieniu tego wymagania oraz zależności (31) wzór (30) zapiszemy w postaci

$$C_{uk} \leq \frac{1}{2} \sqrt{P_d \cdot R_x} \cdot \frac{r}{R_x} \cdot \frac{k_m}{100}. \quad (32)$$

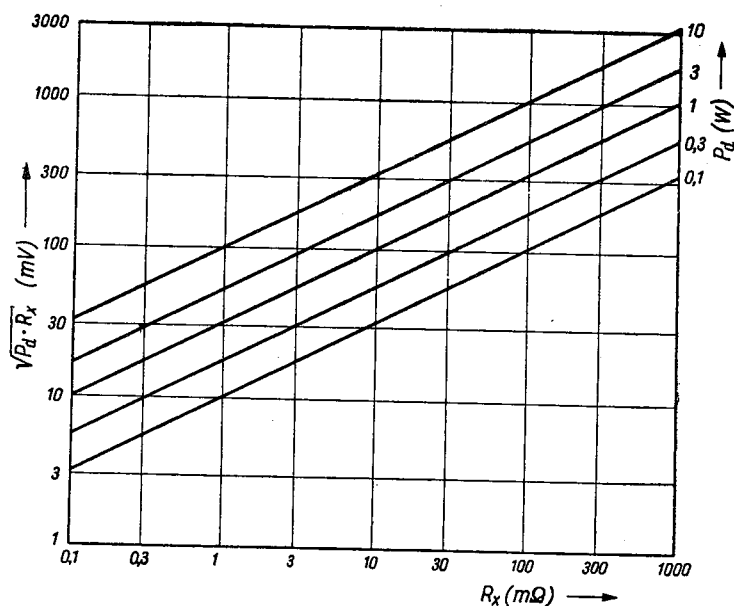
Niżej przytoczone tablice i wykresy mają za zadanie zilustrować wartości $\sqrt{P_d \cdot R_x}$ oraz $(C_{uk})_{\max}$, z jakimi wypadnie się spotkać.

Tablica 2

$$\sqrt{P_d \cdot R_x} = f(P_d; R_x) \text{ [mV] przy } n = 1$$

R_x [mΩ]	P_d [W]				
	0,1	0,3	1	3	10
0,1	3,2	5,5	10	17	32
0,3	5,5	9,5	17	30	55
1,0	10	17	32	55	100
3,0	17	30	55	95	170
10,0	32	55	100	170	320
30,0	55	95	170	300	550
100	100	170	320	550	1000
300	170	300	550	950	1700
1000	320	550	1000	1700	3200

Na podstawie powyższej tablicy sporządzono wykresy przedstawione na rys. 7.



Rys. 7. Dopuszczalny spadek napięcia

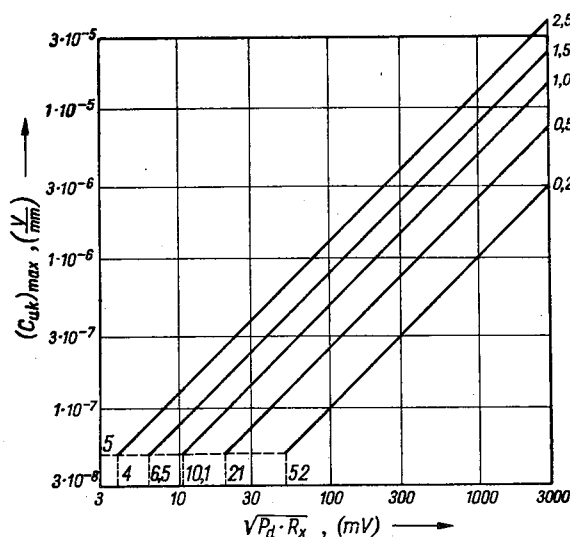
$$\sqrt{P_d R_x} = f(P_d; R_x); \quad n = 1$$

Tablica 3

$$(C_{uk})_{\max} = f(k_m; \sqrt{P_d R_x}) \left[\frac{V}{dz.} \right] \text{ przy } \frac{r}{R_x} = 10^{-3}$$

$\sqrt{P_d R_x}$ [mV]	k_m				
	0,2	0,5	1,0	1,5	2,5
3	$3,0 \cdot 10^{-9}$	$7,5 \cdot 10^{-9}$	$1,5 \cdot 10^{-8}$	$2,3 \cdot 10^{-8}$	$3,8 \cdot 10^{-8}$
5,5	$5,5 \cdot 10^{-9}$	$1,4 \cdot 10^{-8}$	$2,8 \cdot 10^{-8}$	$4,1 \cdot 10^{-8}$	$6,9 \cdot 10^{-8}$
10	$1,0 \cdot 10^{-8}$	$2,5 \cdot 10^{-8}$	$5,0 \cdot 10^{-8}$	$7,5 \cdot 10^{-8}$	$1,3 \cdot 10^{-7}$
17	$1,7 \cdot 10^{-8}$	$4,3 \cdot 10^{-8}$	$8,5 \cdot 10^{-8}$	$1,3 \cdot 10^{-7}$	$2,1 \cdot 10^{-7}$
30	$3,0 \cdot 10^{-8}$	$7,5 \cdot 10^{-8}$	$1,5 \cdot 10^{-7}$	$2,3 \cdot 10^{-7}$	$3,8 \cdot 10^{-7}$
55	$5,5 \cdot 10^{-8}$	$1,4 \cdot 10^{-7}$	$2,8 \cdot 10^{-7}$	$4,1 \cdot 10^{-7}$	$6,9 \cdot 10^{-7}$
100	$1,0 \cdot 10^{-7}$	$2,5 \cdot 10^{-7}$	$5,0 \cdot 10^{-7}$	$7,5 \cdot 10^{-7}$	$1,3 \cdot 10^{-6}$
170	$1,7 \cdot 10^{-7}$	$4,3 \cdot 10^{-7}$	$8,5 \cdot 10^{-7}$	$1,3 \cdot 10^{-6}$	$2,1 \cdot 10^{-6}$
300	$3,0 \cdot 10^{-7}$	$7,5 \cdot 10^{-7}$	$1,5 \cdot 10^{-6}$	$2,3 \cdot 10^{-6}$	$3,8 \cdot 10^{-6}$
550	$5,5 \cdot 10^{-7}$	$1,4 \cdot 10^{-6}$	$2,8 \cdot 10^{-6}$	$4,1 \cdot 10^{-6}$	$6,9 \cdot 10^{-6}$
1000	$1,0 \cdot 10^{-6}$	$2,5 \cdot 10^{-6}$	$5,0 \cdot 10^{-6}$	$7,5 \cdot 10^{-6}$	$1,3 \cdot 10^{-5}$
1700	$1,7 \cdot 10^{-6}$	$4,3 \cdot 10^{-6}$	$8,5 \cdot 10^{-6}$	$1,3 \cdot 10^{-5}$	$2,1 \cdot 10^{-5}$
3000	$3,0 \cdot 10^{-6}$	$7,5 \cdot 10^{-6}$	$1,5 \cdot 10^{-5}$	$2,3 \cdot 10^{-5}$	$3,8 \cdot 10^{-5}$

Powyższą zależność ilustruje rodzina krzywych na rys. 8.



Rys. 8. Stała napięciowa galwanometru

$$(C_{uk})_{\max} = f(k; \sqrt{P_d R_x}); \frac{r}{R_x} = 1 \cdot 10^{-3}$$

8. DOKŁADNOŚĆ POMIARU

Przed wszystkim zestawimy warunki uzyskania dobrej dokładności pomiaru, omówione poprzednio w punktach 6 i 7.

Zastosowane mierniki o klasie dokładności k_m powinny pracować powyżej $2/3$ zakresu. Oporniki mostka powinny być klasy $k_r \leq k_m/5$. Mostek o układzie poziomo-symetrycznym, tj. $n = 1$. Oporność krytyczna całkowita galwanometru powinna spełniać warunek (28), a stała napięciowa przy tłumieniu krytycznym — (32).

Spełnienie wymagania (28) najdogodniej można uzyskać przez odpowiedni dobór współczynnika p . Wypada zauważyć, że zwykle jest $R_1 \ll R_g$, wobec czego $q \approx R_g/R_b$. Korzystając z tego faktu możemy na podstawie (28) napisać warunek dla oporu krytycznego zewnętrznego R_k użytego galwanometru

$$R_k = R_b n \frac{1+p}{1+n} \text{ i stąd obliczyć potrzebną wartość}$$

$$p = \frac{R_k}{R_b} \cdot \frac{1+n}{n} - 1. \quad (33)$$

Przy spełnieniu powyższych wymagań całkowity uchyb poprawki r jest praktycznie równy uchybowi systematycznemu i w przybliżeniu wynosi

$$[\delta_c(r)]^2 = \frac{1}{2} \left(\frac{k_m}{100} \right)^2. \quad (34)$$

Na podstawie definicji (3) całkowity uchyb oporności mierzonej R_x można w przybliżeniu obliczyć jako

$$[\delta_c(R_x)]^2 = [\delta_s(R_n)]^2 + \left[\frac{r}{R_x} \delta_c(r) \right]^2. \quad (35)$$

Drugi składnik powyższego wyrażenia jest pomijalnie mały wobec pierwszego pod warunkiem spełnienia nierówności:

$$\frac{r}{R_x} |\delta_c(r)| \leq \frac{1}{\sqrt{10}} |\delta_s(R_n)|,$$

która po uwzględnieniu zależności (34) przybiera postać

$$k_m \leq 45 \frac{R_x}{r} |\delta_s(R_n)| = 15 \frac{R_x}{r} |\delta_w(R_n)|, \quad (36)$$

gdyż do obliczeń przyjmujemy wartość sprawdzoną (a nie nominalną) oporu wzorcowego R_n , której uchyb wiarygodny jest 3-krotnie większy od średniego.

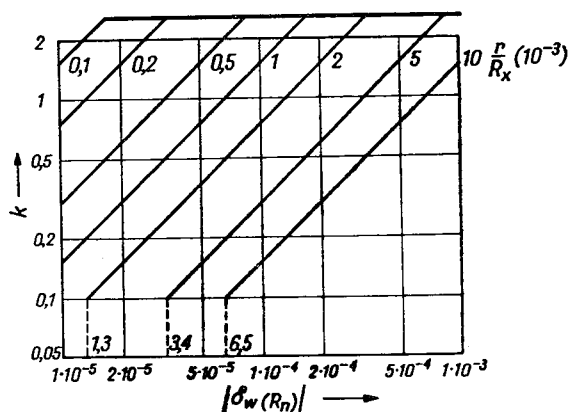
Przy założeniu różnych wartości r/R_x optymalny dobór dokładności elementów układu wg zależności (36) daje następujące wartości liczbowe zestawione w tablicy 4.

Tablica 4

$$(k_m)_{\max} = f\left(\frac{r}{R_x}; \delta_{w(R_n)}\right)$$

$\frac{r}{R_x}$	$\delta_{w(R_n)}$						
	$1 \cdot 10^{-5}$	$2 \cdot 10^{-5}$	$5 \cdot 10^{-5}$	$1 \cdot 10^{-4}$	$2 \cdot 10^{-4}$	$5 \cdot 10^{-4}$	$1 \cdot 10^{-3}$
$1 \cdot 10^{-4}$	1,5	2,5	2,5	2,5	2,5	2,5	2,5
$2 \cdot 10^{-4}$	0,75	1,5	2,5	2,5	2,5	2,5	2,5
$5 \cdot 10^{-4}$	0,3	0,6	1,5	2,5	2,5	2,5	2,5
$1 \cdot 10^{-3}$	0,15	0,3	0,75	1,5	2,5	2,5	2,5
$2 \cdot 10^{-3}$	0,08	0,15	0,38	0,75	1,5	2,5	2,5
$5 \cdot 10^{-3}$	0,0	0,0	0,15	0,3	0,6	1,5	2,5
$1 \cdot 10^{-2}$	0,0	0,0	0,08	0,15	0,3	0,75	1,5

W przytoczonej tablicy wartość 2,5 wpisano wszędzie tam, gdzie wzór (36) daje tę właśnie wartość — lub też większą. To samo dotyczy wykresów powyższej zależności przedstawionych na rys. 9.



Rys. 9. Klasa dokładności mierników

$$k = f\left(\frac{r}{R_x}; |\delta_{w(R_n)}|\right)$$

Przy spełnieniu wszystkich wymagań, podanych krótko w niniejszym punkcie, można przyjąć praktycznie, że uchyb wiarygodny oporu mierzonego jest równy uchybowi opornika wzorcowego, gdyż różnice w uchybach nie powinny przekroczyć ok. 15% uchybu.

9. WYNIK POMIARU PRZYKŁADOWEGO

Celem zilustrowania strony praktycznej przedstawionej metody wykonano pomiar tego samego opornika normalnego opisanym sposobem jak

również (dla porównania) z zastosowaniem podstawienia w zrównoważonym mostku Thomsona.

Jako metodę porównawczą wybrano podstawienie w mostku zrównoważonym, gdyż daje ono dobrą dokładność pomiaru i łatwo pozwala uzyskać daleko idącą zbieżność warunków pomiaru w obu metodach. W obu przypadkach zastosowano wspólny układ pomiarowy, zmontowany na stałe (z wyjątkiem nieznacznych zmian uwarunkowanych różnicami obu metod). Prąd zasilający był również jednakowy w obu pomiarach. Pomiaru wykonano bezpośrednio jedno po drugim, w stałej temperaturze 22,5°C.

Wskaźnikiem zerowym był galwanometr typu M 21/1, nastawiony na maksymalną czułość napięciową. Jako opory R_a i r_a zastosowano oporniki siedmiodekadowe, a jako R_b i r_b — sześciodekadowe klasy 0,02 firmy „Ulrich”. Układ mostka był podwójnie symetryczny, to znaczy $n = p = 1$, a opór R_b wynosił w przybliżeniu 6 000 omów. Pozwoliło to na bardzo płynną regulację układu i dało dobre dopasowanie do oporności krytycznej galwanometru. Jako opornik R_1 zastosowano normalny o wartości 1 Ω . Oporem R_2 w mostku skompensowanym a bocznikiem do obiektu badanego w zrównoważonym był sześciodekadowy opornik klasy 0,02 „Ulrich”. Amperomierz i woltomierz były miernikami klasy 0,2 (magneto-elektryczne) „RFT”.

Obiekt badany stanowił opór normalny 0,1 Ω int. firmy „Wolff”, a opornik pomocniczy R_c — 0,1 Ω „Siemens”. Wszystkie pozostałe oporności, szczególnie wzorzec 0,1 Ω „Schlottheim”, były wyrażone w jednostkach absolutnych.

Pomiary dla każdej z metod wykonano czterokrotnie, zmieniając za każdym razem kierunek prądu. Dobrze dopasowanie układu i wskaźnika równowagi o znacznej czułości (z promieniem świetlnym o długości 4,5 m) zapewniło potrzebną czułość metody. Rozstrojenie mostka o $63 \cdot 10^{-5}$ dawało odchylenie galwanometru o 75 działek. Prąd zasilający wynosił przy tym tylko ok. 1,4 A, co daje obciążenie opornika normalnego w wysokości ca 1/5 W. Dla uzyskania stabilnych warunków nagrzewania prądowego oporniki normalne były stale obciążone; na czas koniecznych zmian w układzie odłączano tylko galwanometr. Dążenie do utrzymania małej wartości oporności połączenia r_o podyktowało układ oporników w obwodzie prądowym w następującej kolejności: R_x , R_c , R_n . Przełączenie się z oporu R_n na R_x wymagało jedynie przeniesienia przewodów z zacisków napięciowych jednego opornika na drugi. Całość układu zajmowała znaczną powierzchnię ze względu na ilość elementów oraz spore rozmiary oporników dekadowych. Zasilanie główne stanowiła bateria akumulatorów 12 V, a źródłem SEM kompensującej było pojedyncze ogniwo ołowiowe 2 V.

Podstawienie w skompensowanym mostku Thomsona dało następujące wyniki:

Wartości stałe: $1 + m = 1,000$, $1 + n = 2,000$, $R_1 = 1 \Omega$

Lp.	I	U	R_2	r/R_n
	A	V	k Ω	10^{-5}
1	+ 1,425	+ 1,97	45,2	61,2
2	- 1,430	- 1,97	42,6	64,7
3	+ 1,435	+ 1,97	44,0	62,4
4	- 1,435	- 1,97	41,9	65,5

Wynik pomiaru $(r/R_n)_{sr} = 63,5 \cdot 10^{-5}$.

Obliczenia wykonano na podstawie wzoru przybliżonego (8).

Wyniki otrzymane przy zastosowaniu podstawienia w zrównoważonym mostku Thomsona kształtują się jak niżej:

Lp.	R	r/R_n
	Ω	10^{-5}
1	159	62,9
2	157	63,7
3	159	62,9
4	157	63,7

Wynik pomiaru $(r/R_n)_{sr} = 63,3 \cdot 10^{-5}$.

W powyższej tablicy R oznacza wartość bocznika, który przyłączony do obiektu badanego sprowadza układ do stanu równowagi. Układ jest również w równowadze przy włączeniu oporu wzorcowego R_n bez bocznika. Obliczenia wykonano na podstawie wzoru

$$r/R_n = \frac{R_n}{R - R_n} \approx \frac{R_n}{R}.$$

Przytoczone wyniki wykazują zasadniczą zgodność obu metod. Rozrzut wyników w mostku skompensowanym jest spowodowany trudnością utrzymania całego rozległego układu w ściśle jednakowej temperaturze. Powstająca SEM pasożytnicza wykazuje trochę większy wpływ na wyniki poszczególnego pomiaru niż w przypadku mostka zrównoważonego. Szybkie powtórzenie odczytu przy zmienionym kierunku prądu pozwala jednak ten wpływ praktycznie zupełnie wyeliminować.

10. ZAKOŃCZENIE

Przedstawione powyżej wywody świadczą wyraźnie, iż podstawienie w skompensowanym mostku Thomsona pozwala osiągnąć dużą dokładność pomiaru, podobnie jak w skompensowanym mostku Wheatstone'a ([2], [3]). Warunki niezbędne do uzyskania znacznej dokładności nie są specjalnie trudne do spełnienia. Jednocześnie sam pomiar nie jest nadmiernie skom-

plikowany i nie pochłania zbyt wiele czasu przy badaniach serii oporników o jednakowej wartości nominalnej. W takich przypadkach pierwszy etap pomiaru wykonuje się jeden raz dla całej serii. Oczywiście wymaga to bardzo dobrej stałości warunków pomiaru.

Przytoczone rozważania pozwalają pozytywnie ocenić wartość proponowanej metody — nie wskazują jednak na jakąś przewagę nad podstawieniem w mostku zrównoważonym. Trzeba jednak zauważyć, że zwykle podstawienie ma mocno ograniczony zakres stosowalności. Wynika to z niżej podanych względów.

Przyłączony bocznik o zmienianej oporności musi pozwolić na regulację w skokach najwyżej po 1%. Wynika stąd, że minimalna oporność bocznika regulowanego skokami po 0,1 Ω powinna wynosić 10 Ω . Mniejsze oporności nie wchodzić również w rachubę ze względu na zbyt duży wówczas wpływ przewodów łączących z oporem bocznikowanym.

Maksymalną zmianę oporności obiektu uzyskujemy przyłączając minimalną wartość bocznika $R_{\min}$. Otrzymana zmiana względna $\varrho_{\max}$ oporności zastępczej w stosunku do oporu bocznikowanego zależy od jego wartości wg niżej podanej tablicy:

Wartość stała: $R_{\min} = 10 \Omega$

$R_x [\Omega]$	$1 \cdot 10^{-1}$	$3 \cdot 10^{-2}$	$1 \cdot 10^{-2}$	$3 \cdot 10^{-3}$	$1 \cdot 10^{-3}$	$3 \cdot 10^{-4}$	$1 \cdot 10^{-4}$
$\varrho_{\max} [\%]$	10	3	1	0,3	0,1	0,03	0,01

Małe opory dokładne mogą (wg wymagań PTR [1]) wykazywać względne odchylenie od wartości nominalnej w wysokości $\pm 1\%$. Widać stąd, że dowolna odmiana podstawienia w mostku zrównoważonym da się zastosować bez ograniczeń przy badaniu oporników o wartości nominalnej ok. 0,01 Ω lub większej. Poniżej tej granicy zwykle podstawienie można zastosować tylko do porównania oporników o rzeczywistych wartościach bardzo bliskich sobie.

Podstawienie w mostku skompensowanym jest wolne od tego ograniczenia, co stanowi jego bardzo istotną zaletę. Warto również podkreślić, że proponowana metoda opiera się na zastosowaniu wyłącznie sprzętu typowego, łatwo dostępnego. Te korzystne cechy niewątpliwie sprzyjać będą jej zastosowaniu.

Politechnika Gdańska
Katedra Miernictwa Elektrycznego

WYKAZ LITERATURY

1. Bekanntmachung über die Beglaubigung elektrischer Präzisionswiderstände durch die Physikalisch-Technische Reichsanstalt vom 24.03.1939.— Amtsblatt der PTR, 15 Reihe, Nr. 3 (1939), s. 53.
2. Sawicki J.: Skompensowany mostek Wheatstone'a. — Arch. Elektrot., t. VII, z. 3, Warszawa 1958.
3. Sawicki J.: Substytucja w skompensowanym mostku Wheatstone'a Arch. Elektrot., t. IX, z. 2, Warszawa 1960.

Е. САВИЦКИ

ЗАМЕЩЕНИЕ В КОМПЕНСИРОВАННОМ МОСТЕ ТОМСОНА

Резюме

Компенсированный мост Томсона отличается от уравновешенного тем, что условие равновесия не точно соблюдено; гальванометр доводится к нулевому отклонению посредством последовательного с ним введения ЭДС, которая известна и регулируется

Измерение по этому методу состоит из двух основных частей. В первом этапе мы вводим сравнительное сопротивление R_n на место где обыкновенно включается измеряемое сопротивление. Мост точно уравновешивается при включенной компенсирующей ЭДС. Точность уравновешения проверяется при прерванном соединении между небольшими сопротивлениями.

Во втором этапе включается измеряемое сопротивление R_x на обыкновенном месте и тогда мост компенсируется, причём все прежде включённые сопротивления остаются без изменений.

Значение испытуемого сопротивления получаем из формулы (3). Точную корректуру r находим при помощи формулы (6). На практике пользуемся приближённым выражением (8). Относительную погрешность приближения δ' определяем по формуле (9). Встречаемые значения этой погрешности видны из таблицы 1-й и чертежа 4-го. Ток питания обозначен I ; компенсирующая ЭДС e находится по формуле (11).

Систематическая погрешность $\delta_{s(r)}$ корректуры определяется уравнением (17), а погрешность от нечувствительности $\delta_{n(r)}$ — (23). Принимая во внимание выражение (24) для эквивалентного сопротивления моста R_M , можно определить последнюю погрешность по формуле (25). В случае моста согласованного с критическим сопротивлением гальванометра, получаем формулу (29).

Когда стрелочные приборы выказывают релятивные отклонения не меньше $2/3$ и класс точности k_r сопротивлений R_1 и R_2 не меньше $1/5$ класса точности k_m измерителей, тогда систематическая погрешность находится по упрощенной формуле (27).

Если условие формулы (30) выполнено, то погрешность от нечувствительности является настолько незначительной по отношению к систематической, что ею можно пренебречь. Наибольшую чувствительность достигаем когда $n = 1$ и при полном использовании нагрузочной способности P_d испытуемого сопротивления. Формула (32) определяет необходимую постоянную по напряжению для гальванометра. Выражение (33) дает коэффициент p как функцию наружного критического сопротивления гальванометра R_k . Условие (36) даёт нам возможность такого подбора точности измерителей, чтобы достоверная погрешность $\delta_{w(R_x)}$ испытуемого сопротивления практически совпадала с достоверной погрешностью $\delta_{w(R_n)}$ сравнительного сопротивления. Относительная разность обоих погрешностей не должна превышать 15% погрешности.

Таблицы 2, 3, 4 а также чертежи 7, 8, 9 иллюстрируют числовые значения для наилучшего подбора элементов схемы с точки зрения чувствительности и точности измерения.

В абзаце 9 описано примерное измерение, выполненное при помощи замещения в компенсированном и уравновешенном мосте. В последем случае мост уравновешен посредством регулируемого шунта к испытываемому сопротивлению. Соответствующие средние результаты: $63,5 \cdot 10^{-5}$ и $63,3 \cdot 10^{-5}$ доказывают нам, что предлагаемый метод даёт хорошую точность.

Описанный метод может быть применяем при сравнении, по отношению 1:1, всяких незначительных четырех-зажимных сопротивлений. Однако, замещение в уравновешенном мосте (посредством регулируемого шунта) может быть применено лишь для сопротивлений не меньше 0,01 Ω . Для меньших сопротивлений шунт недопустимо мал. Важное преимущество компенсированного моста заключается в том, что этот метод пользуется исключительно типичными аппаратами и не требует специальных элементов (на пример: интерполяционные сопротивления).

J. SAWICKI

SUBSTITUTION IN DER KOMPENSIERTEN THOMSON-BRÜCKE

Zusammenfassung

Die kompensierte Thomson-Brücke unterscheidet sich von der abgeglichenen dadurch, dass die Gleichgewichtsbedingung nicht genau erfüllt wird. Das Galvanometer wird vermittle einer in die Reihe eingeschalteten Kompensations-EMK bis zur Nullstellung zugeführt. Der Wert dieser EMK muss bekannt und regulierbar sein.

Die Messung nach dieser Methode besteht aus zwei grundsächlichen Teilen. In der ersten Etappe wird der Vergleichswiderstand R_n in derselben Stelle eingeschaltet, wo gewöhnlich der Prüfling steht. Die Brücke wird genau abgeglichen bei ausgeschalteter Kompensations-EMK. Die Genauigkeit der Abgleichung überprüft man bei unterbrochener Verbindung zwischen den kleinen Widerständen.

In der zweiten Etappe schalten man den Prüfling R_x in die gewöhnliche Stelle ein und dann wird die Brücke kompensiert, wobei die vorher eingestellten Widerstände beibehalten werden.

Der Wert des Prüflings wird aus der Formel (3) berechnet. Der genaue Ausdruck für die Korrektur „ r “ ist durch die Formel (6) festgestellt. Praktisch, wird der angenäherte Ausdruck (8) angewendet. Den relativen Fehler der Annäherung δ' berechnet man aus dem Zusammenhang (9). Die vorkommenden Werte diesen Fehlers sind aus der Tabelle 1 und der Zeichnung 4 ersichtlich. Sie sind vernachlässigbar klein den systematischen Fehler der Korrektur gegenüber.

Der Speisestrom der Brücke wird durch I bezeichnet. Die Kompensations-EMK e berechnet man nach der Formel (11).

Der systematische Fehler der Korrektur $\delta_{s(r)}$ ist durch den Zusammenhang (17) und der Unempfindlichkeitsfehler $\delta_{n(r)}$ — (23) festgestellt. Den Ausdruck für den Ersatzwiderstand der Brücke R_M (24) in betracht nehmend, kann man den letzten Fehler mit der Beziehung (25) ausdrücken. Wird die Brücke dem kritischen Widerstand des Galvanometers angepasst, so erhalten wir den Ausdruck (29).

Im Falle wenn die Zeigerinstrumente relative Ablenkungen nicht kleiner als 2/3 aufweisen und die Genauigkeitsklasse k_r der Widerstände R_1 und R_2 nicht 1/5 der Klasse der Zeigerinstrumente überschreitet — so kann der systematische Fehler aus der vereinfachten Formel (27) berechnet werden.

Wird die Bedingung (30) erfüllt, so ist der Unempfindlichkeitsfehler vernachlässigbar klein dem systematischen gegenüber. Die optimale Empfindlichkeit wird erreicht, wenn $n = 1$ und die Belastbarkeit P_d des Prüflings R_x vollständig ausgenutzt ist. Der Zusammenhang (32) bestimmt die notwendige Spannungskonstante des Galvanometers. Die Formel (33) gibt den Koeffizient p als Funktion des kritischen Aussenwiderstandes R_k des Galvanometers an. Die Bedingung (36) erlaubt uns so die Genauigkeitsklasse der Zeigerinstrumente zu wählen, dass der glaubens-

würdige Fehler $\delta_{w(Rx)}$ der Prüflings praktisch dem glaubenswürdigen Fehler $\delta_{w(Rn)}$ des Vergleichswiderstandes gleich wird. Der relative Unterschied der beiden Fehler darf nicht 15% des Fehlers überschreiten.

Die Tabellen 2, 3, 4 und die Zeichnungen 7, 8, 9 geben die Zahlenwerte für die vorteilhafteste Wahl der Brückenelemente an, welche die Empfindlichkeit und Genauigkeit gewährleistet.

Im Absatz 9 ist eine Beispielmessung beschrieben, die mit der Substitution in der kompensierten und abgeglichenen Brücke durchgeführt wurde. Die Abgleichung im zweiten Falle erfolgte vermittels des bekannten und regulierbaren Nebenschlusses R zu dem Prüfling. Die entsprechenden mittleren Resultate: $63,5 \cdot 10^{-5}$ und $63,3 \cdot 10^{-5}$ beweisen, dass die vorliegende Methode eine gute Genauigkeit gibt.

Die besprochene Methode kann zum Vergleichen im Verhältniss 1:1 allerlei kleiner und vierklemmiger Widerstände angewandt werden. Dagegen ist die Substitution in der abgeglichenen Brücke (vermittels des regulierbaren Nebenschlusses) nur für die Widerstände nicht kleiner als $0,01\Omega$ anwendbar. Für kleinere Widerstände ist der Nebenschluss unzulässigbar klein. Ein wichtiger Vorteil der kompensierten Brücke bildet die Tatsache, dass diese Methode nur typische Apparatur benutzt und keine Sonderausführungen (wie z.B. Interpolationswiderstände) gebraucht.

ZENON KRZYCKI

Sprzęgacz krzyżowy z otworem zapełnionym ferrytem

Rękopis dostarczono 27. 5. 1960

Wyznaczone zostało sprzężenie sprzęgacza krzyżowego mającego otwór wypełniony ferrytem. Rozważania przeprowadzono w oparciu o teorię sprzęgacza z małym otworem Bethego [2], [5] i Stinsona [7]. Przedstawiono porównanie obliczeń teoretycznych z wynikami pomiaru dla otworu zapełnionego częściowo. Wskazano na praktyczne zastosowania rozpatrywanego sprzęgacza — głównie do pomiaru szerokości linii rezonansu ferromagnetycznego i efektywnego spektroskopowego współczynnika rozszczepienia ferrytów.

1. WSTĘP

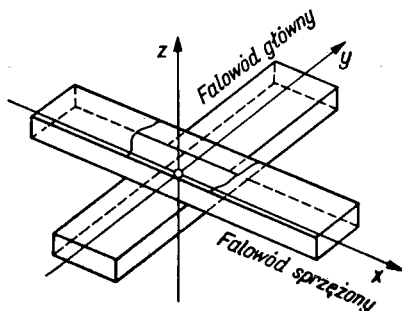
Wśród wielu typów mikrofalowych sprzęgaczy falowodowych sprzęgacz krzyżowy z centralnie położonym otworem wyróżnia się tym, że sprzężenie występuje w nim tylko za pomocą pola elektrycznego [5]. Wprowadzenie do otworu sprzęgającego ferrytu, charakteryzującego się przy podmagnesowaniu stałym polem magnetycznym, tensorową przenikalnością magnetyczną, powoduje wystąpienie także magnetycznego sprzężenia. Sprzężenie elektryczne można w prosty sposób wyeliminować (o czym mowa w dalszej części) i sprzężenie, zależne już tylko od aktualnej przenikalności magnetycznej ferrytu, regulować stałym polem w bardzo dużym przedziale: od zerowego do maksymalnego, występującego przy rezonansie ferromagnetycznym.

Falowodowe sprzęgacze ferrytowe znane są w literaturze już od lat kilku (np. [1], [3] i in.), ale dopiero analiza, podana przez Stinsona [7] w roku 1957, oparta na teorii Bethego [2] dostatecznie dokładnie ujęła zagadnienie.

Celem tego artykułu jest podanie w literaturze krajowej teoretycznej podstawy, na której opiera się pomiar w sprzęgaczu, szerokości linii rezonansu ferromagnetycznego i efektywnego spektroskopowego współczynnika rozszczepienia ferrytów. Układ pomiarowy opisany został przez Stinsona [8] a także przez autora [4].

Artykuł ten oparty jest głównie na pracy Stinsona [7], przy czym wskazano i usunięto pewne merytoryczne nieścisłości tam popełnione, a przez konsekwentnie stosowany rachunek symboliczny uporządkowano także stronę formalną.

Tematem jest wyznaczenie sprzężenia pewnego szczególnego sprzęgacza krzyżowego (rys. 1), utworzonego przez dwa odcinki falowodu prostokątnego, o doskonale przewodzących ściankach połączonych szerszymi bokami prostopadłe względem siebie, mających wspólny, symetrycznie położony mały okrągły otwór wypełniony ferrytem.



Rys. 1. Uproszczony widok sprzęgacza

Rozważania przeprowadzone będą w układzie MKSA dla stanu ustalonego przy sinusoidalnej zależności od czasu.

2. OZNACZENIA

Wprowadzimy następujące oznaczenia:

- d — średnica otworu sprzęgającego,
- t — czas,
- t_0 — grubość ścianki sprzęgającej,
- $\vec{i}_x, \vec{i}_y, \vec{i}_z$ — wersory osi odpowiednio x, y albo z ,
- H_1, E_1 — pole magnetyczne (albo elektryczne) w falowodzie pierwotnym, które istniałoby w środku otworu sprzęgającego, gdyby była tam ścianka,
- $\mathfrak{M}, \gamma$ — wielkości¹⁾ zależne od wielkości i kształtu okna sprzęgającego, równe odpowiednio momentowi dipola magnetycznego okna przypadającemu na jednostkowe pole magnetyczne albo dipola elektrycznego — na jednostkowe pole elektryczne,
- B — indukcja magnetyczna,
- M — magnetyzacja,
- $\vec{I}$ — magnetyczny moment okna,
- Φ — elektryczny moment okna,
- D — indukcja elektryczna,
- 1 — jako indeks — odpowiednie wielkości dla falowodu pierwotnego,

¹⁾ Termin angielski polarizabilities.

- 2 — jako indeks — odpowiednie wielkości dla falowodu wtórnego,
 $\mu_0 = 4 \pi 10^{-7}$ H/m — przenikalność magnetyczna próżni,
 ω — pulsacja,
 λ_0 — długość fali w swobodnej przestrzeni,
 λ_{gr} — długość granicznej fali falowodu,
 λ_f — długość fali w falowodzie,
 g — spektroskopowy współczynnik rozszczepienia,
 $\epsilon_0 = \frac{1}{36 \pi} \cdot 10^{-9}$ F/m — przenikalność elektryczna próżni,
 z_f — impedancja falowa falowodu.

3. MOMENTY DIPOŁOWE OTWORU SPRZĘGAJĄCEGO

Podana przez Bethego [2] teoria dyfrakcji na małych otworach jest bardzo pożyteczna w rozpatrywanym przypadku wyznaczenia sprzężenia przez otwór wypełniony ferrytem. Według niej sprzężenie przez otwór oblicza się określając parę zastępczych dipoli promieniujących: magnetyczny i elektryczny, a następnie wyznacza się pole wzbudzone w falowodzie, do którego dipole te promieniują. Pole wzbudzone określa się z równań wiążących momenty tych dipoli z polami własnych rodzajów falowodu pobudzonego. Momenty ich są proporcjonalne do natężeń pól — odpowiednio magnetycznego albo elektrycznego. Współczynniki proporcjonalności, zależne jedynie od kształtu i wielkości otworu obliczone zostały przez Bethego.

Istnieją równania [2], określające momenty zastępczych dipoli pustego otworu. W zapisie rachunku symbolicznego są one następujące:

$$\vec{\hat{\Pi}}_0 = \mathfrak{M}_1 \hat{H}_{1x} \cdot \vec{\hat{i}}_x + \mathfrak{M}_2 \hat{H}_{1y} \cdot \vec{\hat{i}}_y, \quad (1a)$$

$$\vec{\hat{\Phi}}_0 = \epsilon_0 \cdot \gamma \cdot \vec{\hat{E}}_1 \cdot \vec{\hat{i}}_z, \quad (1b)$$

przy czym wartość chwilowa, np. $\vec{\hat{\Pi}}$ związana jest ze swoją postacią symboliczną $\vec{\hat{\Pi}}$ relacją

$$\vec{\hat{\Pi}} = \text{Re} \left(\vec{\hat{\Pi}} \cdot e^{j\omega t} \right).$$

Dla małego otworu okrągłego o średnicy d jest [2]

$$\mathfrak{M}_1 = \mathfrak{M}_2 = \mathfrak{M} = \frac{d^3}{6} \quad (2a)$$

$$\gamma = \frac{d^3}{12} \quad (2b)$$

Zakłada się [7], że równoważne ze względu na wzbudzenie pola momenty elektryczny $\vec{\hat{\Phi}}_f$ i magnetyczny $\vec{\hat{\Pi}}_f$ ferrytu wypełniającego okno określone

są równaniami podobnego kształtu do (1a) i (1b)

$$\hat{\vec{\Pi}}_f = \mathfrak{M} \cdot \hat{\vec{M}}_{1x} \cdot \vec{i}_x + \mathfrak{M} \hat{M}_{1y} \vec{i}_y, \quad (3a)$$

$$\hat{\Phi}_f = \gamma \cdot \hat{\vec{P}} \cdot \vec{i}_z. \quad (3b)$$

Równania (3a) i (3b) są oczywiste, jeśli zauważymy, że $\mathfrak{M}$ i γ mają wymiar objętości i reprezentują kształt i wielkość otworu sprzęgającego — a więc i ferrytu — oraz że ferryt ma własny moment magnetyczny wynikający z jego magnetyzacji M oraz elektryczny z polaryzacji P .

Uwzględniając znane związki

$$\hat{\vec{B}}_1 = \mu_0 \cdot (\hat{\vec{H}}_1 + \hat{\vec{M}}) \quad (4a)$$

oraz

$$\hat{\vec{D}}_1 = \epsilon_0 \cdot \epsilon \cdot \hat{\vec{E}}_1 \quad (4b)$$

momenty zapełnionego okna wyrażą się następująco:

$$\mu_0 \hat{\vec{\Pi}} = \mathfrak{M} \hat{B}_{1x} \vec{i}_x + \mathfrak{M} \hat{B}_{1y} \vec{i}_y, \quad (5a)$$

$$\hat{\Phi} = \gamma \hat{\vec{D}}_1 \cdot \vec{i}_z. \quad (5b)$$

W przypadku ogólnym (4a) jest zależnością tensorową. Składowe indukcji magnetycznej można więc zapisać następująco¹):

$$\hat{B}_{1x} = \mu_0 (\tilde{\mu} \hat{H}_{1x} - j \tilde{k} \hat{H}_{1y}), \quad (6)$$

$$\hat{B}_{1y} = \mu_0 (j \tilde{k} \cdot \hat{H}_{1x} + \tilde{\mu} \hat{H}_{1y}), \quad (7)$$

gdzie

$\tilde{\mu}$, $\tilde{k}$ — elementy tensora, tzw. zewnętrznej przenikalności magnetycznej [9].

Znając wartości momentów równoważnych otworowi, w falowodzie sprzężonym zostanie obliczone pole powstałe w wyniku promieniowania tych dipoli.

4. WYZNACZENIE WSPÓŁCZYNNIKA SPRZĘŻENIA

Zakładamy następującą postać składowych poprzecznych pola elektrycznego E_{2t} w falowodzie wtórnym, wzbudzonego przez pole E_{1t} istniejące w falowodzie pierwotnym:

¹ Należy podkreślić, że Stinson popełnia w swoim artykule [7] nieścisłość, nie rozróżniając tensora zewnętrznej przenikalności magnetycznej od przenikalności, która odnosi indukcję w ferrycie do istniejącego w nim pola magnetycznego.

dla fali rozchodzącej się w falowodzie 2 w kierunku x

$$\hat{E}_{2t}^- = \sum_n \hat{A}_n \cdot \hat{E}_{n1t} e^{-j\beta_n x}, \quad (8a)$$

dla fali zaś rozchodzącej się w kierunku $-x$

$$\hat{E}_{2t}^+ = \sum_n \hat{B}_n \cdot \hat{E}_{n1t} e^{j\beta_n x}, \quad (8b)$$

gdzie $\hat{A}_n, \hat{B}_n$ — zespolone współczynniki proporcjonalności dla fal dowolnego rodzaju n .

Posługując się równoważnymi dipolami (5a) i (5b) promieniowanie z otworu wyraża się za pomocą własnych funkcji falowodu, do którego promieniowanie następuje.

Korzystając z ortogonalnych właściwości pól w idealnym, jednorodnym falowodzie cylindrycznym Silver [6] wprowadza zależności, które w postaci symbolicznej są następujące:

$$\int_{S_0} (\hat{E}_2^- \times \hat{H}_{n1}^*) \cdot \vec{i}_z \cdot ds = -4 \hat{A}_n \hat{P}_n, \quad (9a)$$

$$\int_{S_0} (\hat{E}_2^+ \times \hat{H}_{n1}^*) \cdot \vec{i}_z \cdot ds = -4 \hat{B}_n \hat{P}_n; \quad (9b)$$

S_0 — pole powierzchni otworu sprzęgającego,

H_2, E_2 — wypadkowe pola wszystkich własnych rodzajów falowodu wtórnego (wzbudzone przez otwór),

H_{n1} — pole w falowodzie pierwotnym, które istniałoby w środku otworu, gdyby była tam pełna ścianka,

$\hat{P}_n$ — zespolona moc fali rodzaju n w falowodzie pierwotnym.

Średnia wartość P_n mocy chwilowej, dla rozpatrywanego przebiegu kształtu $e^{j\omega t}$, wiąże się z mocą zespoloną $\hat{P}_n$ następująco:

$$P_n = \text{Re } \hat{P}_n.$$

Wynika to wprost ze znanej¹⁾ analogicznej zależności dla wektora Poyntinga

$$\vec{S}_n = \text{Re } \hat{S}_n,$$

przy czym zespolony wektor Poyntinga $\hat{S}_n$ fali rodzaju n definiuje się jako

$$\hat{S}_n = \frac{1}{2} \hat{E}_n \times \hat{H}_n^*,$$

gdzie E_n i H_n są amplitudami.

¹ Zob. np. J. A. Stratton: „Elektromagnetic Theory”, Mc Graw-Hill Book Co., 1941, Sec. 2.20.

Wyrażenia (9a) i (9b) zostały obliczone przez Bethego dla pustych otworów. Dla otworu zapełnionego ferrytem podane zostały przez Stinsona [7]. W postaci symbolicznej, po uwzględnieniu (5a) i (5b), przedstawiają się one następująco:

$$\begin{aligned} 4 \hat{P}_n \hat{A}_n &= -j \omega (\mp \mathfrak{M} \cdot \hat{B}_{1y} \hat{H}_{2y}^* + \mathfrak{M} \hat{B}_{1x} \hat{H}_{2x}^* + \gamma \hat{D}_{1z} \hat{E}_{2z}^*, \\ 4 \hat{P}_n \hat{B}_n & \end{aligned} \quad (10)$$

gdzie — odpowiada A_n , + odpowiada B_n .

W wyrażeniu (10) E_1 i H_1 oraz E_2 i H_2 są polami własnych rodzajów, fal o jednostkowych emplitudach, falowodów pierwotnego i sprzężonego, które istniałyby w środku otworu, gdyby tego otworu nie było.

Podstawiając (6) i (7) do (10) otrzymujemy

$$\begin{aligned} 4 \hat{P}_n \hat{A}_n &= -j \omega [\mp \mathfrak{M} \mu_0 (j \tilde{k} \hat{H}_{1x} \hat{H}_{2y}^* + \tilde{\mu} \hat{H}_{1y} \hat{H}_{2x}^*) + \\ 4 \hat{P}_n \hat{B}_n &= -j \omega [\mp \mathfrak{M} \mu_0 (\tilde{\mu} \hat{H}_{1x} \hat{H}_{2x}^* - j \tilde{k} \hat{H}_{1y} \hat{H}_{2z}^*) + \gamma \hat{D}_{1z} \hat{E}_{1z}^*]. \end{aligned} \quad (11)$$

W dalszej części rozpatrzone będzie przypadek sprzężenia dwu jednakowych falowodów prostokątnych, o wymiarach poprzecznych a i b , w których propagują fale w rodzaju podstawowym H_{10} . Jest więc $H_{1x} = H_{2y}$ oraz $H_{2x} = H_{1y}$.

Uwzględniając, że składniki będące iloczynami składowych poprzecznych przez podłużne nie dają przyczynku do mocy przenoszonej przez falowód oraz zmieniając indeks n na aktualny 10 (11) uprości się do postaci

$$\begin{aligned} 4 \hat{P}_{10} \hat{A}_{10} &= -j \omega [j \mu_0 \tilde{k} \mathfrak{M} (\mp H_{1x}^2 - H_{1y}^2) + \varepsilon_0 \cdot \varepsilon \cdot \gamma E_{1z}^2]. \\ 4 \hat{P}_{10} \hat{B}_{10} & \end{aligned} \quad (12)$$

Przyjęto dalej, że falowód wtórny będzie nieskończenie długi, zaś pierwotny — zwarty w odległości l .

Rozkłady pól, dla kierunku propagacji $\pm y$ są następujące:

$$\begin{aligned} \hat{H}_x &= \pm j \hat{H}_0 \sin \frac{\pi}{a} x \cdot e^{\pm j \frac{2\pi}{\lambda_f} y}, \\ \hat{H}_y &= \frac{\lambda_f}{\lambda_{gr}} \hat{H}_0 \cos \frac{\pi}{a} x \cdot e^{\pm j \frac{2\pi}{\lambda_f} y}, \\ \hat{E}_z &= j z_f \hat{H}_0 \cdot \sin \frac{\pi}{a} x \cdot e^{\pm j \frac{2\pi}{\lambda_f} y}, \end{aligned} \quad (13)$$

gdzie:

$$z_f = \sqrt{\frac{\mu_0}{\varepsilon_0}} \cdot \frac{1}{\sqrt{1 - \left(\frac{\lambda_0}{\lambda_{gr}}\right)^2}},$$

a w falowodzie zwartym w odległości l :

$$\begin{aligned}\hat{H}_x &= j 2 \hat{H}_0 \sin \frac{\pi}{a} x \cdot \cos \frac{2\pi}{\lambda_f} (y-l), \\ \hat{H}_y &= 2 \cdot \frac{\lambda_f}{\lambda_{gr}} \cdot \hat{H}_0 \cos \frac{\pi}{a} x \cdot \sin \frac{2\pi}{\lambda_f} (y-l), \\ \hat{E}_z &= 2 z_f \hat{H}_0 \sin \frac{\pi}{a} x \cdot \sin \frac{2\pi}{\lambda_f} (y-l).\end{aligned}\quad (14)$$

W celu wyeliminowania sprzężenia elektrycznego odległość zwarcia l dobiera się tak, ażeby elektryczne pole w środku otworu, tj. dla $y = 0$ było zerowe. Oczywiście będzie to dla

$$l = \frac{m \cdot \lambda_f}{2},$$

gdzie m — liczba całkowita.

Wprowadzając przy tych założeniach (14) do (12) otrzymuje się

$$4 \hat{A}_{10} \hat{P}_{10} = -4 \hat{B}_{10} \hat{P}_{10} = 4 \omega \cdot \mu_0 \cdot \mathfrak{M} \cdot \tilde{k} \cdot H_0^2. \quad (15)$$

Jak nietrudno obliczyć, korzystając z (14), zespolona moc $\hat{P}_{10}$ fali padającej na otwór w falowodzie pierwotnym wynosi

$$\hat{P}_{10} = a \cdot b \cdot z_f \cdot H_0^2. \quad (16)$$

Podstawiając (2a) i (16) do (15) oraz uwzględniając skończoną grubość ścianki sprzęgającej [5], po wykonaniu prostych przekształceń otrzymuje się ostatecznie¹⁾

$$\hat{A}_{10} = -\hat{B}_{10} = \frac{\pi d^3}{3 a \cdot b \cdot \lambda_f} F_H \tilde{k}, \quad (17)$$

gdzie [5]

$$F_H = e^{-2\pi \sqrt{\frac{1}{(1,71 d)^2} - \frac{1}{\lambda_0^2}} \cdot t_0}. \quad (18)$$

W wyrażeniu (18) pominięty jest wpływ właściwości ferrytu na długość fali granicznej otworu.

Sprzężeniem C rozpatrywanego sprzęgacza, w odróżnieniu od powszechnie przyjętej definicji [5], nazwiemy stosunek mocy P_2 w jednym z ramion falowodu sprzężonego do mocy P_1 doprowadzonej do falowodu głównego. Pod pojęciem „moc doprowadzona” należy rozumieć moc fali bieżącej dostarczaną przez generator, podczas gdy fala odbita zostaje wytłumiona. Układowo realizuje się to przez zastosowanie dopasowanego izolatora ferrytowego między generatorem a sprzęgaczem. Wyrażając C

¹⁾ Stinson [7] popełnia tu błąd pisząc, że wielkość $\hat{A}_{10}$, oznaczona przez niego A_n , jest dwukrotnie większa od podanej w wyrażeniu (17). Przyjmuje on bowiem, że A_n zależy od amplitudy pola w falowodzie głównym, co jest sprzeczne z (9a) i (9b).

w decybelach jest

$$C = 10 \log \frac{P_2}{P_1}. \quad (19)$$

Moc w falowodzie sprzężonym pochodzi od aktualnego pola w falowodzie głównym. Równa się więc

$$P_2 = \frac{ab}{4} z_f (2 H_0 \cdot A_{10})^2,$$

zaś moc doprowadzona, zgodnie z wyżej podaną definicją równa jest

$$P_1 = \frac{ab}{4} z_f H_0^2.$$

Sprężenie jest więc równe

$$C = 20 \log 2 A_{10}. \quad (20)$$

Wprowadzając (17) do (20) otrzymuje się¹⁾

$$C = 20 \log \frac{2 \cdot \pi \cdot d^3}{3 ab \lambda_f} \cdot F_H \cdot |\tilde{k}|. \quad (21)$$

Oznaczając

$$C_0 = 20 \log \frac{2 \cdot \pi \cdot d^3}{3 ab \lambda_f} \cdot F_H \quad (22)$$

ostatecznie jest

$$C = C_0 + 20 \log |\tilde{k}|. \quad (23)$$

Tak wyraża się sprzężenie sprzęgacza krzyżowego, którego falowód sprzężony jest obustronnie dopasowany, falowód główny jest zwarty w takiej odległości, że w środku otworu sprzęgającego pole elektryczne jest zerowe, a otwór sprzęgający wypełniony jest ferrytem. Jest ono proporcjonalne do $|\tilde{k}|$; w szczególności osiąga maksimum dla rezonansu ferromagnetycznego, a dla $|\tilde{k}| = 0$ jest zerowe ($-\infty$ dB). Jest rzeczą oczywistą [por. (17)], że rozpatrywany sprzęgacz, według powszechnie przyjętej definicji kierunkowości [5], jest bezkierunkowy.

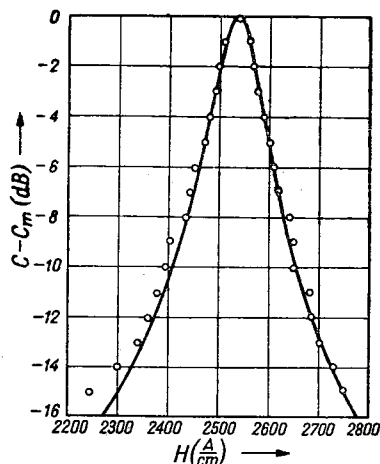
5. ZAKOŃCZENIE

W przypadku tylko częściowego wypełnienia otworu ferrytem zasadnicza postać wyrażenia określającego sprzężenie nie zmieni się. Potwierdza to rys. 2. Linia ciągła przedstawia zależność (23) obliczoną teoretycznie²⁾, naniesione punkty są otrzymane doświadczalnie. Pomiary wykonano na częstotliwości $f = 9285$ MHz w sprzęgaczu krzyżowym, z otworem o średnicy $d = 4$ mm, mającym $C_0 = -41.9$ dB. Kulista próbka

¹⁾ Identyczną postać wzoru (21) Stinson [7] otrzymał pisząc, że sprzężenie wynosi A^2 .

²⁾ Rozważania dotyczące kulistej próbki ferrytowej przedstawione były przez autora na Krajowej Konferencji Ferrytowej w roku 1959 i podane będą osobno.

ferrytowa miała średnicę 0,6 mm i wykonana była z materiału niklowo-cynkowego, wykonanego w Pracowni Materiałów Magnetycznych IPPT



Rys. 2. Porównanie teoretycznej linii rezonansowej $|\tilde{k}|$ z wynikami pomiaru. C_m — maksymalne sprzężenie występujące przy rezonansie próbki

PAN. Jego parametry są następujące: szerokość linii rezonansowej $H = 84$ A/cm, magnetyzacja nasycenia $M = 3940$ A/cm, natężenie pola rezonansowego $H_r = 2540$ A/cm, a efektywny spektroskopowy współczynnik rozszczepienia $g_{ef} = 2.079$.

Różnice między punktami doświadczalnymi a linią teoretyczną są na ogół mniejsze od 1% i zawierają się w dokładności pomiaru stałego pola magnetycznego.

Omawiany sprzęgacz był przeznaczony — jak podano we wstępie — do pomiarów linii rezonansowej ferrytów [4], [8] i pod tym kątem był rozpatrzony. Wykonuje się w nim [4] w Pracowni Materiałów Magnetycznych IPPT PAN pomiary szerokości linii rezonansowej i efektywnego spektroskopowego współczynnika rozszczepienia ferrytów.

Krzyżowy sprzęgacz z otworem zapełnionym ferrytem może być oczywiście zastosowany jako element z magnetycznie kontrolowanym sprzężeniem w szerokich granicach i to zarówno jako element szerokopasmowy, jak też jako element wybitnie selektywny. Szerokość pasma zależy od użytego ferrytu. Może zawierać się w granicach od setnych części procenta do kilkunastu procent. W tym zastosowaniu główny falowód oczywiście nie może być zwarty. W związku z tym sprzężenie sprzęgacza będzie mieszane, elektryczno-magnetyczne, i sprzęgacz będzie kierunkowy. Wartość sprzężenia może być obliczona z wyrażenia (12) i (13) w podobny sposób do przedstawionego w niniejszym artykule.

Autor składa wyrazy podziękowania prof. drowi inż. A. Smolińskiemu za przejrzanie tej pracy i cenne uwagi.

*Instytut Podstawowych Problemów Techniki PAN
Pracownia Materiałów Magnetycznych*

WYKAZ LITERATURY

1. Berk A. D., Strumwasser E.: Ferrite Directional Couplers. — Proc. IRE, vol. 44, October 1956, str. 1439—1446.
2. Bethe H. A.: Theory of Diffraction by Small Holes. — Phys. Rev., vol. 66, No. 7—8, October 1944, str. 163—182.
3. Damon R. W.: Magnetically Controlled Microwave Directional Coupler. — J. Appl. Phys., vol. 26, October 1955, str. 1281—1283.
4. Krzycki Z.: Układ do pomiaru szerokości linii rezonansowej ferrytów w pasmie X. — Zeszyty Problemowe Nauki Polskiej Nr 20, 1960.
5. Montgomery D. C.: Technique of Microwave Measurements. — Mc. Graw Hill Book Co., New York 1947, str. 862.
6. Silver S.: Microwave Antenna Theory and Design.—Mc. Graw Hill Book Co., New York 1947, § 9, 10.
7. Stinson B. C.: Ferrite Line Width Measurements in a Cross-Guide Coupler.—IRE Trans., vol. MTT-6, October 1958, str. 446—450.
8. Stinson D. C.: Coupling Through an Aperture Containing an Anisotropic Ferrite. — IRE Trans., vol. MTT-5, July 1957, str. 184—191.
9. Waldron R. A.: Theory of the Measurement of the Elements of the Permeability Tensor of a Ferrite by Means of Resonant Cavity. — PIEEE, Part B, Supplement 1956, str. 307.

3. КРЖИЦКИ

ПЕРЕКРЕСТНЫЙ ОТВЕТВИТЕЛЬ С ОТВЕРСТИЕМ ЗАПОЛНЕННЫМ ФЕРРИТОМ

Резюме

Определена связь перекрёстного ответвителя имеющего отверстие наполненное ферритом. Рассуждения произведены на основании теории ответвителя с малым отверстием Бетэ (2), (5) и Стинсона (7). Представлено сравнение теоретических расчётов с результатами измерения при заполненном частично отверстием. Показаны практические примечения рассматриваемого ответвителя — в основном для измерения ширины линии ферромагнитного резонанса и эффективного спектроскопического коэффициента расщепления ферритов.

Z. KRZYCKI

CROSS-GUIDE COUPLER WITH AN APERTURE CONTAINING A FERRITE

Summary

The coupling of a cross-guide coupler with an aperture containing a ferrite has been marked out. The considerations have been carried out on the basis of the Bethe's [2], [5] and Stinson's [7] theory of a coupler with a small aperture. The comparison of theoretical computations to the results of the measurement for the aperture filled partially has been presented. Practical uses of the considered coupler—mainly to measure the line width of ferromagnetic resonance and effective spectroscopical coefficient of ferrite cleavage have been pointed out.

PAWEŁ SZULKIN

Wariacyjne metody obliczania nieciągłości w kablach koaksjalnych

Rękopis dostarczono 11. 11. 1960

W pracy zastosowano metodę wariacyjną do rozwiązywania problemu nieciągłości płaskich w kablach koaksjalnych. Wyprowadza się równanie całkowe dla admitancji, spełniające warunki ciągłości składowych E_r i H_φ w aper-
turze i zaniku E_r w pozostałej części płaszczyzny nieciągłości. Powyższe równanie całkowe jest stacjonarne względem małych zmian E_r , daje minimalną wartość Y , jeżeli stosujemy rzeczywisty rozkład pola. Następnie rozwijamy

E_r w pełny zbiór funkcji walcowych z członem podstawowym $\frac{1}{r}$. Współczynniki rozwinięcia wyznacza się dzięki wykorzystaniu własności stacjonarności admitancji.

Zakładając, że mod TEM jest jedynym modem swobodnie rozchodzącym się w kablu, potwierdza się możliwość ścisłej reprezentacji nieciągłości przez pojemność skupioną. Człon główny rozwinięcia pola $\left(\frac{1}{r}\right)$ prowadzi do pojemności zerowego rzędu, która jest zawsze większa od wartości ścisłej. Należy sądzić, że dzięki własnościom stacjonarności rozwiązanie zerowego rzędu będzie wystarczająco dokładne. Wyprowadza się również wzory dla poprawek umożliwiające uzyskanie rozwiązań wyższego rzędu.

Na zakończenie omówiono szereg przypadków szczególnych nieciągłości w kablach.

Ogólnie przyjęto uwzględniać wpływ nieciągłości w kablu koaksjalnym przy pomocy adekwatnej pojemności skupionej, umieszczonej w płaszczyźnie nieciągłości [1]. W pracy tej ustalimy warunki ważności takiego potraktowania problemu i podamy sposób obliczania pojemności równoważnej, oparty na metodzie wariacyjnej [2]. W naszej analizie uwzględnimy szereg przypadków, a wśród nich i najbardziej ogólny — jednoczesnej zmiany średnic obu przewodników kabla.

Zanim przejdziemy do wyznaczenia rozkładu pola w pobliżu nieciągłości kabla koaksjalnego, rozpatrzmy rozwiązania równań Maxwella w układzie współrzędnych walcowych. Zakładając, że ośrodek dielektryczny ($\sigma = 0$) jest izotropowy i jednorodny oraz nie zawiera ładunków i prądów, możemy równania te napisać w postaci

$$\nabla_x \bar{E} + j\omega\mu\bar{H} = 0, \quad (1)$$

$$\nabla_x \bar{H} - j\omega\varepsilon\bar{E} = 0. \quad (2)$$

Tworząc $\nabla x (\nabla x \bar{E})$ i $\nabla x (\nabla x \bar{H})$ i rozwiązując równanie (1) względem $\bar{E}$, zaś równanie (2) względem $\bar{H}$, otrzymujemy

$$\nabla^2 \bar{E} + \omega^2 \varepsilon \mu \bar{E} = 0, \quad (3)$$

$$\nabla^2 \bar{H} + \omega^2 \varepsilon \mu \bar{H} = 0. \quad (4)$$

Rozpatrując jedynie składową z pola mamy zamiast równań wektorowych (3) i (4) zwykle równania falowe skalarne, które można rozwiązać bez trudu. Znając następnie składowe z, obliczamy pozostałe składowe pola przy pomocy równań (1) i (2). Rozwiązania równań (3) i (4) można podzielić na trzy ogólne rodzaje, a mianowicie: 1) mod TEM, który nie posiada składowych pola w kierunku z, 2) mod TM, dla którego wszystkie składowe pola elektromagnetycznego wyprowadza się z E_z , i wreszcie 3) mod TE, dla którego wszystkie składowe wyprowadza się z H_z .

Równania dla składowych z są [3]:

$$\left. \begin{array}{l} \text{(mody TE)} \quad H_z \\ \text{(mody TM)} \quad E_z \end{array} \right\} = R \Phi e^{j\omega t \mp \gamma z}, \quad (5)$$

gdzie

$$\Phi = A \cos n \varphi + B \sin n \varphi, \quad (6)$$

$$R = C J_n(\beta_q r) + D N_n(\beta_q r). \quad (7)$$

A, B, C i D oznaczają dowolne stałe, wielkość zaś β_q jest określona równaniem

$$\beta_q^2 = \omega^2 \varepsilon \mu \gamma_q^2. \quad (8)$$

Równania (1) i (2) można teraz rozwinąć w sposób następujący:

$$\begin{aligned} E_r &= -\frac{1}{\beta_q^2} \left(\gamma \frac{\partial E_z}{\partial r} + \frac{j\omega\mu}{r} \frac{\partial H_z}{\partial \varphi} \right), \\ E_\varphi &= \frac{1}{\beta_q^2} \left(-\frac{\gamma}{r} \frac{\partial E_z}{\partial \varphi} + j\omega\mu \frac{\partial H_z}{\partial r} \right), \\ H_r &= \frac{1}{\beta_q^2} \left(\frac{j\omega\varepsilon}{r} \frac{\partial E_z}{\partial \varphi} - \gamma \frac{\partial H_z}{\partial r} \right), \\ H_\varphi &= -\frac{1}{\beta_q^2} \left(j\omega\varepsilon \frac{\partial E_z}{\partial r} - \frac{\gamma}{r} \frac{\partial H_z}{\partial \varphi} \right). \end{aligned} \quad (9)$$

Warunki brzegowe wymagają, aby styczne pole elektryczne zniknęło na powierzchni przewodnika. Dla modów TM sprowadza się to do warunku $E_z = 0$, gdy $r = a$ i $r = b$, gdzie a i b promienie obu przewodników kabla koaksjalnego. A więc dla modów TM musi być zawsze spełniony warunek

$$\begin{vmatrix} J_n(\beta_q a) & N_n(\beta_q a) \\ J_n(\beta_q b) & N_n(\beta_q b) \end{vmatrix} = 0,$$

czyli

$$J_n(\beta_q a) N_n(\beta_q b) - J_n(\beta_q b) N_n(\beta_q a) = 0, \quad (10)$$

gdzie β_q — pierwiastek równania (10) ($q = 1, 2, \dots$).

Dla modów TE pole styczne elektryczne znika również na powierzchni przewodników. Z równań (9) wynika, że w tym wypadku wymaganie to sprowadza się do zaniku normalnej pochodnej H_z na brzegach. Ostatecznie daje to następujący warunek dla fal TE

$$J'_n(\beta_q a) N'_n(\beta_q b) - J'_n(\beta_q b) N'_n(\beta_q a) = 0, \quad (11)$$

gdzie $J'_n = \frac{\partial J_n}{\partial r}$ itp.

W specjalnym wypadku, gdy $\beta_q = 0$, γ staje się urojone

$$\gamma = \sqrt{-\omega^2 \varepsilon \mu} = j \beta_0 \quad (12)$$

i równanie falowe ma jako rozwiązanie dobrze znany mod TEM, który jako jedyny rozchodzi się swobodnie w praktycznie stosowanych kablach koaksjalnych.

Badając różne struktury fizyczne nieciągłości (rys. 1—4) łatwo się przekonać, że mod TEM nie może sam spełnić wszystkich warunków brzegowych, ponieważ dla $z = 0$ dochodzi wymaganie $E_r = 0$. Należy więc przyjąć, że obok modu TEM powstają mody TM. Ze względu na istnienie symetrii kołowej można jednocześnie położyć $n = 0$ i ograniczyć się do rozpatrywania jedynie modów TM_{0q} , jeżeli modem propagacyjnym jest TEM.

Uwzględniając mody wyższych rzędów, poruszające się naprzód, możemy pola poprzeczne w obszarze (1) napisać w postaci

$$E_r^{(1)} = \frac{1}{r} \frac{V(z)}{\ln \frac{b_1}{a_1}} + \sum_{q=1}^{\infty} A_q^{(1)} Z_1(\beta_q^{(1)} r) e^{\gamma_q z}, \quad (13)$$

$z < 0,$

$$H_\varphi^{(1)} = \frac{I(z)}{2 \pi r} - \sum_{q=1}^{\infty} \frac{j \omega \varepsilon}{\gamma_q^{(1)}} A_q^{(1)} Z_1(\beta_q^{(1)} r) e^{\gamma_q z},$$

gdzie

$$\gamma_q^{(1)} = \beta_q^{(1)} \sqrt{1 - \left(\frac{\beta_0}{\beta_q^{(1)}}\right)^2} = \beta_0 \sqrt{\left(\frac{\beta_q^{(1)}}{\beta_0}\right)^2 - 1}, \quad (14)$$

$Z_1(\beta_q r)$ zaś oznacza następującą kombinację liniową funkcji walcowych

$$Z_1(\beta_q^{(1)} r) = N_0(\beta_q^{(1)} a_1) J_1(\beta_q^{(1)} r) - J_0(\beta_q^{(1)} a_1) N_1(\beta_q^{(1)} r). \quad (15)$$

Warto podkreślić, że ograniczenie się do fali poruszającej się naprzód implikuje, że rozpatrywana nieciągłość jest dostatecznie oddalona od ob-

ciężenia lub innych nieciągłości, aby mody wyższego rzędu były całkowicie stłumione. Wprowadźmy jeszcze liniową kombinację $Z_0(\beta_q^{(1)} r)$ określoną jako

$$Z_0(\beta_q^{(1)} r) = N_0(\beta_q^{(1)} a_1) J_0(\beta_q^{(1)} r) - J_0(\beta_q^{(1)} a_1) N_0(\beta_q^{(1)} r). \quad (16)$$

Przy wyprowadzeniu równań (13) korzystamy z identyczności [4]

$$Z'_0(x) = -Z_1(x) \quad (17)$$

i równania (10) dla modów TM, które w naszym wypadku przyjmuje postać

$$J_0(\beta_q^{(1)} a_1) N_0(\beta_q^{(1)} b_1) - J_0(\beta_q^{(1)} b_1) N_0(\beta_q^{(1)} a_1) = 0. \quad (18)$$

$V(z)$ i $I(z)$ w równaniach (13) reprezentują napięcia i prądy w ujęciu konwencjonalnej teorii linii długiej. Oznaczając przez Y admitancję w punkcie $z = 0$, możemy $V(z)$ i $I(z)$ wyrazić następująco (zakładając nieskończoną przewodność):

$$V(z) = V(\cos \beta_0 z - j Z_0^{(1)} Y \sin \beta_0 z), \quad z < 0, \quad (19)$$

$$I(z) = V \left(Y \cos \beta_0 z - j \frac{1}{Z_0^{(1)}} \sin \beta_0 z \right),$$

gdzie V jest napięciem w $z = 0$, zaś $Z_0^{(1)}$ — impedancją charakterystyczną odcinka (1) kabla, równą

$$Z_0^{(1)} = \frac{1}{2\pi} \sqrt{\frac{\mu}{\varepsilon}} \ln \frac{b_1}{a_1}. \quad (20)$$

Admitancja w $z = 0$ zależy od nieciągłości i admitancji wejściowej odcinka (2) kabla. Ta ostatnia nas nie interesuje, gdyż jej obliczenie jest łatwe na podstawie klasycznej teorii linii długiej; możemy więc po prostu obrać tę admitancję wejściową równą zero. W tych warunkach Y reprezentuje tylko admitancję nieciągłości.

Podstawiając równania (19) do równań (13) otrzymamy ostateczne wyrażenia pola w obszarze (1)

$$E_r^{(1)} = \frac{V(\cos \beta_0 z - j Z_0 Y \sin \beta_0 z)}{r \ln \frac{b_1}{a_1}} + \sum_{q=1}^{\infty} A_q^{(1)} Z_1(\beta_q^{(1)} r) e^{\gamma_q^{(1)} z} \quad (21)$$

$$H_\varphi^{(1)} = \frac{V \left(Y \cos \beta_0 z - j \frac{1}{Z_0^{(1)}} \sin \beta_0 z \right)}{2\pi r} - j\omega\varepsilon \sum_{q=1}^{\infty} \frac{A_q^{(1)}}{\gamma_q^{(1)}} Z_1(\beta_q^{(1)} r) e^{\gamma_q^{(1)} z}.$$

Podobnie dla obszaru (2):

$$E_r^{(2)} = \frac{V \cos \beta_0 z}{r \ln \frac{b_2}{a_2}} + \sum_{q=1}^{\infty} A_q^{(2)} Z_1(\beta_q^{(2)} r) e^{-\gamma_q^{(2)} z}, \quad (22)$$

$$H_\varphi^{(2)} = j \frac{V \sin \beta_0 z}{Z_0^{(2)} 2 \pi r} + j \omega \varepsilon \sum_{q=1}^{\infty} \frac{A_q^{(2)}}{\gamma_q^{(2)}} Z_1(\beta_q^{(2)} r) e^{-\gamma_q^{(2)} z}.$$

Współczynnik $A_q^{(1)}$ wyznaczamy kładąc $z = 0$, mnożąc równanie (21) przez $r Z_1(\beta_n^{(1)} r)$ i całkując względem r w przedziale (a_1, b_1) :

$$\begin{aligned} \int_{a_1}^{b_1} r' E_{r'}^{(1)} Z_1(\beta_n^{(1)} r') dr' &= \frac{V}{\ln \frac{b_1}{a_1}} \int_{a_1}^{b_1} Z_1(\beta_n^{(1)} r') dr' + \\ &+ \sum_{q=1}^{\infty} A_q^{(1)} \int_{a_1}^{b_1} r' Z_1(\beta_n^{(1)} r') Z_1(\beta_q^{(1)} r') dr', \end{aligned} \quad (23)$$

gdzie r' oznacza punkt całkowania ($z = 0$).

Całki występujące w równaniu (23) obliczamy korzystając ze znanych wzorów [4]

$$\int r Z_1^2(\alpha r) dr = \frac{r^2}{2} [Z_1^2(\alpha r) - Z_0(\alpha r) Z_2(\alpha r)], \quad (24)$$

$$\int r Z_1(\alpha r) Z_1(\beta r) dr = \frac{\beta r Z_1(\alpha r) Z_0(\beta r) - \alpha r Z_0(\alpha r) Z_1(\beta r)}{\alpha^2 - \beta^2} \quad (25)$$

$$\alpha \neq \beta$$

$$\int Z_1(\alpha r) dr = -\frac{1}{\alpha} Z_0(\alpha r), \quad (26)$$

$$N_0(x) J_1(x) - J_0(x) N_1(x) = \frac{2}{\pi x}. \quad (27)$$

Łatwo przekonać się, że jeżeli $q \neq n$, prawa strona równania (23) zanika, natomiast dla $q = n$ otrzymujemy

$$A_q^{(1)} = \gamma_q^{(1)} G_q^{(1)} \int_{a_1}^{b_1} r' E_{r'}^{(1)} Z_1(\beta_q^{(1)} r') dr', \quad (28)$$

gdzie

$$G_q^{(1)} = \frac{\pi^2 (\beta_q^{(1)})^2}{2 \gamma_q^{(1)} \left[\frac{J_0^2(\beta_q^{(1)} a_1)}{J_0^2(\beta_q^{(1)} b_1)} - 1 \right]}. \quad (29)$$

Aby wyznaczyć V , całkujemy równanie (21) względem r' w przedziale (a_1, b_1) . Uwzględniając równanie (26) mamy dla $z = 0$

$$V = \int_{a_1}^{b_1} E_r^{(1)} dr' . \quad (30)$$

Jest to oczywiście zgodne z konwencjonalną definicją napięcia przy $z = 0$. Podobne wyrażenia dla $A_q^{(2)}$ i V otrzymamy zmieniając po prostu w równaniach (28) i (30) indeksy⁽¹⁾ na indeksy⁽²⁾. Pozornie może się wydawać, że powyższe wyniki prowadzą do pewnych wątpliwości, ponieważ napięcie jest określone wówczas przez

$$V = \int_{a_1}^{b_1} E_r^{(1)} dr' = \int_{a_2}^{b_2} E_r^{(2)} dr' , \quad (31)$$

$$E_r^{(1)} = E_r^{(2)} = E_r . \quad (32)$$

Oczywiście rezultaty są całkowicie równoważne, ponieważ poprzeczne pole E znika dla $z = 0$ z wyjątkiem apertury, gdzie jest ono ciągłe.

Ciągłość pola magnetycznego wymaga, aby w aperturze $H_\varphi^{(1)} = H_\varphi^{(2)}$. W tych warunkach przyrównanie odpowiednich równań (21) i (22) przy $z = 0$ i wykorzystanie równań (28) i (30) oraz ciągłości pola E_r daje

$$\begin{aligned} \frac{Y \int E_r dr'}{2\pi r} - j\omega\varepsilon \sum_{q=1}^{\infty} G_q^{(1)} Z_1(\beta_q^{(1)} r) \int r' E_r Z_1(\beta_q^{(1)} r') dr' = \\ = j\omega\varepsilon \sum_{q=1}^{\infty} G_q^{(2)} Z_1(\beta_q^{(2)} r) \int r' E_r Z_1(\beta_q^{(2)} r') dr' . \end{aligned} \quad (33)$$

Całkowanie wystarczy ograniczyć do pola apertury, gdyż przy $z = 0$ E_r poza tym polem równa się zeru. Mnożąc to równanie przez rE_r i całkując względem r w aperturze otrzymujemy ostatecznie

$$\begin{aligned} Y = \frac{j2\pi\omega\varepsilon}{\left(\int E_r dr\right)^2} \left\{ \sum_{q=1}^{\infty} G_q^{(1)} \left[\int r E_r Z_1(\beta_q^{(1)} r) dr \right]^2 + \right. \\ \left. + \sum_{q=1}^{\infty} G_q^{(2)} \left[\int r E_r Z_1(\beta_q^{(2)} r) dr \right]^2 \right\} = \\ = \frac{j2\pi\omega\varepsilon}{\left[\int E_r dr\right]^2} \sum_{p=1}^2 \sum_{q=1}^{\infty} G_q^{(p)} \left[\int r E_r Z_1(\beta_q^{(p)} r) dr \right]^2 . \end{aligned} \quad (34)$$

Wykażemy teraz, że to wyrażenie dla Y jest stacjonarne względem małych zmian pola apertury. Z równania (34) wynika bowiem

$$\begin{aligned} \delta Y \left[\int E_r dr \right]^2 + 2Y \int E_r dr' \int \delta E_r dr = \\ = 2 \int \delta E_r dr \left[j2\pi\omega\varepsilon \sum_{p=1}^2 \sum_{q=1}^{\infty} G_q^{(p)} r Z_1(\beta_q^{(p)} r) \int r' E_r Z_1(\beta_q^{(p)} r') dr' \right] , \end{aligned}$$

lecz zgodnie z równaniem (33) wyrażenie w nawiasach równa się po prostu $Y \int E_r dr'$. A więc

$$\delta Y \left[\int E_r dr \right]^2 + 2Y \int E_r dr' \int \delta E_r dr = 2Y \int E_r dr' \int \delta E_r dr,$$

czyli

$$\delta Y = 0$$

i nasze równanie całkowe (34) jest stacjonarne względem wariacji pierwszego rzędu w odniesieniu do ścisłego rozkładu pola.

Można również wykazać, że ścisłe wyrażenie pola apertury E_r prowadzi do bezwzględnego minimum Y . Załóżmy więc, że E_r^0 i Y^0 oznaczają ścisły rozkład pola i odpowiadającą temu admitancję. Niech E_r będzie dowolnym polem próbnym, Y zaś wynikającą stąd admitancją. Ponieważ $G_q^{(p)}$ jest zawsze skończoną wielkością dodatnią, możemy napisać następującą nierówność

$$j 2 \pi \omega \varepsilon \sum_{p=1}^2 \sum_{q=1}^{\infty} G_q^{(p)} \left\{ \int r' Z_1(\beta_q^{(p)} r') \left[\frac{E_r}{\int E_r dr} - \frac{E_r^0}{\int E_r^0 dr} \right] dr' \right\}^2 \geq 0$$

lub

$$j 2 \pi \omega \varepsilon \sum_{p=1}^2 \sum_{q=1}^{\infty} G_q^{(p)} \left\{ \left[\frac{\int E_r r Z_1(\beta_q^{(p)} r) dr}{\int E_r dr} \right]^2 + \left[\frac{\int E_r^0 r Z_1(\beta_q^{(p)} r) dr}{\int E_r^0 dr} \right]^2 + \right. \\ \left. - \frac{2 \int E_r^0 r' Z_1(\beta_q^{(p)} r') dr' \int E_r r Z_1(\beta_q^{(p)} r) dr}{\int E_r^0 dr \int E_r dr} \right\} \geq 0.$$

Zgodnie z równaniem (34) pierwsze dwa człony to Y i Y^0 . A więc

$$Y + Y^0 - j 2 \pi \omega \varepsilon \sum_{p=1}^2 \sum_{q=1}^{\infty} \frac{2 \int E_r dr \left[\frac{G_q^{(p)} r Z_1(\beta_q^{(p)} r) \int E_r^0 r' Z_1(\beta_q^{(p)} r') dr'}{\int E_r^0 dr} \right]}{\int E_r dr} \geq 0.$$

Lecz wyrażenie w nawiasie jest na podstawie równania (33) równe $\frac{Y^0}{j 2 \pi \omega \varepsilon}$, co po podstawieniu daje ostatecznie

$$Y + Y^0 - 2 Y^0 \geq 0$$

i

$$Y \geq Y^0.$$

Widzimy więc, że dowolna funkcja próbna dla E_r daje wartość Y większą niż ścisły rozkład pola. Warto zaznaczyć, że podobne wyrażenie dla $\frac{1}{Y}$ można otrzymać z H_r i w ten sposób otrzymać minimalną i maksymalną

wartość Y dla danego założonego rozkładu pola. Jest oczywiste, że stosując rzeczywisty rozkład pola obie metody dają tę samą wartość Y .

Dla dalszej analizy jest celowe rozwinięcie pola apertury w zbiór pełny funkcji ortogonalnych. Aczkolwiek w zasadzie rozwinięcie takie jest dowolne, w naszym przypadku będzie logiczne obrać w tym celu funkcje walcowe J i N .

Rozpatrzmy więc następujący zbiór

$$E_r = \frac{C_0}{r} + \sum_{n=1}^{\infty} C_n S_1(\beta_n r), \quad (35)$$

gdzie $S_1(\beta_n r)$ jest liniową kombinacją o postaci

$$S_1(\beta_n r) = a J_1(\beta_n r) + b N_1(\beta_n r), \quad (36)$$

zaś a i b dowolne chwilowo nieokreślone stałe. Wprowadzając równanie (35) do równania (36) otrzymujemy

$$Y = j 2 \pi \omega \varepsilon \sum_{p=1}^2 \sum_{q=1}^{\infty} G_q^{(p)} \left\{ \frac{C_0 \int Z_1(\beta_q^{(p)} r) dr + \sum_{n=1}^{\infty} C_n \int r Z_1(\beta_q^{(p)} r) S_1(\beta_n r) dr}{C_0 \int \frac{dr}{r} + \sum_{n=1}^{\infty} C_n \int S_1(\beta_n r) dr} \right\}^2 \quad (37)$$

Ponieważ Y jest niezależne od skali wymiarów pola apertury, wygodnie będzie położyć $C_0 = 1$. Dla zwięzłości napiszemy Y w postaci

$$Y = \frac{j 2 \pi \omega \varepsilon}{\left(\ln \frac{r_2}{r_1} \right)^2} \left\{ \psi + 2 \sum_{n=1}^{\infty} C_n \zeta_n + \sum_{n=1}^{\infty} \sum_{m=1}^{\infty} C_n C_m \xi_{nm} \right\}, \quad (38)$$

gdzie

$$\begin{aligned} \psi &= \sum_{q=1}^{\infty} \sum_{p=1}^2 [P_q^{(p)}]^2, \\ \zeta_n &= \sum_{q=1}^{\infty} \sum_{p=1}^2 P_q^{(p)} P_{qn}^{(p)}, \\ \xi_{nm} &= \sum_{q=1}^{\infty} \sum_{p=1}^2 P_{qn}^{(p)} P_{qm}^{(p)} \end{aligned}$$

$$P_q^{(p)} = \sqrt{G_q^{(p)}} \int_{r_1}^{r_2} Z_1(\beta_q^{(p)} r) dr,$$

$$P_{qn}^{(p)} = \sqrt{G_q^{(p)}} \int_{r_1}^{r_2} r Z_1(\beta_q^{(p)} r) S_1(\beta_n r) dr.$$

Równanie (38) uwzględnia tę okoliczność, że można zawsze tak dobrać $S_1(\beta_n r)$, aby $\int S_1(\beta_n r) dr$ zniknęła na aperturze. Warunek ten będzie spełniony na pewno, jeżeli stałe a i b obrane będą tak, aby

$$S_1(\beta_n r) = N_0(\beta_n r_1) J_1(\beta_n r) - J_0(\beta_n r_1) N_1(\beta_n r), \quad (39)$$

zaś β_n jest określone jako n . pierwiastek równania

$$J_0(\beta_n r_1) N_0(\beta_n r_2) - N_0(\beta_n r_1) J_0(\beta_n r_2) = 0, \quad (40)$$

gdzie r_1 i r_2 są wewnętrznym i zewnętrznym promieniami apertury.

Ponieważ, jak to wykazaliśmy wyżej, Y jest stacjonarne względem małych zmian pola w aperturze, możemy wyznaczyć współczynniki rozwinięcia zakładając $\frac{\partial Y}{\partial C_n} = 0$, co przy wykorzystaniu równania (38) daje

$$\zeta_n + \sum_{m=1}^{\infty} C_m \xi_{nm} = 0. \quad (41)$$

Mnożąc równanie (41) przez C_n i sumując względem n , otrzymujemy

$$\sum_{n=1}^{\infty} C_n \zeta_n + \sum_{m=1}^{\infty} \sum_{n=1}^{\infty} C_n C_m \xi_{nm} = 0, \quad (42)$$

co po podstawieniu do równania (38) prowadzi ostatecznie do

$$Y = \frac{j 2 \pi \omega \varepsilon}{\left(\ln \frac{r_2}{r_1}\right)^2} \left(\psi + \sum_{n=1}^{\infty} C_n \zeta_n \right). \quad (43)$$

W ten sposób matematycznie nasz problem można uważać za rozwiązany, gdyż obliczając C_n z równania (41) i podstawiając do równania (43) mamy ścisłe wyrażenie dla admitancji nieciągłości. Jeżeli mod TEM jest jedynym modem swobodnie rozchodzącym się w kablu, człon w nawiasie w równaniu (43) jest zawsze dodatni i rzeczywisty i wobec tego Y reprezentuje susceptancję skupionej pojemności.

Człon ψ można rozpatrywać jako rozwiązanie zerowego rzędu, stanowiące w większości wypadków zupełnie zadowalającą aproksymację. Stosowanie tej aproksymacji jest właściwie równoważne założeniu $E_r \sim \frac{1}{r}$ w aperturze.

Łatwo się przekonać, że człon zerowego rzędu jest dodatni, podczas gdy pozostałe człony korekcyjne są ujemne. Stąd wynika, że aproksymacja zerowego rzędu jest zawsze większa od wartości rzeczywistej Y . W miarę tego jak zwiększamy ilość członów korekcyjnych, coraz bardziej zbliżamy się do wartości rzeczywistej Y jako granicy dolnej. Ten wynik jest zresztą oczywisty z wariacyjnego punktu widzenia.

Otrzymane wyżej ogólne rozwiązanie, można również wyrazić w postaci macierzowej, co nieraz ułatwia obliczenia. W tym celu wprowadzamy następujące macierze

$$[\xi] = \begin{bmatrix} \xi_{11} & \xi_{12} & \dots & \xi_{1\infty} \\ \xi_{21} & \xi_{22} & \dots & \xi_{2\infty} \\ \dots & \dots & \dots & \dots \\ \xi_{\infty 1} & \xi_{\infty 2} & \dots & \xi_{\infty \infty} \end{bmatrix} \quad (44)$$

$$[\zeta] = \begin{bmatrix} \zeta_1 \\ \zeta_2 \\ \vdots \\ \zeta_{\infty} \end{bmatrix} \quad (45)$$

$$[C] = \begin{bmatrix} C_1 \\ C_2 \\ \vdots \\ C_{\infty} \end{bmatrix} \quad (46)$$

Pozwala to nadać równaniu (41) postać

$$\begin{aligned} [\zeta] + [\xi][C] &= 0, \\ [C] &= -[\xi]^{-1}[\zeta]. \end{aligned} \quad (47)$$

Zamiast równania (43) mamy wówczas

$$Y = \frac{j 2 \pi \omega \varepsilon}{\left(\ln \frac{r_2}{r_1}\right)^2} \{\bar{\psi} + [\zeta]_t [C]\}, \quad (48)$$

gdzie $[\zeta]_t$ jest macierzą transponowaną. Podstawiając dalej równanie (47) do równania (48) otrzymujemy ostatecznie

$$Y = \frac{j 2 \pi \omega \varepsilon}{\left(\ln \frac{r_2}{r_1}\right)^2} \{\bar{\psi} - [\zeta]_t [\xi]^{-1} [\zeta]\}. \quad (49)$$

Poprawka rozwiązania zerowego rzędu zawiera macierz nieskończoną, jednak należy sądzić, że dla większości praktycznie spotykanych nieciągłości w kablu wystarczy uwzględnić jeden lub dwa człony korekcyjne. Dalej posunięte uściślenie wyniku wiąże się z ogromnym wysiłkiem obliczeniowym i chyba nie wnosi żadnych poważnych zmian.

Dotychczas nasza analiza miała charakter stosunkowo ogólny i obejmowała różne rodzaje nieciągłości. Postaramy się teraz zastosować do wypadków konkretnych, pokazanych na rysunkach 1÷4.

A) $a_2 > a_1$; $b_2 < b_1$ (rys. 3).

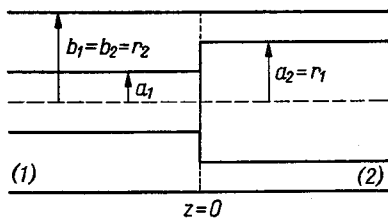
Brzeży apertury są takie same jak dla odcinka (2) kabla. A więc

$$r_1 = a_2, \quad (50)$$

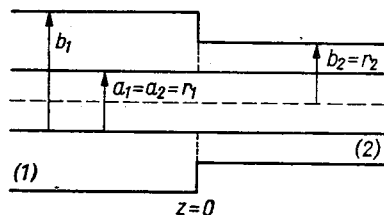
$$r_2 = b_2. \quad (51)$$

Z równań (39) i (40) wynika, że

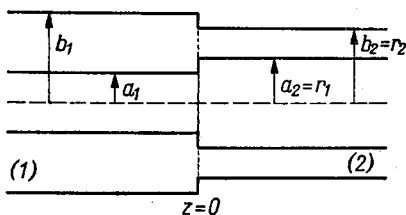
$$S_1(\beta_n r) = Z_1(\beta_n^{(2)} r). \quad (52)$$



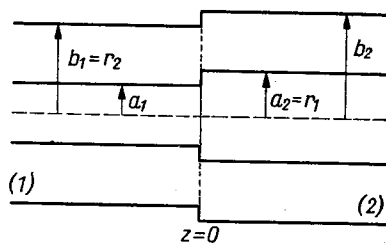
Rys. 1



Rys. 2



Rys. 3



Rys. 4

Nasz problem sprowadza się więc do obliczenia $P_q^{(p)}$ i $P_{qn}^{(p)}$ przy uwzględnieniu równań (50), (51) i (52). Korzystając z równań (24) ÷ (27) oraz (18) otrzymujemy

$$P_q^{(1)} = \frac{\sqrt{G_q^{(1)}}}{\beta_q^{(1)}} [Z_0(\beta_q^{(1)} a_2) - Z_0(\beta_q^{(1)} b_2)],$$

$$P_q^{(2)} = 0, \quad (53)$$

$$P_{qn}^{(1)} = \frac{2\beta_q^{(1)}}{\pi\beta_n^{(2)}} \sqrt{G_q^{(1)}} \left[\frac{Z_0(\beta_q^{(1)} a_2) - \frac{J_0(\beta_n^{(2)} a_2)}{J_0(\beta_n^{(2)} b_2)} Z_0(\beta_q^{(1)} b_2)}{(\beta_q^{(1)})^2 - (\beta_n^{(2)})^2} \right],$$

$$P_{qn}^{(2)} = \delta_{nq} \frac{1}{\gamma_q^{(2)} \sqrt{G_q^{(2)}}}, \quad \begin{array}{ll} \delta_{nq} = 0 & n \neq q, \\ = 1 & n = q, \end{array} \quad (53)$$

$$G_q^{(1)} = \frac{\pi^2 \beta_q^{(1)}}{2} \left[\frac{J_0^2(\beta_q^{(1)} a_1)}{J_0^2(\beta_q^{(1)} b_1)} - 1 \right]^{-1} \frac{1}{\sqrt{1 - \left(\frac{\beta_0}{\beta_q^{(1)}} \right)^2}},$$

$$G_q^{(2)} = \frac{\pi^2 \beta_q^{(2)}}{2} \left[\frac{J_0^2(\beta_q^{(2)} a_2)}{J_0^2(\beta_q^{(2)} b_2)} - 1 \right]^{-1} \frac{1}{\sqrt{1 - \left(\frac{\beta_0}{\beta_q^{(2)}} \right)^2}}.$$

Podstawiając te wyrażenia do równania (49) przy uwzględnieniu równania (38) i przyjmując $\varepsilon = \varepsilon_0 = \frac{1}{36 \pi \cdot 10^9}$ F/m, otrzymujemy następującą aproksymację zerowego rzędu dla równoważnej pojemności nieciągłości

$$C^{(0)} = \frac{1}{18 \cdot 10^9 \left(\ln \frac{b_2}{a_2} \right)^2} \sum_{q=1}^{\infty} \frac{G_q^{(1)}}{(\beta_q^{(1)})^2} [Z_0(\beta_q^{(1)} a_2) - Z_0(\beta_q^{(1)} b_2)]^2 \quad (54)$$

Ponieważ $Z_0(\beta_q^{(1)} r)$ jest liniową kombinacją funkcji J i N [patrz równanie (16)], otrzymane rozwiązanie da się wyznaczyć numerycznie przy pomocy funkcji stabilizowanych. Aczkolwiek rozwiązanie zerowego rzędu zawiera nieskończoną sumę, to jednak zbieżność jest na ogół bardzo dobra i można ograniczyć się do 3—5 pierwszych członów.

Rozwiązania wyższego rzędu otrzymamy z równania (49) w postaci

$$C^{(\infty)} = C^{(0)} - \frac{1}{18 \cdot 10^9 \left(\ln \frac{b_2}{a_2} \right)^2} [\zeta]_t [\xi]^{-1} [\zeta], \quad (55)$$

gdzie elementy macierzy obliczamy przy pomocy równań (53) i (38).

$$\text{B) } a_2 > a_1; \quad b_2 = b_1 \text{ (rys. 1).}$$

Ten rodzaj nieciągłości jest wypadkiem szczególnym sytuacji rozpatrzonej w A). Tym razem mamy

$$r_1 = a_2, \quad (56)$$

$$r_2 = b_2 = b_1 \quad (57)$$

i zamiast równań (53) otrzymujemy

$$P_q^{(1)} = \frac{\sqrt{G_q^{(1)}}}{\beta_q^{(1)}} Z_0(\beta_q^{(1)} a_2),$$

$$P_q^{(2)} = 0, \quad (58)$$

$$P_{qn}^{(1)} = \frac{2 \beta_q^{(1)}}{\pi \beta_n^{(2)}} \sqrt{G_q^{(1)}} \frac{Z_0(\beta_q^{(1)} a_2)}{(\beta_q^{(1)})^2 - (\beta_n^{(2)})^2},$$

$$P_{qn}^{(2)} = \delta_{nq} \frac{1}{\gamma_q^{(2)} \sqrt{G_q^{(2)}}}.$$

Korzystając z równania (54) możemy rozwiązanie zerowego rzędu napisać w postaci

$$C^{(0)} = \frac{1}{18 \cdot 10^9 \left(\ln \frac{b_2}{a_2} \right)^2} \sum_{q=1}^{\infty} \frac{G_q^{(1)}}{(\beta_q^{(1)})^2} Z_0^2(\beta_q^{(1)} a_2). \quad (59)$$

Rozwiązania wyższego rzędu obliczamy podobnie jak w przypadku A) stosując równanie (55).

$$C) \quad a_2 = a_1; \quad b_2 < b_1 \text{ (rys. 2).}$$

Również i tę nieciągłość można traktować jako przypadek szczególny analizy w A) z tym, że

$$r_1 = a_2 = a_1, \quad (60)$$

$$r_2 = b_2. \quad (61)$$

Uwzględniając fakt, że $a_1 = a_2$, mamy zamiast równań (53):

$$P_q^{(1)} = -\frac{\sqrt{G_q^{(1)}}}{\beta_q^{(1)}} Z_0 (\beta_q^{(1)} b_2),$$

$$P_q^{(2)} = 0, \quad (62)$$

$$P_{qn}^{(1)} = -\frac{2 \beta_q^{(1)}}{\pi \beta_q^{(2)}} \sqrt{G_q^{(1)}} \frac{J_0(\beta_n^{(2)} a_2)}{J_0(\beta_n^{(2)} b_2)} \frac{Z_0 (\beta_n^{(1)} b_2)}{(\beta_q^{(1)})^2 - (\beta_q^{(2)})^2}$$

$$P_{qn}^{(2)} = \delta_{nq} \frac{1}{\gamma_q^{(2)} \sqrt{G_q^{(2)}}}.$$

Daje to rozwiązanie zerowego rzędu

$$C^{(0)} = \frac{1}{18 \cdot 10^9 \left(\ln \frac{b_2}{a_1} \right)^2} \sum_{q=1}^{\infty} \frac{G_q^{(1)}}{(\beta_q^{(1)})^2} Z_0^2 (\beta_q^{(1)} b_2), \quad (63)$$

dalsze zaś poprawki można uzyskać przy pomocy równania (55).

$$D) \quad a_2 < a_1; \quad b_2 > b_1 \text{ (rys. 4).}$$

Tym razem rozwiązanie jest nieco bardziej skomplikowane, ponieważ $S_1(\beta_n r)$ nie redukuje się do $Z_1(\beta_n^{(1)} r)$ lub $Z_1(\beta_n^{(2)} r)$. Dzieje się tak dlatego, że apertura nie pokrywa się z brzegami żadnego z obu odcinków kabla w przeciwieństwie do wypadków wyżej rozpatrzonych.

Kładąc

$$r_1 = a_2, \quad (64)$$

$$r_2 = b_1 \quad (65)$$

całkujemy bezpośrednio równania (38) przy wykorzystaniu równań (24)–(27) oraz (18) i (40). W rezultacie otrzymujemy

$$P_q^{(1)} = \frac{\sqrt{G_q^{(1)}}}{\beta_q^{(1)}} Z_0 (\beta_q^{(1)} a_2),$$

$$P_q^{(2)} = -\frac{\sqrt{G_q^{(2)}}}{\beta_q^{(2)}} Z_0 (\beta_q^{(2)} b_1), \quad (66)$$

$$P_{qn}^{(1)} = \sqrt{G_q^{(1)}} \frac{f_q^{(1)}}{\pi \beta_n} \frac{2 Z_0 (\beta_q^{(1)} a_2)}{(\beta_q^{(1)})^2 - (\beta_n)^2},$$

$$P_{qn}^{(2)} = - \frac{\beta_q^{(2)}}{\pi \beta_n} \frac{J_0(\beta_n a_2)}{J_0(\beta_n b_1)} \frac{2 Z_0 (\beta_q^{(2)} b_1)}{(\beta_q^{(2)})^2 - (\beta_n)^2}$$

i rozwiązanie zerowego rzędu

$$C^{(0)} = \frac{1}{18 \cdot 10^9 \left(\ln \frac{b_1}{a_2} \right)^2} \sum_{q=1}^{\infty} \frac{G_q^{(1)}}{(\beta_q^{(1)})^2} Z_0^2 (\beta_q^{(1)} a_2) + \frac{G_q^{(2)}}{(\beta_q^{(2)})^2} Z_0^2 (\beta_q^{(2)} b_1).$$

Podobnie jak poprzednio, rozwiązania wyższego rzędu można uzyskać z równania (55).

WYKAZ LITERATURY

1. Marcuvitz N.: Waveguide Handbook. — Rad. Lab. Ser. vol. 10, New York 1951.
2. Szulkin P.: O pewnych podstawowych zagadnieniach teorii analizy nieciągłości w falowodach. — III. Symp. teorii pola elektromagnetycznego. Szklarska Poręba 1959.
3. Stratton J. A.: Electromagnetic theory. New York 1941.
4. Рызык J. M., Градштейн Т. S.: Таблицы интегралов. Москва 1951.

П. ШУЛЬКИН

ВАРИАЦИОННЫЕ МЕТОДЫ РАСЧЕТА НЕОДНОРОДНОСТИ В КОАКСИАЛЬНЫХ КАБЕЛЯХ

Резюме

В работе применен вариационный метод для решения проблемы плоских неоднородностей в коаксиальных кабелях. Выводится интегральное уравнение полной проводимости, выполняющее требования непрерывности составляющих E_r и H_y в апертуре и зажиме E_r в остальной части плоскости неоднородности. Вышеуказанное интегральное уравнение является стационарным относительно незначительных изменений E_r ; оно даёт минимальное значение Y , если мы при этом действительное распределение поля. Далее развёртываем E_r в полное семейство цилиндрических функций с основным членом $\frac{1}{r}$, где коэффициенты определяются путем использования свойства стационарности полной проводимости.

Предполагая, что тип ТЕМ является единственным типом свободно распространяющейся в кабеле волны, подтверждается возможность точного отображения неоднородности сосредоточенной ёмкостью. Главный член развёртки поля $\frac{1}{r}$ ведёт к ёмкости нулевого порядка, которая всегда превышает точную величину. Следует полагать, что благодаря свойствам стационарности решение нулевого порядка будет достаточно точным. Выводятся также формулы для коррекции дающие возможность получения решений высшего парядка. В заключении описан ряд отдельных неоднородности в кабелях.

P. SZULKIN

VARIATION METHODS OF MEASURING DISCONTINUITY IN COAXIAL CABLES

Summary

The variation method has been applied in the paper to solve the problem of flat discontinuity in coaxial cables. An integral equation has been deducted for an admittance, fulfilling the conditions of constituent discontinuities E_r and H_y in the aperture and disappearance of E_r in the remaining part of discontinuity plain. The above integral equation is stationary in relation to petty changes of E_r and gives the minimum value Y if the real field arrangement is applied. Then E_r is developed into a full collection of cylinder functions with the basic part $\frac{1}{r}$. The development coefficients are marked out thanks to making use of admittance stationariness properties.

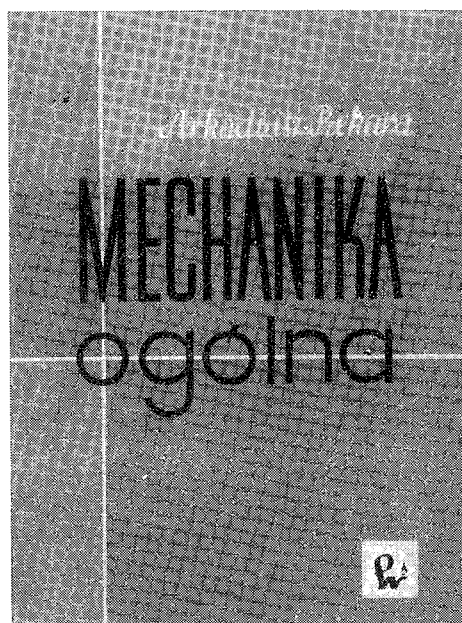
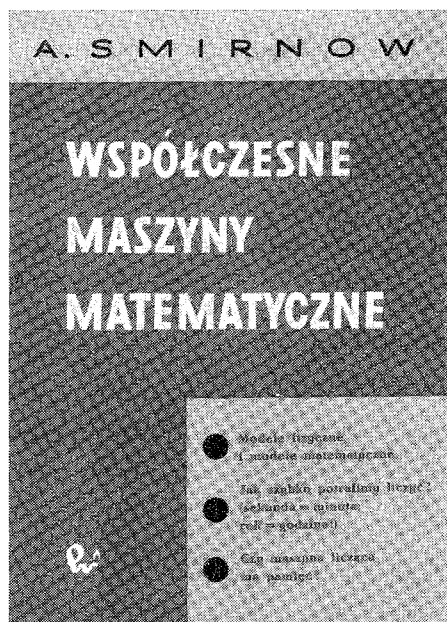
Assuming that the type TEM is the only type spreading freely in the cable, the possibility of exact representation of discontinuity by concentrated capacity.

The main front of the field development $\frac{1}{r}$ leads to capacity of zero order, which is always larger than the exact value. It should be thought that owing to the stationariness properties the solution of the zero order will be sufficiently precise. Formulae are deducted, too, for corrections which make it possible to obtain solutions of a higher order. Some peculiar discontinuity cases in cables have been discussed at the end.

NAJNOWSZE

WYDAWNICTWA

P
W
N



JERZY OWCZAREK

Dokładny pomiar prędkości obrotowej silnika miniaturowego

Rękopis dostarczono 16. 4. 1960

Praca stanowi dalszy ciąg cyklu artykułów dotyczących badań maszyn miniaturowych, zapoczątkowanego w zeszycie 1/1961 roku Archiwum Elektrotechniki opracowaniem autora pt. „Dokładny pomiar mocy pobranej przez silnik miniaturowy z zastosowaniem kompensatora w układzie współrzędnych prostokątnych”. We wstępie tego artykułu uzasadniono potrzebę opracowania metod pomiarów parametrów maszyn miniaturowych, odmiennych od stosowanych przy badaniu większych maszyn.

Niniejszy artykuł omawia pomiar następnego podstawowego parametru maszyn miniaturowych, którym jest prędkość obrotowa. Przeanalizowano w nim zalety i wady stosowanych dotychczas metod z punktu widzenia wymagań występujących przy badaniach rozpatrywanego rodzaju maszyn. Zaproponowano dwie metody pomiaru obrotów nie mające omówionych wad, jedną — opartą na zasadzie heterodynowego pomiaru różnicy częstotliwości sterowanej badaną maszyną i częstotliwości wzorcowej, drugą — liczenia czasu trwania okresu sygnału otrzymywanego z nadajnika fotoelektrycznego. Podano podstawowe założenia do budowy przyrządów wykorzystujących te zasady, których stosowanie pozwoli na uzyskiwanie dokładnych wyników przy prostej obsłudze i bezpośrednim obiektywnym odczycie wyników.

Ważną zaletą przyrządów jest możliwość stosowania ich bez żadnych dodatkowych urządzeń, jak również stosowania ich do badania maszyn wirujących dowolnej wielkości i rodzaju przy zachowaniu tej samej klasy dokładności.

1. WSTĘP

Pomiary prędkości obrotowej maszyn wykonywane są w stanie ustalonym i w stanach nieustalonych. Pomiary w stanach nieustalonych wskutek krótkotrwałości mogą być przeprowadzane tylko przy wykorzystaniu automatycznego zapisu wyników (np. oscylograficznego), których odczyt jest możliwy dopiero po zakończeniu przebiegu nie ustalonego. Podczas większości badań maszyn elektrycznych wymagane jest przeprowadzenie pomiaru ustalonej lub quasi-ustalonej prędkości obrotowej w danym, przeważnie krótkim, czasie, z możliwością natychmiastowego odczytu wyniku.

Większość stosowanych metod pomiarów prędkości obrotowej powoduje pewne obciążenie wału badanego silnika. W przypadku większych

maszyn obciążenie to może być pominięte wobec mocy na wale i pomiar nie nastęrcza trudności. Natomiast w przypadku mniejszych maszyn, jeżeli obciążenie wału miernikiem obrotów jest współmierne z mocą maszyny, pomiar powoduje zmianę warunków jej pracy, przestaje więc być miarodajny. W przypadku badania maszyn miniaturowych zostają więc od razu wyeliminowane wszystkie metody z zastosowaniem mechanicznych obrotomierzy, liczników obrotów i prądniczek tachometrycznych o dostrzegalnej bezwładności układu ruchomego. Pozostają metody powodujące mniejsze obciążenie wału badanej maszyny lub nie powodujące go zupełnie. Są to praktycznie metody częstotliwościowe i stroboskopowe.

2. OMÓWIENIE STOSOWANYCH METOD POMIARU PRĘDKOŚCI OBROTOWEJ MASZYN MINIATUROWYCH

Przy częstotliwościowej metodzie pomiaru prędkości obrotowej wał badanej maszyny steruje nadajnikiem częstotliwości. Nadajnik ten może składać się z obwodu magnetycznego, w którego polu obraca się wirnik sprzężony z badanym wałem, a skonstruowany tak, że jego obrót powoduje zmianę przewodności na drodze strumienia magnetycznego nadajnika. Najprostszy nadajnik częstotliwości tego typu opisany w [1] ma wirnik w postaci elipsy o 2- do 3-krotnej różnicy długości osi osadzony na badanym wale i obracający się w kołowym wycięciu zwory, w której uzwojenie prądu stałego wzbudza strumień magnetyczny. Obrót wirnika powoduje zmianę przewodności magnetycznej drogi strumienia. Strumień pulsuje z częstotliwością proporcjonalną do liczby obrotów badanego wału indukując w uzwojeniu wyjściowym zmienną siłę elektromotoryczną. Częstotliwość wyjściowego napięcia jest mierzona częstotliciemierzem, a na podstawie otrzymanego wyniku oblicza się obroty badanego wału. Wykonanie uzębienia o jednakowej podziałce na cylindrycznym wirniku i kołowym wycięciu zwory, podobnie jak w synchronicznych silnikach reluktancyjnych, zwiększa częstotliwość w stosunku do omówionego najprostszego rozwiązania proporcjonalnie do liczby zębów wirnika. W ten sposób można otrzymać zwiększoną częstotliwość wyjściową przy danej liczbie obrotów wirnika, co pozwala na dopasowanie jej do zakresu miernika lub do częstotliwości wzorcowej stabilizowanego generatora. Wyjściową częstotliwość nadajnika poza bezpośrednim pomiarem można porównywać z regulowaną częstotliwością generatora wzorcowego przy wykorzystaniu metody zerowej [8] lub ze stabilizowaną częstotliwością generatora nieziennej częstotliwości przy wykorzystaniu metody różnicowej [6].

Nadajnikiem częstotliwości może być również zespół składający się z umocowanej na badanym wale tarczy z otworami na obwodzie, przez które oświetlana jest fotokomórka. Częstotliwość impulsów sterowana fotokomórką proporcjonalna do liczby obrotów badanego wału i liczby

otworów na obwodzie tarczy może być wykorzystywana do pomiaru prędkości obrotowej badanego wału jednym z wymienionych sposobów, na przykład z wykorzystaniem metody różnicowej [7].

W przypadku pomiaru częstotliwości wysyłanej przez nadajnik za pomocą częstociomierza dokładność pomiaru określona jest klasą dokładności częstociomierza. Przy porównywaniu częstotliwości wyjściowej nadajnika z częstotliwością generatora wzorcowego i wykorzystaniu metody zerowej, dokładność pomiaru zależy od klasy generatora. Może być ona zwiększona względem poprzedniej, jednak kosztem większego skomplikowania pomiaru. Przy stosowaniu metody różnicowej i generatora stałej stabilnej częstotliwości (np. kwarcowego) można osiągnąć najwyższą dokładność (rzędu $\pm 0,05\%$) w wąskim zakresie częstotliwości (rzędu $\pm 5\%$ względem wzorcowej) [7].

Zaletą wszelkich odmian metody częstotliwościowej jest możliwość pomiarów również w stanach nieustalonych przy stosowaniu oscylograficznego zapisu częstotliwości wyjściowej nadajnika.

Nadajnik częstotliwości stanowi niewielkie obciążenie wału badanego silnika, a metoda częstotliwościowego pomiaru prędkości obrotowej pozwala na osiągnięcie dużej dokładności wyników. Z tego względu metoda jest bardzo wygodna i godna polecenia do pewnych granic minimalnej mocy i maksymalnych obrotów. Jednakże w przypadku maszyn miniaturowych — szczególnie najmniejszych mocy — rzędu ułamków wata i szybkoobrotowych to obciążenie wału nie zawsze może być pominięte, zarówno ze względu na tarcie o powietrze i moment wentylatorowy wirnika lub tarczy nadajnika, jak i ze względu na zmianę momentu bezwładności badanego układu ruchomego, co wpływa szczególnie na wyniki pomiarów podczas stanów nieustalonych.

Metoda stroboskopowa pozwala teoretycznie na pomiar prędkości obrotowej maszyny zupełnie bez obciążenia jej wału. Jednakże w tym przypadku otrzymany wynik pomiaru nie zawsze jest pewny. Stroboskopy używane do pomiaru prędkości obrotowej można podzielić na błyskowe — wyposażone w okresowe migające źródło światła oraz mechaniczne — wyposażone w wirującą przed okularzem tarczę z otworami. W stroboskopach błyskowych zmienna częstotliwość migania światła uzyskiwana jest przez zastosowanie układów lamp elektronowych, oporników i kondensatorów. Elementy tych układów nie zawsze zachowują stałość właściwości fizycznych w czasie, co powoduje obniżenie klasy dokładności przyrządu i konieczność okresowego, stosunkowo kłopotliwego, skalowania. Praktycznie osiągalna dokładność pomiarów waha się w granicach określonych uchybem od 1 do 3%. Stroboskopy mechaniczne są prostsze w budowie i obsłudze i nie wymagają tak częstego skalowania. Przy ich stosowaniu również trudno uzyskać uchyb pomiaru mniejszy od około 1%.

Podstawowymi wadami stroboskopowej metody pomiaru są subiektywność oceny wyniku i niepewność co do jego jednoznaczności. Dla wy-

eliminowania tych wad stosuje się szereg specjalnie opracowanych metod postępowania. Z opublikowanych w literaturze opracowań tych metod za najbardziej wyczerpujące i pełne można uważać [2] i [5].

Wieloznaczność wyniku można wyeliminować za pomocą dwu lub kilkakrotnej obserwacji efektu stroboskopowego, polegającego na pozornym zatrzymaniu się badanego wirującego ciała przy rozmaitych wybranych częstotliwościach i rozwiązanie odpowiedniej liczby równań z tyłomaż niewiadomymi [2]. Wadą tego postępowania jest skomplikowanie pomiaru, wymagającego dodatkowych obliczeń, konieczność utrzymania stałych obrotów badanego obiektu przez dłuższy czas, przeprowadzenia kilku pomiarów i niemożliwość natychmiastowego odczytu wyniku.

W tym samym celu można stosować specjalne tarcze mocowane na badanym wale z nakreślonymi na nich koncentrycznymi przerywanymi pierścieniami o stopniowo zmiennej podziałce przerw. Obserwacja pozornego bezruchu niektórych pierścieni, przy powolnym pozornym ruchu innych (z szybkością poślizgu) pozwala na jednoznaczne określenie wyniku pomiaru. Dokładność określenia prędkości obrotowej badanego ciała może być powiększona przez liczenie obrotów poślizgu odpowiednich pośrednich pierścieni i zastosowanie interpolacji [2]. Wadą tego postępowania jest konieczność dokładnie centrycznego mocowania tarcz na badanym wale i zachowania dostatecznej szerokości poszczególnych wykreślonych na niej pierścieni, co prowadzi do zwiększenia wymiarów tarczy i, wobec tego, obciążenia badanego wału. Dla uniknięcia tej ostatniej wady w [2] proponowana jest modyfikacja postępowania polegająca na mocowaniu na badanym wale małej tarczki sterującej miganiem światła, a tarczy z pierścieniami na pomocniczym silniku o wzorcowej liczbie obrotów. Komplikacja układu i wprowadzenie dodatkowego źródła błędów (silnik wzorcowy) może być celowa, ale tylko w stosunku do silników o mocy większej od rozpatrywanych. Wady tej modyfikacji postępowania przy maszynach miniaturowych zostały omówione przy metodzie częstotliwościowej, do której została ona właściwie sprowadzona.

Pomiary z zastosowaniem stroboskopu mechanicznego omówione są w [5]. Przy wykorzystywaniu różnicowej oceny wyniku pomiaru można dla wyeliminowania subiektywności postrzegania wprowadzić fotokomórkę. W tym przypadku na badany wał trzeba nałożyć tarczę z otworami, a przyrząd rejestrujący dostosować do pomiaru częstotliwości różnicowej względem wzorcowej otrzymywanej jednym z podanych w [5] sposobów. Metoda sprowadza się więc w zasadzie do omówionych poprzednio i jest obciążona tymi samymi wadami.

Ponadto w [5] zaproponowano uniknięcie wieloznaczności odczytu przy pomiarze stroboskopowym przez przeprowadzanie wstępnego pomiaru mniej dokładnym przyrządem rzędu wielkości mierzonej prędkości obrotowej. Jak już wspomniano w niniejszym opracowaniu, takie postę-

powanie nie może być brane pod uwagę przy badaniu maszyn miniaturowych.

Zwiększenie dokładności pomiaru przy stroboskopowej metodzie prowadzi więc w każdym przypadku do konieczności zakładania na wał badanego obiektu takiej lub innej tarczy. Wady takiego postępowania w przypadku pomiarów maszyn miniaturowych zostały już omówione przy metodzie częstotliwościowej.

Tak więc metoda stroboskopowa charakteryzująca się poważnymi zaletami przy badaniu większych maszyn, a w szczególności asynchronicznych (pomiaru poślizgu) i umożliwiająca w tych przypadkach uzyskanie wystarczającej dokładności pomiarów — nie może być zalecana do badań maszyn miniaturowych. Składają się na to — niejednoznaczność wyników przy braku tarcz pomocniczych, a w przypadku ich stosowania, obciążenie wału (momenty hamujący i wentylatorowy, zmiana bezwładności elementu wirującego), trudności mocowania tarcz (dokładna centryczność), długotrwałość i kłopotliwość pomiaru, niemożliwość natychmiastowego odczytu dokładnego wyniku przy większości metod postępowania i subiektywności oceny tego wyniku. Ponadto poważną wadą stroboskopowej metody pomiaru prędkości obrotowej jest niemożliwość jej zastosowania podczas badań stanów nieustalonych maszyn.

Do pomiarów prędkości obrotowej maszyn miniaturowych należy opracować metodę pomiaru nie mającą wad występujących w omówionych metodach stosowanych dotychczas. Powinna ona umożliwiać bezpośredni i obiektywny pomiar prędkości obrotowej z wystarczająco dużą dokładnością, bez obciążania badanej maszyny. Dwie metody częstotliwościowe opracowane przez autora i podane w p. 3 i 4 spełniają te wymagania.

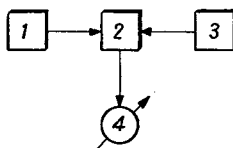
3. METODA OKREŚLANIA PRĘDKOŚCI OBROTOWEJ PRZEZ HETERODYNOWY POMIAR CZĘSTOTLIWOŚCI

Zasada tej metody polega na pomiarze różnicy dwóch niewiele różniących się od siebie częstotliwości — otrzymanej z nadajnika sterowanego badaną maszyną i wzorcowej, otrzymanej z generatora kwarcowego. Przygotowanie badanej maszyny do pomiaru polega wyłącznie na zaczerznięciu części obwodu jej wału lub innego elementu wirującego, przewidzianego do sterowania nadajnika. Jest więc spełnione wymaganie całkowitego braku obciążenia badanej maszyny (brak sprzężenia z jakimkolwiek ruchomym elementem pomiarowym). Pomiar zostaje przeprowadzony przy pomocy przyrządu, którego uproszczony schemat blokowy podany jest na rys. 1.

Nadajnik częstotliwości 1 zawiera fotokomórkę reagującą na natężenie światła odbitego od badanego wirującego elementu oraz żarówkę oświetlającą ten element. Podczas wirowania zaczerznięcie części wału powodu-

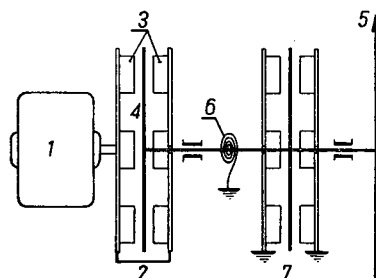
je modulację napięcia w obwodzie fotokomórki z częstotliwością proporcjonalną do liczby jego obrotów.

Modulowane napięcie z nadajnika 1 i napięcie wzorcowe o dużej stabilności częstotliwości z generatora kwarcowego 3 doprowadzane są do



Rys. 1. Schemat blokowy układu do heterodynowego pomiaru częstotliwości

elementu 2. Tu ich częstotliwości zostają sprowadzone do jednakowego rzędu wielkości za pomocą powielaczy lub dzielników częstotliwości i podane do mieszacza. Z mieszacza i układu odpowiednich filtrów otrzymuje



Rys. 2. Przyrząd wskazówkowy

się napięcie o częstotliwości stanowiącej różnicę częstotliwości przekształconych napięć nadajnika i generatora kwarcowego. Poprzez wzmacniacz mocy zasila ono przyrząd wskazówkowy 4.

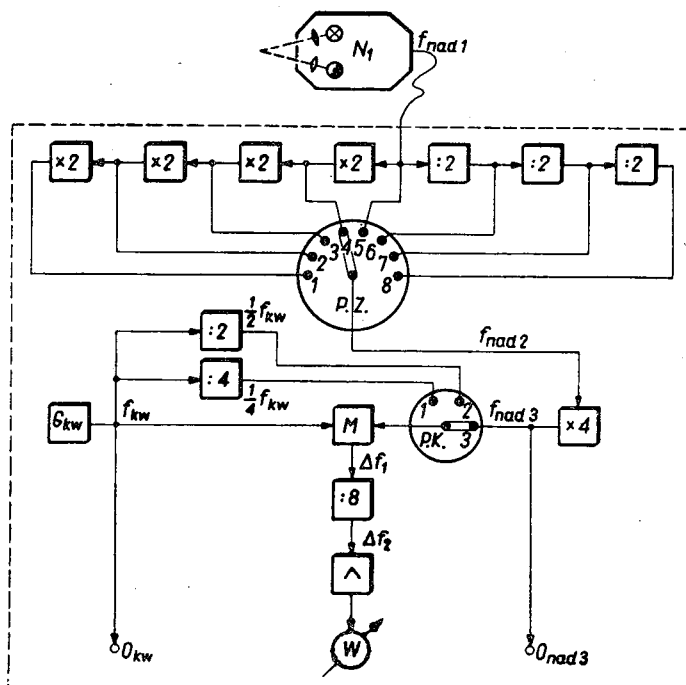
Układ przyrządu wskazówkowego podany jest na rys. 2. Synchroniczny silnik histerezy 1 zasilany przez wzmacniacz mieszacza obraca zespół magnetyczny 2 z prędkością proporcjonalną do częstotliwości różnicowej. W szczelinie między stałymi magnesami 3 zespołu umieszczona jest tarcza 4. Przy obrotach zespołu magnetycznego 2 w tarczy 4 indukują się prądy wirowe, wytwarzające wspólnie z polem wirujących magnesów stałych moment obrotowy porywający tarczę 4. Na wspólnej z nią osi umieszczone są tarcza tłumika magnetycznego układu ruchomego 7 i wskazówka przyrządu 5. Kąt obrotu osi jest ograniczony sprężynką zwrotną 6.

Przemysł krajowy produkuje aktualnie zespoły magnetyczne o podanej konstrukcji, które mogłyby być wykorzystane w całości do budowy omawianego przyrządu.

Stosowanie miernika o podanym układzie z napędem silnikiem synchronicznym pozwala na otrzymanie liniowej skali przyrządu, która może obejmować kąt obrotu wskazówki bliski do 360° . Dzięki temu przy

umiarkowanej długości wskazówki można wykorzystać roboczą długość skali rzędu $500 \div 1000$ mm, co zapewnia jej dużą i równomierną zdolność rozdzielczą. Jako silnik napędowy miernika został zaproponowany silnik histerezy, gdyż może on zadowalająco pracować w szerokim zakresie częstotliwości ($20 \div 600$ Hz) i nie jest wrażliwy na znaczne zmiany napięcia zasilającego. Proponowane rozwiązanie miernika spełnia wymaganie bezpośredniości i obiektywności odczytu.

Schemat blokowy rozwiązania konstrukcyjnego przyrządu z jedną częstotliwością wzorcową podany jest na rys. 3.



Rys. 3. Schemat przyrządu z generatorem na jedną częstotliwość wzorcową

Nadajnik częstotliwości N_1 stanowi oddzielny element konstrukcyjny powiązany z przyrządem za pomocą giętkiego kabla. Umożliwia to niezależenie ustawienia przyrządu od położenia i wymiarów badanej maszyny.

Podstawowymi elementami właściwego przyrządu są powielacze i dzielniki częstotliwości oznaczone w zależności od krotności jako $\times 2$, $\times 4$, $:2$, $:4$, $:8$, generator kwarcowy — G_{kw} , mieszacz M , przełączniki zakresów $P.Z.$ i kontrolny $P.K.$, wzmacniacz $\wedge$, oraz filtry nie uwidocznione na schemacie.

Przykładowy układ przyrządu podany na rys. 3 pozwala na pomiar obrotów od 125 obr/min do 32 000 obr/min przy użyciu ośmiu zakresów pomiarowych podzielonych jak w tabelicy 1. Stała przyrządu podana jest

w założeniu 1000 działkowej skali dla zakresu od 1000 obr/min do 2000 obr/min, jako najwygodniejszej do przeliczeń.

Tablica 1

Położenie P.Z.	Zakres pomiaru obrotów	Stała przrządu
1	125 — 250	1/8
2	250 — 500	1/4
3	500 — 1000	1/2
4	1000 — 2000	1
5	2000 — 4000	2
6	4000 — 8000	4
7	8000 — 16000	8
8	16000 — 32000	16

Przy pomiarze obrotów od 125 do 32 000 obr/min z nadajnika otrzymuje się częstotliwość $f_{nad\ 1}$ od $2 \frac{1}{12}$ do $533 \frac{1}{3}$ Hz. Kaskada zdwajaczy i dzielników częstotliwości sprowadza ją dla poszczególnych zakresów obrotów do wartości $f_{nad\ 2}$ wynoszącej od $33 \frac{1}{3}$ do $66 \frac{2}{3}$ Hz w sposób wynikający wprost ze schematu. Częstotliwość $f_{nad\ 2}$ zwiększona 4-krotnie jest przekształcona na $f_{nad\ 3} = 133 \frac{1}{3}$ do $266 \frac{2}{3}$ Hz dla wszystkich zakresów pomiaru.

Napięcie nadajnika o przekształconej częstotliwości $f_{nad\ 3}$ i napięcie z generatora kwarcowego o częstotliwości wzorcowej $f_{kw} = 500$ Hz podawane są na mieszacz, w którym zostaje wydzielona częstotliwość różnicowa $\Delta f_1 = f_{kw} - f_{nad\ 3}$, wynosząca od $366 \frac{2}{3}$ do $233 \frac{1}{3}$ Hz. Zostaje ona ośmiokrotnie zmniejszona do wartości Δf_2 wynoszącej $45 \frac{5}{6}$ do $29 \frac{1}{6}$ Hz.

Napięcie o częstotliwości Δf_2 poprzez wzmacniacze zasila miernik W o omówionej konstrukcji (rys. 2), na którym odczytuje się mierzoną prędkość obrotową przy uwzględnieniu stałej przrządu, zależnej od położenia podczas pomiaru przełącznika P.Z.

Do skalowania i kontroli miernika przewidziana jest możliwość zasilania drugiego wejścia mieszacza dwoma stabilizowanymi częstotliwościami z generatora kwarcowego. Z rys. 3 widać, że przy ustawieniu przełącznika P.K. w położeniu 1 na mieszacz podawane są częstotliwości 500 i 125 Hz, a w położeniu 2 — 500 i 250 Hz. Otrzymane wzorcowe częstotliwości różnicowe podzielone i wzmocnione sterują miernik, którego wskazówka powinna ustawić się odpowiednio na wartościach 937,5 obr/min i 1875 obr/min. Według tych punktów wykonuje się liniową skalę mierni-

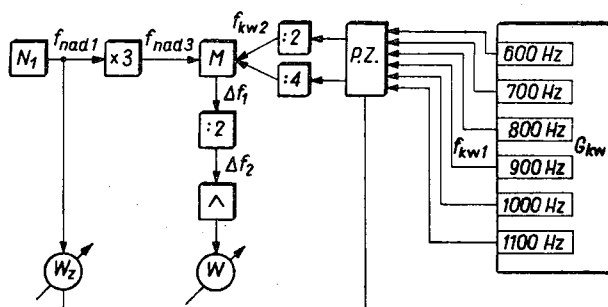
ka o roboczym zakresie od 1000 do 2000 obr/min i przeprowadza ewentualne sprawdzenie ustawienia miernika przed przystąpieniem do pomiarów. Przyrząd jest przygotowany do pomiarów przy ustawieniu przełącznika P.K. w położeniu 3.

Na schemacie blokowym, rys. 3, uwidoczniono również wyprowadzenia O_{kw} i $O_{nad 3}$, umożliwiające oscylografowanie częstotliwości generatora kwarcowego i przekształconej częstotliwości nadajnika przy pomiarach prędkości obrotowej w stanach nieustalonych.

Dokładność pomiaru prędkości obrotowej opisanym przyrządem można określić rozpatrując dokładność poszczególnych jego elementów. Jeżeli wykonać skalę kołową o promieniu rzędu 150 mm, to przy roboczym zakresie 1000 obr/min jednemu obrotowi odpowiada odległość kresek podziałki rzędu 1 mm. Zakładając możliwość dokładności odczytu 0,5 odległości kresek, otrzymamy uchyb odczytu wskazań na skali ok. 0,05%.

Nadajnik z fotokomórką, dzielniki i powielacze częstotliwości, mieszacz, wzmacniacze i filtry nie wprowadzają uchybów częstotliwości. Uchyb stabilizowanej częstotliwości wzorcowej otrzymywanej z generatora kwarcowego (10^{-5}) jest pomijalny wobec uchybu odczytu wskazań na skali. Uchyb miernika przy starannym jego wykonaniu może być również oceniony na około 0,05% [7]. Tak więc otrzyma się uchyb przyrządu rzędu 0,1% dla wszystkich zakresów pomiarowych. Dokładność pomiaru prędkości obrotowej przewyższa uzyskiwaną przy stosowaniu innych metod spełniających warunki braku obciążania badanej maszyny przez przyrząd i obiektywnej oceny wyników pomiaru.

Dokładność pomiaru można zwiększyć przy zastosowaniu generatora kwarcowego pozwalającego na otrzymanie większej liczby wzorcowych częstotliwości. Zasadę postępowania tą drogą wytłumaczymy rozpatru-



Rys. 4. Schemat przyrządu z generatorem na kilka częstotliwości wzorcowych

jąc układ do pomiaru prędkości obrotowej w zakresie 1000—10000 obr/min z generatorem kwarcowym na sześć częstotliwości wzorcowych: 600, 700, 800, 900, 1000 i 1100 Hz. Uproszczony schemat blokowy tego układu jest podany na rys. 4, a podział zakresów w tablicy 2. W układzie znajdują się dwa wskaźniki — W_z i W . Wskaźnik W_z wykonany jako częstotcio-

miernik wskazuje zakresy pomiarowe i jest wyskalowany w tysiącach obrotów na minutę — może więc być stosunkowo mało dokładny. Wskaźnik W_z steruje jednocześnie przełącznik P.Z., włączający w zależności od zakresu tysięcy obrotów odpowiedni zakres generatora częstotliwości wzorcowej i dzielnik napięcia — zgodnie z tablicą 2. Na mierniku W , wykonanym analogicznie jak w poprzednim układzie, odczytujemy dokładną liczbę obrotów na minutę, którą trzeba dodać do wskazań zakresu przyrządu W_z , aby otrzymać mierzoną wartość prędkości obrotowej.

Tablica 2

Lp.	Zakres pomiaru obrotów	$f_{nad\ 3} = 3f_{nad\ 1}$	f_{kw}	Krotność podziału	$f_{kw\ 2}$	$f_1 = f_{nad\ 3} - f_{kw}$
	obr/min	Hz	Hz	—	Hz	Hz
1	1000—2000	50—100	600	4	150	100—50
2	2000—3000	100—150	800	4	200	100—50
3	3000—4000	150—200	1000	4	250	100—50
4	4000—5000	200—250	600	2	300	100—50
5	5000—6000	250—300	700	2	350	100—50
6	6000—7000	300—350	800	2	400	100—50
7	7000—8000	350—400	900	2	450	100—50
8	8000—9000	400—450	1000	2	500	100—50
9	9000—10000	450—500	1100	2	550	100—50

Jeżeli rozpatrzmy dokładność wyników pomiarów otrzymywanych w tym układzie i porównamy ją z poprzednim rozwiązaniem, to stwierdzimy, że dokładność odczytu wskazań rośnie ze wzrostem wartości bezwzględnych mierzonych obrotów, osiągając przy zakresie 10000 obr/min dziesięciokrotnie większą wartość względem zakresu 1000 obr/min. Tak więc przy ostatecznym obliczaniu dokładności przyrządu dla zakresu do 10000 obr/min uchyb spowodowany odczytem można będzie pominąć wobec założonego uprzednio uchybu miernika W , otrzymując uchyb przyrządu rzędu 0,05%.

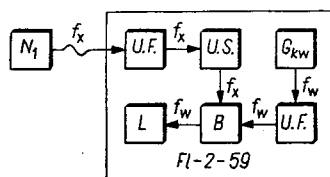
Dalsze zwiększenie dokładności pomiaru może być osiągnięte przez zwięźanie zakresów pomiarowych i stosowanie większej liczby częstotliwości wzorcowych. Liczba niezbędnych częstotliwości wzorcowych będzie wzrastała wolniej niż liczba zakresów pomiarowych.

4. METODA POMIARU PRĘDKOŚCI OBROTOWEJ PRZY POMOCY LICZNIKA IMPULSÓW

Zasada tej metody polega na pomiarze czasu trwania 1 lub 10 okresów napięcia otrzymywanego z nadajnika częstotliwości przy pomocy aktualnie produkowanego typowego krajowego przyrządu — falomierza liczącego F. 1. -2-59 [4] i [9].

Przygotowanie silnika do pomiaru obrotów, zasada działania i konstrukcja nadajnika częstotliwości N_1 z fotokomórką są identyczne z omówionymi w rozdz. 3.

Pomiar czasu trwania 1 lub 10 okresów mierzonego napięcia odbywa się w układzie, którego ideowy schemat blokowy uwidoczniony jest na rys. 5. Napięcie z nadajnika N_1 o mierzoneym czasie trwania okresu podawane jest na wejście falomierza liczącego F.l. Po przekształceniu na impulsy w układzie formującym UF powoduje ono otwarcie układu bramko-



Rys. 5. Schemat układu pomiarowego z falomierzem liczącym

wego B podczas 1 lub 10 pełnych okresów. Przez czas otwarcia układ bramkowy B przepuszcza impulsy o znanej częstotliwości wzorcowej f_w otrzymywanej z generatora kwarcowego. Impulsy o częstotliwości f_w , która powinna znacznie przewyższać częstotliwość mierzonego napięcia, są zliczane przez licznik L . Licznik wskazuje całkowitą liczbę n przepuszczonych przez bramkę impulsów. Stąd czas trwania okresu mierzonego napięcia $t_x = \frac{n}{m \cdot f_w}$, gdzie m liczba mierzonych okresów. Częstotliwość wzor-

cowa f_w jest dobrana tak, aby wynik cyfrowy odczytywany z licznika miał wymiar czasu. Dla otrzymania prędkości obrotowej w obrotach na minutę należy 60 podzielić przez otrzymany czas trwania okresu w sekundach. Dokładność pomiaru zależy od dokładności falomierza, gdyż przeliczenie może nie wprowadzać błędu. Zastosowany w falomierzu licznik nie daje uchybu. Układ otwierający bramkę na czas m okresów również nie daje

uchybu. Uchyb względny wzorca częstotliwości $\frac{\Delta f_w}{f_w} = 10^{-5}$. Uchyb bez-

względny powodowany przez układ bramkowy wynikający z zasady jego działania wynosi $\Delta n = \pm 1$. Całkowity uchyb bezwzględny pomiaru czasu trwania okresu Δt_x po podstawieniu wyżej wymienionych wartości i od-

powiednich przekształceniach wyraża się wzorem: $\Delta t_x = \pm \left(\frac{1}{m f_w} + 10^{-5} t_x + P \right)$, gdzie P pomijalny zazwyczaj błąd pomiaru uwarunkowa-

ny amplitudą i kształtem napięcia sygnału. Jego wartość można obliczyć ze wzoru $P \leq \frac{5 t_x}{m U_x} \cdot 10^{-3}$ [s], gdzie U_x — napięcie sygnału w woltach, a t_x — czas trwania okresu w sekundach. Zakres pomiaru czasu trwania okresu mierzonego napięcia przez falomierz FL-2-59 wynosi od 100 μ sek

do 10^3 sek, odpowiada to zakresowi od 600 000 do 0,06 obr/min, a więc pokrywa z nadwyżką cały spotykany zakres obrotów silników miniaturowych.

W miarę wzrostu prędkości obrotowej przy danej liczbie okresów, których czas trwania jest mierzony, maleje czas pomiaru i wzrasta jego uchyb. Po przeprowadzeniu obliczeń uchybów można stwierdzić, że przy użyciu najmniejszej jednostki pomiarowej falomierza Fl-2-59 wynoszącej $t_w = \frac{1}{f_w} = 1 \mu\text{sek}$ obroty od najwolniejszych do 6000 obr/min mogą być mierzone podczas trwania jednego obrotu, a więc w czasie mniejszym od 1 sek poczynając już od 60 obr/min ze względnym uchybem procentowym nie przekraczającym 0,01%. Przy 60000 obr/min czas trwania pomiaru wynosi już 0,001 sek, a uchyb 0,1%. Dla zwiększenia dokładności można przy tym zakresie prędkości obrotowej mierzyć czas trwania dziesięciu okresów — wtedy czas trwania pomiaru wydłuży się do 0,01 sek, a uchyb zmaleje do 0,01%.

Jeżeli uda się skonstruować dostatecznie prosty i pracujący bez uchybów przelicznik cyfrowy umożliwiający dzielenie 60 przez t (względnie n), to po wprowadzeniu go do schematu falomierza przed licznikiem L , można będzie otrzymywać cyfrowy wynik pomiarów wyrażony bezpośrednio w obrotach na minutę. Licznik niezależnie od tego, czy pokazuje czas trwania okresu, czy liczbę obr/min, może być sprzęgnięty z drukarką, co umożliwia automatyczne utrwalanie wyników pomiarów co 0,2 do 5 sek.

Pomiar prędkości obrotowej drogą bezpośredniego określania częstotliwości sygnału otrzymywanego z nadajnika N_1 możliwy przy pomocy falomierza Fl-2-59 nie może być zalecany, gdyż jest zbyt mało dokładny. Uchyb pomiaru przybiera wartości mniejsze od 0,1% przy obrotach przewyższających 60 000 obr/min, a więc dopiero w pobliżu górnej granicy spotykanych w praktyce prędkości obrotowych silników miniaturowych.

5. ZAKOŃCZENIE

Opracowane metody są przeznaczone do dokładnego pomiaru prędkości obrotowych maszyn miniaturowych, dla których stabilizacja obrotów w zautomatyzowanych układach wymagana jest zazwyczaj z dokładnością określoną maksymalnym dopuszczalnym uchybem 0,3—0,5% [7]. Potrzeba opracowania wynika z nieprzydatności do tego celu metod stosowanych przy badaniach większych maszyn, gdyż dokładniejsze z nich powodują obciążenie badanej maszyny, zaś nie powodujące obciążenia są za mało dokładne.

Obydwie proponowane metody odpowiadają wszystkim podstawowym wymaganiom stawianym przy badaniach maszyn miniaturowych: badany

wał nie jest obciążony podczas pomiaru, przygotowanie maszyny i wykonanie pomiaru jest proste, czas trwania pomiaru — krótki, odczyt — łatwy, a ocena wyniku pomiaru — obiektywna. Dokładność wyników przewyższa osiągalną przy stosowaniu dotychczasowych metod pomiarowych. Układy pomiarowe są tak skonstruowane, że nieuniknione zmiany w czasie fizycznych właściwości niektórych ich elementów nie wpływają na dokładność otrzymywanych wyników, odróżniając je pod tym względem korzystnie od generatorów lampowych regulowanej częstotliwości, stosowanych w stroboskopach błyskowych i przy niektórych dotychczasowych rozwiązaniach metod częstotliwościowych. Zapewniona jest samokontrola przyrządów, pozwalająca uniezależnić wyniki pomiarów od przypadkowych wpływów zewnętrznych, co eliminuje konieczność okresowego sprawdzania lub skalowania według obcego wzorca, gwarantując stałą dokładność wyników pomiarów w czasie eksploatacji przyrządu.

Obydwie metody łączą zalety metod stroboskopowych i częstotliwościowych eliminując ich wady. Mianowicie wykorzystują one możliwości uzyskania dużej dokładności i obiektywności oceny wyników pomiarów właściwe metodom częstotliwościowym, przy braku obciążenia badanego wału — właściwym metodom stroboskopowym. Obiektywność oceny wyników pomiaru daje poważną przewagę omawianym metodom nad zalecanymi na ogół dla silników miniaturowych metodami stroboskopowymi, gdyż w tych ostatnich ocena wyników zależy od obserwatora, wymaga dużej wprawy i naraża na częste okazje do pomyłek [2], [3], [5].

Oryginalność opracowanych metod polega na tym, że zespół wszystkich omówionych zalet nie występuje jednocześnie w żadnej ze stosowanych dotychczas metod pomiarowych.

Z porównania opracowanych metod wynika, że metoda heterodynowego pomiaru częstotliwości może zapewnić otrzymanie mniejszej dokładności. Praktycznie osiągalny uchyb minimalny 0,1% przy skomplikowaniu konstrukcji może być doprowadzony do wartości 0,05%. Jednak ta metoda pozwala na ciągły pomiar prędkości obrotowej w stanach nieustalonych drogą oscylografowania, co stanowi jej przewagę nad metodą z zastosowaniem falomierza liczącego. Ta ostatnia jest dokładniejsza, gdyż uchyb może być łatwo utrzymany mniejszym od 0,01%, ale pomiar nie może być wykonywany w sposób ciągły. Najmniejszy odstęp czasu między dwoma kolejnymi pomiarami wynosi 0,2 sek. Dla jego wykorzystania niezbędne jest oczywiście użycie dodatkowego przyrządu — drukarki przyłączonej do licznika falomierza. Współpraca z drukarką niezależnie od częstości pomiarów stanowi element automatyzacji pomiaru. Dodatkową zaletą jest otrzymywanie wyniku w postaci cyfrowej, co całkowicie eliminuje uchyb odczytu. Czas przeprowadzenia pomiaru w stanach quasi-ustalonych przy metodzie heterodynowej jest określony wyłącznie prędkością odczytu wyniku ze skali przyrządu, a przy metodzie z falomierzem wynosi w zależności od mierzonej prędkości obrotowej od 1 do

0,01 sek. Wynik może więc być uważany w obydwu przypadkach za wartość chwilową prędkości obrotowej w momencie pomiaru. Obydwie metody uzupełniają się w pewnym sensie. Decyzja, którą z nich zastosować, zależy od wymagań stawianych wynikowi pomiaru.

Przy budowie przyrządów dla obydwu opracowanych metod mogą być wykorzystane typowe, produkowane aktualnie elementy, co w dużej mierze ułatwia zadanie.

Dzięki wykonaniu nadajnika w postaci oddzielnego elementu konstrukcyjnego i prostocie przygotowania do pomiaru wirującego elementu badanej maszyny, przyrządy można stosować do pomiaru prędkości obrotowej nie tylko maszyn miniaturowych, dla których są one w zasadzie przeznaczone, lecz również maszyn o dowolnej wielkości i rodzaju, przy zachowaniu wszystkich zalet obydwu metod.

*Politechnika Warszawska
Katedra Maszyn Elektrycznych*

WYKAZ LITERATURY

1. Czeczet Ju. S.: Issledowanie elektriczeskich mikromaszin. — MEI, 1953.
2. Hermann P. K.: Stroboskopische Drehzahlmessungen. — ETZ-A. 1953, Bd. 24.
3. Owczarek J.: Dokumentacja i egzemplarz stroboskopu mechanicznego. — ZME PW, 1958.
4. Porządkowski E.: Falomierz liczący Mod. Fle-155. — Zeszyty Naukowe Politechniki Warszawskiej, Elektryka Nr 14, 1957.
5. Reinhardt F.: Stroboskopische Drehzahlmessungen. — Archiv für technisches Messen, Sept. 1953.
6. Utiamyszew R. I.: Primienienie czastotnogo metoda dla tocznogo izmierienia ugłowych skorostiej. — Izmieritielnaja tiechnika Nr 2/1958.
7. Utiamyszew R. I.: Tocznyj pribor dla izmierienia ugłowej skorosti mikrodwigatielej. — Priborostrojenie 5/1958.
8. Dokumentacja hamownicy IE1 dla maszyn ułamkowej mocy z urządzeniem do pomiaru obrotów.
9. Falomierz liczący, model Fl-2-59. Instrukcja techniczna. Zakład Urządzeń Radiotechnicznych Politechniki Warszawskiej 1959.

Е. ОВЧАРЕК

ТОЧНОЕ ИЗМЕРЕНИЕ СКОРОСТИ ВРАЩЕНИЯ МИКРОДВИГАТЕЛЯ

Резюме

Статья является разработкой метода измерения следующего параметра микромашины после опубликованной в Архивум Электротехники 1/61 статьи об измерениях потребляемой мощности. В статье констатировано, что для измерения скорости вращения микродвигателей можно применять исключительно частотные и стробоскопические методы измерения. Оба эти метода обсуждены с указанием их преимуществ и недостатков.

Далее предложены два оригинальные частотные метода автора, отличающиеся наличием преимуществ предыдущих методов при избежании их недостатков.

Первый метод основан на измерении разности двух близких по значению частот — полученной и датчика, управляемого исследуемым двигателем, и эталонной, полученной с кварцевого генератора частоты. Основными элементами измерительной схемы кроме датчика и кварцевого генератора являются смеситель частоты и измерительный прибор с гистерезисным электродвигателем описанной конструкции. Гетероидное устройство для получения разности частоты базированной на высокостабильной эталонной частоте кварцевого генератора позволяет получить точность измерения в классе 0,1 или еще более точном. Конструкция измерительного прибора с линейной шкалой, обладающей большой разделительной способностью, позволяет производить непосредственную мгновенную считку результатов, оценка которых может производиться объективно. Соответственные выводы напряжений позволяют одновременно осциллографировать частоты датчика и кварцевого генератора, что дает возможность применять этот метод также при измерении вращения в нестационарных режимах.

Второй метод основан на измерении продолжительности 1 или 10 периодов сигнальной частоты, получаемой от фотоэлектрического датчика при помощи типового прибора Fl-2-59 разработанного в кафедре Радиотехнического Оборудования Варшавского Политехнического Института. Основными элементами прибора являются эталонный кварцевый генератор, пропускная схема с импульсным управлением и счетчик дающий цифровой результат измерения. Точность измерения можно получить без затруднений в классе 0,01. Измерения можно вести каждые 0,2 сек, причем необходимо конечно применение автоматического регистрирующего устройства, или без него в произвольно больших интервалах времени.

Датчик частоты, использующий фотоэлемент реагирующий на световой пучек отраженный от вращающегося элемента, позволяет производить измерения при полном отсутствии нагрузки исследуемой машины. Это является главным преимуществом предлагаемых методов измерения. Подготовка исследуемого двигателя к измерениям благодаря принятому решению датчика предельно проста. Время необходимое для измерения скорости вращения — коротко. Считка результатов не представляет затруднений, а их оценка объективна. Точность определения скорости вращения выше получаемой обычно при применении известных методов измерения. Схемы приборов обеспечивают неизменность точности во времени и исключают необходимость перманентных контрольных регулировок приборов. Высокая и постоянная точность измерения представляет собой следующее основное ценное преимущество предлагаемых методов. Совокупность всех перечисленных преимуществ не выступает одновременно ни в одном из известных методов измерения скорости вращения.

В заключении подsumмированы преимущества предлагаемых методов по сравнению с применяемыми и подчеркнута универсальность разработанных приборов заключающаяся в том, что их можно применять не только для измерения скорости вращения микромашин, для которых они в принципе построены, но также для машин произвольных видов и размеров при сохранении всех преимуществ методов.

J. OWCZAREK

ACCURACY IN MEASUREMENT OF MICROMOTOR REVOLUTIONS

Summary

The paper deals with the measurement methods concerned with micromotors and is a continuation of the work set out by the author in *Electrotechnic Archives*, Nr. 1, 1961 on the measurement of the power drawn by such motors.

Actually two groups of measurement methods in measurement of the angular velocity are involved, namely the frequency and the stroboscopic method. Advantages and disadvantages of both methods headed by an analysis are discussed in detail. The author then presents his own two novel frequency methods combining all the advantages of two above mentioned groups of methods and void of all their disadvantages.

The principle of the first novel method consists in measurement of the difference between two slightly diverging frequencies, from the emitter controlled by the examined machine and the quartz oscillator whose frequency is considered as sample frequency.

Besides the emitter and the oscillator there are such essential elements in the network as the frequency mixer and the indicator device comprising hysteresis motor whose design is described. The heterodyne arrangement devised to originate the frequency difference based upon highly stable sample frequency provided by the quartz oscillator, permits to meet the measurement requirements of an accuracy of an order equal or even higher than 0,1. The design of pointer device affords a highly distributive faculty of the line scale enabling, thus to take off direct and instant readings and to appraise the recorded measurement results objectively.

Oscillographic tracing of both the emitter frequency proportional to the measured angular velocity and that of the quartz oscillator considered as time standard may be carried out simultaneously owing to the appropriately made leads. This equally permits to apply this method for the measurements of the angular velocity in transient states.

The principle of the second novel method consists in measurement of cycle time for one or ten cycles of the signal frequency obtained from the photoelectric emitter with assistance of a typical device Fl-2-59 designed by the Chair of Radio-technical Equipment of the Polytechnic University of Warsaw. Essential components of this device consist of the quartz oscillator of sample frequency, gate circuit and a counter recording the measurement results in numerical form.

A degree of measurement accuracy of order 0,01 may be achieved without difficulty. Applying automatically a recording arrangement the readings may be taken every 0,2 sec. or at an arbitrary long space of time if made without it.

The frequency emitter exploiting the performance of photocell sensitive to the intensity of light reflected by revolving element permits to carry out the measurement at no-load run of the examined machine. This is a most outstanding advantage of suggested measurement method. Such an emitter design makes all necessary preparations, which machine measurement involves, extremely simple. The time required to measure the angular velocity is short. The reading of results is easy and entirely independent of the observer's perception. The accuracy is of a higher order than in any other known measurement methods. The equipment systems secure the permanence of accuracy and exclude necessary elsewhere and somewhat troublesome periodic instrument calibrations. In fact a high and constant measurement accuracy may be considered as next of the significant advantages of suggested methods. It is worth while to note that in any method known up to now the entire set of specified advantages could not occur simultaneously.

In concluding part all the advantages of suggested methods compared with those presently used are summed up. The stress is put on their versatility which consists in the possibility of using the designed devices not only for the measurement of the angular velocity of micromotors, for which they have been primarily intended, but also for arbitrary chosen size and type of machine preserving all the advantages of both measuring methods.

FRANCISZEK SONDIJ

Ocena pracy asymetrycznego stalowniczego urządzenia łukowego oraz dobór optymalnych warunków jego pracy

Rękopis dostarczono 30. 3. 1961

W pracy przeprowadzono analizę obwodu trójfazowego asymetrycznego stalowniczego urządzenia łukowego. Analizę przeprowadzono dla stanu symetrii prądowej oraz dla symetrycznego układu napięć zasilających urządzenie. W wyniku przeprowadzonej analizy uzyskano wskaźniki K_g oraz K_e (50), (53) określające stopień asymetrii pieca dla dwóch charakterystycznych stanów jego pracy. Wymienione wskaźniki uzyskano na podstawie pomiarów stałych rozpatrywanego obwodu. Wskaźniki K_g oraz K_e umożliwiają wyznaczenie przy pomocy prostych związków podstawowych wielkości charakteryzujących zarówno pracę pieca, jak i urządzenia dla krańcowych stanów optymalnego zakresu pracy. Wyznaczono również wpływ asymetrii pieca na stopień nierównomierności cieplnego obciążenia wewnętrznej ściany pieca (67), (68), (69) oraz określono górną granicę technicznego wskaźnika asymetrii uwarunkowanego dopuszczalnym stopniem nierównomierności cieplnego obciążenia ściany pieca (70), (71), (72).

Wyprowadzone zależności zostały sprawdzone na drodze pomiarów przeprowadzonych na sześciotonowym łukowym trójfazowym piecu stalowniczym.

1. WSTĘP

Analiza zależności funkcjonalnych pomiędzy elektrycznymi wielkościami występującymi w obwodzie trójfazowego stalowniczego urządzenia łukowego, w wyniku której wyznacza się warunki ekonomicznej pracy pieca, przyjmuje upraszczająco, że urządzenie łukowe jest odbiornikiem symetrycznym, co oznacza, że prądy fazowe urządzenia oraz jego napięcia fazowe są sobie równe.

Poza tym przyjmuje się również upraszczająco, że oporności czynne oraz oporności bierne poszczególnych faz są wielkościami niezależnymi od prądu. Założenia takie pozwalają na rozpatrywanie tylko jednej fazy urządzenia oraz na uogólnienie wniosków wynikających z takiej analizy na pozostałe fazy urządzenia [5], [7].

W rzeczywistych układach stalowniczych urządzeń łukowych zachodzą zawsze odstępstwa od podanych założeń upraszczających, ponieważ urządzenia te są w zasadzie odbiornikami asymetrycznymi.

Asymetria trójfazowych stalowniczych urządzeń łukowych może mieć charakter konstytucjonalny, związany z budową urządzenia, jak również

może wynikać ze sposobu eksploatacji pieca, a w szczególności z nierównych obciążeń poszczególnych faz urządzenia [5].

Asymetria wynikająca z budowy urządzenia jest nazywana asymetrią konstrukcyjną.

Asymetrię konstrukcyjną warunkuje głównie nierówność oporności czynnych oraz nierówność oporności biernych poszczególnych faz toru wielkoprądowego podczas pracy pieca.

Każda z faz toru wielkoprądowego jest przewodem utworzonym z szeregowo połączonych odcinków przewodu sztywnego, odcinka przewodu giętkiego i elektrody połączonej z przewodem sztywnym przy pomocy ochłodzonego uchwytu.

Występujące podczas pracy pieca czynne i bierne oporności poszczególnych faz toru wielkoprądowego różnią się między sobą. Różnice te wynikają przede wszystkim z powodu nierównych sprzężeń magnetycznych pomiędzy poszczególnymi przewodami toru, a szczególnie pomiędzy ich odcinkami giętkimi.

Przy stosowanym powszechnie w praktyce układzie przewodów toru wielkoprądowego, sprzężenia magnetyczne pomiędzy środkową fazą toru i fazami skrajnymi uważa się praktycznie za równe, natomiast sprzężenia te różnią się od sprzężenia magnetycznego skrajnych faz toru. Te różnice w sprzężeniach magnetycznych powodują podczas pracy pieca wzrost oporności czynnej jednej ze skrajnych faz toru i równocześnie zmniejszenie się oporności czynnej pozostałej skrajnej fazy toru, przy czym przyrost oporności jednej fazy jest równy ubytkowi oporności pozostałej fazy skrajnej.

Wpływ na różnice pomiędzy opornościami czynnymi poszczególnych faz toru wielkoprądowego mają poza tym ewentualne różnice w wymiarach i opornościach właściwych przewodów tworzących poszczególne fazy toru oraz ewentualne różnice pomiędzy opornościami styków uchwytów poszczególnych elektrod.

Zmianom ulegają również oporności indukcyjne poszczególnych faz toru.

Powstałe w wyniku różnych sprzężeń magnetycznych poszczególnych faz toru wielkoprądowego przy założeniu toru płaskiego, różnice w opornościach czynnych skrajnych faz toru mogą prowadzić do znacznego przeciążenia jednej ze skrajnych faz urządzenia i do niedociążenia pozostałej fazy skrajnej [12]. Faza skrajna przeciążona w wyniku istnienia nierównych sprzężeń magnetycznych pomiędzy poszczególnymi fazami toru nazywa się fazą mocną, natomiast faza niedociążona fazą słabą.

Powstałe w taki sposób odstępstwa od podanych na wstępie założeń upraszczających, jakie przyjmuje się przy analizie układów trójfazowych urządzeń stalowniczych, powodują, że wyniki uzyskane na drodze analitycznej przy przyjętych uproszczeniach, różnią się od rzeczywistych przebiegów analizowanych wielkości.

Zjawisko asymetrii wpływa niekorzystnie na pracę stalowniczego urządzenia łukowego. Konsekwencje praktyczne asymetrii są różnorodne.

W wyniku powstawania fazy mocnej i fazy słabej występuje — przy zachowaniu symetrii prądowej — asymetria mocy w piecu, przy czym największą moc wykazuje łuk fazy mocnej, natomiast moc najmniejszą ma łuk fazy słabej. W wyniku asymetrii mocy w piecu elektroda fazy mocnej ulega szybszemu ugarowi. Poza tym wynikiem asymetrii mocy jest również asymetryczne cieplne obciążenie wewnętrznych ścian pieca. Nierównomierność cieplnego obciążenia ścian pieca może prowadzić do szybszego zużywania się wyprawy pieca, szczególnie w pobliżu łuku o największej mocy, co z kolei staje się przyczyną częstszych przerw w ruchu, koniecznych do przeprowadzenia remontów, i w konsekwencji prowadzi do zmniejszenia się wydajności pieca [5], [11]. Innym przejawem asymetrii mocy w piecu jest nierównomierne nagrzewanie się wsadu z uwagi na różne warunki pracy występujące w zasięgu poszczególnych elektrod. Taki stan pracy pogarsza równomierność procesu piecowego.

Elektrycznym przejawem asymetrii urządzenia jest również przesunięcie punktu zerowego we wsadzie piecowym w stosunku do punktu zerowego transformatora piecowego. Zjawisko to ma jednak dla pracy pieca znaczenie podrzędne i wyznaczanie wielkości tego przesunięcia jest mało użyteczne [1], [2], [5].

Asymetria konstrukcyjna urządzenia wyrażana jest ilościowo zależnością

$$\Delta X_m = \omega M (1 - C),$$

przy czym

ω — oznacza pulsację prądu zasilającego piec,

M — indukcyjność wzajemną pomiędzy środkowym i skrajnym przewodem toru,

C — wskaźnik asymetrii konstrukcyjnej.

Wskaźnik asymetrii konstrukcyjnej określony jest jako

$$C = \frac{N}{M} < 1.$$

We wzorze tym N oznacza indukcyjność wzajemną pomiędzy skrajnymi przewodami toru wielkoprądowego.

W zależnościach powyższych przyjmuje się, że indukcyjności wzajemne pomiędzy przewodem środkowym toru i jego przewodami skrajnymi są równe i większe od indukcyjności wzajemnej między skrajnymi przewodami toru. Założenie takie jest zgodne z praktyką i uzasadnione teoretycznie [13].

Określony powyżej wskaźnik C ma charakter raczej teoretyczny. Jego wyznaczenie wymaga znajomości wielkości M oraz N . Wielkości te są

praktycznie trudne do zmierzenia, szczególnie w warunkach pracy urządzenia łukowego.

Wielkość ΔX_m jest podstawową wielkością, od której zależy asymetria. Wielkość ta nie stanowi jednak operatywnego wskaźnika oceny pracy pieca z uwzględnieniem jego asymetrii.

Biorąc pod uwagę, że asymetria konstrukcyjna wpływa na takie istotne wielkości i wskaźniki eksploatacyjne, jak rozkład mocy na poszczególne łuki, wydajność i trwałość pieca, celowe jest wprowadzenie operatywnej metody oceny pracy urządzenia łukowego uwzględniającej jego asymetrię. Wprowadzenie takiej oceny pracy urządzenia łukowego umożliwi jego racjonalniejszą eksploatację i da możliwość porównywania urządzeń łukowych pod względem asymetrii.

2. CEL I PLAN PRACY

Celem pracy jest analiza pracy trójfazowego łukowego urządzenia stalowniczego uwzględniająca asymetrię konstrukcyjną urządzenia. W wyniku tej analizy przeprowadzonej dla stanu symetrii prądowej powinno być możliwe:

a) wyznaczenie procentowego rozkładu mocy na poszczególne fazy zarówno w odniesieniu do urządzenia, jak i dla pieca,

b) wyznaczenie warunków optymalnej pracy pieca, a w szczególności natężenia prądu i mocy użytecznej, warunkujących najmniejsze zużycie energii elektrycznej na tonę wsadu, oraz natężenie prądu i mocy pieca, warunkujących największą jego przelotność,

c) wyznaczenie współczynników mocy i sprawności elektrycznej dla optymalnych warunków pracy pieca,

d) ocena stanu termicznego na wewnętrznej ścianie pieca,

e) klasyfikacja urządzeń łukowych na urządzenia, które można by uważać za praktycznie symetryczne, i na urządzenia, których asymetria przekracza granice ustalone normatywnie.

Analiza zostanie przeprowadzona na uproszczonym schemacie trójfazowego obwodu stalowniczego urządzenia łukowego dla warunku symetrii prądowej.

W rozważaniach przyjęto za podstawę stan symetrii prądowej, ponieważ jest to korzystne dla sieci zasilającej urządzenie.

Zastosowane uproszczenie polega na pominięciu strat stanu jałowego transformatora piecowego. Straty te wynoszą przeciętnie mniej niż 1% maksymalnej mocy urządzenia łukowego. Wielkości tego rzędu mieszczą się w granicach uchybów pomiarów, jakie wykonuje się na stalowniczych urządzeniach łukowych, z tego też powodu pominięcie tych strat jest uzasadnione.

W rozważaniach, na podstawie przyjętego schematu uproszczonego, zostaną wyznaczone ogólne wzory na oporności i prądy obwodu oraz warunki powstania symetrii prądowej w obwodzie.