

POLSKA AKADEMIA NAUK
INSTYTUT PODSTAWOWYCH PROBLEMÓW TECHNIKI

ROZPRAWY ELEKTROTECHNICZNE

TOM VI • ZESZYT 1-2

KWARTALNIK

PAŃSTWOWE WYDAWNICTWO NAUKOWE
WARSZAWA 1960

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, Młodowa 10

Nakład 660 (466+194)	Oddano do składania 23.I.1960
Ark. wyd. 8,5, ark. druk. 7,75	Podpisano do druku 3.V.1960
Pap. druk. sat. kl. III, 80 g, 70×100	Druk ukończono w maju 1960
Zamówienie nr 145/60	Cena zł. 50.— C-23

Drukarnia im. Rewolucji Październikowej, Warszawa

MICHAŁ JABŁOŃSKI

Rozwinięcie metody obliczania oporności rozproszenia transformatora z niesymetrycznymi uzwojeniami cylindrycznymi

Rękopis dostarczono 19.12.1958

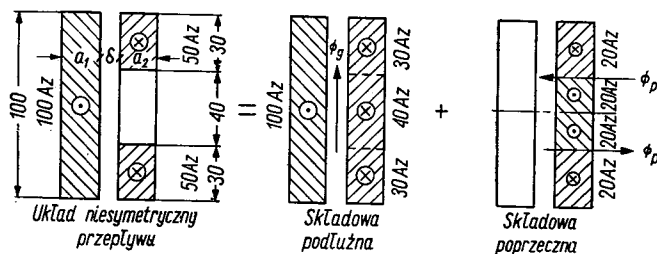
W pracy rozwinięto i rozbudowano opublikowaną już wcześniej [5] metodę obliczania oporności biernej zwarcia transformatora z uzwojeniami cylindrycznymi nawet przy bardzo znacznych asymetriach przestrzennych amperozwojów.

Opracowano wzory na skupianie rozrzedzeń w uzwojeniach, na eliminowanie luk skompensowanych oraz ostatecznie na obliczanie oporności biernej w przypadku jednoczesnego występowania luk brzegowych i środkowych w obu uzwojeniach. Opracowano również sposób obliczania oporności biernej w autotransformatorach wielozaczepowych i w transformatorach wieloprądowych z małą liczbą zwojów.

Rozważania są ilustrowane przykładami obliczeniowymi.

1. ROZWAŻANIA PODSTAWOWE

Obecność niesymetrii w układzie cylindrycznym uzwojeń powoduje znaczne komplikacje w obliczaniu oporności rozproszenia. Spośród kilku



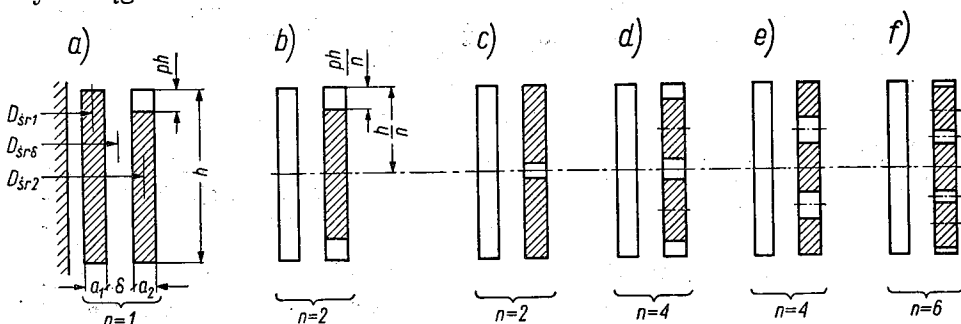
Rys. 1. Przykład podziału amperozwojów niesymetrycznych na składową podłużną i poprzeczną

metod stosowanych w tym przypadku najprostszą, lecz i najmniej dokładną jest metoda amperozwojów podłużnych i poprzecznych [2], [3], [4]. Analiza tej metody i jej założeń pozwala jednak na wyprowadzenie ściś-

lejszych wyrażeń oraz doświadczalnych współczynników poprawkowych, co powoduje bardzo znaczne zwiększenie dokładności rezultatów.

Nie powtarzając tutaj w całości rozumowań opublikowanych już wcześniej [5] przytoczymy w rozdziale pierwszym jedynie zasadę rozkładu amperozwojów na składową podłużną i poprzeczną — w formie przykładu na rys. 1 — oraz podstawowe wyjaśnienia i wzory wynikające ze wspomnianych rozumowań. Rozdziały następne stanowią dalsze rozwinięcie i rozszerzenie zakresu metody.

Na rysunku 2 przedstawiono znaczenie symboli stosowanych w dalszym ciągu.



Rys. 2. Oznaczenia poszczególnych wymiarów i określenie liczby n wzajemnie skompensowanych części elementarnych asymetrii

Oporność rozproszenia dla składowej podłużnej oblicza się znanym wzorem Kappa-Rogowskiego

$$X_{pod} = 7,9 \cdot f \cdot z^2 l_{sr} \frac{\delta'}{h} k_R \cdot 10^{-8} [\Omega], \quad (1)$$

gdzie

$$l_{sr} = \pi \frac{D_{sr1} \frac{a_1}{3} + D_{sr\delta} \cdot \delta + D_{sr2} \frac{a_2}{3}}{\frac{a_1}{3} + \delta + \frac{a_2}{3}} [\text{cm}]; \quad (2)$$

$$\delta' = \frac{a_1 + a_2}{3} + \delta [\text{cm}]; \quad (3)$$

k_R — współczynnik Rogowskiego dla układu składowej podłużnej;

w postaci uproszczonej $k_R = 1 - \frac{a_1 + \delta + a_2}{\pi \cdot h}$;

$\frac{h}{k_R} = h_s$, czyli wysokość rozproszeniowa układu [cm].

Oporność dodatkową X_{dod} wynikającą ze składowej poprzecznej strumienia rozproszenia określimy początkowo dla przypadku, gdy umyślone

uzwojenie krążkowe rozkłada się na jednakowe wzajemnie skompensowane części elementarne, których liczba wynosi n , rys. 2. Liczbą p oznaczmy stopień asymetrii, którego istotę wyjaśniają rys. 2a i 2b. Zjawisko zakrzywiania się linii sił pola magnetycznego i powstający w związku z tym wzajemny wpływ pola podłużnego i poprzecznego uwzględnimy wprowadzając do wyrażenia na X_{dod} wspomniany wyżej współczynnik poprawkowy k_p .

Wzór na oporność dodatkową od składowej poprzecznej dla jednej części elementarnej umyślnego układu krążkowego X'_{dod} analogicznie do wzoru (1) jest:

$$X'_{dod} = 7,9 \cdot f \left(\frac{zp}{n} \right)^2 l_{sr} \frac{l_{wt}}{l_{sr}} \cdot \frac{h}{3n} \cdot \frac{k'_R}{a} k_p \cdot 10^{-8} [\Omega], \quad (4)$$

Dla wszystkich n części elementarnych:

$$X_{dod} = X'_{dod} \cdot n = 7,9 \cdot f \frac{z^2 p^2}{n} l_{sr} \frac{l_{wt}}{l_{sr}} \cdot \frac{h}{3n} \cdot \frac{k'_R}{a} k_p \cdot 10^{-8} [\Omega], \quad (5)$$

gdzie:

$\frac{zp}{n}$ — liczba zwojów przypadająca na każdą lukę uzwojenia asymetrycznego przy równomiernym rozłożeniu amperozwojów na całej wysokości uzwojenia,

$l_{sr} \frac{l_{wt}}{l_{sr}}$ — przejście ze średniej długości obu uzwojeń na własną średnią długość obwodową umyślnego uzwojenia krążkowego l_{wt}

$\frac{h}{3n}$ — zastępcza szczelina rozproszeniowa δ'_p jednego zespołu elementarnego umyślonych uzwojeń krążkowych,

k_p — współczynnik poprawkowy uwzględniający wzajemny wpływ strumienia podłużnego i poprzecznego,

k'_R — współczynnik Rogowskiego dla krążkowego umyślnego układu poprzecznego. Współczynnik ten należy wykorzystać w postaci pełnej [1] przy zaniedbaniu obecności kadzi, a uwzględniając obecność rdzenia.

Oznaczając poszczególne odległości zgodnie z rys. 3 i stosując symbole:

$$s_n = e^{2 \cdot \pi \cdot n \cdot b/h} \quad (6a)$$

oraz

$$m_n = \frac{1 - e^{-\frac{\pi \cdot n \cdot a}{h}}}{\frac{\pi \cdot n \cdot a}{h}} \quad (6b)$$

możemy napisać

$$k'_R = 1 - m_n \left(1 - \frac{m_n \cdot \pi \cdot n \cdot a}{2 \cdot h \cdot s_n} \right). \quad (7)$$

Podstawiając (7) do (5) uzyskujemy:

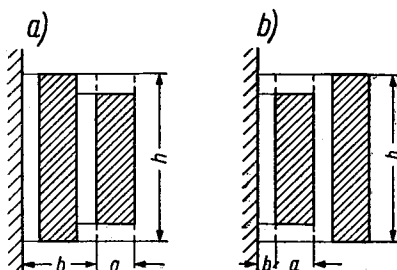
$$\begin{aligned} X_{doa} &= 7,9 \cdot f \cdot z^2 \cdot l_{sr} \cdot 10^{-8} \left\{ \frac{p^2}{n^2} k_p \frac{l_{wl}}{l_{sr}} \cdot \frac{h}{3a} \left[1 - m_n \left(1 - \frac{m_n \cdot \pi \cdot n \cdot a}{2 \cdot h \cdot s_n} \right) \right] \right\} = \\ &= 7,9 \cdot f \cdot z^2 \cdot l_{sr} \cdot \beta \cdot 10^{-8}, \end{aligned} \quad (8)$$

gdzie

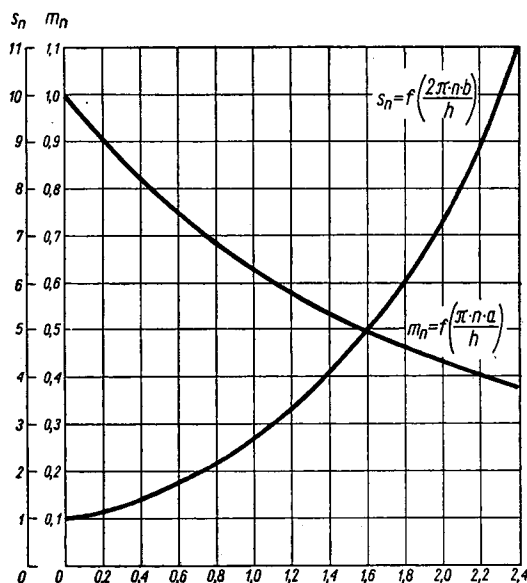
$$\beta = k_p \cdot p^2 \left\{ \frac{1}{n^2} \cdot \frac{l_{wl}}{l_{sr}} \cdot \frac{h}{3a} \left[1 - m_n \left(1 - \frac{m_n \cdot \pi \cdot n \cdot a}{2 \cdot h \cdot s_n} \right) \right] \right\} = k_p \cdot p^2 \cdot \gamma, \quad (9)$$

czyli

$$\gamma = \frac{1}{n^2} \cdot \frac{l_{wl}}{l_{sr}} \cdot \frac{h}{3a} \left[1 - m_n \left(1 - \frac{m_n \cdot \pi \cdot n \cdot a}{2 \cdot h \cdot s_n} \right) \right]. \quad (9a)$$



Rys. 3. Oznaczenia stosowane we wzorach 4÷9



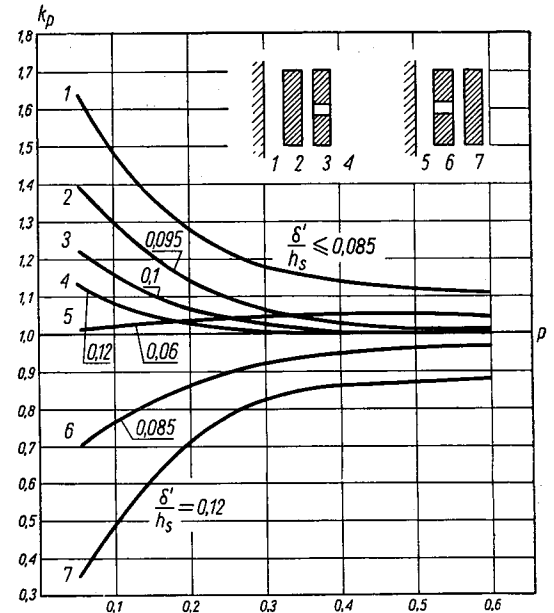
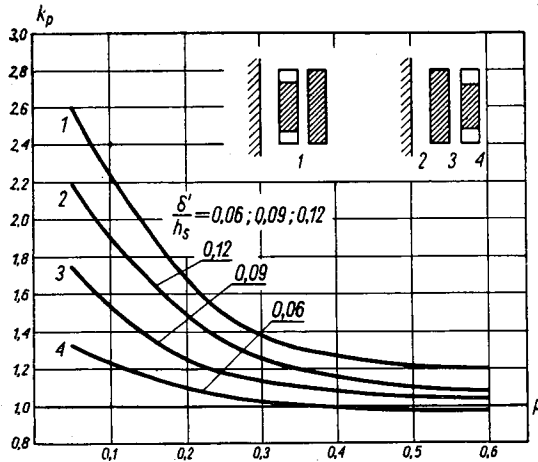
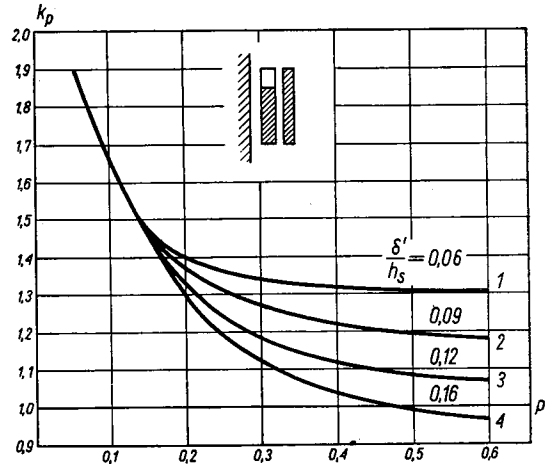
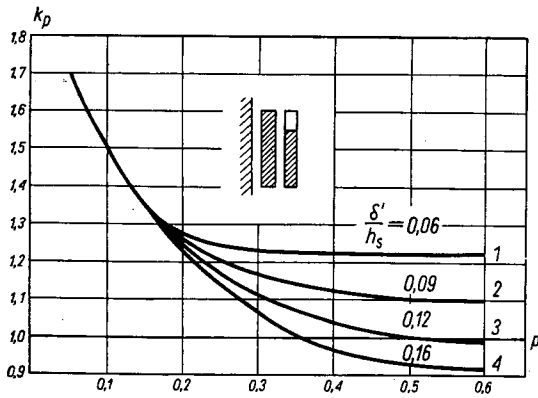
Rys. 4. Przebieg współczynników s_n oraz m_n

Całkowita oporność rozproszenia jest sumą oporności składowych:

$$X = X_{pod} + X_{doa} = 7,9 \cdot f \cdot z^2 \cdot l_{sr} \left(\frac{\delta'}{h_s} + \beta \right) 10^{-8} [\Omega] \quad (10)$$

lub

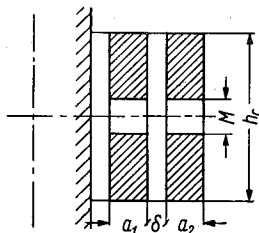
$$u_x = 7,9 \cdot f \cdot \frac{Iz}{e'} l_{sr} \left(\frac{\delta'}{h_s} + \beta \right) 10^{-6} [\%]. \quad (10a)$$



Rys. 5. Przebieg współczynnika poprawkowego k_p w zależności od stopnia asymetrii dla różnych układów asymetrii

Wartość wyrażenia $s_n = f\left(\frac{2 \cdot \pi \cdot n \cdot b}{h}\right)$ oraz $m_n = f\left(\frac{\pi \cdot n \cdot a}{h}\right)$ w zakresie spotykanym w praktyce przedstawia rys. 4. Uzyskaną empirycznie na doświadczalnych transformatorach wartość k_p przedstawiono na rys. 5a, b, c, d.

Do dalszych rozważań potrzebny jest również opublikowany już uprzednio [5] wzór określający wysokość zastępczą układu cylindrycznego



Rys. 6. Przypadek luki skompensowanej w środku uzwojenia

go (z uwzględnieniem współczynnika Rogowskiego) zawierającego dwie jednakowe naprzeciw siebie położone luki w uzwojeniach górnego i dolnego napięcia, rys. 6. Luki takie będziemy nazywać dalej skompensowanymi.

$$h_s = h_c - M + \frac{a_1 + a_2 + \delta}{3} \left(2 - e^{-\frac{3M}{a_1 + a_2 + \delta}} \right). \quad (11)$$

Jeżeli liczba luk skompensowanych o wysokościach poosiowych $M_1, M_2, \dots, M_k$ wynosi k , to wzór (11) przybiera postać w oparciu o poz. [4] wykazu literatury:

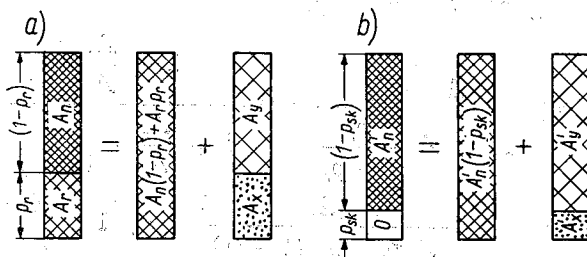
$$h_s = h_c - \left(\sum_{n=1}^{n=k} M_n \right) + \frac{a_1 + a_2 + \delta}{3} \left(k + 1 - \sum_{n=1}^{n=k} e^{-\frac{3 M_n}{a_1 + \delta + a_2}} \right). \quad (11a)$$

2. PRZYGOTOWANIE ZASTĘPCZEGO UKŁADU OBLICZENIOWEGO Z LUKAMI SKUPIONYMI

Luki występujące w uzwojeniach transformatora są bądź rzeczywistymi lukami skupionymi (np. wyłączenie części cewek zaczepekowych lub tzw. odstępy grupowe), bądź rozrzedzeniami na tle wzmacniania izolacji, kompensowania luk zaczepekowych itp. Rozrzedzenia powstają przez zastosowanie zwiększonych odstępów między cewkami lub przez zmniejszenie liczby zwojów w cewkach na skutek pogrubienia izolacji.

Przygotowując zastępczy układ obliczeniowy rozrzedzenia należy skupić i utworzyć zastępcze luki skupione. Przedstawiony na rys. 7a naj-

prostszy układ z rozrzedzeniem ma wymiary przedstawione w liczbach względnych odniesionych do wysokości uzwojenia. Podobnie jak na rys. 1



Rys. 7. Skupianie rozrzedzenia

a) Rozkład amperozwojów z rozrzedzeniem na składową podłużną i poprzeczną, b) Analogiczny rozkład amperozwojów z luką skupioną

układ ten rozkłada się na składową podłużną i poprzeczną. Amperozwoje składowej poprzecznej równoważą się wzajemnie, czyli

$$A_x p_r = A_y (1 - p_r),$$

a superpozycja obciążeń liniowych obu składowych musi dawać rzeczywisty układ wyjściowy. Można więc zgodnie z rys. 7a napisać:

$$[A_n (1 - p_r) + A_r p_r] - A_x = A_r \quad (12)$$

oraz

$$[A_n (1 - p_r) + A_r p_r] + A_y = A_n. \quad (13)$$

Układ ten ma być w pełni równoważny układowi ze skupioną luką przedstawionemu na rys. 7b.

Dla układu z luką skupioną można napisać analogicznie do poprzedniego:

$$A'_x p_{sk} = A'_y (1 - p_{sk})$$

oraz

$$[A'_n (1 - p_{sk}) + 0 \cdot p_{sk}] - A'_x = 0, \quad (14)$$

$$[A'_n (1 - p_{sk}) + 0 \cdot p_{sk}] + A'_y = A'_n. \quad (15)$$

Równoważność obu układów (rys. 7a oraz 7b) występuje wtedy, gdy zarówno ich składowe podłużne, jak i poprzeczne są odpowiednio równoważne, to znaczy

$$A_n (1 - p_r) + A_r p_r = A'_n (1 - p_{sk}) \quad (16)$$

oraz

$$A_x p_r = A_y (1 - p_r) = A'_x p_{sk} = A'_y (1 - p_{sk}). \quad (17)$$

Rozwiązanie otrzymanego układu sześciu równań (12 do 17) z sześcioma niewiadomymi (A_x , A_y , A'_x , A'_y , A'_n , p_{sk}) w zależności od trzech wiadomych (A_n , A_r , p_r) pozwala na uzyskanie równania luki skupionej p_{sk} :

$$p_{sk} = p_r \frac{A_n - A_r}{A_n + A_r \frac{p_r}{1 - p_r}} = p_r \frac{1 - \frac{A_r}{A_n}}{1 + \frac{A_r}{A_n} \cdot \frac{p_r}{1 - p_r}} =$$

$$= p_r \left[1 - \frac{A_r}{A_r p_r + A_n (1 - p_r)} \right] = p_r \left(1 - \frac{A_r}{A_{sr}} \right), \quad (18)$$

gdzie A_{sr} jest średnim obciążeniem liniowym dla łącznej wysokości normalnej części uzwojenia oraz rozpatrywanej części rozrzedzonej. Na ogół wygodniej jest operować bezwzględными wartościami wysokości odpowiednich części uzwojeń, niż określać ich wartości względne (p). W takim przypadku wzór (18) można wyrazić w postaci:

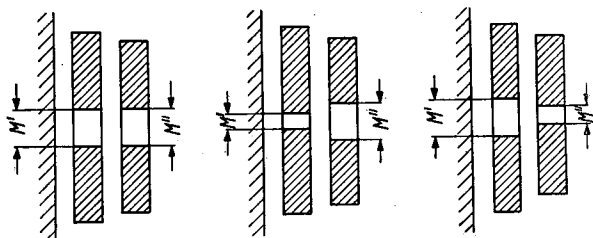
$$M_{sk} = M_r \frac{1 - \frac{A_r}{A_n}}{1 + \frac{A_r}{A_n} \cdot \frac{M_r}{M_n}}. \quad (18a)$$

Tu M_r jest wysokością rozpatrywanej części rozrzedzonej; M_n — to wysokość części normalnej uzwojenia, a M_{sk} — wysokość luki skupionej zastępującej M_r .

W porównaniu z szeroko dotychczas używanym wyrażeniem

$$p_{sk} = p_r \frac{A_n - A_r}{A_n} = p_r \left(1 - \frac{A_r}{A_n} \right)$$

wzór (18) lub (18a) daje mniejszą wartość luki skupionej — tym mniejszą, im rozrzedzenie jest rozciągnięte na większej przestrzeni i im bardziej



Rys. 8. Podstawowe praktyczne przypadki asymetrii

zbliżone są do siebie wartości obciążeń liniowych części rozrzedzonej i normalnej, czyli im stopień rozrzedzenia $1 - \frac{A_r}{A_n}$ jest bliższy 0. Pomiary wy-

konane na modelach rozrzedzeń całkowicie potwierdziły słuszność wzoru (18) i (18a).

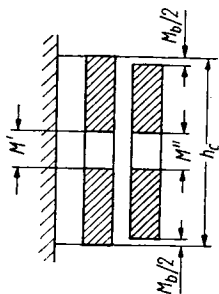
W przypadku kilku rozrzedzeń i luk występujących w obu uzwojeniach rozważanie komplikuje się, jednak — jak wykazało porównanie obliczenia i pomiaru — wystarczy osobne rozważenie obu uzwojeń, a w każdym z nich osobne rozpatrzenie poszczególnych rozrzedzeń.

Najczęściej zastępczy układ ze skupionymi lukami odpowiada jednej z trzech możliwości z rys. 8: $M' = M''$; $M' < M''$; $M' > M''$. Te trzy możliwości, przy jednoczesnym istnieniu luk brzegowych w jednym z uzwojeń, obejmują większość zdarzających się w praktyce asymetrii.

3. ROZPATRYWANIE PRZYPADKÓW $M' \geq M''$

a) $M' = M''$, rys. 9

Przypadek ten odpowiada istnieniu luki skompensowanej (rys. 6);



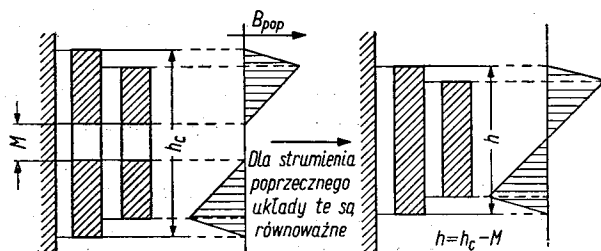
Rys. 9. Przypadek asymetrii brzegowej i luki skompensowanej

oznaczymy dla niego trzy różne wysokości uzwojeń potrzebne podczas rozważań

- h_c — całkowita wysokość uzwojeń po skupieniu ewentualnych rozrzedzeń końcowych,
- h_s — zastępcza wysokość rozproszeniowa obliczana wzorem (11) i wykorzystywana we wzorze (10),
- h — zredukowana wysokość uzwojeń, która powstaje po zredukowaniu w układzie luki skompensowanej (rys. 10) i jest wykorzystywana do określania stopnia asymetrii oraz do wzoru (9).

Wysokość zredukowana służy do wyznaczenia współczynnika β oddającego liczbowo wpływ dodatkowej reaktancji od składowej poprzecznej strumienia rozproszenia.

Rozpatrując przebieg dodatkowego strumienia poprzecznego w przypadku gdy luka skompensowana jest w środku symetrii układu i pomijając zakrzywianie się linii sił, a przez to przechodzenie dodatkowego stru-

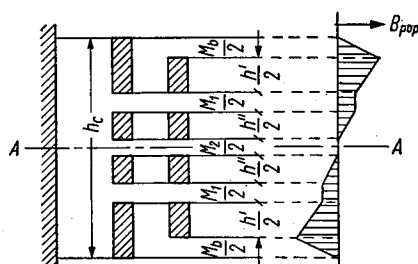


Rys. 10. Przebieg strumienia poprzecznego przy obecności i nieobecności luki skompensowanej w środku uzwojenia

mienia poprzecznego w podłużny, obraz indukcji strumienia poprzecznego można przedstawić jak na rys. 10. Przy takich założeniach indukcja dodatkowego strumienia poprzecznego na przestrzeni luki skompensowanej wynosi 0, co pozwala napisać

$$h = h_c - M. \quad (19)$$

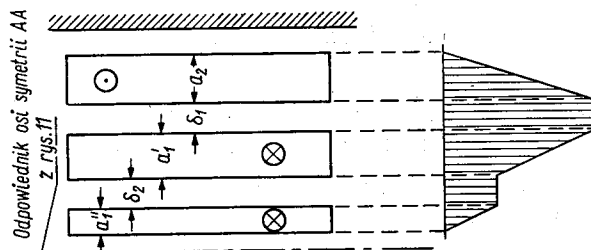
Jeżeli indukcja dodatkowego strumienia poprzecznego w miejscu luki



Rys. 11. Przebieg strumienia poprzecznego przy obecności wielokrotnych luk skompensowanych

skompensowanej jest nierówna 0, określenie zredukowanej wysokości uzwojenia jest znacznie kłopotliwsze.

Przebieg wartości indukcji poprzecznej (do osi symetrii AA) na rys. 11 jest analogiczny z przedstawionym na rys. 12 przebiegiem indukcji podłużnej w uzwojeniach kilkucylindrycznych ze szczelinami poosiowymi (np. cylinder uzwojenia głównego i cylinder uzwojenia regulacyjnego — rys. 12).



Rys. 12. Rozkład strumienia (podstawowego) w uzwojeniu kilkucylindrycznym ze szczelinami poosiowymi

Rysunek 12 jest obrócony o 90° , aby zwiększyć podobieństwo z rys. 11. Wartość połowy zredukowanej wysokości uzwojenia $\frac{h}{2}$ (do osi symetrii AA — rys. 11) podzielona przez 3 odpowiada szczelinie zastępczej dla strumienia poprzecznego δ'_p [patrz objaśnienia do wzoru (5)], czyli

$$\delta'_p = \frac{h}{6} \quad \text{lub} \quad h = 6 \delta'_p.$$

Obliczenie zastępczej szczeliny rozproszeniowej dla układu na rys. 12 przeprowadził prof. E. Jezierski, wyprowadzając wzór¹⁾:

$$\delta' = \frac{a_2}{3} + \delta_1 + \left(1 - c + \frac{c^2}{3}\right) a'_1 + (1 - c)^2 \left(\frac{a''}{3} + \delta_2\right),$$

gdzie

$$c = \frac{a'_1}{a'_1 + a''_1}; \quad (1 - c) = \frac{a''_1}{a'_1 + a''_1}.$$

Ze względu na analogię przebiegu indukcji na rysunkach 11 i 12 można wprowadzić porównanie odpowiednich odcinków (po lewej stronie znaku równości oznaczenia z rys. 12, po prawej — z rys. 11):

$$a_2 = \frac{M_b}{2}; \quad \delta_1 = 0; \quad a'_1 = \frac{h'}{2}; \quad \delta_2 = \frac{M_1}{2}; \quad a''_1 = \frac{h''}{2};$$

po takim podstawieniu $\delta' = \delta'_p$. Wykonanie prostych matematycznych przeróbek upraszczających daje w wyniku wyrażenie dla rys. 11:

$$h = h_c - M_2 - M_1 \left[1 - 3 \left(\frac{h''}{h' + h''} \right)^2 \right]. \quad (20)$$

Jeżeli luka skompensowana M_1 znajduje się daleko od środka symetrii, można uzyskać nawet $h > h_c$.

b) $M' < M''$, rys. 13

W tym przypadku część luki środkowej jest skompensowana, można więc wydzielić

lukę środkową skompensowaną,
lukę środkową nieskompensowaną,
lukę brzegową.

Część skompensowana luki środkowej odpowiada wymiarowi mniejszemu, czyli M' . Różnica $M'' - M' = M_s$ tworzy już lukę środkową nieskompensowaną.

¹⁾ E. Jezierski: „Transformatory. Podstawy teoretyczne”. PWT, Warszawa 1956, str. 88 oraz 89.

$$n = 2 \text{ oraz gdy } \frac{p_m}{p_n} = 1, n = 4.$$

W zakresie pośrednim $0 < \frac{p_m}{p_w} < 1$ nakładają się na siebie dwa układy: I — mający $n = 2$, oraz II — mający $n = 4$. W celu określenia wypadkowego n wydzielamy w obu lukach — środkowej i brzegowej — asymetrię p_m , czyli łącznie $2p_m$, która tworzy układ I ($n = 4$), pozostała w większej luce asymetria $p_w - p_m$ tworzy układ II ($n = 2$). Wypadkowy n uznać można za średnią ważoną obu układów cząstkowych, a więc

$$n = \frac{4 \cdot 2p_m + 2(p_w - p_m)}{2p_m + (p_w - p_m)}. \quad (21)$$

Trzeba tu jednak pamiętać, że wydzielenie z większej luki asymetrii p_m wymaga skupienia tej asymetrii rozrzedzonej na całej rozciągłości luki p_w do rozciągłości mniejszej, wynoszącej p_m . Jak wynika z dyskusji wzoru (18), przy skupieniu asymetrii jej wartość maleje, tak więc, aby otrzymać skupioną lukę p_m , należy z luki p_w wydzielić większą nieco asymetrię p'_m rozrzedzoną na całej rozciągłości p_w . Pozostała część $(p_w - p'_m)$ jest również rozrzedzona na całym wymiarze p_w i po skupieniu ulegnie zmniejszeniu.

W celu uwzględnienia wpływu skupiania luki należy zastosować wzór (18). We wzorze tym p_r odpowiada całemu wymiarowi luki p_w , a wyraz $\left(1 - \frac{A_r}{A_n}\right)$, czyli stopień rozrzedzenia, odpowiada stosunkowi $\frac{p'_m}{p_w}$, czyli stopniowi rozrzedzenia cząstkowej asymetrii p'_m na całym wymiarze p_w .

Stąd:

$$p_m = p_w \frac{1 - \frac{A_r}{A_n}}{1 + \frac{A_r}{A_n} \cdot \frac{p_w}{1 - p_w}} = p_w \frac{\frac{p'_m}{p_w}}{1 + \frac{p_w - p'_m}{p_w} \cdot \frac{p_w}{1 - p_w}};$$

po uproszczeniu otrzymujemy

$$p'_m = \frac{p_m}{1 - (p_w - p_m)} \quad (22)$$

albo

$$p_m = p'_m \frac{1 - p_w}{1 - p'_m}. \quad (22a)$$

Skupiając pozostałą część $(p_w - p'_m)$ należy analogicznie jak we wzorze (22a) pomnożyć ją przez współczynnik $\frac{1 - p_w}{1 - (p_w - p'_m)}$. Ostatecznie wyjściowy wzór (21) zostanie zastąpiony skorygowanym wzorem (23).

$$n = \frac{4 \cdot 2 p_m + 2(p_w - p'_m) \frac{1 - p_w}{1 - (p_w - p'_m)}}{2 p_m + (p_w - p'_m) \frac{1 - p_w}{1 - (p_w - p'_m)}}. \quad (23)$$

Po podstawieniu (22) do (23) i uproszczeniu otrzymujemy:

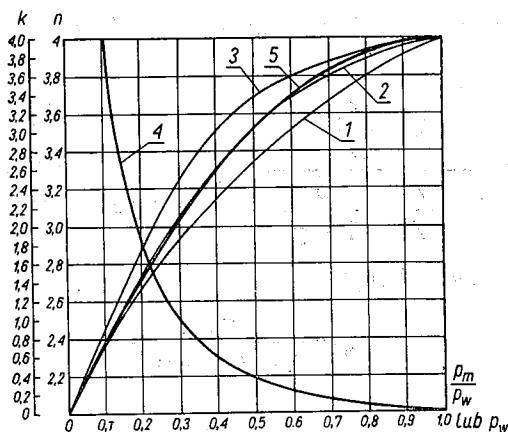
$$n = 2 \frac{4 \left(\frac{p_m}{p_w}\right)^2 + \left(3 \frac{p_m}{p_w} + 1\right) k}{2 \left(\frac{p_m}{p_w}\right)^2 + \left(\frac{p_m}{p_w} + 1\right) k}, \quad (24)$$

gdzie

$$k = \frac{\left(\frac{1}{p_w} - 1\right)^2}{\frac{2}{p_w} - 1}. \quad (25)$$

Wyprowadzone wzory (24) i (25) są skomplikowane, a przez to kłopotliwe do stosowania. Aby otrzymać w oparciu o nie wzór wygodniejszy w użytku, posłużymy się rozumowaniem uproszczonym.

Pamiętając, że $p'_m > p_m$, w członie $(p_w - p_m)$ wzoru (21) powiększymy wyraz p_m stosując przy nim współczynnik α większy od 1 i tak dobrany,



Rys. 14. Przebieg wyrażeń (24), (25) i (26)

$$1 - n = f\left(\frac{p_m}{p_w}\right) \text{ przy } p_w = 0; \quad 2 - n = f\left(\frac{p_m}{p_w}\right) \text{ przy } p_w = 0,29 \ (k = 1);$$

$$3 - n = f\left(\frac{p_m}{p_w}\right) \text{ przy } p_w = 0,5; \quad 4 - k = f(p_w) \text{ [wzór (25)];}$$

$$5 - n = f\left(\frac{p_m}{p_w}\right) \text{ [wzór uproszczony (26)]}$$

aby otrzymać wyrażenie możliwie proste, dające rezultaty zbliżone do wyników wzoru (23), przy asymetriach p_w rzędu $0,25 \div 0,3$. Warunek taki spełnia współczynnik $\alpha = 1 + \frac{p_w - p_m}{2p_w - p_m}$, który zmienia się od 1 przy $p_m = p_w$ (wtedy $p'_m = p_m$) do 1,5 przy $p_m = 0$ ($p'_m = 0$). Wstawiając ten współczynnik do wyrażenia (14) otrzymamy wzór (26):

$$n = \frac{4 \cdot 2 p_m + 2 (p_w - p_m \alpha)}{2 p_m + (p_w - p_m \alpha)} = 2 \left[2 - \left(1 - \frac{p_m}{p_w} \right)^2 \right]. \quad (26)$$

Stosowanie uproszczonego wzoru (26) przy asymetriach mniejszych ($p_w < 0,25$) można uznać za dopuszczalne, gdyż w miarę zmniejszania wartości asymetrii dodatkowa oporność rozproszenia gra coraz mniejszą rolę i nieduży uchyb popełniany przy jej obliczaniu traci stopniowo znaczenie.

W przypadkach $p_w > 0,3$ należy stosować pełny wzór (24). Przebieg wyrażenia (24) dla kilku wartości p_w oraz przebieg wyrażenia (25) przedstawia rys. 14.

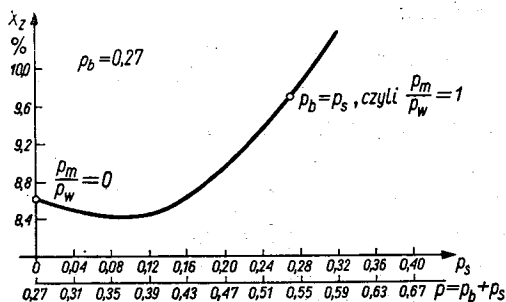
Obliczenie współczynnika β [wzór (9)] wymaga jeszcze określenia współczynnika poprawkowego k_p . Wykresy wartości tego współczynnika dla $n = 2$ podane są w rozdziale 1 na rys. 5c, d. Współczynnik k_p oznaczmy dodatkowo indeksem m lub w , zależnie od tego, czy dotyczy on luki o mniejszym, czy też o większym stopniu asymetrii. Współczynnik k_p zarówno dla luki środkowej, jak i brzegowej należy odczytywać dla całej wartości asymetrii p . Wypadkowy współczynnik k_p jest średnią ważoną z obu odczytanych wartości, to znaczy:

$$k_p = \frac{k_{pm} \cdot p_m + k_{pw} \cdot p_w}{p}.$$

Ponieważ [wzór (9)] $\beta = \gamma \cdot k_p \cdot p^2$, więc ostatecznie:

$$\beta = \gamma \cdot p (k_{pm} \cdot p_m + k_{pw} \cdot p_w). \quad (27)$$

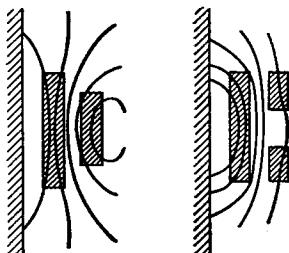
Zarówno obliczenia oporności rozproszenia układu ze stałą asymetrią brze-



Rys. 15. Przebieg wartości oporności biernego rozproszenia dla konkretnego przypadku odpowiadającego rys. 13 przy zmianie wartości luki środkowej

gową, a środkową zmieniającą się od $p_s = 0$ do $p_s = p_b$ [wzory (24), (9), (27)], jak i potwierdzający je pomiar modelowy wykazują występowanie pewnego minimum oporności (rys. 15), to znaczy, że pomimo wzrostu $p = p_b + p_s$ indukcyjność rozproszenia układu początkowo maleje osiagając przy $\frac{p_s}{p_w} \approx 0,3$ swe minimum, a następnie wzrasta. Mamy tu więc zjawisko zbliżone do tego, które występuje w przypadku rozrzedzenia kompensującego łukę zaczepową¹⁾.

Fizyczną interpretację tego zjawiska daje rys. 16 przedstawiający przebieg linii sił przy luce brzegowej i środkowej. Wygięcie linii sił w obu



Rys. 16. Przebieg linii sił rozproszenia przy luce brzegowej i środkowej w uzwojeniu zewnętrznym

przypadkach jest przeciwne, więc jednoczesne występowanie obu luk powoduje „prostowanie się” linii sił, a zatem częściową kompensację strumienia poprzecznego. Z „prostowaniem się” linii sił i zmniejszeniem dodatkowej indukcyjności rozproszenia związane jest oczywiście malenie poosiowych sił elektrodynamicznych. Zjawisko takie zaobserwował Festl [6].

c) $M > M''$, rys. 17.

Podobnie jak w przypadku b) część łuki środkowej jest skompensowana, można więc wydzielić:

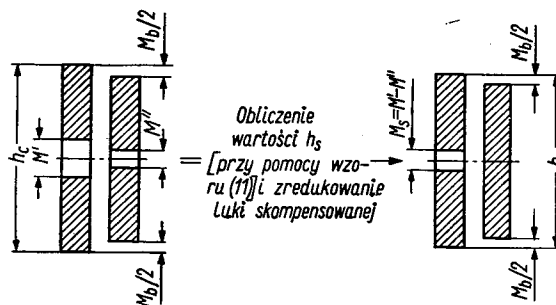
łukę środkową skompensowaną M'' ,

łukę środkową nieskompensowaną $M_s = M' - M''$ w uzwojeniu I,

łukę brzegową M_b w uzwojeniu II.

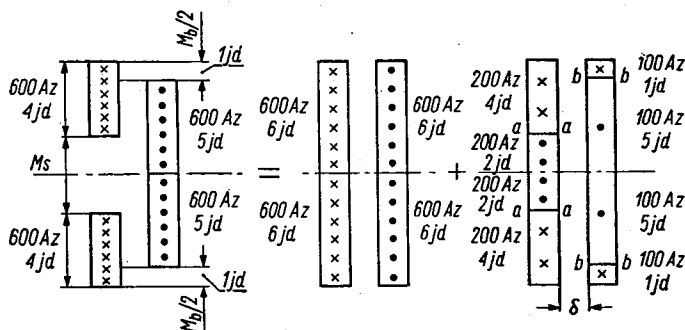
Łukę skompensowaną M'' rozpatrujemy jak w punkcie a), obliczamy rozproszeniową wysokość zastępczą h_s [wzór (11)] oraz wysokość zredukowaną h , w danym przypadku równą $h = h_c - M''$. Otrzymany po tej czynności zredukowany układ uzwojeń, a raczej amperozwojów, należy rozłożyć na składową podłużną i poprzeczną, jak w przykładzie na rys. 18.

¹⁾ Zob. Wykaz literatury [5], rozdział 2.



Rys. 17. Przypadek asymetrii brzegowej i częściowo skompensowanej luki środkowej, dający się sprowadzić do asymetrii brzegowej i środkowej w różnych uzwojeniach

Z rys. 18 widać, że amperozwoje poprzeczne obu uzwojeń wspomagają się wzajemnie.



Rys. 18. Rozkład niesymetrycznych amperozwojów z rys. 17 na składową podłużną i poprzeczną (w pewnym konkretnym przypadku); jd — jednostka długości

Ujęcie matematyczne jest jednak w tym przypadku bardzo trudne, ponieważ:

1. Wartości amperozwojów poprzecznych każdego z uzwojeń różnią się, gdyż są proporcjonalne do odpowiedniego stopnia asymetrii

$$p_b = \frac{M_b}{h} \quad \text{lub} \quad p_s = \frac{M - M'}{h}.$$

2. Osi podziału amperozwojów poprzecznych aa oraz bb są przesunięte względem siebie, co powoduje przechodzenie składowej poprzecznej strumienia w podłużną; osi pokrywają się tylko w jednym szczególnym przypadku, $p_b = p_s = 0,5$.

3. Wzajemne skojarzenie amperozwojów poprzecznych obu uzwojeń ze sobą zależy od odległości δ uzwojeń, a raczej od stosunku $\frac{\delta}{h}$.

4. Wyznaczenie współczynnika Rogowskiego dla tak zestawionego układu jest praktycznie nieosiągalne.

Wszystkie wymienione trudności skłaniają do przybliżonego rozpatrywania tego układu w oparciu o wskazówki doświadczalne. Rozumowanie uproszczone oprzemy na dwóch przypadkach skrajnych:

1. Założymy, że obie asymetrie współdziałają całkowicie i są w pełni ze sobą skojarzone. Współczynnik k_p należy więc określić w poszczególnych uzwojeniach dla całej wartości asymetrii $p_b + p_s$. Współczynnik γ [wzór (9)] oblicza się osobno dla każdego z uzwojeń składowych, a następnie iloczyn $k_p \gamma$ dla całości przyjmuje się jako średnią ważoną z wartości obliczonej dla poszczególnych uzwojeń

$$k_p \cdot \gamma = \frac{k_{ps} \cdot \gamma_s \cdot p_s + k_{pb} \cdot \gamma_b \cdot p_b}{p_b + p_s} \quad (28)$$

Ostatecznie

$$\beta' = k_p \cdot \gamma (p_b + p_s)^2$$

albo

$$\beta' = (k_{ps} \cdot \gamma_s \cdot p_s + k_{pb} \cdot \gamma_b \cdot p_b) (p_s + p_b) \quad (29)$$

2. Założymy, że obie asymetrie działają całkowicie niezależnie od siebie, tak więc:

$$\beta'' = k_{ps} \cdot \gamma_s \cdot p_s^2 + k_{pb} \gamma_b \cdot p_b^2 \approx k_p \cdot \gamma (p_s^2 + p_b^2), \quad (30)$$

gdzie $k_p \cdot \gamma$ określa wzór (28).

Rzeczywista wartość β zawiera się między β' oraz β'' , które — jak można łatwo udowodnić — nie różnią się między sobą więcej niż dwukrotnie. Doświadczenie wskazuje, że wartość β leży bliżej β' niż β'' , wpływa na nią jednak szereg czynników wymienionych na wstępie rozważań. Dla przeciętnie spotykanych kształtów transformatora z uzwojeniami cylindrycznymi można przyjąć, że β jest większe od średniej geometrycznej z β' oraz β'' .

$$\begin{aligned} \beta &> \sqrt{k_p \cdot \gamma (p_s + p_b)^2 k_p \cdot \gamma (p_s^2 + p_b^2)} = \\ &= (k_{ps} \cdot \gamma_s \cdot p_s + k_{pb} \cdot \gamma_b \cdot p_b) \sqrt{p_s^2 + p_b^2} = \beta''' \end{aligned} \quad (31)$$

Ponieważ jednocześnie $\beta < \beta'$, w oparciu o szereg pomiarów przyjęto

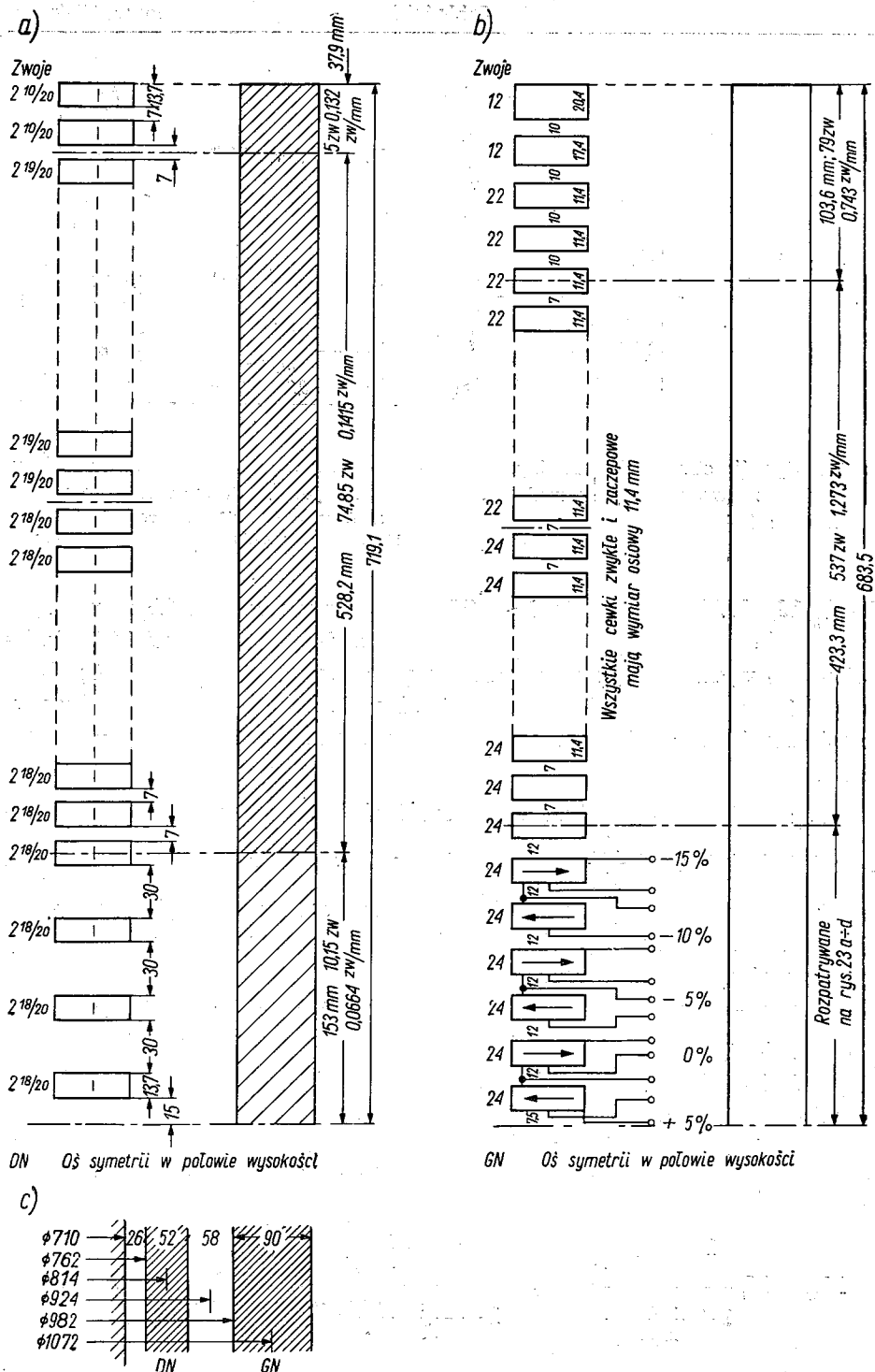
$$\beta = \sqrt{\beta' \cdot \beta'''} \quad (32)$$

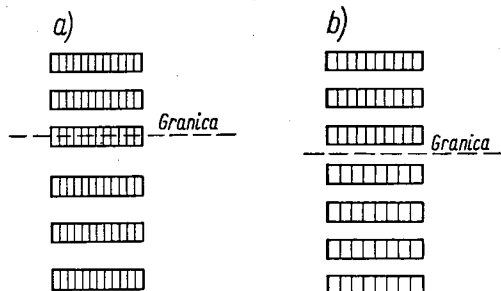
czyli

$$\beta = (k_{ps} \cdot \gamma_s \cdot p_s + k_{pb} \cdot \gamma_b \cdot p_b) \sqrt[4]{(p_s + p_b)^2 (p_s^2 + p_b^2)}.$$

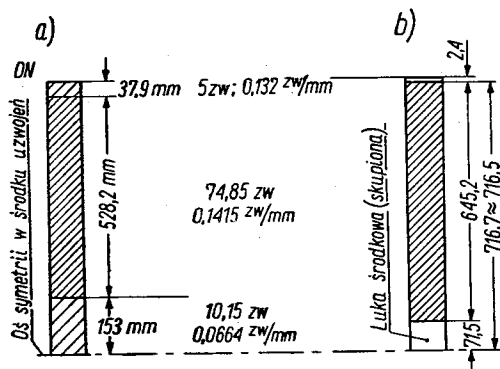
d) Przykład obliczeniowy 1 z komentarzami.

Dany jest transformator o mocy 31,5 MVA i przekładni





Rys. 20. Sposób określania granicy między amperozwojami: a) normalnymi i b) rozrzedzonymi



Rys. 21. Zastąpienie rozrzedzeń w uzwojeniu DN lukami skupionymi
a) Rzeczywisty układ z rozrzedzeniami,
b) Zastępczy układ z lukami skupionymi

$117 \pm \frac{5\%}{-15\%} / 16,5 \text{ kV } \nabla \Delta$ (uzwojenie GN ma zaczepty co $1,67\%$); napięcie zwojowe $e' = 91,8 \text{ V/zw}$; całkowita liczba zwojów GN wynosi 772; liczba zwojów DN wynosi 180; prąd znamionowy fazowy DN wynosi 636 A; układ uzwojeń i zaczeptów przedstawia rys. 19. Należy obliczyć i porównać z wynikami pomiaru napięcia rozproszenia na zaczeptach $+5$; 0 ; -5 oraz -15% .

Przy skupianiu rozrzedzeń w uzwojeniach [wzór (18a)] zamiast stosunku obciążeń liniowych $\frac{A_r}{A_n}$ wykorzystamy stosunek „zagęszczeń zwojów”, czyli liczb zwojów przypadających na 1 cm wysokości w części rozrzedzonej i normalnej. Granicę między różnymi zagęszczeniami przyjmujemy wg następujących kryteriów:

a) jeżeli cewki są jednakowe, odstęp zaś między cewkami różne (rys. 20a), to granica przechodzi przez środek cewki granicznej;

b) jeżeli odstęp są jednakowe, a cewki różne, to granica przechodzi przez środek odstępu granicznego (rys. 20b).

Skupianie rozrzedzeń w uzwojeniu DN przedstawiono na rys. 21. Wychoząc z rys. 21a skupiamy rozrzedzenie brzegowe [wzór (18a)].

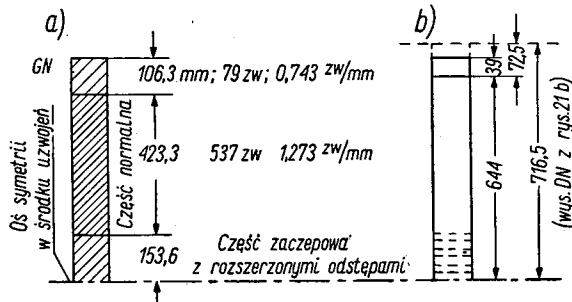
$$M_{sk} = M_r \frac{1 - \frac{A_r}{A_n}}{1 + \frac{A_r M_r}{A_n M_n}} = 37,9 \frac{1 - \frac{0,132}{0,1415}}{1 + \frac{0,132 \cdot 37,9}{0,1415 \cdot 528,2}} = 2,35 \approx 2,4 \text{ mm}.$$

Analogicznie dla rozrzedzenia środkowego otrzymujemy

$$M_{sk} = 153 \frac{1 - \frac{0,0664}{0,1415}}{1 + \frac{0,0664 \cdot 153}{0,1415 \cdot 528,2}} = 71,5 \text{ mm.}$$

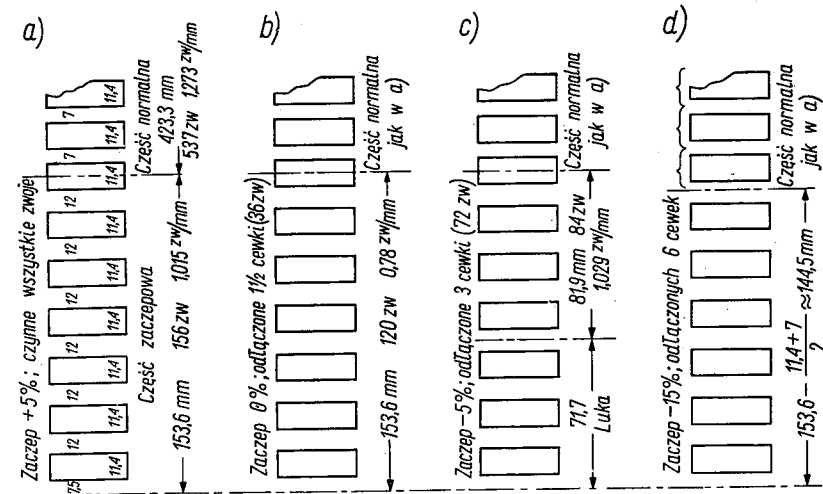
Wymiary uzwojenia po skupieniu przedstawia rys. 21b.

Skupienie rozrzedzenia brzegowego GN przedstawia rys. 22.



Rys. 22. Zastąpienie rozrzedzenia brzegowego w uzwojeniu GN przez lukę skupioną

a) Układ z rozrzedzeniem, b) Zastępcza luka skupiona



Rys. 23. Określenie zastępczej luki środkowej dla poszczególnych zaczepów w uzwojeniu GN

a) Zaczep +5%, b) Zaczep 0%, c) Zaczep -5%, d) Zaczep -15%

$$M_{sk} = 106,5 \frac{1 - \frac{0,745}{1,273}}{1 + \frac{0,743}{1,273} \cdot \frac{106,5}{423,3}} = 38,7 \approx 39 \text{ mm.}$$

Skupienie rozrzedzenia środkowego GN i uwzględnienie luki zaczepowej dla rozpatrywanych czterech zaczepów wykonamy w oparciu o rys. 23.

a) Zaczep + 5‰

$$M_{sk} = 153,6 \frac{1 - \frac{1,015}{1,273}}{1 + \frac{1,015}{1,273} \cdot \frac{153,6}{423,3}} = 24,2 \text{ mm} \approx 24 \text{ mm}.$$

b) Zaczep 0‰

$$M_{sk} = 153,6 \frac{1 - \frac{0,78}{1,273}}{1 + \frac{0,78}{1,273} \cdot \frac{153,6}{423,3}} = 48,5 \text{ mm}.$$

c) Zaczep - 5‰

$$M_{sk} = 71,7 + 81,9 \frac{1 - \frac{1,029}{1,273}}{1 + \frac{1,029}{1,273} \cdot \frac{81,9}{423,3}} = 85,2 \text{ mm} \approx 85 \text{ mm}.$$

d) Zaczep - 15‰; czysta luka skupiona $M_{sk} = 144,5 \text{ mm}$.

Dla obliczania oporności rozproszenia [wzór (10a)] wyznaczmy wartość δ' , l_{sr} [wzór (2)], l_{wlDN} , l_{wlGN} (wg rys. 19c).

$$\delta' = \frac{a_1}{3} + \delta + \frac{a_2}{3} = \frac{5,2}{3} + 5,8 + \frac{9,0}{3} = 10,55 \text{ cm};$$

$$l_{wlDN} = \pi \cdot 81,4 = 255,5 \text{ cm}; \quad l_{wlGN} = \pi \cdot 107,2 = 337 \text{ cm};$$

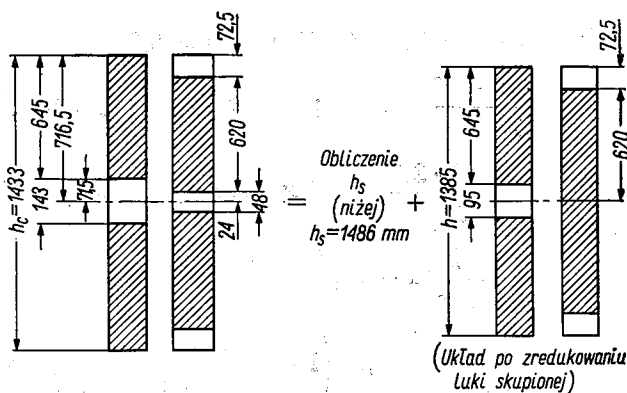
$$l_{sr} = \pi \frac{81,4 \frac{5,2}{3} + 92,4 \cdot 5,8 + 107,2 \frac{9,0}{3}}{10,55} = 298 \text{ cm}.$$

Napięcie rozproszenia w procentach obliczamy przy pomocy wzoru (10a):

$$\begin{aligned} u_x &= 7,9 \cdot f \frac{Iz}{e'} l_{sr} \cdot 10^{-6} \left(\frac{\delta'}{h_s} + \beta \right) = \\ &= 7,9 \cdot 50 \frac{636 \cdot 180}{91,8} 298 \cdot 10^{-6} \left(\frac{10,55}{h_s} + \beta \right) = 146,5 \left(\frac{10,55}{h_s} + \beta \right). \end{aligned}$$

Obliczenie h_s , β i ostatecznie u_x przeprowadzimy osobno dla poszczególnych zaczepów.

a) Zaczep + 5‰, rys. 24, przypadek odpowiadający rysunkowi 17.



Rys. 24. Układ wymiarowy uzwojenia na zaczepie +5% odpowiadający rys. 17

1) Luka skompensowana występuje tu w środku symetrii uzwojeń i wynosi $M = 48$ mm. Wysokość rozproszeniową h_s obliczamy stosując wzór (11), a wysokość zredukowaną h stosując wzór (19).

$$h_s = h_c - M + \frac{a_1 + a_2 + \delta}{3} \left(2 - e^{-\frac{3M}{a_1 + \delta + a_2}} \right) =$$

$$= 1433 - 48 + \frac{200}{3} \left(2 - e^{-\frac{3 \cdot 48}{200}} \right) = 1486 \text{ mm} = 148,6 \text{ cm};$$

$$h = h_c - M = 1433 - 48 = 1385 \text{ mm} = 138,5 \text{ cm}.$$

2) Stopień asymetrii; asymetria brzegowa w uzwojeniu GN

$$p_b = \frac{2 \cdot 72,5}{1385} = 0,1050;$$

asymetria środkowa w uzwojeniu DN

$$p_s = \frac{95}{1385} = 0,0685;$$

asymetria całkowita $p = p_b + p_s = 0,1735$.

3) Współczynnik poprawkowy k_p każdego z uzwojeń odczytujemy z rys. 5 dla całej wartości p .

Parametr $\frac{\delta'}{h_s} = \frac{10,55}{148,6} = 0,071$; dla asymetrii brzegowej w uzwojeniu zewnętrznym interpolujemy współczynnik pomiędzy krzywych 3 i 4 na rys. 5c; $k_{pb} = 1,2$; dla asymetrii środkowej w uzwojeniu wewnętrznym interpolujemy współczynnik pomiędzy krzywych 5 i 6 na rys. 5d; $k_{ps} = 0,95$.

4) Liczba n dla obu uzwojeń, w których występuje asymetria, rys. 24, wynosi 2.

5) Obliczenie γ przeprowadzamy dla każdego z uzwojeń stosując wzór (9):

uzwojenie wewnętrzne DN z luką środkową; $a = 5,2$ cm

$$\frac{\pi \cdot n \cdot a}{h} = \frac{\pi \cdot 2 \cdot 5,2}{138,5} = 0,236; \quad m_n \text{ z rys. 4 wynosi } 0,89;$$

$$\frac{2 \cdot \pi \cdot n \cdot b}{h} = \frac{2 \cdot \pi \cdot 2 \cdot 2,6}{138,5} = 0,236; \quad s_n \text{ z rys. 4 wynosi } 1,27;$$

$$\gamma_s = \frac{1}{n^2} \cdot \frac{l_{wi}}{l_{sr}} \cdot \frac{h}{3a} \left[1 - m_n \left(1 - \frac{m_n \cdot n \cdot \pi \cdot a}{2h \cdot s_n} \right) \right] =$$

$$= \frac{1}{4} \cdot \frac{255,5}{298} \cdot \frac{138,5}{3 \cdot 5,2} \left[1 - 0,89 \left(1 - \frac{0,89 \cdot 2 \cdot \pi \cdot 5,2}{2 \cdot 138,5 \cdot 1,27} \right) \right] = 0,348;$$

uzwojenie zewnętrzne GN z luką brzegową; $a = 9,0$ cm

$$\frac{\pi \cdot n \cdot a}{h} = \frac{\pi \cdot 2 \cdot 9,0}{138,5} = 0,409; \quad m_n = 0,822 \text{ (z rys. 4);}$$

$$\frac{2 \cdot \pi \cdot n \cdot b}{h} = \frac{2 \cdot \pi \cdot 2 \cdot 13,6}{138,5} = 1,235; \quad s_n = 3,48 \text{ (z rys. 4);}$$

$$\gamma_b = \frac{1}{4} \cdot \frac{337}{298} \cdot \frac{138,5}{3 \cdot 9} \left[1 - 0,822 \left(1 - \frac{0,822 \cdot 2 \cdot \pi \cdot 9}{2 \cdot 138,5 \cdot 3 \cdot 48} \right) \right] = 0,316.$$

Mając już stopnie asymetrii, liczbę γ i współczynnik poprawkowy k_p znajdujemy β przy pomocy wzoru (32):

$$\beta = (0,95 \cdot 0,348 \cdot 0,0685 + 1,2 \cdot 0,316 \cdot 0,105) \sqrt[4]{0,1735^2 (0,105^2 + 0,0685^2)} =$$

$$= 0,00925.$$

Napięcie rozproszenia wynosi ostatecznie:

$$u_x = 146,5 \left(\frac{10,55}{148,6} + 0,00925 \right) = 11,7 \, \%.$$

b) Zaczep 0, rys. 25, jak w przypadku poprzednim odpowiada rysunkowi 17.

Prowadząc obliczenie analogiczne do poprzedniego otrzymujemy:

$h_s = 145,4$ cm; $h = 133,6$ cm; $p_b = 0,1087$; $p_s = 0,0345$; $p = 0,1432$;

$k_{pb} = 1,27$; $k_{ps} = 0,94$;

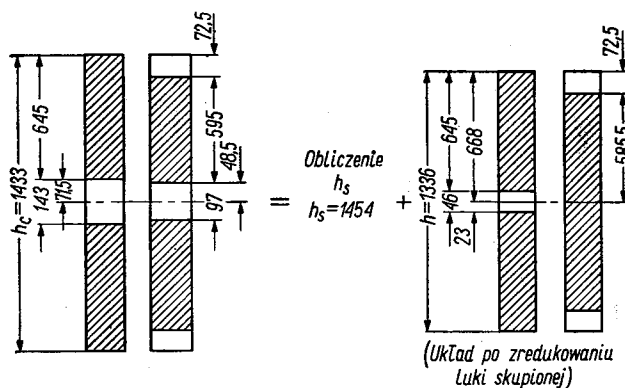
dla GN z luką brzegową $m_n = 0,815$; $s_n = 3,65$; $\gamma_b = 0,312$;

dla DN z luką środkową $m_n = 0,887$; $s_n = 1,29$; $\gamma_s = 0,344$,

$$\beta = (0,94 \cdot 0,344 \cdot 0,0345 +$$

$$+ 1,27 \cdot 0,312 \cdot 0,1087) \sqrt[4]{0,1432^2 (0,1087^2 + 0,0345^2)} = 0,0069$$

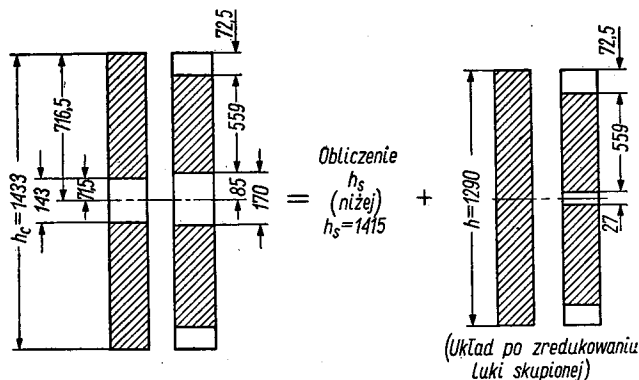
$$u_x = 146,5 \left(\frac{10,55}{145,4} + 0,0069 \right) = 11,52 \, \%.$$



Rys. 25. Układ wymiarowy uzwojenia na zaczepie 0% odpowiadający rys. 17

c) Zaczep -5% , rys. 26, przypadek odpowiadający rys. 13.

1) Luka skompensowana występuje tu w środku asymetrii uzwojeń

Rys. 26. Układ wymiarowy uzwojeń na zaczepie -5% odpowiadający rys. 13

1) wynosi 143 mm. Wysokość rozproszeniową h_s obliczamy przy pomocy wzoru (11), a wysokość zredukowaną h przy pomocy wzoru (12).

$$h_s = h_c - M + \frac{a_1 + a_2 + \delta}{3} \left(2 - e^{-\frac{3M}{a_1 + a_2 + \delta}} \right) = 1433 - 143 +$$

$$+ \frac{200}{3} \left(2 - e^{-\frac{3 \cdot 143}{200}} \right) = 1415 \text{ mm} = 141,5 \text{ cm};$$

$$h = h_c - M = 1433 - 143 = 1290 \text{ mm} = 129 \text{ cm}.$$

2) Stopień asymetrii; całość asymetrii występuje w uzwojeniu GN;

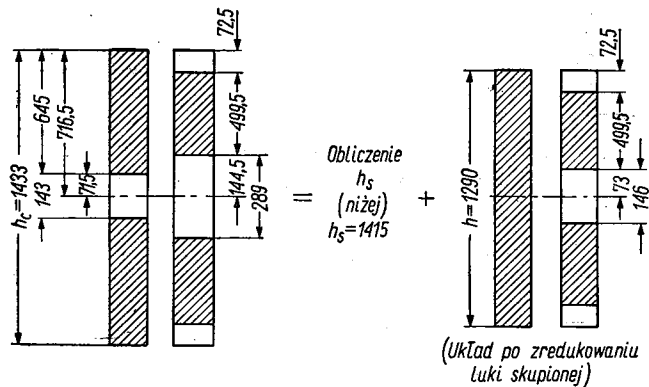
$$\text{asymetria brzegowa } p_b = \frac{2 \cdot 72,5}{1290} = 0,1125;$$

$$\text{asymetria środkowa } p_s = \frac{27}{1290} = 0,0209;$$

asymetria całkowita $p = p_b + p_s = 0,1334$.

Asymetria brzegowa jest większa — p_w , a środkowa mniejsza — p_m .

3) Współczynniki poprawkowe k_p odczytujemy z rys. 5 dla całej wartości $p = 0,1334$. Parametr $\frac{\delta'}{h_s} = \frac{10,55}{141,5} = 0,0745$; dla asymetrii brzegowej współczynnik k_p interpolujemy pomiędzy krzywych 3 i 4 na rys. 5c;



Rys. 27. Układ wymiarowy uzwojeń na zaczepie —15% odpowiadający rys. 13

$k_{pb} = 1,3$; dla asymetrii środkowej odczytujemy współczynnik z krzywej 1 na rys. 5d; $k_{ps} = 1,37$.

4) Liczba n — określamy ją wzorem uproszczonym (26)

$$n = 2 \left[2 - \left(1 - \frac{p_m}{p_w} \right)^2 \right] = 2 \left[2 - \left(1 - \frac{0,0209}{0,1125} \right)^2 \right] = 2,676.$$

5) Obliczenia γ dokonujemy stosując wzór (9):

$$\frac{\pi \cdot n \cdot a}{h} = \frac{\pi \cdot 2,676 \cdot 9,0}{129,0} = 0,586; \quad m_n \text{ z rys. 4 wynosi } 0,76;$$

$$\frac{2 \cdot \pi \cdot n \cdot b}{h} = \frac{2 \cdot \pi \cdot 2,676 \cdot 13,6}{129,0} = 1,775; \quad s_n \text{ z rys. 4 wynosi } 5,90;$$

$$\begin{aligned} \gamma &= \frac{1}{n^2} \cdot \frac{l_{wt}}{l_{sr}} \cdot \frac{h}{3a} \left[1 - m_n \left(1 - \frac{m_n \cdot n \cdot \pi \cdot a}{2h \cdot s_n} \right) \right] = \\ &= \frac{1}{2,676^2} \cdot \frac{337}{298} \cdot \frac{129}{3 \cdot 9} \left[1 - 0,76 \left(1 - \frac{0,76 \cdot 2,676 \cdot \pi \cdot 9}{2 \cdot 129 \cdot 5,90} \right) \right] = 0,204. \end{aligned}$$

Mając stopnie asymetrii, γ oraz k_p znajdujemy β przy pomocy wzoru (27)

$$\begin{aligned} \beta &= \gamma \cdot p (k_{pm} \cdot p_m + k_{pw} \cdot p_w) = 0,204 \cdot 0,1334 (1,3 \cdot 0,1125 + 1,37 \cdot 0,0209) = \\ &= 0,00475. \end{aligned}$$

Napięcie rozproszenia wynosi ostatecznie:

$$u_x = 146,5 \left(\frac{10,55}{141,5} + 0,00475 \right) = 11,6 \, \%$$

d) Zaczep — 15%, rys. 27, przypadek odpowiadający rys. 13.

1) Wysokość rozproszeniowa h_s [wzór (11)] i wysokość zredukowana h , wzór (19); luka skompensowana wynosi tutaj, jak w przypadku poprzednim, 143 mm, obliczenie daje więc ten sam rezultat: $h_s = 141,5$ cm; $h = 129$ cm.

2) Stopień asymetrii; jak widać z rys. 27

$$p_b \approx p_s = \frac{146}{1290} = 0,113; \quad p = p_b + p_s = 0,226.$$

3) Współczynnik poprawkowy k_p ; ze względu na to, że $p_b \approx p_s$, współczynnik k_p przyjmujemy jako średnią z wartości odczytanych dla luki środkowej i brzegowej. Parametr $\frac{\delta'}{h_s} = \frac{10,55}{141,5} = 0,0745$; wartości współczynników odczytujemy jak w punkcie poprzednim:

$$k_{pb} = 1,15; \quad k_{ps} = 1,25; \quad k_p = \frac{1,25 + 1,15}{2} = 1,2.$$

4) Liczba n ; ponieważ $p_b \approx p_s$, więc $n = 4$.

5) Obliczenie, wzór (9)

$$\frac{\pi \cdot n \cdot a}{h} = \frac{\pi \cdot 4 \cdot 9}{129} = 0,877; \quad m_n = 0,645;$$

$$\frac{2 \cdot \pi \cdot n \cdot b}{h} = \frac{2 \cdot \pi \cdot 4 \cdot 13,6}{129} = 2,55; \quad s_n = 14,1;$$

$$\gamma = \frac{1}{4^2} \cdot \frac{337}{298} \cdot \frac{129}{3 \cdot 9} \left[1 - 0,645 \left(1 - \frac{0,645 \cdot 4 \cdot \pi \cdot 9}{2 \cdot 129 \cdot 14,1} \right) \right] = 0,124.$$

Obliczenie β przeprowadzamy stosując wzór (9):

$$\beta = \gamma \cdot p^2 \cdot k_p = 0,124 \cdot 0,226^2 \cdot 1,2 = 0,0076.$$

Napięcie rozproszenia wynosi ostatecznie:

$$u_x = 146,5 \left(\frac{10,55}{141,5} + 0,0076 \right) = 12,02 \, \%.$$

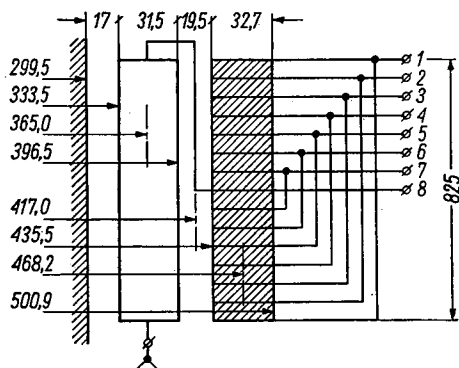
Zestawienie wyników obliczenia i pomiaru przeprowadzonego na stacji prób po wykonaniu transformatora wskazuje na zupełną zgodność rezultatów (tablica).

Tablica 1
Tablica wyników obliczeń i pomiarów

Zaczep	u_x obliczone	u_x pomierzone
—	%	%
+ 5 %	11,70	11,74
0 %	11,52	11,55
— 5 %	11,65	11,70
— 15 %	12,02	12,10

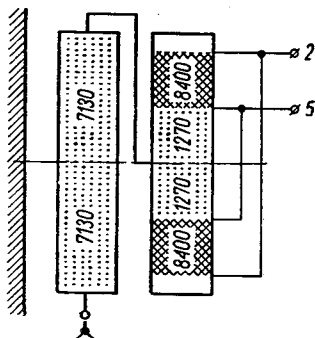
4. PRZYPADEK SZCZEGÓLNY UKŁADU AUTOTRANSFORMATOROWEGO

Szczególnym przypadkiem niejednakowych wartości asymetrii środkowej i brzegowej jest autotransformator przedstawiony na rys. 28. Przy



Rys. 28. Układ wielozaczepowego autotransformatora rozważanego w rozdziale 4

obliczaniu napięcia rozproszenia dla dowolnej kombinacji zacze-
pów (w danym przypadku GN — 1, 2, 3, 4 i DN — 5, 6, 7, 8) należy tu brać pod uwagę nie tylko luki brzegowe powstające przy pracy na niższych zacze-



Rys. 29. Przestrzenne rozmieszczenie amperozwojów przy pracy autotransformatora z rys. 28 na zacze-
pach 2 i 5

pach GN (2, 3, 4), lecz również „wtrącone” amperozwoje przeciwne w części środkowej uzwojenia zewnętrznego przy pracy na wyższych zaczepekach DN (5, 6, 7). Konkretny przykład przestrzennego rozmieszczenia amperozwojów przy pracy na zaczepekach 2 i 5 przedstawia rys. 29. Podkreślić trzeba, że przedstawione na rys. 28 rozwiązanie autotransformatora ma bardzo poważną wadę — możliwość tworzenia wielkich sił dynamicznych poosiowych podczas zwarć.

Rozkład amperozwojów tego układu niesymetrycznego na składową podłużną i poprzeczną wykazuje w zastępczym uzwojeniu krążkowym układu poprzecznego większe obciążenie liniowe na przestrzeni odpowiadającej luce środkowej niż luce brzegowej. Obliczymy współczynnik zwiększenia obciążenia liniowego łuki środkowej z amperozwojami wtrąconymi w stosunku do obciążenia łuki brzegowej „czystej”.

W luce czystej obciążenie liniowe składowej poprzecznej musi być równe i przeciwnie skierowane do obciążenia liniowego składowej podłużnej, natomiast w luce z amperozwojami wtrąconymi musi być ono większe o obciążenie liniowe amperozwojów wtrąconych. Prąd płynący w uzwojeniu wewnętrznym i w części p_s uzwojenia zewnętrznego (łuka środkowa) wynosi I_w . Symetryczny podłużny odpowiednik amperozwojów uzwojenia wewnętrznego wynosi na przestrzeni łuki środkowej $I_w z_w p_s$.

Współczynnik κ powiększenia obciążenia liniowego w luce środkowej wynosi

$$\kappa = \frac{I_w \cdot z_w \cdot p_s + I_w \cdot z_z \cdot p_s \frac{h_w}{h_z}}{I_w \cdot z_w \cdot p_s} = 1 + \frac{z_z \cdot h_w}{z_w \cdot h_z}, \quad (33)$$

gdzie

h_w — wysokość uzwojenia wewnętrznego,

h_z — wysokość uzwojenia zewnętrznego,

z_w — liczba zwojów uzwojenia wewnętrznego,

z_z — liczba zwojów uzwojenia zewnętrznego.

Współczynnik κ jest, jak widać, stały niezależnie od wartości p_s , czyli niezależnie od zaczepek.

Obliczenie dodatkowego napięcia rozproszenia wymaga najpierw sprowadzenia układu amperozwojów poprzecznych zastępujących obie łuki, środkową i brzegową, do tej samej wartości obciążenia liniowego. Obciążenie liniowe dla łuki środkowej jest κ -krotnie większe niż dla brzegowej. Sprowadzając całość do większego obciążenia liniowego zmniejszymy lukę brzegową κ razy. Po dokonaniu sprowadzenia oblicza się wartość β [wzór (9)] podanym poprzednio sposobem. Dla uzyskania oporności rozproszenia [wzory (10) lub (10a)] należy jednak wartość β pomnożyć przez współczynnik κ , w tym bowiem stosunku większe są sprowadzone ampe-

rozwoje układu poprzecznego od amperozwojów z przypadku zwykłej, „czystej” luki.

Obliczenie napięcia rozproszenia dla dowolnego zacze- pu najdogodniej przeprowadza się systemem transformatorowym przy założeniu stałej mo- cy własnej, którą nazywamy mocą odniesienia P_{wo} . Moc taka występuje rzeczywiście przy pracy na zacze- pach skrajnych (1 i 8 na rys. 28), przy czym uzwojeniem podstawowym jest całe uzwojenie wewnętrzne, a do- dawczym całe uzwojenie zewnętrzne.

Przeliczenie napięcia rozproszenia u_x na odpowiednią moc przechod- nią P w układzie autotransformatorowym (u_x) nie przedstawia już żad- nych trudności:

$$P_w = P \frac{\vartheta - 1}{\vartheta};$$

$$u'_x = u_x \frac{P_w}{P_{wo}};$$

$$u_{xa} = u'_x \frac{\vartheta - 1}{\vartheta};$$

a więc

$$u_{xa} = u_x \frac{P}{P_{wo}} \left(\frac{\vartheta - 1}{\vartheta} \right)^2 = u_x \frac{P}{P_{wo}} \left(\frac{U_{GN} - U_{DN}}{U_{GN}} \right)^2. \quad (34)$$

Przykład obliczeniowy 2

Przeprowadzimy obliczenie napięcia rozproszenia autotransformatora przedstawionego na rys. 28 dla zacze- pów 2—5 (rys. 29). Dane autotrans- formatora: Moc przechodnia dla dowolnej pary zacze- pów 6300 kVA; prze- kładnia 38—36,4 — 34,8 — 33,2/31,6 — 30 — 28,4 — 26,8 kV.

Moc własna odniesienia

$$P_{wo} = P_p \frac{U_{GN} - U_{DN}}{U_{GN}} = 6300 \frac{38 - 26,8}{38} = 1860 \text{ kVA}.$$

Płynący przy tym prąd w uzwojeniu wewnętrznym

$$I_w = \frac{1860 \text{ kVA}}{\sqrt{3} \cdot 26,8 \text{ kV}} = 40 \text{ A}.$$

Liczba zwojów w uzwojeniu wewnętrznym $z_w = 944$ zwojów.

Liczba zwojów w uzwojeniu zewnętrznym $z_z = 392$ zwojów (2 grupy rów- noległe), zacze- py co 56 zwojów. W obu uzwojeniach zastosowano równo- mierne rozłożenie zwojów na całej wysokości.

Stosunek $\frac{\delta'}{h_s}$ wynosi dla rozpatrywanego autotransformatora 0,048, śred-

nia długość obwodowa uzwojeń [obliczona przy użyciu wzoru (2)] $l_{sr} = 131$ cm; długość obwodowa uzwojenia zewnętrznego $l_{wt} = 147$ cm. Amperozwoje uzwojenia wewnętrznego odpowiadające mocy własnej odniesienia

$$I_w \cdot z_w = 40 \cdot 944 = 37700 \text{ Az}.$$

Napięcie zwojowe

$$e' = \frac{26800 \text{ V}}{\sqrt{3} \cdot 944 \text{ zw}} = 16,4 \text{ V/zw}.$$

Napięcie rozproszenia obliczamy przy pomocy wzoru (10a):

$$u_x = 7,9 \cdot 50 \frac{37700}{16,4} 131 (0,048 + \beta) = 119 (0,048 + \beta).$$

Obliczenie współczynnika β przeprowadzimy dla przypadku zacze-
pów 2—5.

Zaczepty dzielą obie grupy równoległe na 7 części. Rozpiętość jednej części wyrażona liczbą względną wynosi więc 0,143;

$$\text{luka brzegowa (zaczepek 2) wynosi } \frac{1}{7} = 0,143,$$

$$\text{luka środkowa (zaczepek 5) wynosi } \frac{3}{7} = 0,429.$$

Aby sprowadzić całość do amperozwojów luki środkowej, stosujemy wzór (33)

$$z = 1 + \frac{z_z}{z_w} \cdot \frac{h_w}{h_z} = 1 + \frac{392}{944} \cdot \frac{825}{825} = 1,415.$$

Sprowadzona luka brzegowa (czyli tutaj luka mniejsza):

$$p_m = \frac{0,143}{1,415} = 0,101.$$

Luka większa (środkowa) $p_w = 0,429$, a więc $p = p_m + p_w = 0,530$. Dla znalezienia β [wzór (27)] obliczymy γ [wzór (9a)]; w tym celu wyznaczmy poszczególne czynniki tego wyrażenia.

n obliczamy stosując nieuproszczony wzór (24) i (25) (gdyż $p_w > 0,3$):

$$k = \frac{\left(\frac{1}{p_w} - 1\right)^2}{\frac{2}{p_w} - 1} = \frac{\left(\frac{1}{0,429} - 1\right)^2}{\frac{2}{0,429} - 1} = 0,49;$$

$$\frac{p_m}{p_w} = 0,235; \left(\frac{p_m}{p_w}\right)^2 = 0,0554;$$

$$n = 2 \frac{4 \cdot 0,0554 + (3 \cdot 0,235 + 1) 0,49}{2 \cdot 0,0554 + (0,235 + 1) 0,49} = 2,95;$$

$$\frac{\pi \cdot n \cdot a}{h} = \frac{\pi \cdot 2,95 \cdot 3,27}{82,5} = 0,368; m_n \text{ (rys. 4) wynosi } 0,839;$$

$$\frac{2 \cdot \pi \cdot n \cdot b}{h} = \frac{2 \cdot \pi \cdot 2,95 \cdot 6,8}{82,5} = 1,52; s_n \text{ (rys.4) wynosi } 4,63;$$

$$\gamma = \frac{1}{2,95^2} \cdot \frac{147}{131} \cdot \frac{82,5}{3 \cdot 3,27} \left[1 - 0,839 \left(1 - \frac{0,839 \cdot 2,95 \cdot \pi \cdot 3,27}{2 \cdot 82,5 \cdot 4,52} \right) \right] = 0,2055.$$

β znajdujemy przy pomocy wzoru (27) i w tym celu określamy współczynniki poprawkowe k_p dla łuki brzegowej i środkowej, odczytując je z rys. 5c i 5d dla całej wartości $p = 0,53$.

$k_{pm} = 1; k_{pw} = 1,11$; dla zastosowania wyniku we wzorze (10a) mnożymy uzyskaną wartość β przez współczynnik $\kappa = 1,415$, a więc

$$\begin{aligned} \kappa \cdot \beta &= \kappa \cdot \gamma \cdot p (k_{pm} \cdot p_m + k_{pw} \cdot p_w) = \\ &= 1,415 \cdot 0,2055 \cdot 0,53 (1 \cdot 0,101 + 1,11 \cdot 0,429) = 0,089. \end{aligned}$$

Ostatecznie dla zaczeppów 2—5, czyli $\vartheta = 36,4/31,6$ kV

$$u_x = 119 (0,048 + 0,089) = 16,3 \%.$$

Pomiar wykazał przy tym (po przeliczeniu na moc własną odniesienia) wartość $u_x = 16,2\%$.

Tablica 2

Porównanie wyników obliczeń oraz wyników pomiaru wykonanych dla poszczególnych par zaczeppów

Para zaczeppów	—	1—8	1—7	1—6	1—5	2—8	2—7	2—6	2—5
Obliczone u_x	%	5,7	7,4	11,4	18,1	6,7	7,5	10,9	16,3
Pomierzone u_x	%	5,8	7,5	11,9	17,9	6,6	7,2	10,8	16,2
Uchyb	%	—2	—1,5	—4,2	+1	+1,5	+4	+1	+0,6
Para zaczeppów	—	3—8	3—7	3—6	3—5	4—8	4—7	4—6	4—5
Obliczone u_x	%	9,3	8,9	11,9	17,0	13,6	11,4	14,2	19,3
Pomierzone u_x	%	9,2	8,6	11,4	16,4	13,9	12,0	14,0	19,1
Uchyb	%	+1	+3,5	+4,4	+3,5	—2,2	—5	+1,5	+1

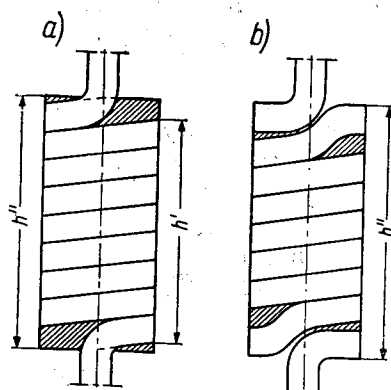
Przeliczenia napięcia rozproszenia na układ autotransformatorowy (u_{xa}) przy stałej mocy przechodniej $P = 6300$ kVA dokonujemy stosując wzór (34):

$$u_{xa} = u_x \cdot \frac{P}{P_{wo}} \left(\frac{U_{GN} - U_{DN}}{U_{GN}} \right)^2 = 16,3 \cdot \frac{6300}{1860} \left(\frac{36,4 - 31,6}{36,4} \right)^2 = 0,96 \%.$$

Zestawione w tablicy 2 porównanie wyników obliczeń wykonanych dla poszczególnych par zaczerwów oraz wyników pomiaru, po sprowadzeniu ich do mocy własnej odniesienia w układzie transformatorowym, potwierdza słuszność zaproponowanej powyżej metody obliczania napięcia rozproszenia.

5. TRANSFORMATOR WIELKOPRĄDOWY Z MAŁĄ LICZBĄ ZWOJÓW

W wielkoprądowych transformatorach z małą liczbą zwojów powstają zasadniczo dwa zagadnienia: zagadnienie luk brzegowych oraz zagadnienie konturu zewnętrznego, którego jednak rozpatrywać tutaj nie będziemy¹⁾. Zagadnienie luk brzegowych rozpatrzmy dla obu alternatyw kon-



Rys. 30. Luki brzegowe w wielkoprądowych transformatorach z małą liczbą zwojów
a) Przypadek spirali ściślej, b) Przypadek wstawki wyrównawczej pod I zwojem

strukcyjnych pokazanych na rys. 30. W przypadku przedstawionym na rys. 30a oznaczmy dwie różne wysokości uzwojenia wielkoprądowego:

h' — całkowita wysokość rozproseniowa uzwojenia odpowiadająca liczbie zwojów z ;

h'' — całkowita wysokość konstrukcyjna uzwojenia odpowiadająca liczbie zwojów $(z + 1)$.

¹⁾ Duże znaczenie konturu zewnętrznego występuje szczególnie w transformatorze zgrzewarkowym, w którym liczba zwojów roboczych wynosi 1 lub 2 (a więc $z^2 = 1$ lub 4); kontur zewnętrzny tworzy praktycznie również 1 zwoj (czysto rozproseniowy). Znaczenie konturu zewnętrznego w miarę wzrostu liczby zwojów roboczych szybko maleje.

Jeżeli wysokość h' uzwojenia wielkoprądowego jest równa wysokości uzwojenia przeciwnego, to uznać można, że niesymetria nie występuje. Jeśli natomiast wysokość h'' równa się wysokości uzwojenia przeciwnego, wtedy występuje niesymetria brzegowa:

$$p_b = \frac{h'' - h'}{h''} = \frac{(z+1) - z}{z+1} = \frac{1}{z+1} \quad (35)$$

przy $n = 2$.

Nierównomierność w rozkładzie niesymetrii wokół obwodu uzwojenia można zaniedbać, gdyż jak wskazuje praktyka błąd popełniany przez to nie ma większego znaczenia.

W przypadku przedstawionym na rys. 30b występuje w zasadzie tylko całkowita wysokość konstrukcyjna uzwojeń h'' , mamy jednak na obu końcach uzwojenia bezpośrednio za pierwszym zwojem luki nierównomiernie rozłożone wokół obwodu uzwojenia. Podobnie jak poprzednio zastępujemy je lukami równomiernymi o średniej wartości wynoszącej po $1/2$ wysokości zwoju i traktujemy skrajny zwój i występującą przy nim lukę jako rozrzedzenie.

Stopień rozrzedzenia $\left(1 - \frac{A_r}{A_n}\right)$ wynosi tu $1 - \frac{1}{1,5} = 0,333$, a wyrażona w liczbach względnych obustronna przestrzeń, na której występuje rozrzedzenie, wynosi $p_r = \frac{1,5 + 1,5}{z+1} = \frac{3}{z+1}$. Rozrzedzenie to można zastąpić luką skupioną stosując wzór (7):

$$p_b = p_r \frac{1 - \frac{A_r}{A_n}}{1 + \frac{A_r}{A_n} \cdot \frac{p_r}{1 - p_r}} = \frac{3}{z+1} \cdot \frac{0,333}{1 + 0,667 \cdot \frac{3}{z-2}} = \frac{z-2}{z(z+1)} \quad (36)$$

W tym przypadku również $n = 2$.

Porównując uzyskane wzory (35) i (36) można stwierdzić, że w przypadku, gdy uzwojenie przeciwne ma wysokość h'' , konstrukcja przedstawiona na rys. 30b łagodzi wpływ niesymetrii brzegowych, lecz ich nie likwiduje.

Przykład obliczeniowy 3.

W transformatorze z uzwojeniem spiralnym 8-zwojowym i wysokością uzwojenia przeciwnego równą wysokości h'' spirali otrzymujemy dla przypadku jak na rys. 30a [wzór (35)]:

$$p_b = \frac{1}{8+1} = 0,111;$$

natomiast dla przypadku jak na rys. 30b skupiona luka brzegowa wynosi [wzór (36)]:

$$p_b = \frac{8 - 2}{8(8 + 1)} = 0,0833.$$

Zastosowanie konstrukcji z rys. 30b zmniejsza więc asymetrię brzegową w tym transformatorze o 25%.

Obliczenia oporności rozproszenia dokonuje się w obu tych przypadkach omówionym w poprzednim rozdziale sposobem.

*Politechnika Łódzka
Katedra Maszyn Elektrycznych*

WYKAZ LITERATURY

1. Rogowski W.: Über das Streufeld und den Streuinduktionskoeffizienten eines Transformators mit Scheibenwicklung und geteilten Endspulen. — Dissertation. Berlin 1908.
2. Heiles F.: Zusätzliche Streuung in Transformatoren. — ETZ, 1933, str. 861.
3. Stephens H. O.: Transformer Reactance and Losses with Nonuniform Windings. — AJEE Trans., 1934, str. 346.
4. Knaack W., Schwaab H.: Zusätzliche Streuung bei Transformatoren. — A.F.E., 1938, str. 470 oraz uzupełnienie na str. 626.
5. Jabłoński Michał: Obliczanie oporności indukcyjnej zwarcia transformatora z niesymetrycznymi uzwojeniami cylindrycznymi. — Rozprawy Elektrotechniczne, 1956, t. I, z. 4, str. 255—273.
6. Festl E.: Ermittlung der Stromkräfte von Transformatoren und ihre Kontrolle auf dem Prüfstand. — ETZ, 1951, str. 173.
7. Jezierski E.: Transformatory — Podstawy teoretyczne. PWT, Warszawa 1956.

М. ЯВЛОНЬСКИ

РАЗВИТИЕ МЕТОДА РАСЧЁТА СОПРОТИВЛЕНИЯ РАССЕЙНИЯ ТРАНСФОРМАТОРА С НЕСИММЕТРИЧЕСКИМИ ЦИЛИНДРИЧЕСКИМИ ОВМОТКАМИ

Резюме

На основании метода продольных и поперечных ампервитков (рис. 1) выведены формулы, а также приведены экспериментальные корректировочные коэффициенты дающие возможность расчёта пассивного сопротивления рассеяния трансформатора даже при весьма значительной асимметрии распре-

ления ампервитков на стержне (напр. в стержневых индукционных печах для плавки металлов). В начальных рассуждениях по этому вопросу, появившихся уже в печати (Л. 5), были разработаны несложные примеры несимметрических распределений, которые могли быть разложены на одинаковые элементы, число которых равно n (рис. 2); разработан также способ учёта перерыва скомпенсированного перерывом или же разрешением в противоположной обмотке (рис. 6). В настоящей статье приведены в сокращении результаты раньше произведенных разработок и в дальнейшем выведена формула необходимая для замещения разрешений сосредоточенными перерывами (18 и 18а). Эта формула отличается от до сих пор применяемой формулы:

до сих пор применяемое выражение $M_{sk} = M_r \left(1 - \frac{A_r}{A_n} \right)$,

выведенное в работе выражение: $M_{sk} = M_r \left(\frac{1 - \frac{A_r}{A_n}}{1 + \frac{M_r}{M_n} \cdot \frac{A_r}{A_n}} \right)$,

где:

M_{sk} — размер вдоль оси сосредоточенного перерыва,

M_r — размер вдоль оси разрешения,

A_r — линейная нагрузка в зоне разрешения,

A_n — линейная нагрузка в нормальной зоне.

Ряд экспериментов оправдал правильность выведенной формулы.

После замещения разрешений сосредоточенными перерывами получается расчётный вид обмотки. В наиболее часто встречаемых практических случаях это один из видов обмотки представленных на рисунке 8. [при чём береговые перерывы могут находиться во внешней или же внутренней обмотке].

В случае большего количества перерывов в обмотке, то есть при $n > 4$ влияние асимметрии очень сокращается и можно без опасения значительных погрешностей упростить схему (по примеру рис. 2) и применить выражение 9 и 10 или 10а, принимая при том корректировочный коэффициент $k_p = 1,2$.

Схемы из рис. 8 были детально обсуждены и были выведены соответственные формулы для расчёта сопротивления рассеяния в таких случаях; тогда находят прежде всего эквивалентную высоту рассеяния h_s , учитывая наличие скомпенсированных перерывов в обмотке (формула 11 или 11а) и в дальнейшем исключают скомпенсированный перерыв из схемы (формула 19 или 20), получая сокращённую высоту обмоток h и окончательный эквивалентный вид схемы (рисунки 10, 13, 17).

Для дальнейшего произведения расчёта в случае представленном на рис. 13 следует определить соответственное значение „ n “, находящееся между 2 и 4 (формула 21 и 25 или формула 26), а также следует определить соответственное значение корректировочного коэффициента k_p . С этой целью вместо формулы 9 применяется формула 27 и окончательное значение сопротивления рассеяния получается обычным путём по формуле 10 или 10а.

В расчёте случая представленного на рис. 17 значение сопротивления находится между результатами двух расчётов:

I-го полагая, что асимметрии в обеих обмотках тесно содействуют,

и II-го полагая, что асимметрии в обеих обмотках независимы друг от друга.

Приведён, основанный на ряде измерений, способ определения соответственного значения коэффициента β (формула 32), нужного для формулы 10 или 10а.

С целью пояснения обсуждаемых случаев произведён примерный расчёт конкретного трансформатора: результаты расчётов хорошо совпадают с результатами измерений.

В главе 4 разработан способ учёта сопротивления рассеяния для любой пары зажимов многоотводного автотрансформатора, представленного на рис. 28; этот способ иллюстрирован примером расчёта.

Во главе 5 разработан способ расчёта сопротивления рассеяния для трансформаторов со значительными токами с небольшим количеством витков z для двух типовых конструктивных вариантов представленных на рис. 30. Самая малая пространственная асимметрия ампервитков, в случае из рисунка 30а, выступает тогда, когда высота рассеяния h' равна высоте противоположной обмотки. Но если конструктивная высота h'' равна высоте противоположной обмотки, тогда выступает асимметрия, степень которой p составляет (формула 35):

$$p = \frac{1}{z + 1}.$$

Для схемы из рис. 30b степень асимметрии в аналогичном случае сокращается к значению (формула 36):

$$p = \frac{1}{z + 1} \cdot \frac{z - 2}{z}.$$

После определения степени береговой асимметрии можно начертить эквивалентную схему обмоток и рассчитать сопротивление рассеяния по методу изложенному в предыдущих главах работы.

M. JABŁŃSKI

DÉDUCTION DE LA MÉTHODE DU CALCUL DE RÉACTANCE DE DISPERSION DU TRANSFORMATEUR AVEC LES ENROULEMENTS CYLINDRIQUES ASSYMETRIQUES

Resumé

En s'appuyant sur la méthode des ampères-tours longitudinaux et transversaux (fig. 1), on a déduit des formules et on a trouvé des coefficients expérimentaux qui permettent de calculer la réactance de fuites magnétiques des transformateurs même en cas d'une très grande assymétrie d'ampères-tours sur la colonne (p. ex. dans les fours à induction à noyau).

Dans les considérations initiales, qui ont été publiées plus tôt [5], on a discuté les cas simples de dispositions assymétriques, qui peuvent être partagées en „ n ” éléments identiques (fig. 2); on y a déduit aussi la méthode de prise en considération d'une lacune compensée par une autre lacune dans le second enroulement (fig. 6).

Dans cet article-ci on a cité les résultats des considérations préalables; puis on a déduit la formule nécessaire pour remplacer des raréfactions d'ampères-tours par des lacunes comprimées (form. 18 et 18a). Cette formule se distingue de la formule utilisée jusqu'à maintenant:

la formule ancienne $M_{sk} = M_r \left(1 - \frac{A_r}{A_n} \right)$,

la formule nouvelle $M_{sk} = M_r \left(\frac{1 - \frac{A_r}{A_n}}{1 + \frac{M_r}{M_n} \cdot \frac{A_r}{A_n}} \right)$

où

M_{sk} — la mesure axiale de la lacune comprimée,

M_r — la mesure axiale de la raréfaction d'ampères-tours,

A_r — ampères-tours par cm de hauteurs dans le rayon de la raréfaction,

A_n — ampères-tours par cm de hauteurs dans le rayon normal.

Toute une série d'essais a confirmé cette formule. Après avoir remplacé les raréfactions par des lacunes comprimées, on obtient la forme d'enroulement pour le calcul.

Pratiquement, le plus souvent c'est une des formes représentées sur la fig. 8 (les lacunes de côté peuvent être situées dans l'enroulement intérieur, ou bien extérieur).

En cas d'un plus grand nombre de lacunes dans l'enroulement, c'est-à-dire ayant $n > 4$, l'influence de l'assymétrie diminue tellement, que sans crainte d'une erreur on peut simplifier la disposition à une forme simple (selon fig. 2) et appliquer les formules 9 et 10 ou 10a en utilisant le coefficient de correction $k_p = 1,2$.

Les dispositions de la figure 8 ont été discutées en détail et on a déduit des formules nécessaires pour le calcul de la réactance de dispersion en ces cas.

Tout d'abord il faut trouver la hauteur suppléante de dispersion h_s ayant égard à la présence des lacunes compensées dans les enroulements (form. 11 ou 11a); puis les lacunes compensées sont éliminées de la disposition (form. 19 ou 20) et on obtient la hauteur réduite de l'enroulement h et la finale forme suppléante de la disposition (fig. 10, 13, 17).

Pour le calcul suivant du cas représenté sur la fig. 13 il faut définir la valeur „ n ” qui est comprise entre les limites 2 et 4 (formules 24 et 25 ou formule 26) et il faut trouver la propre valeur du coefficient de correction k_p .

Pour cela, au lieu de la formule 9 on applique la formule 27 et la valeur définitive de la réactance de dispersion peut être calculée comme d'habitude avec la formule 10 ou 10a.

Pour le calcul suivant du cas de la fig. 17 la valeur de la réactance est comprise entre les résultats de deux calculs:

- I — partant du principe, que les assymétries dans les deux enroulements collaborent strictement;
- II — partant du principe, que les assymétries dans les deux enroulements sont complètement indépendantes.

Pour établir la valeur convenable de la réactance il faut trouver le coefficient β nécessaire pour la formule 10 ou 10a. On a donné le moyen empirique pour établir ce coefficient résultant du calcul I et II (form. 32).

Comme illustration de ces cas discutés on a donné un exemple du calcul pour un certain transformateur. Les résultats des calculs s'accordent bien avec la mesure.

Dans le chapitre 4 on a présenté la méthode de calculer la réactance de dispersion pour une paire quelconque de prises d'autotransformateur de réglage représenté sur la fig. 28; ce cas a été illustré aussi par le calcul.

Dans le chapitre 5 on a déduit la méthode de calculer la réactance de fuites pour le transformateur à très grand courant et avec un petit nombre de spires z en cas de deux solutions de construction (fig. 30a et 30b).

La moindre assymétrie spatiale d'ampères-tours en cas représenté sur la fig. 30a a lieu, quand la hauteur suppléante de dispersion h' est égale à la hauteur du second enroulement.

Quand la hauteur de construction „ h ” est égale à la hauteur du second enroulement, c'est une certaine assymétrie, qui se forme ayant degré p (form. 35)

$$p = \frac{1}{z + 1}.$$

Pour la disposition de la fig. 30b le degré d'assymétrie analogue diminue jusqu'à la valeur (form. 36)

$$p = \frac{1}{z + 1} \cdot \frac{z - 2}{z}.$$

Après la définition du degré d'assymétrie de côté on peut dessiner la forme suppléante de l'enroulement et calculer la réactance avec la méthode donnée dans les chapitres précédents.

M. JABŁOŃSKI

DEVELOPMENT OF THE METHOD FOR CALCULATING THE STRAY REACTANCE OF ASYMMETRIC CYLINDRICAL TRANSFORMER WINDINGS

Summary

In terms of the method of longitudinal and transversal ampere-turns, formulae are evolved and experimental correction factors are presented, which make it possible to calculate the stray reactance in transformers, even in the case of large asymmetry of the ampere-turn distribution on the column (e.g. core type induction furnaces for melting metals). Former considerations, already published (L 5) concerned cases of simple asymmetric systems, which can be resolved into identical asymmetry elements, the number of which is „ n ” (Fig. 2). A method is also elaborated taking into consideration a gap compensated by a gap or thinning in the ampere turn distribution of the second winding (Fig. 6). The author presents in brief results of former considerations, and then derives a formula (18) and (18a) necessary for replacing thinnings by lumped gaps. This formula differs from the expression used till now:

$$\text{the former expression: } M_{sk} = M_r \left(1 - \frac{A_r}{A_n} \right),$$

the expression introduced in the paper: $M_{sk} = M_r \left(\frac{1 - \frac{A_r}{A_n}}{1 + \frac{M_r}{M_n} \cdot \frac{A_r}{A_n}} \right)$,

where:

M_{sk} — axial dimension of lumped gap,

M_r — axial dimension of thinning,

A_r — linear load in thinned area,

A_n — linear load in normal area.

A series of experiments have confirmed the correctness of the deduced expression.

After replacing the thinned areas by lumped gaps, a winding form ready for calculation is obtained. In most practical cases it has one of the forms shown in Fig. 8 (the border gaps may be in the external or infernal windings).

In the case of a larger number of gaps, i.e. when $n > 4$, the asymmetry effect is markedly decreased and the system may be reduced to a simple form (according to Fig. 2) and the expressions (9) and (10) or (10a) with the correction factor $k_p = 1,2$ may be used without fear of making greater errors.

The systems of Fig. 8 are considered in detail and suitable formulae for calculating the stray reactance in these cases are introduced; first the equivalent stray height h_s is found, taking into account the presence of compensated gaps in the winding [formulae (11) or (11a)], then the compensated gap is eliminated from the system [formula (19) or (20)], and the reduced height of windings h and the final equivalent form of the system are obtained (Figs. 10, 13 and 17).

In the further course of calculation of the case shown in Fig. 13, the suitable value of n comprised between 2 and 4 must be determined [formulae (24) and (25) or (26)] and a suitable value of the correction factor k_p should be established. For this purpose formula (27) is used instead of formula (9) and the final value of the stray reactance is obtained as usually from formula (10) or (10a).

The value of the reactance calculated for the case shown in Fig. 17, is comprised between the results of two calculations:

I — assuming that the asymmetries in both windings interact closely,

II — assuming that the asymmetries in both windings are mutually independent.

A method of determining a suitable value of the coefficient β [formula (32)], necessary for the formula (10) or (10a), based on a series of measurements is given.

In order to illustrate the discussed cases the numerical calculation of a concrete transformer is carried out; the results are in good agreement with the results of measurements.

In chapter 4 a method of calculating the stray reactance of a freely chosen pair of terminals of a multi-tap autotransformer shown in Fig. 28 is elaborated and illustrated on an example.

A method of calculating the stray reactance of high-current transformers with a small number of turns „ z ” for two typical designs shown in Fig. 30 is presented in chapter 5. The minimum spatial asymmetry of ampere turns for the case of Fig. 30a occurs, when the stray height h' is equal to the height of the second winding. If, however, the construction height h'' is equal to the height of the second winding, an asymmetry occurs, the degree p of which is [formula (35)]:

$$p = \frac{1}{z + 1}.$$

For the system shown in Fig. 30b the asymmetry degree for an analogical case is reduced to the following value [formula (36)]

$$p = \frac{1}{z + 1} \cdot \frac{z - 2}{z}.$$

After establishing the degree of the border asymmetry, an equivalent picture of the windings may be drawn and the stray reactance can be calculated according to the method presented in the former chapters of this paper.

629.135.15:621.316.9

JERZY BADER, JANUSZ LECH JAKUBOWSKI, JERZY ZIELIŃSKI

Zagrożenie piorunowe szybowców

Rękopis dostarczono 11. 3. 1959

W artykule omówiono mechanizm powstawania ładunków na szybowcach, a mianowicie drogą influencji i drogą ładowania przez tarcie i przez rozbryzgiwanie kropeł wody na powierzchni szybowca. Ponadto zestawiono parametry pioruna trafiającego w szybowiec. Na podstawie analizy przydatności metod modelowania, polegających na odtworzeniu czoła wyładowania piorunowego za pomocą elektrody ostrzowej i zastąpieniu szybowca modelem w skali 1:20, wykonano serię pomiarów posługując się generatorem udarowym Instytutu Elektrotechniki na napięcie 2800 kV i o energii 32 kW. Badania te miały na celu wybranie miejsc dla założenia zwodów na szybowcu oraz określenie ich przestrzeni chronionych. Badania powyższe umożliwiły wyciągnięcie wniosków co do rozwiązania ochrony odgromowej i jej skuteczności na konkretnym szybowcu typu „Bocian”.

1. WSTĘP

Sport szybowcowy w okresie powojennym osiągnął w Polsce niezwykle wysoki stopień rozwoju. Ilustruje to choćby zwiększenie sprzętu — ze 120 szybowców w 1948 do 710 w r. 1958 — i jego unowocześnienie. Jest to zarówno sport wyczynowy (uzyskanie przez Polaków 34 odznak z trzema diamentami oraz zdobycie Mistrzostwa Świata w klasie Standart), jak i sport już w pewnej mierze masowy (2820 pilotów).

Rozwój szybownictwa zwrócił uwagę na niebezpieczeństwo lotów pod chmurami burzowymi i w tych chmurach, lotów bardzo atrakcyjnych dla szybowników, gdyż występujące tam wznoszące się prądy powietrzne umożliwiają im osiągnięcie dużych wysokości. Kroniki polskiego szybownictwa niestety mają już do zanotowania dwa wypadki śmierci załogi i zniszczenia szybowców wywołane uderzeniem pioruna.

Opisany stan zagrożenia skłonił Ligę Przyjaciół Żołnierza a następnie Aeroklub Polskiej Rzeczypospolitej Ludowej do powierzenia Katedrze Wysokich Napięć Politechniki Warszawskiej sprawy ochrony szybowców

i pilotów od skutków uderzeń pioruna. Wynikiem dwuletnich prac laboratoryjnych na ten temat będą dwa artykuły w „Rozprawach Elektrotechnicznych”: „Zagrożenie piorunowe szybowców” i „Ochrona odgromowa szybowców”. Prace badawcze biegnące dalej dotyczą rozładowywania ulotowego samolotów i szybowców oraz ochrony radiostacji.

Wszelkie rozważania ochrony odgromowej szybowców muszą wyjść z założenia, że uderzenie pioruna w szybowiec jest zjawiskiem zasadniczo nie do uniknięcia. Można wprawdzie zmniejszyć zagrożenie trafienia pioruna w szybowiec, mianowicie przez usuwanie z niego ładunków związanych z dużymi natężeniami pola przy szybowcu, ale są przypadki, gdy ten zabieg nie jest skuteczny. Toteż ochrona szybowca liczy się z trafieniem przez piorun i dąży do tego, aby piorun uderzył w specjalnie konstrukcyjnie dostosowane miejsce szybowca oraz aby prąd pioruna przepływał przez szybowiec nie wywołując rażenia pilota i uszkodzeń konstrukcji (wypalenie dziur, rozszczepianie tworzywa, stopienie łożysk i przez to unieruchomienie sterów, uszkodzenie instalacji elektrycznej i nawigacyjnej).

Zagadnienie ochrony szybowców jest zasadniczo zbliżone do zagadnienia ochrony samolotów, jednak różnice w obu przypadkach są bardzo istotne; samoloty posiadają zwykle pełną obudowę metalową, stanowiącą prawie doskonałą klatkę Faradaya. Przy trafieniu ich przez piorun, prądy pioruna nie przenikają w ogóle do wnętrza lub przenikają w niewielkim stopniu (trafienie w nieodpowiednio zabezpieczoną antenę). Trafienie do wnętrza przez okno jest niezwykle mało prawdopodobne. Antycypując dalsze rozważania, można na wstępie już stwierdzić, że koncepcja ochrony szybowców polega na możliwym upodobnieniu ich pod względem osłony do samolotu metalowego, a więc na założeniu odpowiedniej sieci zwodów, stanowiącej częściową klatkę Faradaya.

Niniejszy artykuł ujmuje tylko część omówionego wyżej tematu, mianowicie omawia zagrożenie piorunowe szybowców, w szczególności warunki powstawania dużych natężeń pól elektrycznych przy szybowcu i określanie wybiorności pioruna, to jest wyznaczanie miejsc predestynowanych do trafienia przez piorun. Znajomość tych miejsc pozwala na prawidłowe zaprojektowanie ochrony.

Prace Katedry Wysokich Napięć nie ograniczyły się do badań podstawowych, ale objęły również wytyczne projektu instalacji piorunochronnej szybowca typu „Bocian”. Wytyczne te zostaną tutaj również omówione szkicowo, jako ilustracja praktycznego zastosowania badań laboratoryjnych.

Rozdziały 2, 3 i 4 niniejszego artykułu mają charakter sprawozdawczy i zawierają krytyczny przegląd literatury z nielicznymi uogólnieniami własnymi. Rozdziały 5 i 6 opisują badania własne autorów.

2. POWSTAWANIE ŁADUNKÓW ELEKTRYCZNYCH NA SZYBOWCACH

Podczas lotu szybowca występuje szereg zjawisk powodujących wystąpienie na jego powierzchni dużego natężenia pola elektrycznego K , a więc dużej gęstości powierzchniowej ładunku elektrycznego (zgodnie z zależnością $K = \frac{\sigma}{\varepsilon}$, gdzie ε — przenikalność dielektryczna ośrodka).

W dalszym ciągu niniejszego artykułu będą wielokrotnie rozważane albo wartości natężeń pola przy powierzchniach metalowych samolotu i szybowca, albo gęstości ładunku na tych powierzchniach. W rozważaniach tych nie należy zapominać, że wartości liczbowe wzmiankowanych natężeń pola i gęstości ładunku różnią się stałym współczynnikiem proporcjonalności.

Jedno ze zjawisk powodujących powstawanie ładunku na szybowcach występuje jedynie przy przelotach przez obszary burzowe (obszary o dużym natężeniu pola elektrycznego), inne mogą zachodzić niezależnie od istnienia obszarów burzowych. Ładunek powierzchniowy związany z pierwszym wzmiankowanym zjawiskiem wynika z rozdziału ładunków w częściach metalowych szybowca pod wpływem zewnętrznego pola elektrycznego, przy sumarycznym ładunku szybowca równym zeru, co określamy terminem influencji. Druga grupa zjawisk prowadzi do ładunku powierzchniowego przez wprowadzenie do szybowca ładunków dodatkowych z zewnątrz. W tym przypadku jego sumaryczny ładunek jest różny od zera.

Najpierw zajmiemy się wprowadzeniem do szybowca dodatkowych ładunków z zewnątrz. Ma ono znaczenie praktyczne, gdyż występuje stosunkowo często. Zjawiska te były dość szeroko badane; niżej podajemy próbę ułożenia przejrzystego wzoru, nie spotykanego w literaturze, pozwalającego na łatwe określanie roli poszczególnych parametrów.

Gromadzenie się ładunków na szybowcu związane jest z dopływem ładunków z zewnątrz, czyli z prądem elektrycznym ładującym. Szybowiec posiada pojemność własną¹⁾, ładowanie go prowadzi więc do powstawania na nim potencjału elektrycznego. Przy dostatecznie dużym potencjale zjawia się na elementach szybowca o dużej krzywiznie ulot i powstaje prąd rozładowujący szybowiec (prąd upływu). Prąd ten rośnie począwszy od wartości zerowej w miarę wzrostu potencjału szybowca. Gdy prąd ładujący zrówna się z rozładowującym (zakładamy, że prąd ładujący jest stały), potencjał szybowca osiąga maksimum.

¹⁾ Pojemność własna — pojemność określona w założeniu, że druga elektroda jest w nieskończoności.

Prąd upływu (ulotu) można wyrazić przybliżonym wzorem, opartym o wzór I. S. Townsenda [1], stosowany dla układu walców koncentrycznych:

$$i_u = a(V - V_0) V, \quad (1)$$

gdzie:

- a — współczynnik skuteczności rozładowywania,
- V_0 — potencjał początkowy jonizacji (potencjał jonizacji),
- V — potencjał szybowca; $V \geq V_0$.

Zakładając, zgodnie z podanym wyżej równaniem, że prąd upływu i_u równa się prądowi ładowania i_l , otrzymamy:

$$i_l = a(V - V_0) V,$$

oraz

$$V = \frac{V_0}{2} + \sqrt{\frac{V_0^2}{4} + \frac{i_l}{a}} \quad (2)$$

przy $i_l > 0$.

Ze wzoru tego widać, iż potencjał szybowca będzie tym mniejszy, im mniejszy będzie potencjał jonizacji, im współczynnik skuteczności rozładowywania będzie większy i im prąd ładujący będzie mniejszy. Wartości potencjału jonizacji szybowca i współczynnika skuteczności odprowadzania ładunków zależą głównie od konstrukcji szybowca. Wartości prądu ładującego szybowiec zależą głównie od czynników atmosferycznych oraz od prędkości lotu szybowca.

Natężenie pola elektrycznego na powierzchni szybowca, przy pominięciu jonizacji powietrza otaczającego szybowiec i uwzględnieniu szeregu założeń upraszczających, można wyrazić ogólnie wzorem:

$$K = b \cdot V, \quad (3)$$

gdzie współczynnik b jest funkcją wymiarów geometrycznych. Podstawiając wyrażenie (3) do wzoru (2) otrzymamy dla $K > K_0$ wzór

$$K = \frac{K_0}{2} + \sqrt{\frac{K_0^2}{4} + \frac{i_l b^2}{a}}, \quad (4)$$

gdzie:

- K_0 — początkowe natężenie jonizacji, zależne od kształtu szybowca i rozpatrywanego miejsca na szybowcu,
- b — współczynnik zależny od tych samych parametrów.

Tak więc dla danego szybowca i dla określonego na nim miejsca natężenie pola elektrycznego jest funkcją prądu ładującego.

Przy dużych prądach ładujących można się spodziewać, że natężenie pola elektrycznego na powierzchni szybowca będzie mniejsze, niż wynika

to z zależności (4). Przy takich prądach potencjał szybowca osiąga tak wielkie wartości, że powstaje bardzo intensywny ulot, a związana z nim jonizacja powietrza powoduje zmniejszenie współczynnika b we wzorze (3).

Rozpatrzmy obecnie przyczyny powstawania prądu ładującego i związane z nimi wartości tego prądu.

W wyniku obszernych badań [2], [3], [5] stwierdzono następujące przyczyny gromadzenia się ładunku elektrycznego na szybowcach:

a) bezpośrednie zetknięcie się szybowca z obszarem ładunków przestrzennych,

b) rozbryzgiwanie krople wody przy zetknięciu się ich z powierzchnią szybowca,

c) tarcie (stykanie się) szybowca z suchymi cząstkami śniegu i lodu.

W samolotach występują ponadto dodatkowe przyczyny gromadzenia się ładunku elektrycznego, a mianowicie wyrzucanie gazów spalinowych oraz ruch śmigła [2]. Nie odgrywają one zresztą poważniejszej roli.

Omówimy najpierw bezpośrednie zetknięcie się szybowca z obszarem ładunków przestrzennych. Mechanizm wymiany ładunków między szybowcem a otoczeniem jest w tym przypadku dość złożony. Można go w pewnym zakresie porównać do mechanizmu przenoszenia ładunków występującego w generatorze Van de Graaffa. L. P. Harrison [3] podaje, że ładowanie tego typu występuje zwłaszcza tam, gdzie prądy wznoszące niosą ze sobą ładunki elektryczne, np. przy przelotach nad miastami. Wydaje się więc, że to zjawisko może często występować na szybowcach, które odbywają loty w tak zwanych kominach (w obszarach, w których występują prądy wznoszące). Ponieważ prądy wznoszące związane są przeważnie z małą gęstością ładunków elektrycznych, występuje mały prąd ładujący, a więc mały potencjał szybowca.

Zjawisko ładowania się mas metalowych wskutek rozbijania się o nie krople wody zostało wyjaśnione przez P. Lenarda [4]. Wyjaśnił on wprawdzie elektryzowanie się krople deszczu przy rozrywaniu ich przez silny wiatr, ale mechanizm przy rozrywaniu się krople trafiających w masy metalowe jest analogiczny. Ładowanie następuje mianowicie wskutek oddzielania się krople warstwy powierzchniowej. Ta warstwa ładuje się ujemnie w stosunku do wewnętrznej warstwy, pozostającej na obiekcie metalowym. Obiektowi temu zostaje przekazany zatem ładunek dodatni. Prąd ładujący jest w tym przypadku uzależniony od intensywności i rodzaju deszczu, od prędkości wiatru i prędkości szybowca. Jest on większy od prądu ładującego występującego przy bezpośrednim zetknięciu się szybowca z chmurą ładunków, natomiast jest mniejszy od prądu ładującego występującego przy tarcia szybowca o cząstki stałe.

Ilość ładunku elektrycznego gromadzącego się pod wpływem tarcia

szybowca o suche cząstki śniegu i lodu zależy od stanu powierzchni szybowca, gęstości padania śniegu lub cząstek lodu, kąta padania w stosunku do powierzchni szybowca, prędkości szybowca względem cząstek opadu oraz temperatury otoczenia. Na podstawie badań [5] został opracowany wzór, pozwalający na obliczanie prądu ładowania:

$$i_l = c \cdot v^n, \quad (5)$$

gdzie:

- c — stała zależna od konstrukcji i wymiarów szybowca,
- v — prędkość szybowca,
- n — stała zależna od rodzaju powierzchni, jej wielkości i temperatury otoczenia; dla większych powierzchni jest ona rzędu 3 i wykazuje maksimum przy temperaturze ok. -9°C .

W wyżej wspomnianych badaniach zarejestrowano dla samolotu metalowego o dużej powierzchni (miarą jej jest pojemność własna 780 pF) przy słabym opadzie śnieżnym prąd ładujący 0,10 mA, zaś przy średnim opadzie — 0,155 mA. W szybowcach ze względu na mniejsze wymiary (pojemność własna rzędu 200 pF) należy się spodziewać mniejszych prądów ładujących związanych ze zjawiskiem tarcia. Trzeba jednak podkreślić, że przytoczone wartości prądu ładującego nie są największe. Większe wartości występują przy bardzo intensywnych opadach, zwłaszcza cząstek lodu.

Zjawisko gromadzenia się ładunku pod wpływem tarcia nie doprowadza do największej gęstości ładunku powierzchniowego, jednak jest prawdopodobnie najbardziej częstą przyczyną ładowania szybowca. Największe gęstości są związane ze zjawiskiem influencji.

Zjawisko influencji polega na rozdzieleniu ładunków wewnętrznych w naładowanej masie metalowej umieszczonej w polu elektrycznym. Prowadzi ono, jak wspomniano, do największych gęstości ładunku powierzchniowego, to znaczy do największego natężenia pola elektrycznego na powierzchni szybowca. Trzeba jednak podkreślić, że obszary w atmosferze o dużym natężeniu pola elektrycznego, w których influencja jest silna, nie są na ogół rozległe i szybowiec rzadko może trafić na te obszary. Tak więc duże natężenie pola elektrycznego na powierzchni szybowca związane z influencją występuje raczej rzadko.

Natężenie pola elektrycznego na powierzchni szybowca K jest w omawianym przypadku wprost proporcjonalne do natężenia pola elektrycznego K_z obszaru, przez który szybowiec przelatuje:

$$K = K_z \cdot d, \quad (6)$$

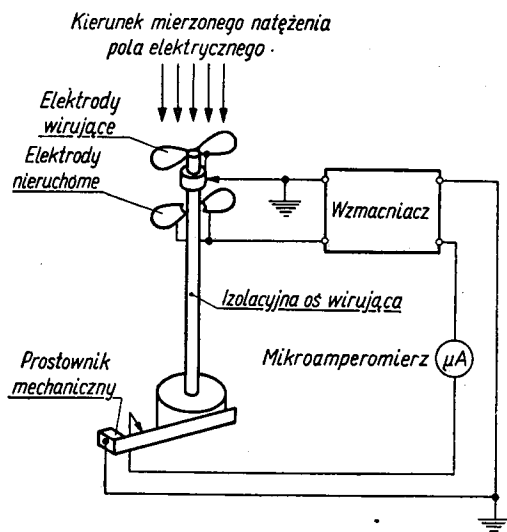
gdzie d — współczynnik zależny od kształtu szybowca, jego położenia w stosunku do linii pola elektrycznego oraz od miejsca na szybowcu. Za-

leżność (6), analogicznie jak (4), ulegnie zmianie w przypadku uwzględnienia silnej jonizacji powietrza otaczającego szybowiec.

Wartości natężenia pola elektrycznego na powierzchni szybowca, prądu ładującego i prądu upływu, ładunku gromadzącego się na szybowcu, poziomu zakłóceń radiowych oraz wielkości meteorologicznych były szczegółowo badane zarówno w czasie lotów, jak i w warunkach laboratoryjnych. Badania wielkości meteorologicznych i poziomu zakłóceń radiowych przeprowadza się przy pomocy powszechnie stosowanych metod. Do pomiaru pozostałych wielkości stosuje się metody specjalne, które tutaj zostaną omówione w sposób ogólny.

Prąd dopływający do szybowca można mierzyć wstawiając między część metalową odizolowaną od niego a jego główną masę odpowiedni przyrząd pomiarowy. Udział prądu ładującego, wywołanego naładowanymi cząstkami, można określić przy pomocy elektrody wykonanej w postaci umieszczonego w pobliżu powierzchni samolotu pierścienia, przez który przelatują cząstki naładowane [5].

Prąd upływu mierzy się za pomocą mierników prądu włączonych między samolot a umieszczone na nim rozładowywacze — specjalne ostrzowe elektrody o początkowym potencjale jonizacji znacznie niższym od początkowego potencjału jonizacji innych części.

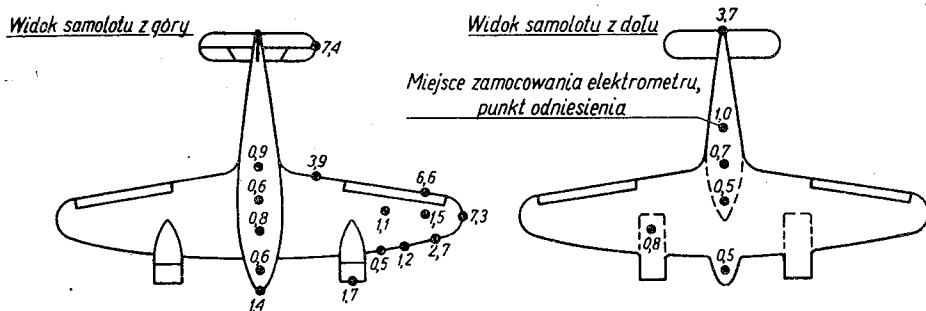


Rys. 1. Układ do pomiaru natężenia pola elektrycznego w pobliżu samolotu lub szybowca

Pomiar natężenia pola elektrycznego w pobliżu powierzchni szybowca przeprowadza się za pomocą elektrometru wirującego. Schemat elektrometru stosowanego do tego celu przez R. C. Waddela i współpracowników [6] uwidoczniiony jest na rys. 1. Elektrody wirujące połączone z masą

szybowca, obracają się wraz z osią wykonaną z materiału izolacyjnego, napędzaną przez mały silniczek z prędkością 1000 obrotów na minutę. W niewielkiej odległości od tych elektrod umieszczone są elektrody nieruchome, odizolowane od masy szybowca. Między elektrodami nieruchomymi a masą szybowca powstaje napięcie zmienne proporcjonalne do natężenia pola elektrycznego. Napięcie to mierzy się przy pomocy wzmacniacza połączanego z mikroamperomierzem i prostownikiem mechanicznym umieszczonym na osi silniczka. Metoda powyższa nie nadaje się do określenia natężenia pola elektrycznego w miejscach o dużej krzywiznie. W przypadku szybowców zastrzeżenie to nie jest ważne, gdyż poza rozładowywaczami i anteną w konstrukcji szybowca unika się silnie zakrzywionych elementów przewodzących.

W celu porównywania wartości natężenia pola elektrycznego występującego w różnych punktach szybowca przeprowadza się w zdefiniowanych warunkach pomiary w tych punktach. Na podstawie tych pomiarów



Rys. 2. Współczynniki wyrażające wartość natężenia pola elektrycznego w poszczególnych punktach szybowca w stosunku do punktu odniesienia

ustala się współczynniki wyrażające wartość natężenia pola elektrycznego w poszczególnych punktach szybowca w stosunku do punktu odniesienia. Za taki punkt odniesienia przyjmuje się zwykle punkt pod kadłubem. Na rys. 2 podane są współczynniki określone przez R. Guna i J. P. Parkera [7] dla samolotu metalowego. Przy pomiarach natężenia pola elektrycznego w warunkach eksploatacyjnych elektrometr umieszcza się w punkcie odniesienia.

Pomiar ładunku zgromadzonego na samolocie przeprowadza się za pomocą urządzenia pozwalającego sztucznie go ładować. Urządzenie takie, opracowane przez Naval Research Laboratory [6] polega na ładowaniu na samolocie kropel wody ładunkiem elektrycznym, a następnie na wydaleniu ich z samolotu. Krople unosząc ze sobą ładunek, ładują samolot do biegunowości przeciwnej niż biegunowość, do której były naładowane. Urządzenie takie zasilano napięciem 10 kV; przy użyciu 340 litrów wody na godzinę pozwalało ono naładować samolot do potencjału około 500 kV.

Przejdziemy teraz do omówienia najważniejszych wyników badań przeprowadzonych w czasie lotów [5].

Ze względu na ciężar przyrządów pomiarowych i znaczne miejsce potrzebne do ich zainstalowania, badania wyżej omówionych parametrów przeprowadza się przeważnie na samolotach zamiast na szybowcach. Wiele wyników uzyskanych przy tych badaniach ma zastosowanie bezpośrednio do szybowców, a wiele można przenieść drogą interpolacji. Natężenie pola elektrycznego w punkcie odniesienia samolotu wyniosło [5]:

przy rozbryzgiwaniu na powierzchni samolotu kropel deszczu	0,05 kV/cm
przy tarciu samolotu o suche cząstki śniegu lub lodu	0,2 kV/cm
przy przelotach w silnym polu elektrycznym	3,4 kV/cm

W ostatnim przypadku zarejestrowano prądy upływu z samolotu dochodzące do 10 mA. Pojemność samolotu, na którym przeprowadzono badania, była około 4 razy większa od pojemności przeciętnego szybowca. Można więc spodziewać się, że prąd upływu z szybowca, wywołany polem elektrycznym, będzie ze względu na mniejsze wymiary szybowca również około 4 razy mniejszy od analogicznego prądu samolotu. Tak więc na podstawie badań przeprowadzonych na samolotach maksymalny prąd upływu szybowca można oszacować na około $\frac{10}{4}$ mA = 2,5 mA.

Duże natężenie pola elektrycznego na powierzchni szybowca jest przyczyną powstawania przy jego powierzchni wyładowań niezupełnych: świetlających, snopiących, a nawet piorunowych wyładowań wstępnych (liderów). Według obserwacji przeprowadzonych na samolotach metalowych [3] wyładowania osiągają długość do 4,5 m. Wyładowania wstępne w znacznym stopniu zwiększają prawdopodobieństwo uderzenia pioruna w szybowiec, a wyładowania niezupełne wszelkich typów wytwarzają w jego otoczeniu tak duży poziom szumów, że komunikacja radiowa może stać się niemożliwa. Dlatego jednym z podstawowych zagadnień ochrony odgromowej szybowca jest zmniejszenie natężenia pola elektrycznego na jego powierzchni.

Zmniejszenie natężenia pola elektrycznego uzyskuje się przez zainstalowanie na szybowcu tak zwanych rozładowywaczy z materiałów półprzewodzących o dużym współczynniku skuteczności odprowadzania ładunków oraz małym początkowym natężeniu jonizacji [patrz wzór (4)]. Rozładowywacze te posiadają stosunkowo dużą oporność własną, a więc odprowadzając prąd nie wywołują szumów szkodliwych dla radiokomunikacji.

Przedstawiony w niniejszym rozdziale przegląd zagadnień związanych z powstawaniem dużych natężeń pola przy szybowcach i samolotach jest materiałem wyjściowym zarówno do badań przeprowadzonych przez au-

torów w zakresie wybiórczości piorunów w odniesieniu do szybowców, jak i badań rozładowywaczy.

3. PIORUN I JEGO PARAMETRY

Badanie piorunów trafiających w samoloty i szybowce nastęrcza znacznie więcej trudności niż badanie piorunów doziemnych. Te trudności wypływają z faktu przebywania w samolocie lub szybowcu obsługi, dla której piorun lub towarzyszące mu zjawiska mogą być groźne. Z tego względu na ogół w przypadku samolotów unika się lotów w obszarach burzowych.

Metody stosowane do określania parametrów piorunów trafiających w samoloty i szybowce są znacznie mniej liczne niż metody stosowane do badań piorunów doziemnych. W zasadzie ograniczają się one do zestawiania statystyk częstości uderzeń piorunów w samoloty i szybowce, pomiaru wartości szczytowej prądu pioruna oraz pomiaru ładunku pioruna. Jak wykazały statystyki, znaczny odsetek piorunów trafia w anteny umieszczone na samolocie lub szybowcu. J. M. Bryant i współpracownicy [8], biorąc to pod uwagę opracowali specjalny układ ochrony przepięciowej radiostacji, w którym przewidziano jednocześnie sztabki magnetyczne [9] pozwalające na przeprowadzenie pomiaru wartości szczytowej prądu pioruna. Wartości ładunków elektrycznych przenoszonych przez pioruny trafiające w samoloty i szybowce określa się przeważnie na podstawie skutków wywołanych przez te pioruny, np. na podstawie wytopienia kul iskiernika umieszczonego w wyżej wspomnianym układzie J. M. Bryanta lub na podstawie wytopienia łożysk części ruchomych (np. sterów i lotek).

Opierając się na zebranych materiale statystycznym A. C. Van Dorsten [10] stwierdza, że prawdopodobieństwo uderzenia pioruna w samolot wzrasta w miarę powiększania jego wymiarów oraz wzrostu prędkości lotu. Pewne średnie dane statystyczne podaje J. H. Hagenguth [11]. Według tych danych na 1363 doświadczalne ekspedycje burzowe (loty w obszarach burzowych) stwierdzono 21 uderzeń piorunów w samolot. Stanowi to 1,5% ogólnej liczby ekspedycji burzowych.

Energetycy, jak wiadomo, zebrali obszerne materiały dotyczące trafiaiania piorunów w ziemię i obiekty energetyczne. Materiały te postaramy się wykorzystać do przeprowadzenia porównania między częstością występowania piorunów uderzających w samoloty i piorunów doziemnych (rejestrowanych przez energetyków). Przy założeniu, że jedna ekspedycja burzowa odpowiada jednemu dniu burzowemu, i przyjmując średnią powierzchnię samolotu równą 500 m^2 — otrzymujemy roczną liczbę uderzeń piorunów w samoloty, przypadającą na 1 km^2 powierzchni i na je-

den dzień burzowy, równą 30. Analogiczna liczba dla piorunów doziemnych wynosi 0,01—0,10 [9]. Z tego porównania wynika, że częstość uderzeń piorunów w samoloty jest $3 \times (10^2 \text{ do } 10^3)$ większa, niż częstość piorunów doziemnych. To porównanie wyraźnie wskazuje, że przy uderzeniach piorunów w samoloty wchodzi w grę pewne dodatkowe czynniki. Do takich czynników można by zaliczyć:

a) większą częstość występowania piorunów na dużych wysokościach (poza piorunami doziemnymi — pioruny między chmurami),

b) lokalne zwiększanie natężeń pola elektrycznego przez samoloty i szybowce znajdujące się w obszarach burzowych.

Pierwszy z tych czynników wydaje się odgrywać mniejszą rolę. Podstawowe znaczenie natomiast przypisywane jest czynnikowi drugiemu, jako przyczynie zapoczątkowywania pioruna przez samolot lub szybowiec [3], [12]. E. M. Workman i współpracownicy szacują liczbę piorunów trafiających samolot skutkiem znalezienia się jego na drodze pioruna tylko na 0,1% ogólnej liczby piorunów trafiających w samolot; tak więc w przeważającej większości przypadków przyczynę wywołującą piorun stanowi sam samolot lub szybowiec. Potwierdzeniem tej tezy może być zestawienie badań E. J. Workmana i R. E. Holzera [12] oraz badań R. G. Strimmela i współpracowników [5]. Według pierwszych samolot jest najbardziej narażony na uderzenia pioruna przy lotach między izotermami -5°C do -25°C , wg drugich w niskich temperaturach, a zwłaszcza — przy temperaturze -7°C ... -9°C gromadzi się na samolocie wielki dodatkowy ładunek elektryczny. W tych warunkach występują więc na powierzchni samolotu duże natężenia pola elektrycznego, to znaczy istnieją korzystne warunki do zapoczątkowania pioruna przez samolot.

Ochrona odgromowa szybowca sprowadza się, jak już wspomniano, do zmniejszenia prawdopodobieństwa uderzenia pioruna, z drugiej strony — do uodpornienia szybowca na działanie pioruna, które może być termiczne, elektryczne i magnetyczne.

Działanie pioruna na szybowiec zależy od parametrów elektrycznych pioruna. Na podstawie pomiarów przeprowadzonych na samolotach podczas lotów J. H. Hagenguth [11] podaje następujące wartości:

dla pioruna między chmurą a ziemią — wartość szczytowa prądu do 100 kA,

dla pioruna między chmurą:

i chmurą — wartość szczytowa prądu do 10 kA,

a ładunek przenoszony do 500 C.

Nieco inne dane na podstawie dwuletnich badań piorunów uderzających w anteny samolotów podaje Lightning and Transient Research Institute: najczęściej występująca wartość szczytowa prądu pioruna rzędu 5 kA, ładunek przenoszony rzędu 100 C.

Można przypuszczać, iż przytoczone wyżej wartości ładunku przenozonego przez piorun uderzający w samolot występują rzadko. Brak jest danych w tej sprawie w literaturze, jednak za przyjęciem przez różnych autorów podanego wyżej przypuszczenia świadczy fakt uwzględnienia w ich obliczeniach mniejszych wartości ładunku. Tak np. J. M. Bryant i współpracownicy [8] proponują — do celów badań samolotów — znormalizowanie wartości ładunku probierczego udaru prądowego na poziomie 25 C. Stwierdzają oni konieczność powiększenia tej wartości ładunku tylko przy niektórych badaniach.

To krótkie zestawienie parametrów piorunów uderzających w samoloty wykazuje, iż ich znajomość jest jeszcze stosunkowo niewielka.

4. OKREŚLENIE WYBIORCZOŚCI PIORUNA NA MODELACH

Określanie miejsc, w które może trafić piorun, nosi nazwę określania wybiorczości pioruna. Badania takie, w odniesieniu do szybowców, nie są możliwe w warunkach naturalnych w czasie burz, a to ze względu na rzadkość uderzeń pioruna w szybowce i na niebezpieczeństwo dla załogi szybowca. Inaczej sprawa przedstawia się w stosunku do samolotów, z reguły pokrytych metalem, stanowiącym doskonałą osłonę dla personelu. Prądy przy uderzeniach pioruna w samolot mogą dostać się do wnętrza tylko wtedy, gdy piorun trafi w antenę (w razie braku odpowiedniego zabezpieczenia) lub w otwór okienny, co jest niezwykle mało prawdopodobne. Przeprowadzono szereg bardzo interesujących eksperymentów, których celem było ustalenie zagrożenia od uderzenia piorunów samolotów wojskowych. Częściowo opublikowane są wspomniane już wyniki badań amerykańskich [11], omawiające tysiące lotów w obszarach burzowych. Materiał uzyskany w ten sposób okazał się jednak bardzo skąpy wobec rzadkości trafień pioruna w samolot.

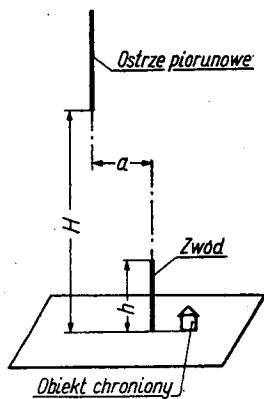
Obserwacje miejsc uderzeń pioruna w samolot w czasie lotów w chmurach burzowych i pod chmurami oraz ustalenie miejsc trafienia na podstawie wytopień metalu stanowią cenny materiał porównawczy, nie dający się jednak bezpośrednio zastosować do szybowców. Wynika to z innych kształtów i wymiarów geometrycznych szybowca i z innego rozkładu części metalowych.

Badania wybiorczości pioruna odnośnie szybowców w warunkach naturalnych są dużo trudniejsze niż w przypadku celów naziemnych trafianych przez piorun. Cele takie, jako nieruchome, „oczekują” na trafienie w czasie każdej burzy przechodzącej nad okolicą, podczas gdy szybowce muszą dopiero unieść się w chmury lub w ich pobliże, aby znaleźć się w warunkach zagrożenia, występujących w eksploatacji. Jednak i w przypadku celów naziemnych bezpośrednia obserwacja zwodów i obiektów

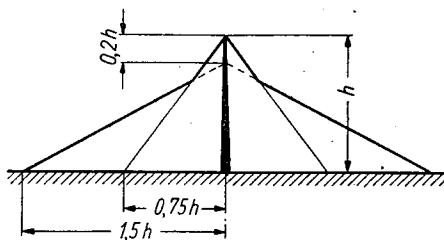
chronionych daje bardzo skąpy materiał — ze względu na niezwykle rzadkość trafień. Nic więc dziwnego, że dla tego przypadku powstała myśl badań laboratoryjnych na modelach w zmniejszonej skali. Takie badania modelowe były prowadzone bardzo intensywnie w ciągu ostatnich lat 25 i mają obszerną literaturę (vide np. [13], [14], [15], [16]).

Przez analogię nasuwa się myśl zastosowania badań modelowych do określania wybiórczości pioruna w przypadku szybowców. Z góry można przewidzieć, że ocena wartości tej metody powinna być bardzo ostrożna. Już badania wybiórczości przy celach naziemnych podlegają ostrej krytyce — przypadek szybowca jest, jak zobaczymy, jeszcze trudniejszy pod względem doboru parametrów przy modelowaniu.

Większość uwag krytycznych dotyczących modelowania układu w przypadku trafienia celów naziemnych odnosi się i do modelowania



Rys. 3. Modelowanie odległości ostrza piorunowego przy badaniu celów naziemnych: h — wysokość zwodu, H — odległość od ziemi ostrza piorunowego



Rys. 4. Przestrzeń chroniona zwodu prętowego według „Wskazówek” radzieckich

układu z szybowcem. Dlatego celowe jest krótkie przedstawienie tych uwag. Cele naziemne, jak zwody piorunowe, mające przyjąć uderzenie pioruna (piorunochron, przewód odgromowy linii energetycznych) i obiekty chronione, jak budynki, linie i stacje elektroenergetyczne — modeluje się w pewnej skali, np. 1 : 50. Według obecnych poglądów nie odtwarza się natomiast na modelu chmury, a tylko zastępuje ostrzem czoło odgórnego wyładowania wstępnego pioruna (lidera), przesuwającego się w naturze w kierunku od chmury do ziemi. To tak zwane ostrze piorunowe łączy się w czasie badań z generatorem, dającym udary napięciowe rzędu setek ... tysięcy kV. Swego czasu w literaturze rozgorzała ożywiona dyskusja nad wysokością umieszczenia tego ostrza nad płytą imitującą powierzchnię ziemi (rys. 3). Jako parametr porównawczy wybrano stosu-

nek wysokości tego ostrza H do wysokości pionowego pręta, imitującego zwód piorunowy h , a więc stosunek $\frac{H}{h}$. Pierwotną tendencją był stosunek $\frac{H}{h}$ bardzo wysoki, np. równy 10; uzasadniano to lepszym imitowaniem warunków naturalnych, w których czoło lidera, powstające w chmurze, zbliża się do danego obiektu z dużej wysokości. Sprawa jednak nie jest tak prosta; pokazała to analiza fotograficznych zdjęć wyładowań w układach modelowych. Wynikało z niej, że na spotkanie wyładowania odgórnego, wychodzącego z ostrza piorunowego, idą powstające przy ostrzach uziemionych (zwody, obiekty chronione lub ich części) wyładowania oddolne, skierowane w górę. Wyładowania te są przy tym na modelach często dużo dłuższe (z uwzględnieniem skali) niż w naturze. Takie wyładowania oddolne wychodzące z ostrza-zwodu (piorunochronu) łącząc się z wyładowaniem z ostrza piorunowego powodują iskrę — odpowiednik pioruna. Jest jasne, że długie wyładowania oddolne z piorunochronu zwierają część odstępu ostrze piorunowe-zwód i ułatwiają uderzenie pioruna w zwód, a przez to zmniejszają zagrożenie obiektów chronionych znajdujących się w pobliżu piorunochronu. Tak zwana przestrzeń chroniona (rys. 4), to jest część przestrzeni, w której mogą się znajdować obiekty uziemione, nie narażone na uderzenie pioruna, jest więc — w przypadkach długich wyładowań oddolnych z piorunochronu — większa niż bez nich. W naturze wyładowania oddolne z obiektów uziemionych mają długość — jak wykazały obserwacje — do 20 m, a więc zwierają bardzo małą część przestrzeni, którą przebiega iskra pioruna. Przy badaniach modelowych, zwłaszcza gdy $\frac{H}{h}$ jest duże (np. 10 i więcej), wyładowania oddolne, idące w górę, mogą natomiast osiągnąć długość połowy przerwy iskrowej, a nawet większą, zanim spotkają się z wyładowaniem odgórnym, idącym w dół.

Jak wykazały fotografie wyładowań, przy próbach modelowych długość wyładowań oddolnych jest specjalnie duża w przypadku ujemnej biegunowości ostrza piorunowego. Pioruny doziemne wychodzące z ujemnie naładowanej części chmury są, jak wiadomo, najczęstsze, stanowią one do 90% wszystkich uderzeń w ziemię [9]. Dlatego też przypadek modelowania, w którym ostrze piorunowe ma biegunowość ujemną, jest szczególnie interesujący. Niestety, przy ostrzu piorunowym ujemnym wyładowania wstępne oddolne przekraczają często połowę długości przerwy iskrowej, zanim spotkają się z wyładowaniem odgórnym. Dyskwalifikuje to całkowicie takie badania. Pierwszy z analizy tego stanu rzeczy wyciągnął wnioski A. A. Akopian [15], który zaleca stosowanie do badań zawsze biegunowości dodatniej ostrza piorunowego, również wówczas,

gdy modeluje się piorun, bijący z ujemnie naładowanej części chmury. Zasada ta została później przyjęta ogólnie.

Stosowanie się do wytycznej Akopiana nie jest oczywiście wystarczające. Różni autorzy posługują się poza tym dodatkowym kryterium, zwłaszcza dotyczącym stosunku $\frac{H}{h}$. Na przykład R. H. Golde [16] twierdzi:

„Jedną z dróg znalezienia prawidłowej wysokości górnej elektrody przy próbach modelowych byłoby takie zredukowanie tej wysokości, żeby dla zwodu prętowego stosunek ochrony był około 1”¹⁾.

Autorzy niniejszego artykułu są zdania, że jako kryterium należy uznać zasadę bardziej ogólną, wynikającą wprost z podanej wyżej analizy, mianowicie przyjęcie takiego układu elektrod, aby wyładowania oddolne — przy kontroli fotograficznej — miały długość wyładowań oddolnych w naturze, a więc wymiary do 20 m (w skali), a ponadto, aby zwierały tylko niewielką część przerwy iskrowej, np. $\frac{1}{5}$ ²⁾. Takie sformułowanie ma jeszcze tę zaletę, że można je zastosować również do prób z szybowcami.

W przypadku szybowców modelowanie jest bardziej złożone niż w przypadku obiektów naziemnych. Szybowce w powietrzu znajdują się nad ziemią, to też pierwszą myślą jest odtworzenie na modelu w skali odległości szybowiec-ziemia. Jest to jednak w praktyce laboratoryjnej niemożliwe, gdyż prowadziłoby do stosowania albo tak wysokich napięć, jakimi technika laboratoryjna jeszcze nie dysponuje, albo do tak miniaturowych modeli szybowców, że dokładne odtworzenie ich budowy i dokładne obserwacje wyładowań byłyby pod znakiem zapytania.

Biorąc pod uwagę, że o trafieniu w daną część szybowca decydują lokalne pola elektryczne, oraz uwzględniając, że wykładnikiem struktury pola jest długość wyładowań wychodzących z szybowca (zarówno w górę, jak w dół i na boki) — można uznać długość tych wyładowań za kryterium prawidłowości modelowania.

Według L. P. Harrisona [3] obserwowane długości wyładowań wstępnych z samolotu w naturze są rzędu 5 m. Dotyczy to przypadku, gdy piorun nie uderza w samolot. Bezpośrednio przed uderzeniem można się spodziewać długości wyładowań wstępnych z szybowca takich, jak z obiektów naziemnych, a więc do 20 m. W przypadku skali modeli 1 : 50 oznacza to długość wyładowań na modelu do 40 cm.

¹⁾ Stosunek ochronny — stosunek promienia koła ograniczającego przestrzeń chronioną na powierzchni ziemi do wysokości zwodu.

²⁾ Przerwa iskrowa — odstęp między ostrzem piorunowym i ostrzem naziemnym (porównaj rys. 3).

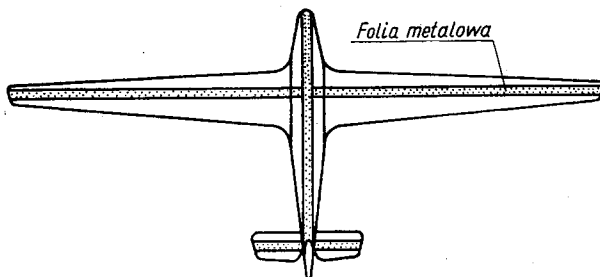
Powyższe rozważania dotyczą modelowania układów, w których szybowiec jest trafiony zarówno przez piorun przebiegający między chmurą a ziemią, jak i przebiegający między dwoma chmurami.

Aby określić długość wyładowań wstępnych, można w praktyce laboratoryjnej oczywiście ucinać udary napięciowe specjalnym zewnętrznym iskiernikiem. Prościej jest jednak określać długość tych wyładowań na fotografiach z przeskokami iskrowymi. Iskra na modelu, podobnie jak iskra w iskierniku zewnętrznym, ucinając udar napięciowy, zmniejsza napięcie do małej wartości i uniemożliwia dalszy rozwój tych wyładowań wstępnych, które nie utworzyły iskry.

5. OKREŚLANIE WYBIORCZOŚCI PIORUNA NA MODELACH UPROSZCZONYCH

Dotychczasowe rozważania miały na celu przedstawienie krytyczne stanu zagadnienia i ustalenie wytycznych do wykonanych w Katedrze Wysokich Napięć badań nad polskimi szybowcami. Teraz zajmiemy się szczegółowo tymi badaniami w zakresie określenia wybiorności piorunów, ich metodyką i wynikami.

Celem badań wybiorności pioruna na uproszczonych modelach szybowca było określenie miejsc szybowca najbardziej narażonych na ude-

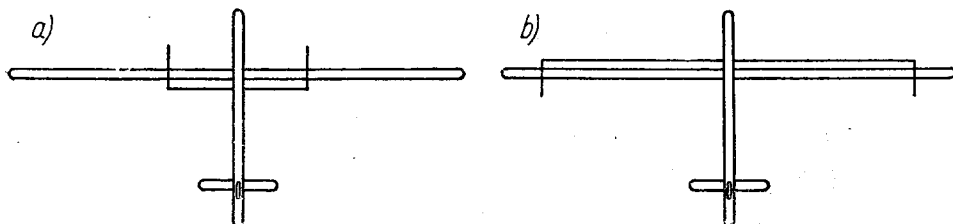


Rys. 5. Uproszczony drewniany model szybowca stosowany do badań wybiorności pioruna

wienie pioruna oraz wyznaczenie przybliżonych stref chronionych zwodów, umieszczonych w miejscach najbardziej narażonych na uderzenie pioruna. Do badań tych użyto dwu modeli szybowca. Pierwszy, wykonany z drewna, oklejony paskami folii aluminiowej w sposób imitujący rozmieszczenie części metalowych na szybowcu (rys. 5) służył do określenia miejsc najbardziej narażonych na wyładowania. Drugi, wykonany z drutu miedzianego, wyposażony w ostrza o regulowanej długości i o zmiennym ustawieniu wzdłuż skrzydeł (rys. 6) służył do określenia przybliżonej strefy chronionej zwodów na szybowcu. Obydwa modele posiadały te sa-

me wymiary gabarytowe i były zmniejszone w stosunku do rzeczywistego szybowca w skali około 1 : 50.

Zgodnie z rozważaniami przeprowadzonymi w rozdziale 4 przyjęto, że najbardziej narażone na uderzenia pioruna są te miejsca szybowca, w któ-

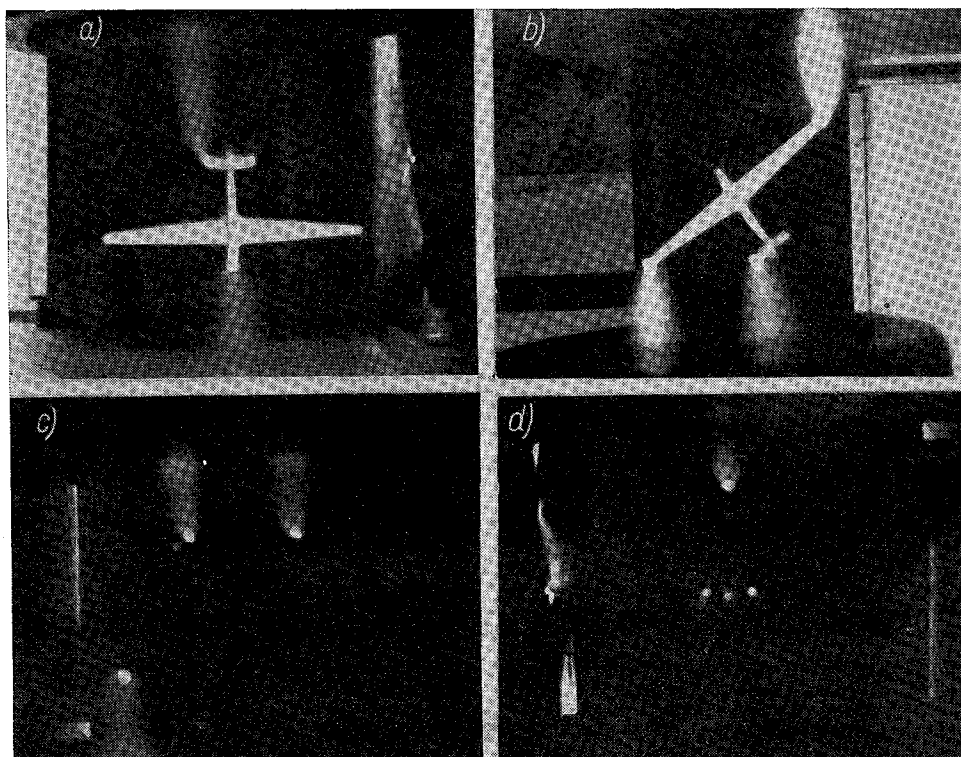


Rys. 6. Model szybowca wykonany z drutu miedzianego z regulowanymi ostrzami do badań przestrzeni chronionej: a) ostrze w małej odległości od kadłuba, b) ostrze w dużej odległości od kadłuba

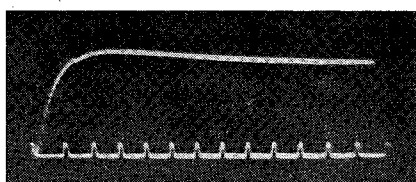
rych w polu elektrycznym pod wpływem dużych nateżeń występują wyładowania niezupełne (świetenia, snopienia). Biorąc pod uwagę, że ścisła rejestracja fotograficzna wyładowań niezupełnych przy napięciu udarowym jest utrudniona, we wstępnych badaniach zastosowano napięcie zmienne o częstotliwości 50 Hz. Przy napięciu takim wyładowania niezupełne można utrzymywać przez czas dowolnie długi dla przeprowadzenia obserwacji i rejestracji fotograficznej. Napięcie to, którego źródłem był transformator probierczy 300 kV, wytwarzało w płaskim układzie elektrod pole elektryczne w przybliżeniu jednostajne. Po wprowadzeniu między elektrody modelu szybowca, w niektórych punktach na powierzchni modelu zjawiają się punkty świetlne, świadczące o występowaniu wyładowań niezupełnych. W punktach tych nateżenia są specjalnie duże. Przez zmiany położenia szybowca w stosunku do elektrod można w ten sposób określić wszystkie miejsca szybowca o dużym nateżeniu pola (rys. 7).

Badania powyższe miały charakter orientacyjny, toteż ograniczono je do przypadku z szybowcem nienaładowanym. Innymi słowy — uwzględniono tylko przypadek influencji.

Punkty o największym nateżeniu pola, to końce skrzydeł, dziób oraz krawędzie sterów wysokościowych i steru kierunkowego. Sprawdzenia prawdopodobieństwa uderzeń pioruna w poszczególne punkty dokonano przy napięciu udarowym, wytwarzanym przez generator udarowy Instytutu Elektrotechniki MPC o danych znamionowych 2800 kV, 32 kWs. Kształt napięcia udarowego wytwarzanego w generatorze podano na rys. 8. Modele szybowca umieszczono w przerwie między elektrodą ostrzową a uziemioną płytą metalową. Elektroda ostrzowa, tak zwane ostrze piorunowe, odwzorowywała czoło zbliżającego się do szybowca wstępnego wyładowania piorunowego. Płyta metalowa służyła do wytworzenia pola



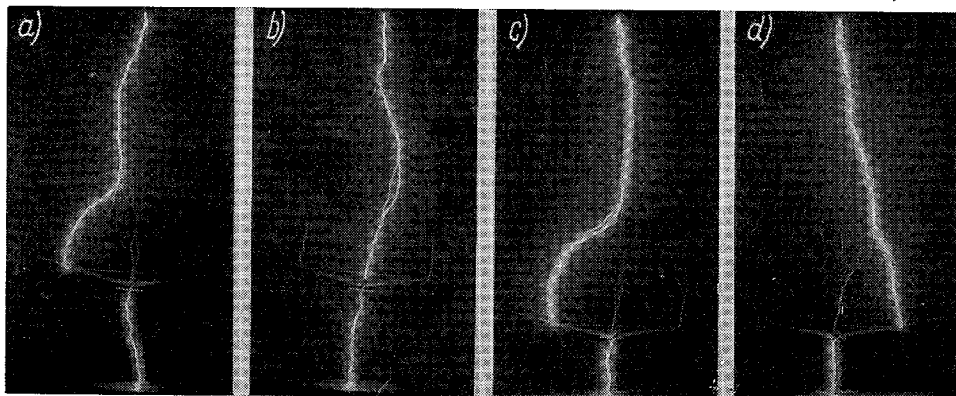
Rys. 7. Wyładowania niezupełne z modelu szybowca przy napięciu zmiennym. Zdjęcia a), b) — wykonane w świetle dziennym, zdjęcia c), d) — wykonane dla tych samych położeń modelu, lecz w pomieszczeniu zaciemnionym



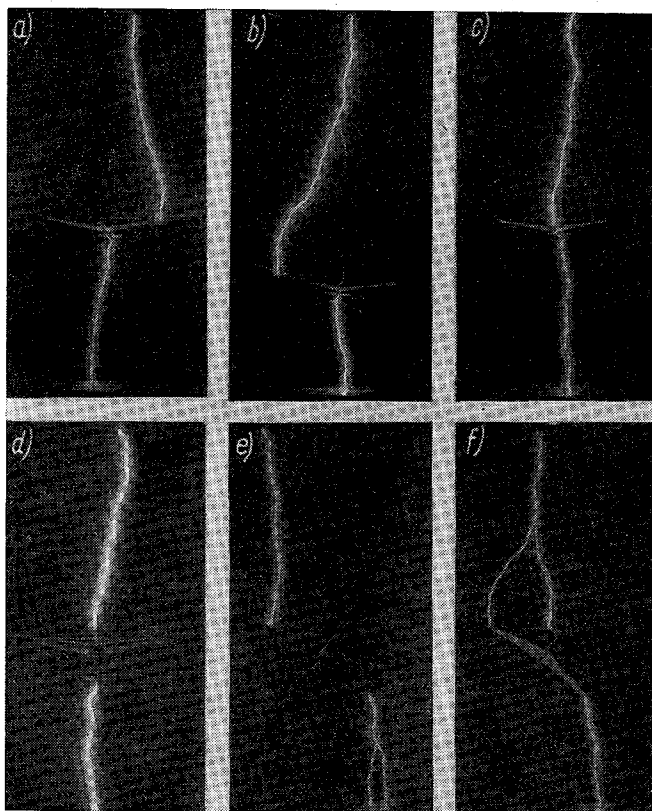
Rys. 8. Oscylogram napięcia wytwarzanego w generatorze udaru 2800 kV. Skalowanie co 1 mikrosekundę

elektrycznego, jakie panowałoby między chmurami lub między chmurą a ziemią. Wymiary płyty 300×400 cm były tak dobrane w stosunku do odstępów między elektrodami (ok. 1,5 m), aby na krawędziach płyty nie występowały wyładowania niezupełne, mogące wpłynąć na zmianę kierunku wyładowania z ostrza piorunowego.

Zgodnie z rozważaniami poprzedniego rozdziału stosowano do badań udary o biegunowości dodatniej. Wartość szczytowa napięcia udarowego



Rys. 9. Wyładowania udarowe w uproszczony model szybowca (wg rys. 5). Długość wyładowań wstępnych z szybowca nie przekracza w skali ok. 10 m (rozpiętość skrzydeł szybowca 18 m). Wyładowania te są widoczne jako cienkie nitki wychodzące z końców skrzydeł i ogona



Rys. 10. Przykłady wyładowań udarowych w uproszczony model szybowca: a), b), e) — wyładowania w skrzydło, c), d) — wyładowania w dziób, f) — wyładowanie podwójne — w skrzydło i ogon

była rzędu 1 miliona woltów, co zapewniało uzyskanie wyładowania przy każdym uderze.

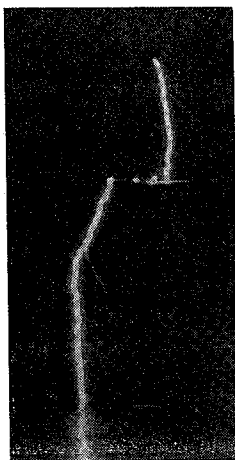
Przy dobieraniu wysokości zawieszenia szybowca w przerwie między-elektrodowej stosowano się do kryterium ustalonego w rozdziale 4. Wysokość tę określono w ten sposób, aby wyładowania wstępne wychodzące z szybowca nie przekraczały w skali długości 20 m, a więc największej długości występującej w naturze w przypadku obiektów uziemionych (por. rys. 9).

Odległość modelu szybowca od płyty uziemionej przyjmowano jak największą z dwóch przyczyn:

- aby pole elektryczne wokół szybowca było możliwie takie, jak w przypadku szybowca wysoko nad ziemią,

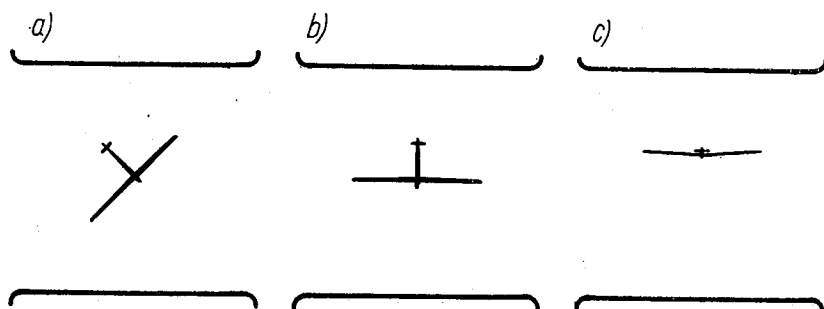
- aby szybowiec nie uzyskiwał potencjału ziemi skutkiem dojścia wyładowania wstępnego z szybowca do ziemi wcześniej, nim spotkają się w przestrzeni wyładowania ostrze piorunowe-szybowiec. Typowe fotografie wyładowań przedstawiono na rys. 10.

Zgodnie z przypuszczeniami najbardziej narażone na wyładowania piorunowe są wysunięte punkty modelu, jak końce skrzydeł, sterów i dziób

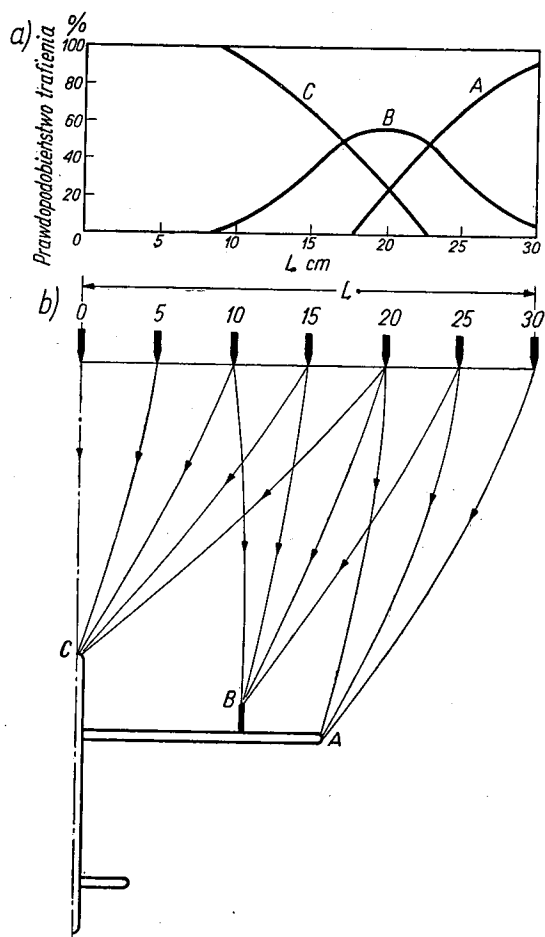


Rys. 11. Przypadek wyładowania udarowego w skrzydło, w miejsce, gdzie występowało „iskrzenie”

kadłuba. Istotną rolę grają przy tym wyładowania wstępne, wychodzące z szybowca. Bardzo dobrze ilustruje to rys. 11. W tym przypadku przerwa w połączeniach folii metalowej na skrzydle była przyczyną powstawania wyładowań niezupełnych — „iskrzenia” w miejscu przerwania. Wyładowanie piorunowe, przed tym trafiające w koniec skrzydła, zostało skierowane w miejsce przerwania folii. Należy stąd wyciągnąć praktyczny wniosek: aby instalacja odgromowa szybowca była skuteczna, powinna być wykonana bez przerw w połączeniach wywołujących iskrzenia.

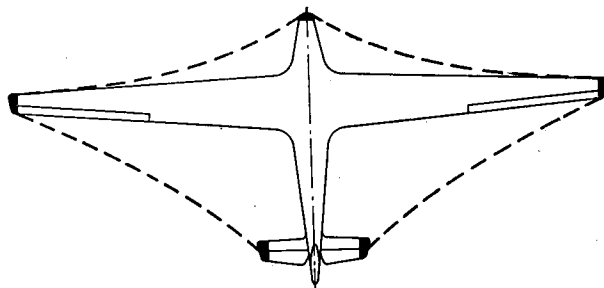


Rys. 12. Położenie modelu szybowca przy określaniu przestrzeni chronionych



Rys. 13. Przykład wyznaczania krzywej prawdopodobieństwa uderzenia pioruna w poszczególne punkty modelu. A — koniec skrzydła, B — ostrze służące do określania przestrzeni chronionej, C — dziób

Badania przestrzeni chronionych na szybowcu przeprowadzono na modelu uproszczonym z rys. 6. Na $1/3$ i $2/3$ długości skrzydła modelu umocowano ostrze, które stopniowo wysuwano. Badanie przeprowadzono w różnych położeniach modelu względem ostrza piorunowego, pokazanych na rys. 12. Przy dobieraniu tych położen załozono, że szybowiec w chmurach burzowych może być trafiony przez wyładowania piorunowe przychodzące praktycznie z każdej strony. Ostrze piorunowe przesuwano podczas pomiarów w płaszczyźnie pionowej poprzecznej do osi szybowca w odstępach 5 cm. Dla każdego położenia wyznaczano prawdopodobieństwo uderzenia pioruna w ostrze na modelu oraz w inne punkty modelu, wytwarzając w układzie 20 uderów i notując miejsca trafione. W ten sposób wykreślono krzywe prawdopodobieństwa uderzenia pioruna w różne punkty szybowca przy różnych położeniach szybowca względem kierunku rozwoju kanału pioruna. Przykładowe krzywe przedstawiono na rys. 13. Wyszukano również takie największe długości ostrzy, wysuwanych z przodu skrzydła (w położeniu szybowca z rys. 6a) lub z tyłu skrzydła (w położeniu szybowca z rys. 6b), przy których nie występo-



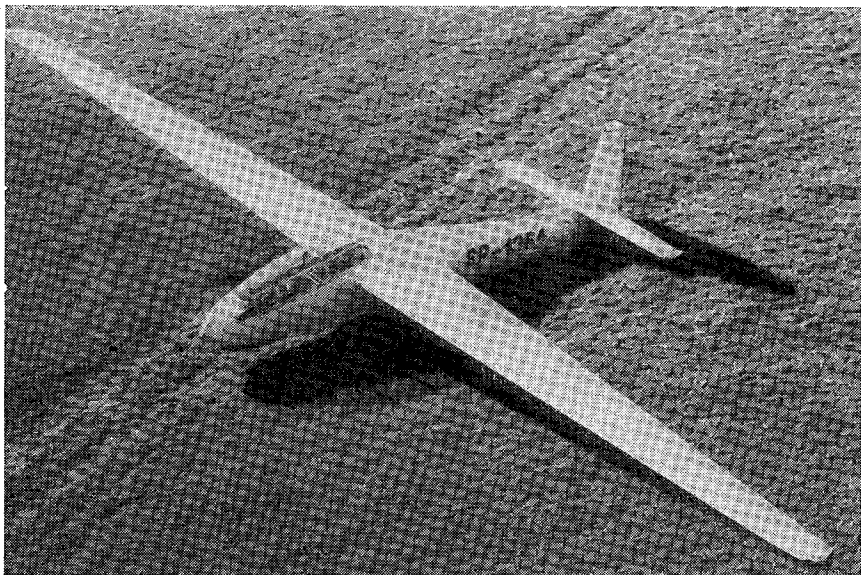
Rys. 14. Zarys przestrzeni chronionej przez zwody na uproszczonym modelu szybowca. Zaczerniono miejsca umieszczenia zwodów

wały ani razu trafienia pioruna w ostrze. Z otrzymanych w ten sposób punktów w otoczeniu modelu szybowca wyznaczono przybliżony zarys przestrzeni chronionej (rys. 14).

Wyniki badań pozwoliły określić miejsca najbardziej narażone na uderzenia pioruna, a więc miejsca, w których powinny być zainstalowane zwody piorunochronne. Określony w przybliżeniu zarys stref przestrzeni chronionej pozwala spodziewać się, że uderzenie pioruna w nie zaopatrzone w zwody miejsca szybowca, jest raczej mało prawdopodobne, o ile nie „wystają” one z przestrzeni chronionej.

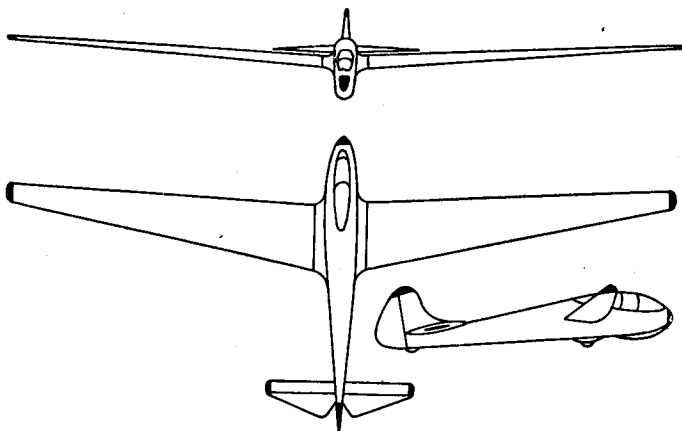
6. OKREŚLANIE WYBIORCZOŚCI PIORUNA NA MODELU SZYBOWCA TYPU „BOCIAN”

Praktycznym celem badań wyborczości pioruna przeprowadzanych na modelach szybowców było uzyskanie podstaw doświadczalnych do za-



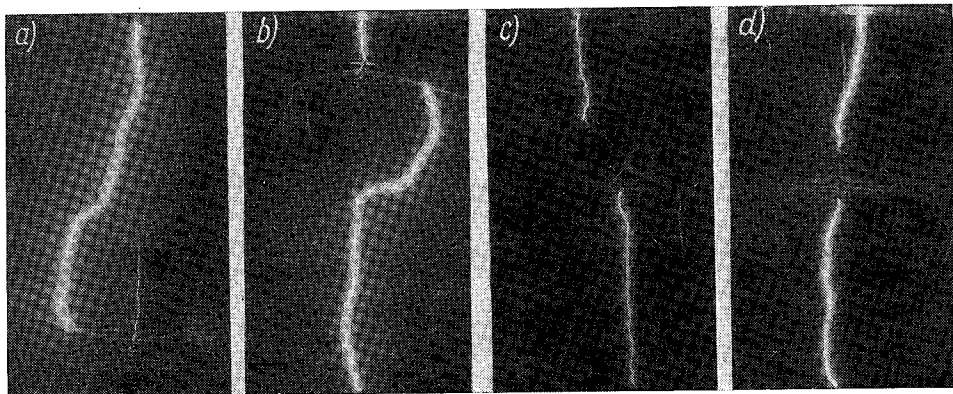
Rys. 15. Szybowiec dwumiejscowy typu „Bocian” SZD-9b, dla którego projektowano ochronę odgromową

projektowania skutecznej ochrony szybowca dwumiejscowego produkcji polskiej typu „Bocian” SZD-9b. Fotografię szybowca tego typu przedstawiono na rys. 15. Szybowiec ten jest skonstruowany do lotów o charakterze wyczynowym, służy jednak również do celów szkoleniowych ze wzglę-

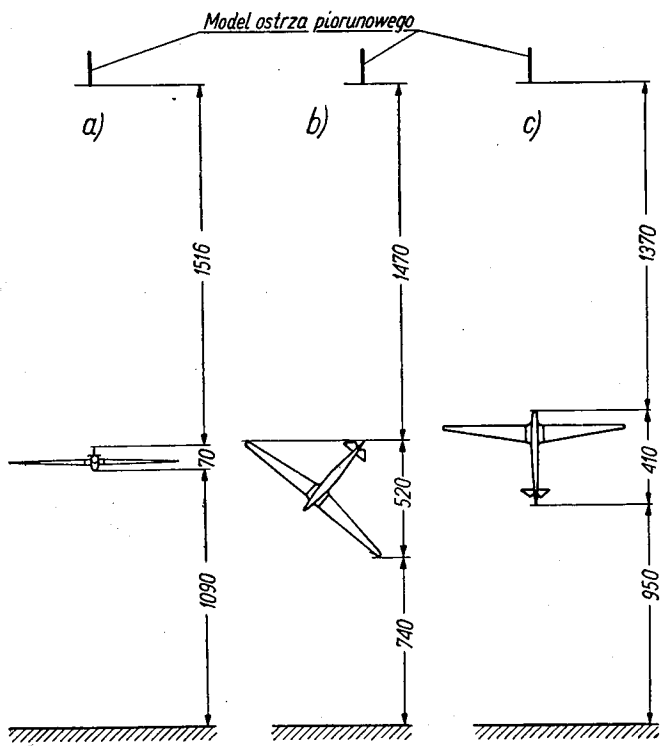


Rys. 16. Rozmieszczenie zwodów (miejsca zaczernione) przy badaniach wylądowań w model szybowca typu „Bocian”

du na możliwość kierowania nim z dwóch stanowisk (elewa i instruktora). Na tym typie szybowca uzyskano większość rekordów Polski dla szybowców dwumiejscowych.



Rys. 17. Wyladowanie niezupełne na modelu szybowca „Bocian” wg rys. 16 przy różnych wysokościach zawieszenia modelu nad płytą. Długość wyladowania jest większa w przypadku mniejszej odległości modelu od płyty. Najdłuższe wyladowania (rys. 17a) dotyczą modelu połączonego metalicznie z uziemioną płytą



Rys. 18. Zawieszenie modelu szybowca typu „Bocian” do sprawdzania skuteczności zwodów piorunochronnych

Przed przystąpieniem do badań zaprojektowano wstępnie rozmieszczenie zwodów na szybowcu w oparciu o omówione w rozdziale poprzednim badania na modelach uproszczonych. Jak wynika z tych badań, konieczne jest umieszczenie zwodów na dziobie szybowca, na końcach skrzydeł i na sterach. Ponieważ w tych miejscach stosuje się ze względów konstrukcyjnych (wzmocnienie mechaniczne) kołpaki metalowe, zaproponowano użycie ich jako zwodów odpowiednio połączonych przewodami odprowadzającymi.

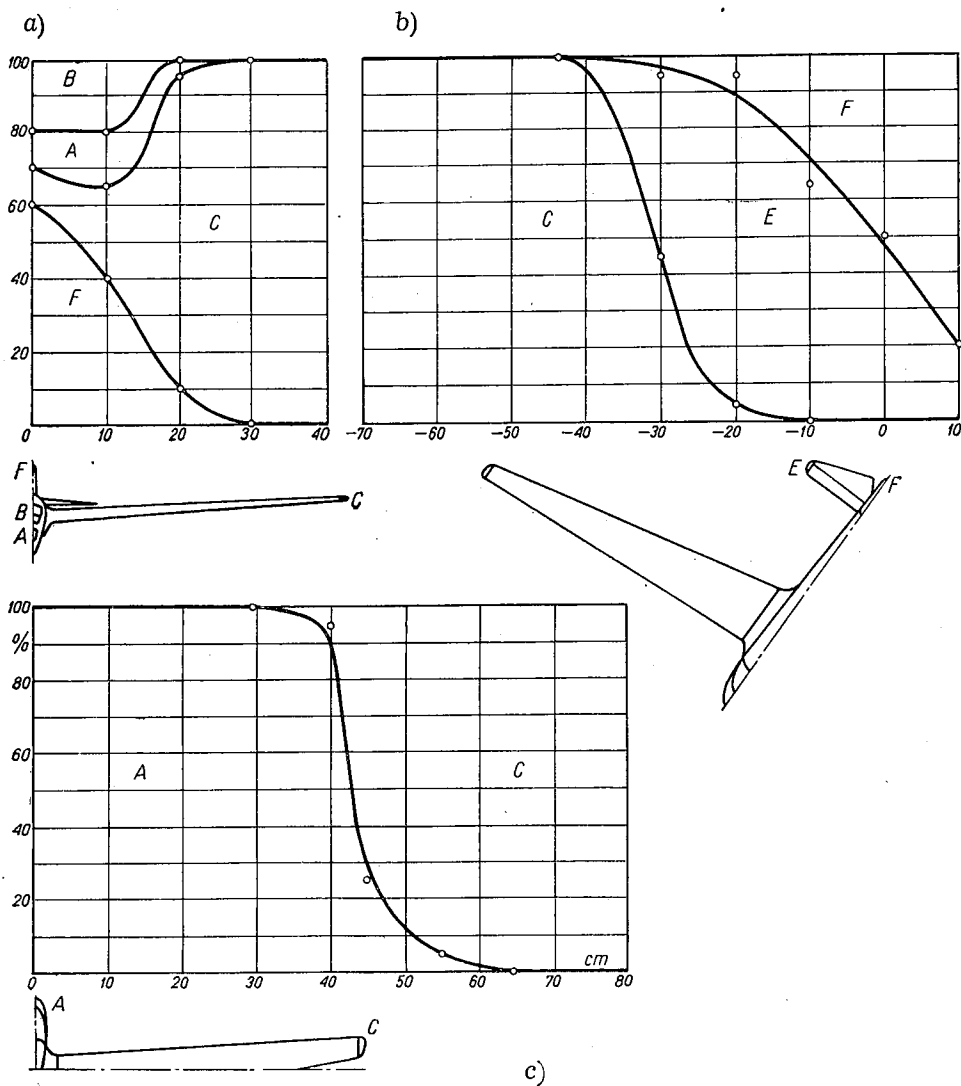
W celu zbadania skuteczności ochrony przygotowano do badań drewniany model szybowca w skali 1 : 20 z zainstalowanymi na nim zwodami, połączonymi przewodami odprowadzającymi, oraz z odwzorowanymi ważniejszymi elementami metalowymi na powierzchni szybowca (por. rys. 16). Badania przeprowadzono w układach analogicznych do omówionych w poprzednim rozdziale. Wobec większych wymiarów modelu przyjęto większy odstęp między ostrzem piorunowym a płytą uziemioną (ok. 2,7 m) i większe napięcie udarowe powodujące wyładowanie (ok. 1,6 miliona woltów).

Stosownie do przyjętego w rozdziale 4 kryterium, wyładowania wstępne z szybowca (w kierunku czoła wyładowania pioruna) nie powinny przekraczać $\frac{1}{5}$ długości przerwy iskrowej, a bezwzględna długość ich dla przyjętej skali 1 : 20 nie powinna przekraczać 100 cm. Jednak, jak wynika z przyjętych wymiarów, decydujący jest warunek pierwszy, według którego długość wyładowań wstępnych nie powinna przekraczać 27—45 cm, w zależności od wysokości zawieszenia modelu nad płytą (porównaj rys. 18).

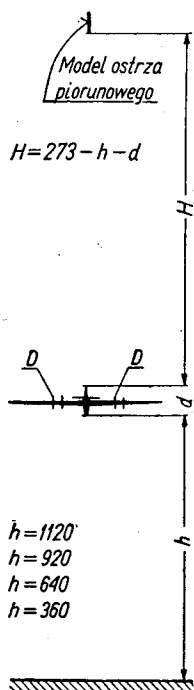
W rzeczywistości ten stosunek długości wyładowań do długości przerwy, jak wynika z fotografii na rys. 17, nie zawsze był zachowany i w niektórych przypadkach zwiększał się do $\frac{1}{4}$. Przy jednakowych odstępach ostrza piorunowego od płyty, w miarę oddalania szybowca od płyty zmniejszała się długość wyładowań wstępnych.

Badania wykonano dla szeregu położeń szybowca przedstawionych na rys. 18. Uzyskane krzywe prawdopodobieństwa trafienia w poszczególne miejsca szybowca przedstawiono na rys. 19. Z krzywych tych wynika, że istnieje prawdopodobieństwo wyładowania w szczyt kabiny (punkt B), jakkolwiek jest ono bardzo niewielkie. Na tej podstawie wysunięto przypuszczenie, że może istnieć także niewielkie prawdopodobieństwo wyładowania w hamulce aerodynamiczne, które wysuwa się ze skrzydeł w pewnych warunkach meteorologicznych. W związku z tym przeprowadzono dodatkową specjalną serię badań mającą na celu określenie prawdopodobieństwa wyładowań w hamulce i wabinę.

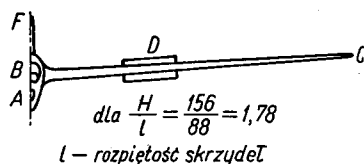
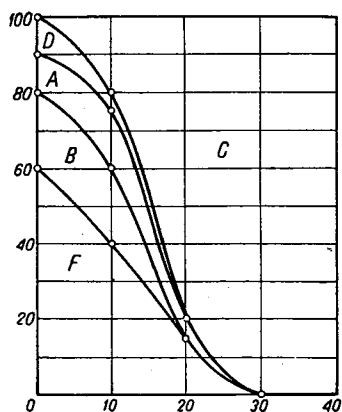
W trakcie badań stwierdzono, że wysunięcie hamulców (rys. 20) na wysokość 12 cm ponad powierzchnię skrzydeł (6 mm na modelu) wpro-



Rys. 19. Krzywe prawdopodobieństwa wylądowań w poszczególne miejsca modelu szybowca typu „Bocian”: a) dla zawieszenia wg rys. 18a, b) dla zawieszenia wg rys. 18b, c) dla zawieszenia wg rys. 18c. Oznaczenia miejsc na szybowcu: A — zwód na dziobie szybowca, B — wierzchołek kabiny, C — zwód na końcu skrzydła, E — zwód na stateczniku poziomym, F — zwód na stateczniku pionowym



Rys. 20. Zawieszenie modelu z wysuniętymi hamulcami (D)



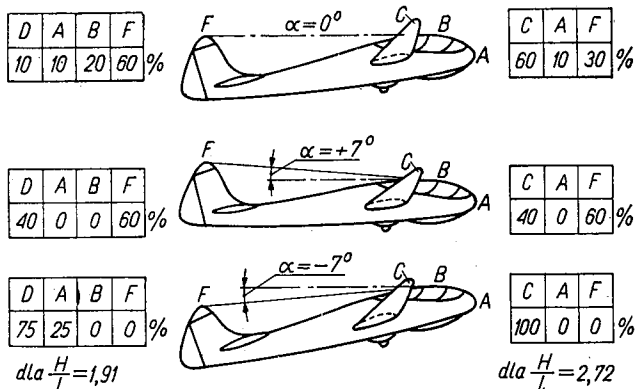
Rys. 21. Krzywe prawdopodobieństwa wyładowań w szybowiec z wysuniętymi hamulcami (D). Oznaczenia jak na rys. 19

wadza pewne prawdopodobieństwo uderzenia pioruna w hamulce (patrz rys. 21). Jednocześnie ulega zmniejszeniu prawdopodobieństwo uderzenia w koniec skrzydła. Natomiast na prawdopodobieństwo uderzenia pioruna w kabinę i wysunięte stery hamulców praktycznie nie wpływa.

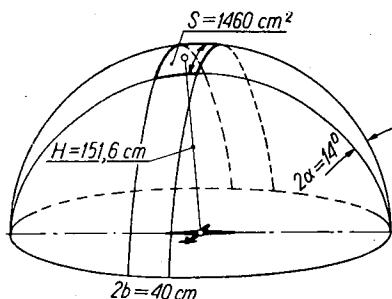
W wyniku specjalnych badań prawdopodobieństwa wyładowania w kabinę modelu stwierdzono, że nawet małe wychylenie osi modelu szybowca z płaszczyzny poziomej (pochylenie go w przód lub w tył) powoduje znaczne zwiększenie liczby wyładowań w dziób, ogon lub końce skrzydeł, a praktycznie zmniejszenie do zera prawdopodobieństwa uderzenia w kabinę. W związku z tym określono doświadczalnie kąt nachylenia, przy którym uderzenie pioruna w kabinę jest już bardzo mało prawdopodobne (rys. 22). Kąt ten wynosi $\alpha = \pm 7^\circ$.

W oparciu o wyniki badań przedstawione na rys. 19 i 22 spróbowano oszacować prawdopodobieństwo uderzenia pioruna w kabinę we wszystkich położeniach szybowca w stosunku do kierunku uderzenia pioruna. W tym celu otoczono go umyślną powierzchnią kulistą o promieniu równym odległości ostrza piorunowego od modelu (rys. 23). Na powierzchni

kulistej wydzielono powierzchnię S , określającą kierunek przestrzeni, z którego możliwe jest dojście lidera, prowadzące do wyładowania w kabine. Przyjmując założenie, że każdy kierunek zbliżania się lidera jest jednakowo prawdopodobny, wyznaczono prawdopodobieństwo nadejścia go



Rys. 22. Wpływ nachylenia osi podłużnej modelu szybowca na prawdopodobieństwo trafienia wyładowania w kabinę



Rys. 23. Szkic do określenia prawdopodobieństwa trafienia wyładowania w kabinę. Powierzchnia S przedstawia miejsca umieszczenia ostrza piorunowego, z których prawdopodobne jest trafienie wyładowania w wierzchołek kabiny

z wyróżnionego kierunku określonego powierzchnią S . Stosunek powierzchni S do powierzchni sferycznej odpowiadającej kątowi przestrzennemu pełnemu wynosi:

$$K = \frac{S}{4\pi \cdot H^2} = 0,005.$$

Biorąc pod uwagę, że średnie prawdopodobieństwo trafienia w kabinę wyładowania nadchodzącego z wyróżnionej powierzchni wynosi około 10%, można spodziewać się, że jedynie jedno na około 2000 wyładowań piorunowych w szybowiec dosięgnie kabiny.

Dla zorientowania się w zagrożeniu pilota tymi uderzeniami można oprzeć się na eksperymentalnych badaniach amerykańskich [11] opisanych w rozdziale 3. Z badań tych wynika, że przeciętnie na 65 lotów w chmurach burzowych przypada 1 uderzenie pioruna w samolot.

Jeśli to samo prawdopodobieństwo przyjąć dla naszych szybowców i naszych warunków klimatycznych, wówczas uderzenia pioruna w kabinę, które — mimo ochrony — mogłyby spowodować przynajmniej czasową niezdolność pilota do kierowania szybowcem, mogłyby zdarzyć się niezmiernie rzadko, jeden raz na około 130 tysięcy lotów w chmurach burzowych.

7. WNIOSKI

— Ze wzoru J. S. Townsenda dla ulotu wynika, że istnieje zależność wiążąca potencjał szybowca z prądem ulotu i potencjałem początkowym jonizacji.

— Ustalono jako kryterium modelowania wyładowań piorunowych w szybowce nieprzekraczanie przez wyładowanie wstępne z modelu długości 20 m (w skali), przy czym wyładowania te nie powinny zwierać więcej niż $1/4 \dots 1/5$ przerwy iskrowej.

— Na podstawie wyznaczenia przestrzeni chronionej stwierdzono, że zwody umieszczone na sterach, na końcach skrzydeł i na dziobie chronią wszystkie części szybowca typu „Bocian” z wyjątkiem wysuniętych hamulców i wierzchołka kabiny. Prawdopodobieństwo uderzenia pioruna w wierzchołek kabiny określone doświadczalnie (na modelu) wydaje się niezwykle małe (jedno uderzenie na 130000 lotów w chmurach burzowych).

Instytut Elektrotechniki

WYKAZ LITERATURY

1. Townsend J. A.: Die Ionisation der Gase (Artykuł w Handbuch der Radiologie Marxa, t. I, 1920).
2. Kimer D., Mc Gee J. W.: Static Radio Interference. — Proceedings of the IRE Vol. 34, maj 1946, s. 234.
3. Harrison L. P.: Lightning Discharges to Aircraft and Associated Meteorological Conditions. Washington D. C., maj 1946.
4. Lenard P.: Über Wasserfallelektrizität und über Oberflächenbeschaffenheit der Flüssigkeiten. — An. der Phys. t. 47, 1915, s. 463.
5. Stimmel R. G., Rogers E. H., Waterfall F. E., Gun R.: Electrification of Aircraft Flying in Precipitation Areas. — Proceedings of the IRE, vol. 34, kwiecień 1946, s. 167.

6. Waddel R. C., Drutowski R. C., Blat W. N.: Aircraft Instrumentation for Precipitation-Static Research. Proceedings of the IRE, vol. 34, kwiecień 1946, s. 161.
7. Gun R., Parker J. T., Mc Gee J. W.: The High-Voltage Characteristics of Aircraft in Flight. — Proceeding of the IRE vol. 34, maj 1946, s. 241.
8. Bryant J. M., Newman M. M., Robb J. D.: Aircraft Protection from Thunderstorm. Discharges to antennas. — A.I.E.E. Trans., vol. 72, II. 1953, s. 248.
9. Jakubowski J. L.: Technika wysokich napięć, 1951.
10. Van Dorsten A. L.: Blitzschlag in Flugzeugen. — Philips Technische Rundschau, 1955, 16, s. 290.
11. Hagenguth J. H.: Lightning Stroke Damage to Aircraft. — A.I.E.E. Trans., vol. 68, 1949, s. 1036.
12. Workman E. J., Holzer R. E.: The Electrical Structure of Thunderstorms. Washington D. C., listopad 1942. (National Advisory Committee for Aeronautics).
13. Jakubowski J. L.: Wpływ światła na strefę chronioną piorunochronów prętowych. — Archiwum Elektr. 1953, s. 165.
14. Jakubowski J. L.: Badania modelowe strefy chronionej piorunochronów prętowych przy świetleniach zwodu. — Archiwum Elektr. 1957, s. 9.
15. Akopian A. A.: Recherches de laboratoire sur les zones protégées par des parafoudres à tiges multiples. C.I.R.R.E. 1937, t. II, s. 328.
16. Golde R. H.: Surge Phenomena, Londyn 1941, wyd. Electrical Research Association.

Е. БАДЭР, Я. Л. ЯКУБОВСКИ, Е. ЗЕЛИНСКИ

ОПАСНОСТЬ ДЛЯ ПЛАНЁРОВ ОТ ГРОЗОВЫХ РАЗРЯДОВ

Резюме

В статье обсуждается опасность, которую составляют для планёров грозовые разряды; в отдельности обсуждаются условия возникновения значительных напряжённостей электрических полей, а также локализация мест предопределённых для поражения. Определение этих мест даёт возможность правильного проектирования защитной установки. В работе приведены также основные положения проекта громоотводной установки планёра типа „Боцян”. Существуют две основные причины возникновения значительной напряжённости электрического поля при поверхности планёра. Одной из них является накопление зарядов на планёре во время полёта — выступающая тогда напряжённость поля представлена на основании формулы Тоунсенда зависимостью (4). Эта зависимость определяет связь между напряжённостью поля и током заряжающим планёр. Далее авторами обсуждён вопрос возникновения тока заряжающего планёр. На основании материалов из технической печати доказано, что вторая причина, то есть индукционное распределение зарядов в планёре вызванное внешним электрическим полем, приводит к самым значительным напряжённостям полей вблизи планёра.

В работе представлена сводка известных из печати результатов исследований параметров грозových разрядов поражающих самолёты и планёры. При определении мест планёра подвергнутых опасности попадания применён метод

моделирования (планёр был воспроизведён в соответственном масштабе — рис. 5, 6 и 12). В качестве критерии моделирования грозовых разрядов поражающих планёры, по аналогии к моделированию наземных объектов, принято, что предварительные разряды из модели не должны превосходить 20 м (в масштабе), причём они не должны замыкать более $1/4 \dots 1/5$ искрового перерыва (см. рис. 9 и 17).

С целью предварительного определения мест образования значительной напряжённости поля произведены моделировочные исследования в условиях медленно изменяющегося напряжения. Не полные разряды из модели представлены на рис. 7.

Лабораторные исследования в условиях ударного напряжения были произведены с помощью ударного генератора в 2,8 мкВ, 32 кВтсек Электротехнического Института в Варшаве. На рис. 10 представлены примеры разрядов поражающих разные места модели. Полученное во-время исследований очертание защищённого пространства представлено на рис. 14. На основании результатов предварительных исследований составлен проект расположения отводов (рис. 16) на изготовленном двухместном планёре польского производства типа „Бочан” (рис. 15). Для разных положений модели по отношению к направлению разряда (рис. 18) определена кривая вероятности попадания разряда в отдельные места планёра вышеуказанного типа (рис. 19). В результате исследований доказано, что отводы расположенные на рулях, на концах крыльев и на носовой части защищают все элементы планёра типа „Бочан”, исключая выдвинутые аэродинамические тормоза и верхушку кабины. Вероятность удара в верхушку кабины была определена экспериментально на модели (рис. 23) и является весьма незначительной (соответствует одному удару на 130 000 полётов в грозовых ту-чах). Тормоза следует полагать отводами и соответственно соединять с защитной установкой.

Дальнейшие работы Кафедры Высоких Напряжений Варшавского Политехнического Института по противоударной защите планёров будут изложены в следующей статье.

J. BADER, J. L. JAKUBOWSKI, J. ZIELIŃSKI

LIGHTNING DANGER FOR GLIDERS

Summary

In the present paper the danger of lightning to which a glider may be exposed is discussed. Particular attention is given to the conditions in which great strengths of the electric field occur and the parts of the glider most liable to be struck by lightning are determined. The knowledge of these points permits to design a correct lightning protection installation. The authors also give guiding principles for a lightning protection installation for gliders of the „Bocian” type.

There are two causes of occurrence of high strengths of the electric field on the glider surface. The first is the accumulation of the electric charge during flight — the strengths of the electric field in this case are expressed by formula (4) in terms of Townsend's equation. This formula gives the relation between the strength of the electric field and the current loading the glider. The authors further

consider the causes of the appearance of the current loading the glider. The second cause, as results from technical literature, is an influence distribution of electric charges in the glider due to an external electric field. This fact causes the appearance of the greatest electric field strengths in the neighbourhood of the glider.

The paper presents the results of research on the parameters of lightning discharges to aeroplanes and gliders, known in the technical literature. In order to determine the parts of the glider most liable to be struck by lightning, a model method was used (the model scale according to Figs. 5, 6 and 12). The adopted criterion of modelling lightning discharges to gliders, analogically to that used for modelling objects on ground level, is that initial discharges from the glider may not exceed 20 m (in scale) and should not short-circuit more than $1/4$ to $1/5$ of the discharge clearance (see Figs. 9 and 17). For a preliminary determination of points of great field strengths, model measurements were carried out at a voltage of industrial frequency. The partial discharges from the model are shown in Fig. 7.

Laboratory experiments with impulse voltage were performed by aid of a surge generator 2,8 MV, 32 kW at the Electrotechnical Institute of Warsaw. Fig. 10 shows examples of discharges to different parts of the model. The outline of the protected area obtained during these measurements is presented by Fig. 14. On the basis of the results obtained in preliminary measurements the arrangement of lightning leads (Fig. 16) was designed for a two-seat glider of Polish construction of the „Bocian” type (Fig. 15). Probability curves (Fig. 19) of discharges to particular parts of the glider were plotted for the model of the glider of „Bocian” type at various positions of the glider in relation to the direction of discharge (Fig. 18). The investigations showed that the lightning leads placed on the empennage and ailerons on the ends of the wings and on the nose protect all parts of the glider except the protruding aerodynamic brakes (Fig. 21) and the cabin top. The probability of lightning striking the cabin top as established experimentally on the model is extremely small (corresponding to 1 stroke in 130.000 flights in storm clouds). The brakes should be considered as leads and should be connected to the lightning protection installation.

Further results of work on lightning protection of gliders at the High Voltage Department of the Warsaw Polytechnic will be published in a subsequent paper.

ZBIGNIEW KACZKOWSKI, ADAM SMOLIŃSKI

Ferryty magnetostrykcyjne i ich zastosowanie

Rękopis dostarczono 19.11.1959

Podano podstawowe definicje wielkości określających własności magnetostrykcyjne materiału. Na podstawie danych zagranicznych przeprowadzono porównanie własności magnetycznych, magnetostrykcyjnych i fizycznych materiałów metalicznych (tablica 1) i ferrytów (tablica 2). W tablicach 3 i 4 podano własności ferrytów produkowanych specjalnie do celów magnetostrykcyjnych. W dalszej części pracy omówiono wymagania stawiane ferrytom stosowanym w przetwornikach ultradźwiękowych (nadajnikach i odbiornikach) oraz urządzeniach radiotechnicznych (filtrach, stabilizatorach częstotliwości i liniach opóźniających). W zakończeniu podano wykaz bibliografii dotyczącej magnetostrykcji ferrytów i dziedzin z nią związanych.

1. WSTĘP

Odkryte przez Joule'a w 1842 r. [1], [2] ¹⁾ zjawisko zmiany wymiarów geometrycznych materiałów magnetycznych pod wpływem przyłożonego pola magnetycznego wraz z odkrytym przez Villariego w 1865 r. [3] zjawiskiem zmiany własności magnetycznych pod wpływem wprowadzonych naprężeń znajduje, począwszy od lat dwudziestych naszego stulecia [4], [A], [5], [6], [7], [8], coraz szersze zastosowanie w przetwornikach ultradźwiękowych [B], [C], [D], [9] oraz urządzeniach telekomunikacyjnych [B], [C], [E], [10], [11], [12], [13], [14], [15].

Dziedzina magnetyzmu obejmująca zjawiska związane z odwracalnymi zmianami magnetomechanicznymi, a więc związane ze zmianą kształtu, wymiarów geometrycznych i modułu sprężystości pod wpływem przyłożonego pola magnetycznego oraz ze zmianą własności magnetycznych pod wpływem zmian naprężeń, nosi nazwę piezomagnetyzmu i jest częścią

¹⁾ Wykaz literatury składa się z dwu części: w pierwszej, oznaczonej liczbami kolejnymi, podano publikacje zamieszczone w czasopiśmie, w drugiej, oznaczonej dużymi literami — literaturę książkową.

ogólniejszej dziedziny obejmującej również zmiany nieodwracalne, a noszące nazwę magnetostrykcji.

Odpowiednie materiały, w których te zjawiska występują, i przetworniki na nich oparte noszą nazwę piezomagnetycznych lub magnetostrykcyjnych.

W chwili obecnej technika dysponuje następującymi materiałami magnetostrykcyjnymi:

1. metale [F], [G]: żelazo, kobalt, nikiel;
2. stopy binarne [F], [G], [H]: Fe-Ni (permaloje), Al-Fe (alfery [16] i alfenole [17]) i Ni-Co [18];
3. stopy ternarne [F], [G], [H]: Fe-Co-Cr (hiperco) [19], Fe-Co-V (permendury wanadowe) i Fe-Ni-Cr (elinwary);
4. ferryty [I], [J], [20]: żelazowe (magnetyt), niklowe, litowe, cynkowe, niklowo-cynkowe, niklowo-kobaltowe, niklowo-miedziowe i bardziej złożone;
5. materiały proszkowe [64].

W zastosowaniach technicznych z metali używa się jedynie niklu; żelazo i kobalt nie znajdują zastosowań praktycznych.

Magnetostrykcyjne przetworniki ultradźwiękowe [B], [C], [D], [E], [K], [21], [22] znalazły zastosowanie w morskich urządzeniach hydroakustycznych (echosondy, hydrolokatory [L]), w badaniach materiałowych (betonoscopia, defektoscopia), przy spawaniu aluminium, w urządzeniach medycznych (w terapii, dentyście, przy usuwaniu nowotworów, kamieni żółciowych itp.), przy wierceniach w szkłe i innych kruchych materiałach [23], przy wygładzaniu powierzchni, przy usuwaniu zanieczyszczeń z mechanizmów, przy przyspieszaniu procesów technologicznych itp.

Materiały magnetostrykcyjne znalazły również zastosowanie w urządzeniach telekomunikacyjnych takich, jak filtry magnetostrykcyjne i elektromechaniczne, stabilizatory częstotliwości [B], [C], [D], [E], [12], [13], [14], [15], [24] oraz magnetostrykcyjne linie opóźniające [10], [11], [25].

We wszystkich wymienionych urządzeniach projektowanych pierwotnie przy użyciu magnetycznych materiałów metalicznych coraz częściej, począwszy od 1949 r. [12] stosuje się ferryty, które ze względu na dużą oporność właściwą i inne korzystne właściwości są dużo odpowiedniejsze od materiałów metalicznych przy pracy przy wielkich częstotliwościach. Poważną wadą ferrytów jest ich kruchość i mała indukcja nasycenia, lecz niskie koszty produkcji i inne równorzędne, a często i lepsze własności magnetostrykcyjne i magnetyczne (w porównaniu z materiałami metalicznymi) przemawiają za stosowaniem ferrytów we wszelkiego rodzaju przetwornikach magnetostrykcyjnych, w których wymienione wady nie mają decydującego znaczenia.

2. WIELKOŚCI OKREŚLAJĄCE WŁASNOŚCI MATERIAŁÓW MAGNETOSTRYKCYJNYCH

Obok wielkości charakteryzujących materiał pod względem magnetycznym i elektrycznym, takich jak indukcja nasycenia B_s , remanencja B_r (pozostałość magnetyczna), natężenie powściągające H_c , przenikalność początkowa μ_p , przenikalność maksymalna μ_m , przenikalność odwracalna μ (μ_0), temperatura Curie T_c , oporność właściwa ϱ , dobroć elektryczna Q_e , tangens kąta strat $\tan \delta$ (lub stratność) — w technice wykorzystującej materiały magnetostrykcyjne niezbędna jest znajomość wielkości określających materiał pod względem magnetostrykcyjnym.

Do tych wielkości wchodzi zarówno znane wielkości fizyczne, takie jak gęstość γ (lub ciężar właściwy), moduł Younga E (lub jego odwrotność — współczynnik sprężystości α [26]), moduł sztywności G , dobroć mechaniczna Q_m , częstotliwość rezonansu mechanicznego f_r (zależna od wymiarów i kształtu próbki i właściwości materiałowych), szybkość rozchodzenia się dźwięku w materiale v , porowatość p , wytrzymałość mechaniczna, jak również wielkości ściśle związane z magnetostrykcją, jak współczynnik magnetostrykcji lub — krócej — magnetostrykcja λ (oznaczana też przez ε , np. [I], [20]), przy czym podawana jest zwykle magnetostrykcja nasycenia λ_s , współczynnik sprzężenia magnetomechanicznego k , przy czym podaje się zwykle wartości dla remanencji k_r oraz maksymalne k_m , czułość magnetostrykcyjna Δ oraz spotykana w literaturze japońskiej [27], [28] magnetostrykcyjność Γ , i inne nie uporządkowane jeszcze określenia.

Przed podaniem definicji określających wielkości magnetostrykcyjne zostanie pokrótce omówione samo zjawisko magnetostrykcji.

O własnościach magnetycznych stanowią nieskompensowane momenty magnetyczne spinów elektronowych z warstwy 3d (dla Fe, Ni, Co, Mn). Obszar, w którym nieskompensowane magnetycznie spiny poszczególnych atomów są równoległe, nosi nazwę domeny [F], [L], [M], [N].

Pod wpływem odpowiednio silnego zewnętrznego pola magnetycznego w jednodomenowym kryształ nie tylko nastąpi obrót wektorów magnetyzacji (równoległe do kierunku przyłożonego pola), ale równocześnie ulegnie zmianie kształt kryształu. W materiale polikrystalicznym deformacji elementarnego kryształu w pewnym stopniu sprzeciwią się kryształy sąsiednie i wtedy w kryształ powstaną naprężenia wewnętrzne. Odwrotnie, pod wpływem sił zewnętrznych, obok odkształceń elastycznych, nastąpi zmiana orientacji spinów w kryształ, co pociąga za sobą zmianę własności magnetycznych. Ze względu na anizotropię kryształu na wielkość tych zmian ma wpływ wzajemne położenie kierunku przyłożonego pola i osi kryształu.

Rozpatrując zagadnienie makroskopowo można stwierdzić, że pod wpływem pola magnetycznego materiał wydłuża się w jego kierunku, równocześnie kurcząc się w kierunkach poprzecznych (magnetostrykcja wzdluzna dodatnia) lub kurczy się w kierunku pola zwiększając swój przekrój poprzeczny (λ wzdluzna ujemna).

Obok zjawiska Joule'a i Villariego znane są zjawiska pochodne mające z nim ścisły związek. I tak zjawisko Guillemina (1846 r.) [29], [30] polega na tym, że zginany (lub zgięty) pręt wykonany z materiału magnetycznego dąży do wyprostowania pod wpływem wzdluznego pola magnetycznego.

Zjawisko Wiedemanna (1862 r.) [31], [F], [30], [M] dotyczy zachowania się zamocowanego na jednym końcu pręta wykonanego z materiału magnetycznego, który jest magnesowany wzdluznie przez zewnętrzne pole i równocześnie obwodowo przez pole pochodzące od przepływającego wzdluz osi pręta prądu elektrycznego. Pod wpływem działania tych pól nastąpi skręt wolnego końca.

Magnetostrykcję określa statyczny współczynnik magnetostrykcji λ obrazujący względne zmiany wymiarów próbki, który jest podawany jako funkcja natężenia pola H [F], [G], [M], [26], [30], indukcji B [32] lub magnetyzacji J [30], [33]

$$\lambda = \frac{\Delta l}{l}. \quad (1)$$

Współczynnik λ przyjmuje znak (+) w przypadku rozciągania i (−) w przypadku kurczenia się próbki pod wpływem wzdluznego pola magnetycznego.

Należy podkreślić, że magnetostrykcja λ jest funkcją parzystą wielkości magnetycznych (tzn. przy zmianie znaku B lub H λ nie zmienia znaku).

Przy badaniach kierunkowych magnetostrykcyjnych, przy współczynniku λ umieszcza się odpowiednio indeks równoległości $\lambda_{||}$ lub prostopadłości $\lambda_{\perp}$ oznaczający kierunek pomiaru w stosunku do przyłożonego pola (przy innych kierunkach oznacza się wartość kątową np. λ_{45}). Przy badaniach kryształów umieszcza się jako indeks nazwę osi np. λ_{100} , λ_{111} , λ_{110} . Wartość współczynnika λ przy nasyceniu nosi nazwę magnetostrykcji nasycenia λ_s .

Przy pomiarach $\lambda = f(H)$ lub $\lambda = f(J)$ występuje podobnie jak przy pomiarach magnetycznych $B = f(H)$ zjawisko histerezy [F], [G], [26], [30], [32].

W pobliżu nasycenia obok magnetostrykcji liniowej występuje magnetostrykcja objętościowa ω [F], [G], [30]

$$\omega = \frac{dV}{V} = \frac{dl}{l} + \frac{2dr}{r}. \quad (2)$$

Zjawisko magnetostrykcji obok konkretnych zastosowań technicznych ma również bardzo duże znaczenie teoretyczne, gdyż z nim są związane inne wielkości magnetyczne. Zagadnieniami teoretycznymi magnetostrykcji zajmowało się wielu naukowców, między innymi Akułow [34], Heisenberg [35], Wonsowski [36], Kittel [37], Kwiek [38] i Diakow [39].

Przy projektowaniu przetworników magnetostrykcyjnych ważne jest zagadnienie, jaka część energii magnetycznej może być zamieniona na mechaniczną (lub odwrotnie). Stosunek energii magnetycznej przetworzonej na mechaniczną E_w do całej energii magnetycznej E określa się jako kwadrat współczynnika sprzężenia magnetomechanicznego [26], [20]:

$$k^2 = \frac{E_w}{E} = \frac{\mu_\sigma - \mu_\lambda}{\mu_\sigma} = \frac{\frac{1}{E_H} - \frac{1}{E_B}}{\frac{1}{E_H}}, \quad (3)$$

gdzie:

- μ_σ — przenikalność odwracalna próbki swobodnej,
- μ_λ — przenikalność odwracalna próbki zamocowanej,
- E_H — moduł Younga przy stałym natężeniu pola,
- E_B — moduł Younga przy stałej indukcji.

Współczynnik k podaje się jako wielkość niemianowaną lub rzadziej w % [28], [40], [41].

Podobnie jak przenikalność przyrostowa μ_Δ jest tangensem kąta nachylenia pętłki magnesowania przy pracy z podkładem stałego pola magnetycznego

$$\mu_{\Delta\sigma} = \left(\frac{\Delta B}{\Delta H} \right)_\sigma, \quad (4)$$

czułość naprężeń Δ jest tangensem kąta nachylenia pętłki zmiennej na tle stałej polaryzacji

$$\Delta_\Delta = \left(\frac{\Delta B}{\Delta \sigma} \right)_H = 4\pi \left(\frac{\Delta \lambda}{\Delta H} \right)_\sigma. \quad (5)$$

Wielkości te w granicznym przypadku, gdy $\Delta H \rightarrow 0$, stają się odwracalne (μ , Δ) [26], [32].

Magnetostrykcyjność Γ [27], [28], [40] oznacza iloczyn pochodnej cząstkowej magnetostrykcji względem magnetyzacji i modułu Younga

$$\Gamma = \left(\frac{\partial \lambda}{\partial J} \right)_\sigma E_j. \quad (6)$$

Spotyka się często określenie stałej magnetostrykcji [J], [17], [18], [65]

$$h = \left(\frac{\partial \sigma}{\partial B} \right)_\lambda, \quad (7)$$

przy czym zamiast h w przytoczonych publikacjach używa się oznaczenia λ (a zamiast $\lambda - \varepsilon$).

3. UKŁADY RÓWNAŃ MAGNETOSTRYKCYJNYCH

Przy matematycznym opisie magnetostrykcji przyjmując, że proces przebiega adiabatycznie, otrzymuje się zespół czterech wielkości (magnetycznych: H i B lub J oraz mechanicznych: λ i σ) wzajemnie ze sobą powiązanych. Zależnie od przyjętych par wielkości zależnych i niezależnych (przy czym do każdej pary wchodzi jedna wielkość magnetyczna i jedna mechaniczna) można otrzymać cztery warianty równań opisujących własności magnetostrykcyjne materiału.

Równania są słuszne dla małych amplitud zmiennych, kiedy: 1) mają miejsce przebiegi odwracalne, 2) praca odbywa się dużo poniżej rezonansu mechanicznego oraz 3) nie występuje depolaryzacja.

W zakresie tej pracy nie leży omawianie wszystkich wariantów równań magnetostrykcyjnych; podany zostanie tutaj jedynie układ tylko najbardziej usprawiedliwiony fizycznie, w którym przyjęto jako wielkości pierwotne natężenie pola H i naprężenia σ , a jako wielkości wtórne indukcję B i wydłużenie względne λ [26].

Różniczki zupełne indukcji i wydłużenia względnego przyjmują postać:

$$dB = \left(\frac{\partial B}{\partial \sigma} \right)_H d\sigma + \left(\frac{\partial B}{\partial H} \right)_\sigma dH, \quad (8)$$

$$d\lambda = \left(\frac{\partial \lambda}{\partial \sigma} \right)_H d\sigma + \left(\frac{\partial \lambda}{\partial H} \right)_\sigma dH, \quad (9)$$

gdzie:

$$\left(\frac{\partial B}{\partial \sigma} \right)_H = A \text{ — czułość naprężeń,}$$

$$\left(\frac{\partial B}{\partial H} \right)_\sigma = \mu_\sigma \text{ — przenikalność odwracalna próbki swobodnej,}$$

$$\left(\frac{\partial \lambda}{\partial \sigma} \right)_H = \frac{1}{E_H} \text{ — współczynnik sprężystości, odwrotność dynamicznego modułu Younga przy stałym natężeniu pola.}$$

Z I i II zasady termodynamiki wynika tożsamość [42], [43], [F], [N]:

$$\left(\frac{\partial B}{\partial \sigma} \right)_H = 4\pi \left(\frac{\partial \lambda}{\partial H} \right)_\sigma = A. \quad (10)$$

Po odpowiednich podstawieniach otrzymuje się następujący układ równań magnetostrykcyjnych:

$$dB = A d\sigma + \mu_\sigma dH, \quad (11)$$

$$d\lambda = \frac{1}{E_H} d\sigma + \frac{A}{4\pi} dH. \quad (12)$$

Ze względu na dotychczas stosowany powszechnie w Polsce w dziedzinie magnetyzmu układ jednostek CGSM w pracy niniejszej również przyjęto ten układ, mimo że w publikacjach zagranicznych z dziedziny magnetostrykcji coraz częściej spotyka się jednostki układu zracjonalizowanego MKSA zalecanego przez normy na materiały piezoelektryczne [44].

Należy podkreślić, że występuje wielka dowolność w oznaczeniu i definiowaniu współczynników związanych z magnetostrykcją, dlatego przed porównaniem wartości liczbowych należy zapoznać się z definicją wielkości i stosowanym w danej publikacji układem jednostek (np. wielkości odpowiadające A mogą być mniejsze od $A 4\pi$ razy). Współczynniki występujące w równaniach magnetostrykcji i współczynnik k powiązane są następującą zależnością:

$$k^2 = \frac{A^2 E_H}{4\pi\mu_\sigma}. \quad (13)$$

Innym powszechnie stosowanym układem równań jest układ przyjmujący jako zmienne niezależne λ i B [18], [19], [J]; wtedy:

$$k^2 = \frac{4\pi h^2 \mu_\lambda}{E_B}. \quad (14)$$

4. FERRYTY A MATERIAŁY METALICZNE

Przetworniki ultradźwiękowe pracują w zakresie częstotliwości do kilkudziesięciu MHz. Zależnie od zastosowań wymagane są różne własności magnetostrykcyjne, magnetyczne, elektryczne i fizyczne materiału. Wymagania te dla poszczególnych typów przetworników zostaną omówione w dalszej części pracy. Tu przedstawione zostaną tylko wymagania ogólne.

O własnościach magnetostrykcyjnych decyduje współczynnik wydłużenia względnego λ . Zarówno stosowane materiały metaliczne, jak i ferryty mają statyczne współczynniki wydłużenia względnego dla nasycenia λ_s tego samego rzędu. Dla najlepszych stosowanych materiałów metalicznych współczynnik λ_s dochodzi do $+70 \cdot 10^{-6}$ (2 V permendur), a nawet do $78 \cdot 10^{-6}$ (stop 35% Fe, 57% Co, 8% Cr [53]), dla ferrytów zaś do $+41 \cdot 10^{-6}$ (ferryt żelazowy — magnetyt). W grupie materiałów o ujemnym współczynniku magnetostrykcji nikiel ($-45 \cdot 10^{-6}$ [I]) i ferryt nikłowy

Magnetyczne, magnetostrykcyjne i fizyczne właściwości materiałów metalicznych i Ferebee [17], [49], Huertera i Bolta [D], Poppera [I],

Badana wielkość		Sym-bol	Jednostka	Nikiel		Permaloj 45	
Skład			% masy	99 99,9	[F] [20]	45 Ni 55 Fe	
Indukcja nasycenia		B_s	kGs	5,33 6,3	[17] [20]	16	[F],[33]
Przenikalność początkowa		μ_p	Gs/Oe	110	[F]	2500	[F]
Przenikalność maksymalna		μ_m	Gs/Oe	600 2500	[F] [17]	25000	[F]
Gęstość		γ	g/cm ³	8,7 8,9	[D] [F],[20]	8,17 8,25	[F] [D]
Magnetostrykcja nasycenia		λ_s	10 ⁻⁶	-33 -45	[33] [I]	+27	[33]
Temperatura Curie		$T_c(\Theta)$	°C	358	[F]	440	[33]
Oporność właściwa		ϱ	$\mu\Omega$ cm	7	[F]	45 70	[F] [D]
Wielkości magneto-mechan. przy B_r	Współczynnik sprzężenia magnetomechan.	k_r	—	0,14 0,23	[20] [18]	0,124 0,17	[D] [32]
	Czułość naprężeń	Δ_r	$\frac{\mu \text{ Gs cm}^2}{\text{dyn}}$	-1,5 -7,2	[20] [18]	4,0 6,0	[33]
	Przenikalność odwracalna	μ_r	Gs/Oe	18 160	[20] [18]	200 350	[33]
	Prędkość dźwięku	v	km/s	4,75 4,9	[20] [17]	4,1	[33]
	Moduł Younga	E_{Hr}	10 ¹² dyn/cm ²	2,0 2,13	[20] [D]	1,38	[D]
Wielkości magneto-mechan. dla k_m	Współczynnik sprzężenia magnetomechan.	k_m	—	0,30 0,31	[18] [17]	0,35	[49]
	Czułość naprężeń	Δ_{km}	$\frac{\mu \text{ Gs cm}^2}{\text{dyn}}$	-6,1	[18]		
	Przenikalność odwracalna	μ_{km}	Gs/Oe	22 70	[49] [18]	100	[49]

Uwaga. Wyniki podane w tablicy zależą od technologii, stąd duże rozbieżności w niektórych z różnych źródeł nie uwzględniając różnic technologicznych.

Tablica 1

wg Bozortha [F], van der Burgta [20], [32], [33], Clarka [18], Davisa Reinbotha [M], Sussmana i Ehrlicha [19].

2V Permendur	Hiperco	Stop 4 Co-Ni	Stop 18 Co-Ni	Alfer 13	Alfer 12 (alfenol)
2V (1,7V [20]) 49 Co 49 Fe	35 Co 0,5 Cr 64,5 Fe	4 Co 96 Ni	18 Co 82 Ni	13 Al 87 Fe	12,5 Al 87,5 Fe
24 [F],[20]	24,2 [F]	6,8 [32]	8,0 [18]	12,2 [17] 13 [32]	13,2 [17]
800 [F]	650 [F]	660 [18]	1450 [18]		
4500 [F]	10000 [M]	~3000 [18]	~4000 [18]	13200 [17]	10800 [17]
8,2 [33]	8,0 [F] 8,1 [20]	8,9 [18]	8,9 [18]	6,3 [17] 6,7 [33]	6,6 [17]
+70 [20],[33]	+55 [20]	-31 [32]		+40 [32]	+35
980 [20],[33]	970 [F] 980 [M]	410 [18]	570 [18]	500 [32]	~500
25 40 [33]	20 [F] 25 [20]	10 [32]		90 [32]	~90
0,13 0,26 [32]	0,14 [19]	0,34 [18]	0,23 [18]	0,19 0,26 [32]	
3,7 [20]	2,6 [19] 3,8 [20]	-14,8 [18]	-12,1 [18]	8,0 11 [33]	
53 [22]	123 [19]	250 [18]	390 [18]	200 350 [33]	
5,18 [33]	5,1 [20]			4,5 [17] 4,75 [33]	4,5 [17]
2,2 [20]	2,1 [20]	1,5 ÷ 1,9 [18]	1,7 [18]	1,5 [17]	1,5 [17]
	0,17 [19]	0,51 (0,59) [18]	0,42 (0,46) [18]	28 [17]	0,30 [17]
	3,6 [19]	13,5 [18]	12,7 [18]		
	75 [19]	130 [18]	150 [18]	80 [17]	100 [17]

pozycjach. Należy również zaznaczyć, że wielkości dotyczące danego materiału czerpano

Tablica 2
Magnetyczne, magnetostrykcyjne i fizyczne własności ferrytów wg van der Burgta [20], [32], [33] i Poppera [1]

Badana wielkość	Symbol	Jednostka	Ferryt żelazowy (Magnetyt)	Ferryt barowy	Ferryt kobaltowy	Ferryt litowy	Ferryt manganowy	Ferryt miedziowy	Ferryt niklowy
Skład	—	—	Fe_3O_4	$\text{BaFe}_{12}\text{O}_{13}$	CoFe_2O_4	$\text{Li}_{0,5}\text{Fe}_{2,5}\text{O}_4$	MnFe_2O_4	CuFe_2O_4	NiFe_2O_4
Indukcja nasycenia	B_s	kGs	~ 6 [33]	$\sim 4,8$	5 [32]	3,9 [32]	5,2	1,7	3,3 [32]
Gęstość	γ	g/cm^3	5,24 [33]	4,6 [1]	5,29 [33]				
Magnetostrykcja nasycenia	λ_s	10^{-6}	+40 [33]	-5 [1]	-200 [32]	-8 [32]	-5	-10	-27 [32]
Temperatura Curie	T_c	$^{\circ}\text{C}$	580 [33]	450 [1]	510 [32]	640 [32]	300	455	590 [32]
Oporność właściwa	ρ	Ωcm	10^{-2} [33]	$35 \cdot 10^6$ [1]	10^6 [32]	10^6 [32]			
Współczynnik sprężenia magnetomechan.	k_r	—	—	0,001 [20]	—	0,03 0,05 [32]			0,14 0,20 [32]

($-27 \cdot 10^{-6}$) mają duże wartości tego współczynnika. Oczywiście nie są to szczytowe osiągnięcia. Istnieją stopy platyny z żelazem (63% Fe, 37% Pt) o współczynniku $\lambda_s = 170 \cdot 10^{-6}$ [45] i żelaza z kobaltem (30% Fe, 70% Co) o $\lambda_s = 130 \cdot 10^{-6}$ [53], ale ze względu na duże koszty produkcji nie znalazły powszechniejszego zastosowania. W ferrytach kobaltowych λ_s jest jeszcze większe i dochodzi do $-210 \cdot 10^{-6}$. Należy dodać, że największą magnetostrykcję nasycenia ferrytów $\lambda_s = +870 \cdot 10^{-6}$ [46] pomierzono w monokryształach ferrytu kobaltowego dla kierunku [100], przy czym pole było przyłożone prostopadłe do tego kierunku. Mimo tak dużych współczynników λ_s ferryty kobaltowe nie znalazły praktycznego zastosowania, gdyż należą do materiałów magnetycznie twardych. Największy współczynnik $\lambda_s = 2750 \cdot 10^{-6}$ pomierzono dla materiału proszkowego o zawartości 75% Fe [64].

W tablicy 1 podano współczynnik magnetostrykcji nasycenia i inne wielkości magnetostrykcyjne materiałów metalicznych, powszechnie stosowanych w przetwornikach ultradźwiękowych. Biorąc pod uwagę możliwości krajowe pierwszeństwo należy oddać permalojowi o zawartości 45% Ni przed niklem i alferem (o ile zostanie opracowana produkcja alferu, materiał ten ze względu na niski koszt, wystarczające zasoby krajowe surowców i dobre własności magnetostrykcyjne powinien wyprzeć pozostałe materiały metaliczne).

W tablicach 2 i 3 podano odpowiednie własności magnetostrykcyjne ferrytów spotykanych na rynku światowym. W chwili obecnej w przetwornikach ultradźwiękowych stosuje się najczęściej ferryty niklowe lub niklowo-cynkowe o małej zawartości cynku (0,05 w stosunku molowym). Opracowano cały szereg specjalnych ferrytów z dodatkiem kobaltu [47], [48], [49] lub miedzi [40], mających znacznie lepsze własności magnetostrykcyjne, przewyższające nawet własności najlepszych materiałów metalicznych.

Dobre materiały magnetostrykcyjne mają współczynnik sprzężenia magnetomechanicznego k_r większy od 0,1 a dochodzący do 0,34 (stop 4-Co-Ni [18]). Na wartość k duży wpływ ma sposób obróbki termicznej. Obecnie dla ferrytów niklowych uzyskuje się współczynnik $k_r = 0,15 \div 0,20$ [32], [50], [51], przy czym k_m dochodzi do 0,21 [50]. Ferryty Ni-Cu [40] mają $k_m = 0,216$, a ferryty niklowo-kobaltowe $k_r = 0,25$ [32]. Dla ferrytów niklowo-kobaltowo-żelazowych k_m dochodzi do 0,37 [49]. Wartości współczynnika k ferrytów zależą od porowatości p . Im większa porowatość, tym niższa wartość sprzężenia magnetomechanicznego.

Ostatnio opracowane ferryty magnetostrykcyjne odznaczają się bardzo małą porowatością ($0,01 \div 0,035$) i bardzo dobrymi własnościami magnetostrykcyjnymi [65].

Magnetyczne, magnetostrykcyjne i fizyczne własności ferrytów Ni-Zn

Wielkość badana		Symbol	Jednostka	[20]	[20]	[20]
Fe ₂ O ₃				50	50	50
NiO			mole w %	15	17	17,5
ZnO				35	33	32,5
Indukcja nasycenia		B_s	kGs	2,6	3,3	3,5
Przenikalność początkowa		μ_p	Gs/Oe	1360	890	775
Gęstość		γ	g/cm ³	5,06	4,90	4,90
Magnetostrykcja nasycenia		λ_s	10 ⁻⁶	-2	-4	-4,5
Temperatura Curie		T_c	°C	85		
Oporność właściwa		ϱ	Ωcm	$>10^4$	$>10^4$	$>10^4$
Wielkości magneto mechaniczne dla B_r	Współczyn. sprzęż. magnetomechan.	k_r	—	0,034	0,067	0,066
	Czułość naprężeń	$-\lambda_r$	$\frac{\mu\text{Gs cm}^2}{\text{dyn}}$	3,2	5,0	4,5
	Przenikalność odwracalna	μ_r	Gs/Oe	1210	735	600
	Prędkość dźwięku	v	km/s	5,97	5,79	5,79
	Moduł Younga	E_{Hr}	10 ¹² dyn/cm ²	1,8	1,63	1,63
	Dobroć mechan.	Q	—	2780	1490	1420
Wielkości magneto mech. drgań skrętnych przy B_{opt}	Współczyn. sprzęż. magnetomechan.	k_{mt} (k_m)*	—	0,063	0,089	0,08
	Czułość naprężeń	$-\lambda_t$	$\frac{\mu\text{G scm}^2}{\text{dyn}}$	7,9	7,9	7,3
	Przenikalność odwracalna	μ_t	Gs/Oe	1880	394	418
	Dobroć mechan.	Q_t	—	1740	2000	1640

* Wartość maksymalna współczynnika sprzężenia magnetomechanicznego dla drgań wzdłuż

Tablica 3

wg van Burgta [20] (wiersze górne) i Golaminy [50] (wiersze dolne)

IVA [20]	IVB [20] [50]	[50]	IVC [20]	[50]	[50]	IVD [20] [50]	[50]	[50]	IVE [20] [50]
50	50	50	50	50	50	50	50	50	50
18	25	30	32	35	37,5	40	42,5	45	50
32	25	20	18	15	12,5	10	7,5	5	0
3,7	4,4		4,7			4,2			2,6
705	219		168			85			24,2
	330	221		160	141	117	97	49	44
4,90	4,85		4,85			4,76			4,20
	5,27	5,25		5,23	5,23	5,20	5,20	5,21	5,20
-5	-11		-16			-21			-26
	-9	-13		-16	-18	-20	-21	-23	-26
200	290		410			490			595
>10 ⁵	>10 ⁵		>10 ⁵			>10 ⁵			>10 ⁵
0,073	0,082		0,093			0,112			0,08
	0,10	0,11		0,11	0,115	0,125	0,13	0,15	0,15
4,5	2,8	2	2,5			2,0			1,2
	4,3	3,15		3,3	3,1	3,1	3,1	2,1	2,2
492	140		86			36			19
	290	165		127	103	85	76	35	34
5,76	5,67		5,57			5,37			4,71
	5,94	5,93		5,90	5,89	5,86	5,83	5,8	5,79
1,62	1,53		1,50			1,37			0,93
	1,85	1,84		1,82	1,81	1,78	1,77	1,75	1,75
1260 (14500)	920		1430			2150			2040
0,0115	0,11		0,105			0,13			0,09
0,089*	0,14*	0,15*		0,165*	0,165*	0,17*	0,17*	0,14*	0,21*
11,2	5,7		4,4			3,7			2,1
470	128		85			35			16
2450	3150		4250			5650			4450

nych.

Największą wartość współczynnika k_m , jeśli chodzi o ferryty, podaje firma Kearfott (k_m torroidu wynosi 0,40) dla ferrytu o nazwie firmowej N-51 [66]. Największą wartość k_m pomierzono dla stopu 4 Co-Ni [18] ($k_m = 0,59$).

Trzecią wielkością magnetostrykcyjną jest czułość naprężeń Δ . Największa czułość ferrytów przy remanencji Δ_r wynosi około $5 \cdot 10^{-6}$ Gs cm²/dynę (dla ferrytu Ni-Zn [20]), natomiast największa czułość Δ_r dla materiałów metalicznych wynosi 14,8 μ Gs cm²/dynę (dla stopu 4% Co-Ni [18]).

Materiały metaliczne mają stosunkowo dużą indukcję nasycenia ($8 \div 24$ kGs) w porównaniu z ferrytami, których B_s zawiera się w granicach $2 \div 5$ kGs. Mała indukcja nasycenia stanowi poważną przeszkodę w stosowaniu ferrytów na przetworniki magnetostrykcyjne dużej mocy, np. na nadajniki ultradźwiękowe.

Bardzo ważną zaletą ferrytów jest ich duża oporność właściwa, często o 10 lub 11 rzędów wielkości wyższa (z wyjątkiem magnetytu) od oporności materiałów metalicznych. Pozwala to na stosowanie przetworników w postaci litej, co ułatwia produkcję i polepsza sprawność.

Wypalanie w wysokich temperaturach (1400°C) polepszające własności magnetostrykcyjne i zmniejszające porowatość obniża oporność właściwą [54], natomiast dodanie kobaltu lub manganu zwiększa tę wielkość [55]. Optymalna ilość dodatków zależy od technologii. Można również polepszyć własności magnetostrykcyjne i zmniejszyć porowatość przez zastąpienie części niklu przez miedź w ferrytach Ni-Zn [40].

Przetworniki z materiałów metalicznych ze względu na duże straty na prądy wirowe muszą być wykonane z blaszek, co jest powodem dodatkowych strat na tarcie. Mimo dzielenia na blaszki możliwość użycia metalicznych przetworników ultradźwiękowych w zakresie większych częstotliwości jest ograniczona górną częstotliwością graniczną, która (zależnie od grubości blaszki) pozwala na pracę w zakresie do 100 kHz. Ostatnio coraz powszechniej stosowane stopy typu alfer dzięki nieco większej oporności właściwej ($\rho = 9 \cdot 10^{-5}$ Ω cm [17]) mogą pracować do 300 kHz.

Należy zaznaczyć, że największa sprawność przetworników metalicznych wynosi około 75% dla najmniejszych częstotliwości (do 15 kHz). W zakresie większych częstotliwości sprawność spada (np. dla 175 kHz $\eta_{ea} = 16\%$, a dla 500 kHz $\eta_{ea} = 8\%$) [52]. W tym zakresie częstotliwości przetworniki pracujące na rdzeniach ferrytowych osiągają sprawność 90%. Zastosowanie ferrytów w przetwornikach magnetostrykcyjnych pozwala na pokrycie zakresu częstotliwości do kilku MHz.

Ferryty magnetostrykcyjne wyróżniają się dużą dobrocią mechaniczną Q_m . Przeciętnie dobroć ferrytów zawiera się w granicach od 2000 do 8000 [47]. W szczególnych przypadkach dobroć ta jest dużo większa, np.

Tablica 4

Własności ostatnio opracowanych ferrytów magnetostrykcyjnych [65], [66], [69], [70]

Wielkość	Symbol	Jednostka	Nazwa firmowa ferrytów					
			Kearfott		Curtis Wright	Ferroxcube		
			N-50	N-51		7A 1	7A 2	7B
Indukcja nasycenia	B_s	kGs	2,175	2,54	2,64	3,20 3,27	3,20 3,27	3,18
Przenikalność początkowa	μ_p	Gs/Oe	34	91		30 40	45 60	
Gęstość	γ	g/cm ³	5,1	5,1	5,2	5,2 5,3	5,2 5,3	5,1
Magnetostrykcja nasycenia	λ_s	10 ⁻⁶		-30		-26 -28	-26 -28	-27
Temperatura Curie	T_c	°C	>500	>500	>500	530	530	590
Oporność właściwa	ϱ	Ω cm		10 ^{2-10⁴}		≥100	≥100	>10 ⁴
Maksymalna wartość współczynnika sprzężenia magnetomech.	k_m	—	0,17* 0,27*	0,40	0,22 0,27	0,25 0,32	0,21 0,26	0,21 0,23
Moduł Younga przy k_m	E_B	$\frac{10^{12} \text{ dyn}}{\text{cm}^2}$	1,42* 1,65*	1,45* 1,60*	1,67	1,68 1,75	1,73 1,78	1,7
Czułość naprężeń przy k_m	A_m	$\frac{\mu \text{Gs dyn}}{\text{cm}^2}$				-2,8 -4,4	-1,6 -2,9	
Stała magnetostrykcji przy k_m	h_k	$\frac{10^4 \text{ dyn}}{\text{Gs cm}^2}$	0,64* -1,64*	0,75* -4,5*	-1,4	-2,2 -2,4	-2,3 -2,8	
Przenikalność odwracalna przy k_m	μ_k	Gs/Oe				15 25	8 15	
Pozostałość magnetyczna	B_r	kGs	1,1	1,75	1,47	1,1 1,6	1,5 1,7	1,6 1,9
Współczynnik k przy B_r	k_r	—		0,33		0,15 0,20	0,15 0,19	0,20 0,22
Czułość naprężeń przy B_r	A_r	$\frac{\mu \text{Gs dyn}}{\text{cm}^2}$				-2,3 -3,8	-2,2 -3,7	-2,7 -3,0
Stała magnetostrykcji przy B_r	h_r	$\frac{10^4 \text{ dyn}}{\text{Gs cm}^2}$				-0,9 -1,2	-0,9 -1,1	-1,5 -1,7
Przenikalność odwracalna przy B_r	μ_r	Gs/Oe	81	330	45	30 50	30 55	20
Natężenie powściągające	H_c	Oe	3,5	1,5		3,1 6,3	2,5 6,0	
Dobroć mechaniczna	Q_m	—	2800	60* 260*	800 2000	2500 5000	2500 5000	3000 4000

* Wielkości te przy innych niezmiennych były podawane w kilku wersjach przytaczanych przy okazji omawiania różnych zagadnień. W niektórych przypadkach prawdopodobnie popełniono omyłki drukarskie, gdyż np. we wszystkich wersjach prospektów $H_c = 35$ Oe dla N-50, a z wykresu wynika $H_c = 3,5$ Oe.

dla ferrytu Ni-Zn typu 4A dobroć dochodzi do 14 500 [20], a dla niklowo-miedziowego do 20 000 [40]. Dobroć materiałów metalicznych, takich jak nikiel lub permalój, zawiera się zazwyczaj w granicach od 25 do 200; natomiast dobroć stopu specjalnego w nazwie firmowej NiSpanC dochodzi do 16 000.

Ferryty wyróżniają się dużą przenikalnością odwracalną dochodzącą do 1200 Gs/Oe, podczas gdy magnetostrykcyjne materiały metaliczne mają przenikalność odwracalną dochodzącą do 400 Gs/Oe (18 Co-Ni [18]).

Temperatura Curie ferrytów jest tego samego rzędu co materiałów metalicznych z wyjątkiem ferrytu Ni-Zn o dużej zawartości cynku, którego T_c spada do 85°C.

Prędkość rozchodzenia się fal dźwiękowych zarówno w materiałach metalicznych, jak i w ferrytach jest tej samej wielkości i wynosi od 4000 do 6000 m/s.

Gęstość ferrytów γ , między innymi ze względu na porowatość, jest niższa niż metali i zwykle nie przekracza 5,3 g/cm³.

Należy zaznaczyć, że materiały metaliczne odznaczają się znacznie większą wytrzymałością mechaniczną. Ferryty ze względu na porowatość dochodzącą w materiałach handlowych do 0,25 są kruche i znoszą znacznie mniejsze moce. Natomiast rdzenie ferrytowe mają dużą odporność na korozję i mogą pracować w wyższych temperaturach (tabl. 2 — ferryty niklowe i litowe).

W tablicy 4 podano własności ostatnio opracowanych ferrytów magnetostrykcyjnych [65], [66], [69], [70]. Własności ich wskazują, że są to ferryty niklowe z niewielkimi domieszkami innych pierwiastków, takich jak miedź, cynk, kobalt, żelazo, mangan lub magnez. Ferryty te przeważnie odznaczają się małą porowatością ($p = 0,01 \div 0,04$) [65], [67] i bardzo dobrymi własnościami magnetostrykcyjnymi. Współczynnik sprzężenia k dochodzi do 0,40, dobroć mechaniczna $Q_m = 5000$, a indukcja nasycenia przewyższa 3kGs. Własności te dowodzą, że technika w tej dziedzinie posuwa się szybko naprzód i że wiele razy zajdzie potrzeba uzupełniania podanych w tablicach zestawień.

5. WYMAGANIA STAWIANE FERRYTOM STOSOWANYM W NADAJNIKACH ULTRADŹWIEKOWYCH

Wiele parametrów przetwornika ultradźwiękowego zależy od jego kształtu i wymiarów. W technice ultradźwiękowej stosuje się kilkanaście odmian vibratorów magnetostrykcyjnych [9], lecz można je sprowadzić do dwóch form podstawowych. Jedna grupa vibratorów to vibratory oparte na zjawisku Joule'a, a więc vibratory prętowe i okienne. Drugą grupę stanowią vibratory pierścieniowe oparte na zjawisku Joule'a lub Wiedemanna.

Od długości wibratora (lub średnicy) i masy drgającej zależy częstotliwość rezonansu mechanicznego przetwornika. Wywołane magnetostrycją odkształcenia mechaniczne zostają dla częstotliwości rezonansowej wielokrotnie wzmacnione (Q_m razy) i w ten sposób powstają silne drgania ultradźwiękowe. Aby ferryty były dobrymi materiałami na nadajniki magnetostrykcyjne, muszą mieć duży współczynnik magnetostrykcji λ , duży współczynnik sprzężenia magnetomechanicznego k i dużą dobroć (o ile praca ma odbywać się przy jednej częstotliwości). Ponadto ferryty stosowane w nadajnikach muszą się wyróżniać dużą wytrzymałością mechaniczną, gdyż od niej uzależnione jest maksymalne dopuszczalne natężenie promieniowania nadajnika (moc na jednostkę powierzchni promieniującej).

Częstotliwość rezonansową określa się z następującej zależności

$$f_r = \frac{1}{b} \sqrt{\frac{E_H}{\gamma}} \quad [\text{Hz}], \quad (15)$$

gdzie:

E_H — moduł Younga przy stałym natężeniu pola w dynach na cm^2 ,

γ — gęstość w g/cm^3 ,

b — współczynnik zależny od kształtu próbki; dla pręta $b_p = 2l$, gdzie l długość pręta w cm, a dla toroidu $b_t = \pi d_s$, gdzie d_s średnica toroidu w cm.

W przypadku drgań skrętnych do wzoru zamiast E wchodzi G — moduł sztywności.

Od nadajników ferrytowych wymagana jest duża sprawność elektroakustyczna η_{ea} będąca iloczynem sprawności elektromechanicznej η_{em} i mechanoakustycznej η_{ma} .

Zagadnienia związane ze sposobem wyznaczania sprawności (np. z kół reaktancji lub admitancji) były przedmiotem kilku publikacji [33], [56], [57]; na szczegółowe omówienie tych zagadnień brak jest tutaj miejsca. Wystarczy tylko powiedzieć, że sprawność elektroakustyczna wzrasta ze wzrostem kwadratu współczynnika sprzężenia magnetomechanicznego k , z dobrocią mechaniczną Q_m i elektryczną Q_e przetwornika.

Bradfield [67] podaje następujący przybliżony wzór na optymalną sprawność

$$\eta_{ea} \approx 1 - \frac{2}{\sqrt{k^2 Q_e \cdot Q_m}}. \quad (16)$$

Przyjmując $k = 0,2 \div 0,3$, $Q_e = 40 \div 300$, $Q_m = 20 \div 5000$ w praktyce sprawność zawierać się będzie w granicach od 65% do ponad 98%.

Wysoka sprawność nadajników jest pożądana nie tylko z punktu widzenia najmniejszych strat energetycznych. Ferryty należące do półprze-

wodników elektrycznych również nie są dobrymi przewodnikami ciepła i grzanie wibratora prowadzi do powstawania cieplnych naprężeń wewnątrz rdzenia, co obniża jego wytrzymałość mechaniczną.

Sprawność nadajnika maleje ze wzrostem częstotliwości na skutek działania prądów wirowych i szkodliwych sprzężeń. Ponadto sprawność η_{ea} nadajnika zależy od jego kształtu, sposobu przygotowania rdzenia (nadajniki na rdzeniach blaszkowych pewien procent energii tracą na tarcie mechaniczne) i od mocy promieniowanej.

Ferryt w czasie drgań poddawany jest naprężeniom rozciągającym i ściskającym (w przypadku drgań skrętnych — ścinającym). Dopuszczalna moc nadajnika ograniczona jest więc przez mechaniczną wytrzymałość materiału. Przy jednokierunkowych naprężeniach wytrzymałość ferrytów na ściskanie jest o rząd wielkości większa od wytrzymałości na rozciąganie [33] i dlatego wytrzymałość na rozciąganie jest wielkością określającą ich przydatność w zastosowaniach na nadajniki, gdyż współczynnik magnetostrykcji statycznej i dobroć pozwalają na pracę przy większych mocach.

Pęknięcie ferrytu zaczyna się w najsłabszych punktach próbki, wynikających z niejednorodności materiału. Ze wzrostem objętości materiału zwiększa się prawdopodobieństwo istnienia takich punktów. W wibratorach prętowych wykonanych z ferrytu o strukturze jednorodnej maksimum naprężeń przypada w okolicy środka wibratora i tam nastąpi pęknięcie.

Wytrzymałość mechaniczna ferrytu jest ściśle związana z jego porowatością. Van der Burgt [32] podaje następujący eksperymentalny wzór na maksymalne dopuszczalne natężenie promieniowania dla prętowego przetwornika wykonanego z ferrytu Ni-Zn drgającego wzdłużnie, gdzie powierzchnią promieniującą jest przekrój poprzeczny pręta:

$$I = 15 (1 - 2p) \quad [\text{W/cm}^2], \quad (17)$$

gdzie p — porowatość

$$p = \frac{\gamma_0 - \gamma}{\gamma_0}. \quad (18)$$

Przy przetwornikach typu jednookiennego powierzchnia promieniująca jest większa od przekroju poprzecznego magnetostryktora i we wzorze należy wprowadzić współczynnik $2/3$, a dla typu wielookiennego $2/3 \div 1/2$.

Dla przetworników prętowych $1/2 \gamma$ o dobroci mechanicznej $Q_m > 10$ natężenie promieniowania można określić [67] z zależności:

$$I = \frac{\pi}{4} \cdot \frac{\sigma_{\max}^2 10^{-9}}{Q_m \varrho v} = \frac{\pi}{4} \cdot \frac{\varepsilon_{\max}^2}{Q_m} \varrho \cdot v^3 \cdot 10^{-1} \quad [\text{W/cm}^2], \quad (19)$$

gdzie:

σ_{max} — maksymalne dopuszczalne naprężenia rozciągające w dynach/cm²,

Q_m — dobroć mechaniczna,

ρ — gęstość ferrytu w g/cm³,

v — prędkość dźwięku w m/s,

ε_{max} — amplituda maksymalnych odkształceń.

Van der Burgt [33] podaje, że ferryty wytrzymują naprężenia rozciągające dochodzące do 400 kG/cm², z czego wynikałoby $I \approx 150$ W/cm². Bradfield i Kikuchi stwierdzili, że ferryty pękają przy naprężeniach 110÷120 kG/cm² ($\approx 10^{10}$ dyn/cm²) [67]. Krytyczna wartość dla przetworników wykonanych na ferrytach 7A [65] oceniana jest na 90 do 110 kG/cm² dla większych rdzeni i 150÷250 kG/cm² dla mniejszych. Wytrzymałość mechaniczna ferrytów na rozciąganie jest od 3 do 15 razy mniejsza od materiałów metalicznych i dlatego w ferrytach stosowanych w nadajnikach wielkością ograniczającą jest ich wytrzymałość, natomiast w materiałach metalicznych o ich przydatności i o możliwości uzyskania maksymalnych natężeń decyduje ich współczynnik magnetostrykcji statycznej i dobroć mechaniczna.

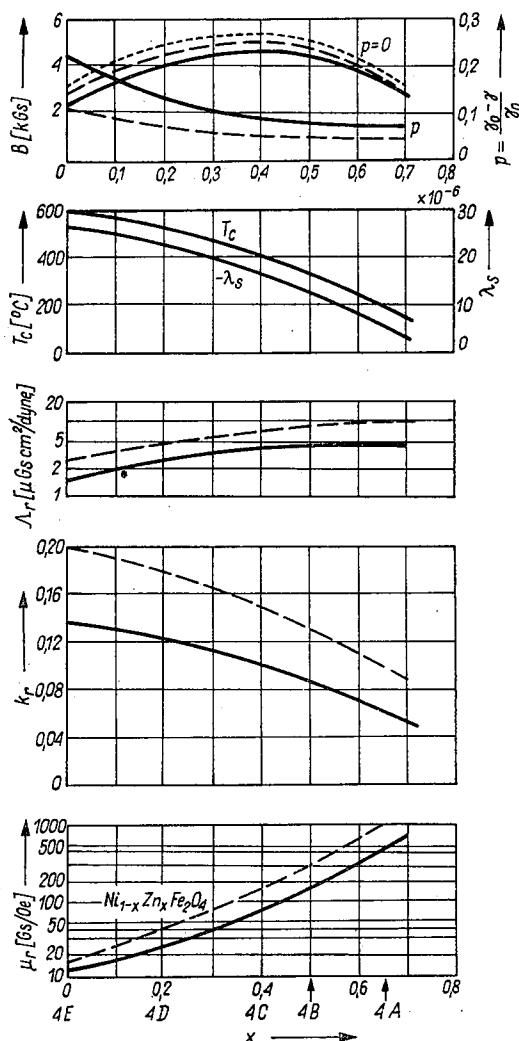
Z tych powodów dla ferrytów praktycznie przyjmuje się, że dla rezonansu stosunek amplitudy drgań do odpowiednich wymiarów magnetostryktora $Q_m \lambda$ nie powinien być większy od $1 \div 1,5 \cdot 10^{-4}$ [32], [67], czyli 0,1‰.

Zależnie od dobroci, która dla ferrytów w wodzie wynosi 15÷45, uwzględniając wytrzymałość mechaniczną ferrytów, można określić wymagania na współczynnik magnetostrykcji nasycenia. Przyjawszy, że praca odbywa się przy polaryzacji stałej i że dla punktu pracy współczynnik $\lambda = 1/3 \lambda_s$, wartość λ_s powinna wynosić około 10^{-5} ($10 \cdot 10^{-6}$).

Wiele spośród ferrytów dostępnych na rynku, a produkowanych nie do celów magnetostrykcyjnych spełnia ten warunek, mając niekiedy znacznie większą magnetostrykcję nasycenia, lecz ze względu na porowatość dochodzącą do 0,25 dopuszczalne natężenie promieniowania jak i współczynnik sprzężenia magnetomechanicznego są stosunkowo niskie. W materiałach produkowanych do celów magnetostrykcyjnych porowatość można zmniejszyć o połowę.

Na rys. 1 podano wg van der Burgt [32], [33], [47] odpowiednie wykresy przenikalności, czułości i współczynnika k dla remanencji oraz porowatość w funkcji zawartości cynku w ferrytach Ni-Zn stosowanych w nadajnikach magnetostrykcyjnych. Krzywe ciągłe dotyczą materiałów handlowych, a krzywe kreskowe dotyczą materiałów produkowanych dla celów magnetostrykcyjnych, w których zmniejszono porowatość. Pod strzałkami u dołu podano nazwy handlowe materiałów. Zmniejszenie porowatości po-

prawia o kilkadziesiąt procent wytrzymałość mechaniczną i wpływa na polepszenie innych własności magnetostrykcyjnych. Dopuszczalne natężenie promieniowania dla ferrytów magnetostrykcyjnych o małej porowatości przy 50 kHz wynosi 10 W/cm^2 [32], podczas gdy dla ferrytów z produkcji masowej (o dużej porowatości) dopuszczalne natężenie jest mniejsze od



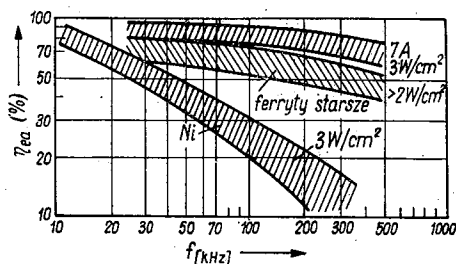
Rys. 1. Porowatość p , indukcja nasycenia B_s , temperatura Curie T_c , magnetostrykcja nasycenia λ_s , czułość Δr , współczynnik sprzężenia magnetomechanicznego k_r i przenikalność odwracalna μ_r dla remanencji materiałów handlowych — krzywe ciągłe — i specjalnych — krzywe przerywane — ferrytów $\text{Ni}_{1-x}\text{Zn}_x\text{Fe}_2\text{O}_4$ wg van der Burgha [33]

8 W/cm² [58]. Ferryty Ni-Zn odznaczają się dużą stabilnością. W czasie 5-godzinnej pracy nie zarejestrowano żadnych wahań mocy promieniowanej [59].

Przy pracy impulsowej, gdzie niebezpieczeństwo powstawania naprężeń termicznych jest b. małe, dopuszczalne natężenie dla promieniowania w wodzie dochodzi do 20 W/cm² [32]. Przy większych objętościach przetwornika prawdopodobieństwo pęknięcia jest większe i dlatego dopuszczalne natężenia należy dzielić przez 2÷3.

Jednym z najważniejszych zadań nadajnika jest przekształcenie jak największej energii elektrycznej w akustyczną. Dlatego mimo trudności konstrukcyjnych stosuje się podkład pola stałego (wstępne podmagnesowanie), aby punkt pracy nadajnika znajdował się w optymalnym położeniu, dla którego występuje maksimum współczynnika sprzężenia k . W nadajnikach opartych na drganiach wzdłużnych optimum wypada przy $0,5 B_s$, natomiast w nadajnikach o drganiach skręcających optymalna indukcja wynosi $0,7 B_s$ [20]. Często wykonuje się ferryty w postaci półtwardej o dużym B_r i praca odbywa się w tym punkcie. Jeżeli są trudności ze stosowaniem obwodu z prądem stałym, punkt pracy przesuwa się za pomocą zewnętrznego magnesu lub przez wstawienie w obwód magnetyczny płytki z ferrytu magnetycznie twardego (o grubości około 1 mm) [41], [60]. W przetwornikach typu okiennego płytę czołową wykonuje się zazwyczaj z ferrytu niemagnetostrykcyjnego o dobrych własnościach magnetycznych (duża B_s i μ). Słupy natomiast wykonane są z ferrytu magnetostrykcyjnego. Klejenie odbywa się zwykle za pomocą aralditu [33].

Rys. 2. Sprawność elektroakustyczna przetworników wykonanych z ferrytów „starszych” [32] i ferrytów typu 7A [65] w porównaniu z η_{ea} przetworników niklowych



Sprawność nadajników ferrytowych waha się w granicach 40÷90% [33], [41], [57], [58], [59]. Najlepsza uzyskana sprawność wynosi 93% przy pracy przy 74 kHz [60], a van der Burgt dopuszcza możliwość uzyskania sprawności większej niż 98%. Na rys. 2 podano wykres $\eta_{ea} = f(f)$ przetworników ferrytowych opracowanych przez van der Burgt.

6. WYMAGANIA STAWIANE FERRYTOM STOSOWANYM W ODBIORNIKACH ULTRADŹWIEKOWYCH

Magnetostrykcyjne przetworniki odbiorcze [32], [33], [60], [61] podobnie jak i nadawcze mogą być typu okiennego lub pierścieniowego. Jeżeli wy-

magana jest ostra kierunkowość przetwornika, używa się rdzeni tłokowych lub rurkowych. Odbiorniki pracujące przy jednej częstotliwości mają tak dobrane wymiary, że częstotliwość rezonansu mechanicznego odpowiada częstotliwości pracy. Przy pracy w szerokim zakresie częstotliwości częstotliwość rezonansu zarówno mechanicznego, jak i elektrycznego powinna znajdować się poza górną częstotliwością pracy.

Zasada pracy odbiorników oparta jest na zjawisku Villariego i pochodnych.

W odbiornikach ultradźwiękowych wielkością charakteryzującą ich czułość jest stosunek amplitudy napięcia wzbudzonego (otwartego obwodu) do amplitudy ciśnienia akustycznego.

$$\frac{\bar{U}}{\sigma} = \frac{2 \pi f n \bar{B} q \cdot 10^{-8}}{\sigma} \sim f n q \Delta. \quad (20)$$

Gdy indukcyjność magnetostryktora L jest stała i indukcyjność rozproszenia jest do zanedbania, czułość odbiornika jest proporcjonalna do współczynnika sprzężenia magnetomechanicznego k [32], który wskazuje, ile mocy mechanicznej zostaje zamienionej na magnetyczną

$$\frac{\bar{U}}{\sigma} \sim f q \frac{\Delta}{\sqrt{\mu}} \sim f q k. \quad (21)$$

W przetwornikach rurkowych występuje duża indukcyjność rozproszenia i wtedy słuszniej jest porównywać czułość Δ .

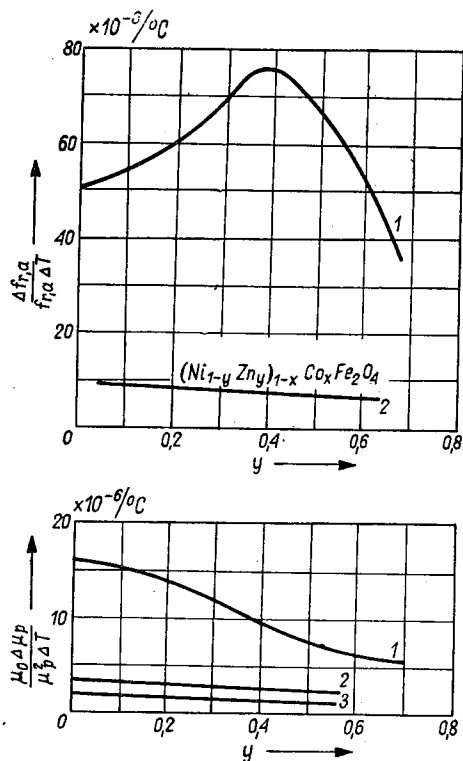
Magnetostryktory odbiorników można wykonywać z tych samych materiałów co przetworniki nadawcze. Najczęściej stosowane są ferryty Ni-Zn, przy czym największą czułość $\Delta_r = 4,5 \div 5 \mu\text{Gs cm}^2/\text{dynę}$ wykazują ferryty o zawartości 33÷32% moli ZnO, lecz ich współczynniki sprzężenia $k < 0,1$. Prace Golaminy [50] wykazały, że przy 15÷25% moli ZnO można osiągnąć czułość $\Delta = 3,3 \div 4,3 \mu\text{Gs cm}^2/\text{dynę}$ przy $k_r = 0,10 \div 0,11$ ($k_m = 0,14 \div 0,165$).

W urządzeniach hydroakustycznych przeważnie ten sam przetwornik spełnia rolę nadajnika i odbiornika i wtedy nacisk należy kłaść na to, by przetwornik był dobrym nadajnikiem, a więc aby miał dużą sprawność, natomiast czułość ze względu na szумы ośrodka wodnego (silniki okrętowe, inne urządzenia hydroakustyczne itp.) nie ma podstawowego znaczenia.

7. FERRYTY MAGNETOSTRYKCYJNE STOSOWANE W URZĄDZENIACH TELEKOMUNIKACYJNYCH

Ferryty magnetostrykcyjne znalazły zastosowanie w takich urządzeniach telekomunikacyjnych, jak stabilizatory częstotliwości [12], w filtrach magnetostrykcyjnych [13], [14], [15] i liniach opóźniających [E].

Wymagania stawiane materiałom stosowanym na filtry i stabilizatory są podobne, będą więc omawiane łącznie. Ferryty w urządzeniach tego typu mogą pracować w zakresie od około kilku kHz do około kilku MHz [14]. Zasada pracy filtru może być oparta na drganiach wzdłużnych lub skręcających. Jako materiału używa się ferrytów Ni-Zn (B1 dla drgań wzdłużnych i B4 dla drgań skręcających) lub ferrytów specjalnych niklowo-cynkowych z domieszką kobaltu [14], [47]. Optymalny punkt pracy dla k i stałej magnetostrykcji h wypada przy $B_0 = 0,6 B_s$ [14], co zwykle jest spełnione przy pracy przy pozostałości magnetycznej. Obok wymagań takich, jak duża dobroć i wysoki współczynnik sprzężenia magnetomechanicznego k , materiały stosowane w filtrach muszą wyróżniać się dużą stałością własności magnetycznych (μ), mechanicznych (f_r i f_a) oraz magneto-



Rys. 3. Zmiana współczynników temperaturowych częstotliwości rezonansowych i antyrezonansowych oraz względnej przenikalności początkowej w zależności od zmiany zawartości cynku w ferrytach niklowo-cynkowo-kobaltowych $(\text{Ni}_{1-y}\text{Zn}_y)_{1-x}\text{Co}_x\text{Fe}_2\text{O}_4$ wg van der Burgt [47]

Krzywe: 1 — ferryty handlowe, 2 — ferryty Ni-Zn-Co o minimalnym współczynniku temperaturowym f_r i f_a , 3 — ferryty Ni-Zn-Co o minimalnym współczynniku temperaturowym μ . Pomiary wykonano dla zakresu typowych temperatur pracy ($20 \div 50^\circ\text{C}$)

mechanicznych (k) w zależności od zmian temperatury. Osiąga się to przez dodanie do ferrytów magnetostrykcyjnych niewielkich domieszek kobaltu [47], [14], [15]. Zagadnieniem temperaturowym zajmowano się w wielu pracach [47], [62]. Na rys. 3 podano wg van der Burgta [47] zmianę współczynników temperaturowych w zależności od składu w ferrytach Ni-Zn i ferrytach Ni-Zn-Co. Dla ferrytów Ni-Zn produkowanych masowo, do których użycia używa się niklu zanieczyszczonego kobaltom ($\approx 0,5\%$), zmiana częstotliwości rezonansowej i antyrezonansowej w zakresie od 20 do 50°C wynosi $0,1\text{--}0,25\%$ (około $0,005\%$ na 1°C). Dla ferrytów Ni-Zn oczyszczonych z domieszek kobaltu zmiana ta wynosi około $0,3\%$ [33]. W ferrytach Ni-Zn-Co opracowanych do zastosowań w filtrach zmiana ta w podanym wyżej zakresie temperatury wynosi $0,03\%$ [33], co wystarcza do praktycznych zastosowań. Przenikalność przy tym składzie również ulega małym zmianom. Dobroć ferrytów jest wystarczająco duża, zawiera się zwykle pomiędzy (2000—8000), dochodząc nawet do 20 000 [40].

Ferryty Ni-Zn z dodatkiem kobaltu odznaczają się również lepszym współczynnikiem sprzężenia k . Ferryty magnetostrykcyjne stosowane w liniach opóźniających powinny charakteryzować się dużym współczynnikiem sprzężenia magnetomechanicznego, dużą dobrocią, dużą stałością parametrów magnetostrykcyjnych przy zmianach temperatury (przynajmniej w zakresie $20\text{--}50^\circ\text{C}$) oraz dużą jednorodnością struktury [E], [10], [11].

Streszczenie tego artykułu referowano na Konferencji Ferrytowej Polskiej Akademii Nauk w maju 1959 r., kładąc nacisk na omówienie własności ferrytów [68].

8. WNIOSKI

Ferryty magnetostrykcyjne mogą z powodzeniem konkurować ze wszystkimi typami materiałów metalicznych w dziedzinie zastosowań w przetwornikach ultradźwiękowych małej i średniej mocy. Należy dążyć do opracowania ferrytów o dużej wytrzymałości mechanicznej (małej porowatości) i dobrych własnościach magnetostrykcyjnych, które pozwolą zastosować ferryty w przetwornikach magnetostrykcyjnych dużej mocy.

Dzięki zastosowaniu ferrytów w przetwornikach magnetostrykcyjnych istnieje możliwość podwyższenia górnej częstotliwości pracy przetworników ze 100 kHz do kilku MHz, co pozwoli na zastąpienie w wielu zastosowaniach znacznie droższych przetworników piezoelektrycznych.

Po niezbyt udanych z powodu niskiej dobroci próbach zastosowań przetworników metalicznych w filtrach magnetostrykcyjnych otwierają się w tej dziedzinie szerokie możliwości przed ferrytami posiadającymi

dużą dobroć mechaniczną i elektryczną. Należy zwrócić uwagę na opracowanie ferrytów o własnościach magnetycznych i magnetostrykcyjnych niezależnych od zmian temperatury przynajmniej w zakresie temperatur od 20°C do 50°C. Ferryty tego typu będą również odpowiednie do pracy w stabilizatorach częstotliwości.

*Instytut Podstawowych Problemów Techniki PAN
Zakład Elektroniki*

WYKAZ LITERATURY

Czasopisma

1. Joule J. P.: On a New Class of Magnetic Forces. — *Am. Electr. Magn. Chem.* 8, 219-24, 1842.
2. Joule J. P.: On the Effects of Magnetism upon the Dimensions of Iron and Steel Bars. — *Phil. Mag. series III*, 30, 76—87, 225—241, 1867.
3. Villari E.: Change of Magnetization by Tension and by Electric Current. — *Ann. Phys.-Chem. Lpz.* 126, 87 — 122, 1865.
4. Kennelly A. E., Pierce G. W.: *Proc. Amer. Acad. of Arts and Science*, 48, 113, 1921.
5. Pierce G. W.: Magnetostriction Oscillators, An Application of Magnetostriction to the Control of Frequency of Audio and Radio Electric Oscillations, to the Production of Sound and to the Measurement of Elastic Constants of Metals. — *Proc. Amer. Acad. of Arts and Science*, 63, 1—47, 1928.
6. Black K. C.: A Dynamic Study of Magnetostriction. — *Proc. Amer. Acad. of Arts and Science*, 63, 49, 1928.
7. Pierce G. W.: Magnetostriction Oscillators. — *PIRE*, 17, 42, 1929.
8. Vincent L. H.: *Proc. Phys. Soc. London*, 41, 476, 1929.
9. Bradfield G.: Ultrasonic Electro-Acoustics. — *Acustica*, 4, 171—181, 1954.
10. Bradburd E.: Magnetostrictive Delay-Line. — *Electrical Communications*, 28, 46 (March 1951).
11. Scarrot G. G., Naylor R.: Wire Type Acoustic Delay-Lines for Digital Storage. — *Proc. Inst. Electr. Eng.*, 103 B, Supplement No 3, 1956.
12. Roberts V. V., Burns L. L.: Mechanical Filters for Radio Frequencies. — *RCA Review*, 10, 348, 1949.
13. Diethelm C. W.: Ferrite als magnetostruktive Resonatoren und deren Anwendung als Elemente elektrischer Filter. — *Technische Mitteilungen PTT*, 29, 8, 281, 1951.
14. Thiele A. P.: Narrow-Band Magnetostrictive Filters. — *Electronic and Radio Engineer*, November, 402—411, 1958.
15. Schweizerhof S.: Über Ferrite für magnetostruktive Schwinger in Filterkreisen. — *Nachrichten Technische Zeitschrift*, 4, 179—185, 1958.
16. Honda, K., Masumato M., Shirakawa Y., Kobayoshi T.: On the Magnetostriction of Iron-Aluminium Alloys and New Alloy-Alfer. — *Science Reports of Research Institutes Tohoku University*, 1A 341—7, Dec. 1949.

17. Davis C. M., Ferebee S. F.: Dynamic Magnetostrictive Properties of Alfenol. — The Journal of the Acoustical Society of America, 28, 2, 286—90, 1956.
18. Clark C. A.: The Dynamic Magnetostriction of Nickel-Cobalt Alloys. — British Journal of Applied Physics, 7, 10, 355—360, 1956.
19. Sussman H., Erlich S. L.: Evaluation on Magnetostrictive Properties of Hyperco. — Journal Acoust. Soc. Amer. 22, 4, 499, 1950.
20. Burgt, v. d., C. M.: Dynamical Physical Parameters of the Magnetostrictive Excitation of Extensional and Torsional Vibrations in Ferrites. — Philips Research Reports, 8, 91—133, 1953.
21. Crawford A. E.: Progress in High Power Ultrasonics, Part I and II. — British Communications and Electronics, 44—54 August (nr 8) 1955 i 56 September 1955.
22. Crawford A. E.: Application of Low Power Ultrasonics. — British Communications and Electronics, V. 3, nr 2, 64—69, February 1956.
23. Neppiras E. A., Foskett R. D.: Ultrasonic Machining, Part I i Part II. — Philips Technical Review, V. 18, nr 11, 323—34 i nr 12, 368—79, 1956/57.
24. Smoliński A.: Filtry mechaniczne. — Postępy Telekomunikacji 1959.
25. Epstein H., Stram O.: Rev. Scient. Instr. 24, 231—2, March 1953.
26. Smoliński A.: Zjawisko magnetostrykcji i materiały magnetostrykcyjne. — Rozprawy Elektrotechniczne, Tom V, z. 2, 211—37, 1959.
27. Kikuchi Y.: On the Magnetostrictivity. — Nippon Electrical Communication Engineering, 30, 91, Oct. 1942.
28. Kikuchi Y., Tsuya W., Shimizu H., i inni: Study on Ferrites for the Use in Magnetostriction Vibrator, Part I Ni-Zn Ferrite. — Sci. Rep. RITU, B (Electr.-Comm.) V. 7, nr 1, 1—7, June 1955.
29. Guillemin A.: Change of Elasticity of Soft Iron. — Comp. rend. 22, 264—265, 432—3, 1846.
30. Lee E. W.: Magnetostriction and Magnetomechanical Effect. — Reports on Progress in Physics, V. 18, 184—229, 1955.
31. Wiedemann G.: Magnetic Investigation. — Pogg. Ann. 117, 193—217, 1862.
32. Burgt, v. d., C. M.: Ferroxcube Material for Piezomagnetic Vibrators. — Philips Technical Review 18, 10, 285—98, 1956/7.
33. Burgt, v. d., C. M.: Performance of Ceramic Ferrite Resonators as Transducers and Filter Elements. — The Journal of the Acoustical Society of America V. 28, nr 6, 1020—32, Nov. 1956.
34. Akulow N. S.: Magnetostriktion in Eisenkristallen. — Z. Physik, 52, 389—405, 1928.
35. Heisenberg W.: Theorie der Magnetostriktion und der Magnetisierungs-kurve. — Z. Physik, 69, 287—97, 1931.
36. Wonsowski S. W.: On Quantum Theory of Ferromagnetic Single Crystals Magnetostriction. — Zurnal Fiziceski SSSR, 3, 181—90, 1940.
37. Kittel C.: Physical Theory of Ferromagnetic Domains. — Review Modern Physics, 21, 541—83, 1949.
38. Kwiek M.: Zjawiska magnetostrykcji i jego związki z termodynamiką oraz teorią sprężystości. — Sesja Naukowa Elektroniki Ciała Stałego PAN, czerwiec 1954.

39. Diakow G. P.: K teorii magnitostrikcji w izotropnych materiałach. — Wiestnik Moskowskiego Uniwersiteta, Nr 3, 75—8, 1957 i prace poprzednie (DAN SSSR 68, 33, 19, 49).
40. Kikuchi Y., Tsuya N., Shimizu H. i inni: Study on Ferrites for the Use in Magnetostriction Vibrator, Part II Ni-Cu ferrite, Sci. Rep. RITU, B (Electr. Comm.), V. 7, No 3, 171-8, Dec. 1955.
41. Kikuchi Y.: Performance of Magnetostrictive Transducers Made of Aluminium — Iron Alloy or Nickel-Copper Ferrite. — The Journal of the Acoustical Society of America, V. 29, No 5, 569—73, May 1957.
42. Stoner E. C.: Phil. Mag., 19, 565, 1935.
43. Stoner E. C.: Phil. Mag., 23, 833, 1937.
44. Standards on Piezoelectric Crystal. — Proc. Inst. Radio Engineers, 37, 1378—95, 1949.
45. Kausmann A., Rittberg G.: Magnetische Eigenschaften von Fe-Pt Legierungen. — Ann. Phys., Leipzig, 7, 173—181, 1950.
46. Bozorth R. M., Tilden E. F., Williams A. J.: Anisotropy and Magnetostriction of Some Ferrites. — Physical Rev., 99, 6, 1788—98, 1955.
47. Burgt, v.d., C.M.: Controlled Crystal Anisotropy and Controlled Temperature Dependence of the Permeability and Elasticity of Various Cobalt-Substituted Ferrites. — Philips Research Reports, V. 12, No 2, 97—122, 1957.
48. Bickford L. R., Pappis J., Stull J. L.: Magnetostriction and Permeability of Magnetite and Cobalt-Substituted Magnetite. — Physical Review 99, 1210-4, 1955.
49. Ferebee S. F., Davis C. M.: Effect of Divalent Ion Substitutions on the Magnetomechanical Properties of Nickel Ferrite. — The Journal of the Acoustical Society of America, V. 30, No 8, 1958.
50. Golamina I. P.: Zastosowanie ferrytów jako materiału na przetworniki elektroakustyczne. — Konferencja Naukowa na temat Przetworników Elektroakustycznych, Krynica, wrzesień, 1958 (materiały z konferencji w druku — w języku agnielskim).
- 50a. Golamina I. P.: Pewne zagadnienia technologii produkcji ferrytów do celów akustycznych, Konferencja Naukowa na temat Przetworników Elektroakustycznych, Krynica, wrzesień, 1958 (materiały z konferencji w druku — w języku angielskim).
51. Kaczkowski Z., Wadas R.: Własności magnetostrykcyjne ferrytów niklowo-kobaltowych. — Zeszyty Problemowe Nauki Polskiej, Nr 20, 1960.
52. Thiede H.: Funk und Ton, 4, 1, 32, 1951.
53. Nesbitt E. A.: Journal Applied Physics, 21, 879, 1950.
54. Uitert L. G.: Dc Resistivity in the Nickel and Nickel-Zinc Ferrite System. — The Journal of Chemical Physics, V. 23, 10, 1883-7, 1955.
55. Uitert L. G.: High Resistivity Nickel Ferrites — the Effect of Minor Additions of Magnese or Cobalt. — The Journal of Chemical Physics, V. 24, 2, 1956.
56. Nodtvedt H.: Some Remarks on the Analysis of the Magnetostrictive Transducers. — Acustica, V. 4, 432-8, 1954.
57. Thiede H.: Untersuchungen an Ferriten auf ihre Eignung als Flüssigkeitsschallwandler. — Acustica, V. 4, 532—536, 1954.
58. Enz V.: Erzeugung von Ultraschall mit Ferriten, Dissertation ETH Zürich, s. 45, 1955.

59. Golamina I. P.: Ultrazwukowej izluczatiel iz fierrita nikiela. — Akusticheskij žurnal, t. II, z. 2, s. 225, 1956.
60. Kikuchi Y., Shimizu H., Terajima M.: Magnetostrictive Ultrasonic Transducers Made of Ferrites. — Sci. Rep. RITU, B — V. 7, 1, 9—15, 1955.
61. Golamina I. P., Sokołow A. D., Czułkowska W. J.: Ispytaniye optimalnykh ultrazwukowych prijemnikow iz fierritow. — Akusticheskij žurnal, t. III, z. 3, 288, 1957.
62. Zimmermann A.: Untersuchung des thermischen Ausdehnungskoeffizienten β von Ferriten im Hinblick auf magnetostruktive Effekte. — Nachrichtentechnik, 8, 342-6, 8, 1958.
63. Rot H. D., McDonald J.: Ceramic Ultrasonic Magnetostrictive Transducer. — Journal of American Ceramic Society, V. 30, No 1, 1—5, 1957.
64. Henkel O.: Ein hochelastischer magnetostruktiver Werkstoff. — Nachrichtentechnik, Heft 8, 346—350, 1957.
65. Burgt, v. d., C. M.: Ferroxcube 7A1 and 7A2, New Ceramic Piezomagnetic Materials for Ultrasonic Power Transducers. — Matronics, no 15, 273—304 (sept.) 1958.
66. Kearfott: Prospekty firmowe: 4(15)58, 6(15)59, 8(1)59.
67. Bradfield G.: Practical Applications and Measurements on Ferrite Electromechanical Transducers. — The Proceedings of the Institution of Electrical Engineers. Part B, Supplement No 7, V. 104, 558—63, 1957.
68. Kaczkowski Z., Smoliński A.: Ferryty magnetostrykcyjne. — Zeszyty Problemowe Nauki Polskiej, nr 20, 1960 r.
69. Burgt v. d. C. M.: Ferroxcube 7A1 und 7A2, neue piezomagnetische keramische Werkstoffe für Ultraschallschwinger höherer Leistung. — Valvo Berichte BV, H1, 1—33, Ap. 1959.
70. Burgt v. d. C. M.: Special Piezomagnetic Ferrites and their Applications in Mechanical Band Pass Filters and High-Power Ultrasonics. — Electronic Technology (w druku), 1960.

Książki

- [A] Konnelly A. E.: Electrical Vibration Instruments. Macmillan Company, New York 1923.
- [B] Mason W. P.: Electromechanical Transducers and Wave Filters, v. Nostrand, New York 1946.
- [C] Hunt F. V.: Electroacoustics. Harvard Univer. Press, Cambridge 1954.
- [D] Hueter T. F., Bolt R. H.: Sonics. J. Wiley, New York 1955.
- [E] Mason W. P.: Physical Acoustics and the Properties of Solids. v. Nostrand, New York 1958.
- [F] Bozorth R. M.: Ferromagnetism, III ed. v. Nostrand, New York 1955.
- [G] Biełow K. P.: Uprugniye, tiepłowyje i elektricheskije jawlenja w ferromagnitnykh metalach. GITTL, Moskwa 1951.
- [H] Snoek J. L.: New Developments in Ferromagnetic Materials. Elsevier, Amsterdam 1949.
- [I] Smith C. G.: Magnetic Alloys and Ferrites. Magnetostrictive Materials, 184, M. C. Say, London 1954.
- [J] Popper P.: Soft Magnetic Materials for Telecommunications. The Magnetostriction of Ferrites, 322, Richards, Lynch, London 1953.
- [K] Bergman L.: Der Ultraschall. Hirzel Verlag, Zürich, 6th ed. 1954.

- [L] Horton J. W.: Fundamentals of Sonar. U. S. Naval Institute, Annapolis (Maryland) 1957.
- [E] Stewart K. H.: Ferromagnetic Domains. Cambridge University Press, 1954.
- [M] Reinboth H.: Technologie und Anwendung magnetischer Werkstoffe. VEB Verlag Technik, Berlin 1958.
- [N] Bates L. F.: Modern Magnetism, III ed. Cambridge University Press, 1951.

З. КАЧКОВСКИ, А. СМОЛИНСКИ

МАГНИТОСТРИКЦИОННЫЕ ФЕРРИТЫ И ИХ ПРИМЕНЕНИЕ

Резюме

Приведены основные определения величин касающихся магнитострикционных свойств материала. На основании заграничных данных произведено сравнение магнитных, магнитострикционных и физических свойств металлических материалов (таблица 1) и ферритов (таблица 2). В таблицах 3 и 4 приведены свойства ферритов производимых специально для магнитострикционных целей. В дальнейшей части работы обсуждены требования к ферритам применяемым в ультразвуковых преобразователях (излучателях и приёмниках), а также в радиоустройствах (фильтрах, стабилизаторах частоты и линиях задержки). В заключении приведена литература по магнитострикции ферритов и связанных с ней областях.

Z. KACZKOWSKI, A. SMOLINSKI

MAGNETOSTRICTIVE FERRITES AND THEIR APPLICATIONS

Summary

Principal definitions for the magnitudes determining magnetostrictive properties of the material are given. On the basis of data obtained from the foreign literature, there is a comparison carried out between magnetic magnetostrictive and physical properties of metallic materials (Table 1) and ferrites (Table 2). Tables 3 and 4 presents the properties of ferrites produced especially for magnetostrictive purposes. In the subsequent part of the paper, the authors discuss the requirements to be met by ferrites applied to ultrasonic transducers (emitters and receivers) and radiotechnical equipments (filters, frequency stabilizers and delay lines). In conclusion, there is an index of bibliographical items concerning ferrite magnetostriction and related fields.

621.314.7

TADEUSZ JANICKI

Zagadnienia termiczne w tranzystorach mocy

Rękopis dostarczono 24. 10. 1958

Artykuł poniższy omawia zależność temperatury przejścia $p-n$ kolektora i wynikającej stąd dopuszczalnej mocy strat od oporności cieplnej tranzystora. Pokazano wpływ konstrukcji i technologii tranzystora na wartość tej oporności oraz podano sposób jej pomiaru.

WSTĘP

Wzrost temperatury tranzystora wywołany wskutek traconej w nim energii elektrycznej powoduje znaczne zwiększenie się koncentracji nośników generowanych termicznie. Konsekwencją tego zjawiska jest wzrost prądu wstecznego kolektora I_{k_0} , niekontrolowanego przez sygnał wejściowy, oraz zmniejszenie się oporności i dopuszczalnego napięcia wstecznego kolektora.

W temperaturach powyżej 100°C wpływ generacji termicznej nośników jest tak silny, że german o oporności kilku Ωcm wykazuje już właściwości przewodnictwa samoistnego, powodującego zanik nieliniowych właściwości przejścia $p-n$.

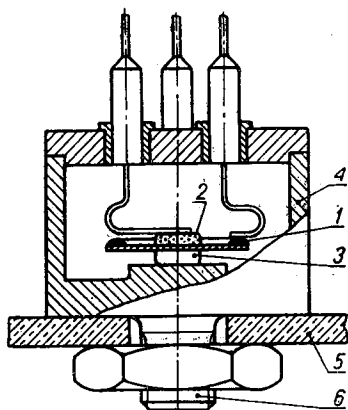
Wpływ temperatury przejścia $p-n$ na wielkość jego parametrów elektrycznych jest tym mniejszy, im mniejszą oporność właściwą ma german użyty do konstrukcji tranzystora. Z drugiej strony zbyt mała oporność właściwa nie pozwala na uzyskanie dostatecznie wysokich dopuszczalnych napięć wstecznych, dlatego zwykle nie stosuje się germanu o oporności mniejszej od $1\ \Omega\text{cm}$.

Wysoka temperatura tranzystora powoduje także wzrost naprężeń mechanicznych występujących między płytką germanową a materiałami elektrod, których współczynniki rozszerzalności liniowej znacznie różnią się od współczynnika rozszerzalności liniowej germanu (ind, cyna, ołów, stopy ind-gal, ołów-antymon).

Biorąc pod uwagę powyższe czynniki, zazwyczaj przyjmuje się, że maksymalna temperatura złącza w tranzystorach germanowych (temperatura pracy złącza) nie powinna przekraczać wartości $+85^{\circ}\text{C}$.

1. ODPROWADZENIE CIEPŁA W TRANZYSTORACH MOCY

Na rysunku 1 przedstawiono przekrój typowego tranzystora mocy, przykręconego do metalowej płyty, która ma za zadanie ułatwienie przekazywania ciepła wydzielanego w tranzystorze do otaczającego powietrza.



Rys. 1. Przekrój typowego tranzystora mocy.

1 — baza, 2 — emiter, 3 — kolektor, 4 — oprawka, 5 — płyta chłodząca, 6 — trzpień umożliwiający przykręcenie tranzystora do płyty

Płytę chłodzącą, do której jest przymocowany tranzystor, będziemy nazywali również radiatorem.

Energia cieplna wydzielana w złączu emitera, kolektora i w obszarze bazy jest doprowadzona do podstawki tranzystora za pośrednictwem miedzianego trzpienia, do którego jest przylutowany kolektor, a następnie do płyty, którą może być również metalowa podstawa lub obudowa urządzenia.

Dla prawidłowo wykonanego tranzystora mocy, o dużym współczynniku α i małej oporności bazy, można pominąć straty mocy w złączu emitera i w bazie wobec strat w kolektorze, dlatego też w dalszych rozważaniach przez moc strat tranzystora będziemy rozumieli moc traconą w złączu kolektora.

Odprowadzenie ciepła od kolektora do podstawki odbywa się dzięki zjawisku przewodzenia, natomiast wymiana ciepła między płytą chłodzącą a otaczającym ją powietrzem następuje wskutek konwekcji i promieniowania.

Prawo przepływu strumienia ciepłego przez przewodnik o długości l da się zapisać w stanie ustalonym jako:

$$\Delta T = P \cdot c \frac{l}{kA}, \quad (1)$$

gdzie

ΔT — różnica temperatur powstała na końcach przewodnika przy przepływie strumienia ciepłego wytworzonego przez moc P [W],

c — elektryczny równoważnik ciepła = 0,239 cal/W·sec,

k — przewodność cieplna,

A — przekrój poprzeczny przewodnika [cm²].

Oznaczając wyrażenie $c \frac{l}{kA}$ przez R otrzymamy:

$$\Delta T = PR \text{ [}^\circ\text{C]}. \quad (1a)$$

Współczynnik R reprezentuje tu oporność termiczną przewodnika i jest wyrażony w stopniach Celsjusza na wat.

W przypadku swobodnej konwekcji moc ciepłą, którą oddaje do otoczenia płyta metalowa o temperaturze T , większej od temperatury otoczenia, można przedstawić przy pomocy zależności:

$$P = \frac{\Delta T}{c \frac{1}{k_k A}}, \quad (2)$$

gdzie:

P — moc oddawana do otoczenia [W],

ΔT — różnica między temperaturą płyty a temperaturą otoczenia [°C],

k_k — współczynnik konwekcji,

A — powierzchnia płyty [cm²].

Wyrażenie $c \frac{1}{k_k A}$ nazwiemy opornością termiczną konwekcji i oznaczymy symbolem R_{konw} . Współczynnik ten podobnie jak współczynnik R z zależności (1a) będziemy wyrażali w stopniach Celsjusza na wat.

Drugim zjawiskiem, dzięki któremu płyta chłodząca oddaje ciepło do otoczenia, jest promieniowanie. W przypadku gdy płyta ta znajduje się w wolnej przestrzeni lub gdy jest ona dostatecznie oddalona od innych elementów urządzenia, promieniowana przez nią moc wyraża się wzorem:

$$P = \varepsilon \sigma A (T_1^4 - T_0^4) \text{ [W]}, \quad (3)$$

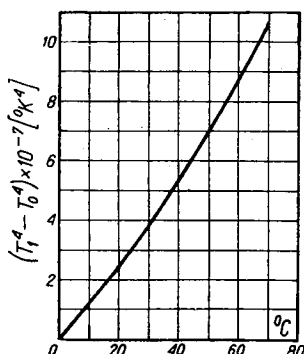
gdzie:

- ε — emisyjność całkowita płyty chłodzącej,
- σ — stała promieniowania wynosząca $5,67 \cdot 10^{-12} \text{ W/cm}^2 \cdot \text{K}^4$,
- A — powierzchnia płyty [cm^2],
- T_1 — temperatura płyty [$^\circ\text{K}$],
- T_0 — temperatura otoczenia [$^\circ\text{K}$].

Wyrażenie $\frac{1}{\varepsilon \sigma A}$ nazwiemy opornością cieplną promieniowania i oznaczymy symbolem R_{prom} . Zależność (3) przybierze wtedy postać:

$$P = \frac{(T_1^4 - T_0^4)}{R_{prom}} [\text{W}]. \quad (3a)$$

Udział promieniowania w ogólnym bilansie odprowadzenia ciepła od płyty chłodzącej jest znaczny i nie może być zaniedbywany. Należy zwrócić uwagę, że w zależności (3) występuje różnica czwartych potęg temperatur wyrażonych w stopniach Kelwina. Na rysunku 2 przedstawiono



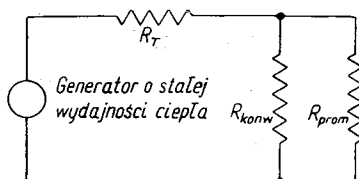
Rys. 2. Wartość $(T_1^4 - T_0^4)$ [$^\circ\text{K}^4$] w funkcji wzrostu temperatury płyty chłodzącej temperaturę otoczenia wynoszącą 25°C

przebieg wartości $(T_1^4 - T_0^4)$ w funkcji wzrostu temperatury płyty ponad temperaturę otoczenia wynoszącą 25°C . Dla różnicy temperatur między płytą chłodzącą i temperaturą otoczenia równej 40°C wyrażenie $(T_1^4 - T_0^4)$, jak to wynika z rysunku, osiąga wartość około $5 \cdot 10^{-7} \cdot ^\circ\text{K}^4$ (temperatura płyty chłodzącej wynosi w tym przypadku 65°C).

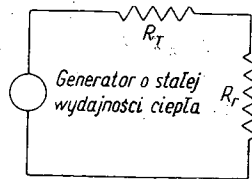
2. CZYNNIKI WPŁYWAJĄCE NA TEMPERATURĘ TRANZYSTORA

Przez analogię do obwodu elektrycznego możemy przedstawić obwód cieplny tranzystora w postaci generatora o stałej wydajności ciepła i kombinacji oporności termicznych tranzystora i płyty chłodzącej. Schemat

zastępczy takiego obwodu jest pokazany na rys. 3a. Oporności termiczne R_{konw} i R_{prom} jest bardzo trudno oddzielić od siebie, dlatego też zastąpimy je przez wypadkową oporność cieplną „płyta chłodząca-powietrze”, wynikającą z jednoczesnego odprowadzania ciepła przez zjawisko konwekcji



Rys. 3a. Schemat zastępczy obwodu cieplnego tranzystora wraz z płytą chłodzącą



Rys. 3b. Modyfikacja obwodu cieplnego z rys. 3a.

i promieniowania. Oporność tę nazwiemy opornością termiczną radiatora i oznaczymy ją symbolem R_r .

Schemat zastępczy obwodu cieplnego tranzystora przekształci się teraz do postaci pokazanej na rys. 3b. Posługując się nim rozważymy wpływ wielkości oporności termicznych i mocy strat na wartość temperatury tranzystora.

Ciepło wydzielane w złączu kolektora powoduje wzrost temperatury tranzystora w stosunku do temperatury otoczenia o wartość proporcjonalną do traconej w nim mocy. Współczynnikiem proporcjonalności jest oporność termiczna, jaką napotyka strumień cieplny na swej drodze od kolektora do powietrza otaczającego radiator, które w przypadku tranzystorów jest najczęściej spotykanym odbiornikiem ciepła (chłodzenia wodnego nie stosuje się).

Zgodnie ze schematem przedstawionym na rys. 3b całkowita oporność cieplna będzie się składała z dwóch części, mianowicie z oporności cieplnej między złączem kolektora i podstawą oprawki, oznaczonej na schemacie symbolem R_T oraz z oporności radiatora R_r . Różnica temperatur między złączem kolektora i oprawką tranzystora wywołana wskutek przepływu strumienia cieplnego wytworzonego przez moc strat w kolektorze o wielkości P_k watów, wyniesie zgodnie z (1a):

$$\Delta T_1 = P_k \cdot R_T \text{ [}^\circ\text{C]},$$

różnica zaś temperatur między płytą chłodzącą, do której przykręcony jest tranzystor, i temperaturą otoczenia wyniesie odpowiednio:

$$\Delta T_2 = P_k \cdot R_r.$$

Zaniedbując oporność cieplną styku: oprawka tranzystora-płyta chłodząca, możemy napisać wyrażenie na całkowitą różnicę temperatur między złączem kolektora a temperaturą otoczenia:

$$\Delta T = \Delta T_1 + \Delta T_2 = P_k (R_T + R_r). \quad (4)$$

Oznaczając temperaturę złącza kolektora przez T_k , temperaturę otoczenia zaś przez T_o otrzymamy z wyrażenia (4)

$$T_k = P_k(R_T + R_r) + T_o. \quad (1)$$

W przypadku gdy napięcie i prąd kolektora zmieniają się w czasie w sposób okresowy, przyjmuje się, że temperatura złącza kolektora w stanie ustalonym jest proporcjonalna do wartości średniej mocy strat w kolektorze, określonej jako:

$$P_{k\text{sr}} = \frac{1}{t_2 - t_1} \int_{t_1}^{t_2} u_k \cdot i_k dt,$$

gdzie u_k oraz i_k są wartościami chwilowymi napięcia i prądu kolektora.

Zależność (5) pozwala na określenie maksymalnej mocy strat, którą możemy wydzielić w tranzystorze przy zadanych warunkach chłodzenia (R_T , R_r i T_o) nie przekraczając dopuszczalnej temperatury złącza kolektora:

$$P_{k\text{max}} = \frac{T_{k\text{max}} - T_o}{R_T + R_r}, \quad (6)$$

gdzie $T_{k\text{max}}$ jest dopuszczalną temperaturą złącza kolektora wyrażoną w stopniach Celsjusza.

Przyjmując, że maksymalna temperatura kolektora w tranzystorach germanowych nie powinna przekraczać 85°C mamy:

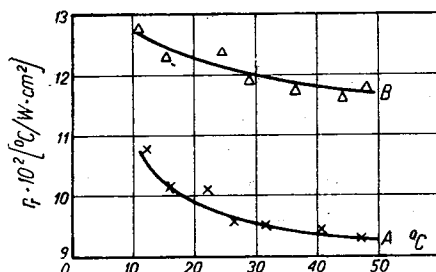
$$P_{k\text{max}} = \frac{85 - T_o}{R_T + R_r}. \quad (6a)$$

Jak wynika z zależności (6) lub (6a) maksymalna moc strat w kolektorze dla danej temperatury otoczenia T_o jest ograniczona wielkością oporności termicznej tranzystora i radiatora. Oporność R_T nie zależy od temperatury otoczenia i od mocy strat w kolektorze. Jest ona zależna jedynie od konstrukcji i technologii tranzystora i jest stałą fizyczną konkretnego egzemplarza (zmiany przewodności cieplnej materiału w zakresie temperatur pracy tranzystora są do pominięcia). Oporność ta dla tranzystorów dużej mocy, rzędu kilkudziesięciu watów, wynosi zwykle 0,3÷1,5°C/W, dla tranzystorów zaś o mocy kilkunastu watów zawiera się w granicach od 3 do 5°C/W. Oporność cieplna radiatora zależy w znacznym stopniu od różnicy temperatur pomiędzy temperaturą płyty chłodzącej i temperaturą otoczenia oraz od stanu powierzchni płyty chłodzącej. Dla różnicy temperatur $\Delta T = 20^\circ\text{C}$ oporność ta przeliczona na jednostkę powierzchni ma wartość około $1,2 \cdot 10^3 \text{ }^\circ\text{C/Wcm}^2$ w przypadku płyty miedzianej o wypolerowanej powierzchni i około $8 \cdot 10^2 \text{ }^\circ\text{C/Wcm}^2$ w przypadku płyty o powierzchni piaskowanej, a następnie czernionej.

Na rysunku 4 przedstawiono zależność oporności termicznej radiatora r_r , przeliczonej na jednostkę powierzchni, od różnicy temperatur między płytą chłodzącą i temperaturą otoczenia. Krzywa A pokazuje tę zależność dla płyty o powierzchni piaskowanej, zaś krzywa B — dla płyty o powierzchni polerowanej. Jak widać z zamieszczonych wykresów, opor-

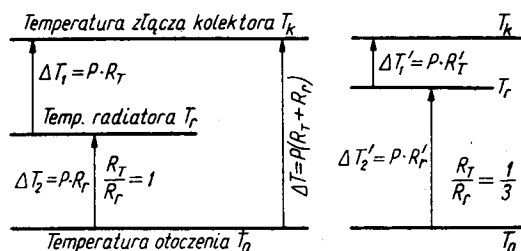
Rys. 4. Przebieg wartości oporności termicznej radiatora, przeliczonej na jednostkę powierzchni, w funkcji wzrostu temperatury płyty chłodzącej ponad temperaturę otoczenia

A — płyta miedziana piaskowana, B — polerowana; $T_0 = 22,5^\circ\text{C}$



ność radiatora jest tym mniejsza, im większa jest różnica temperatur między radiatorem i otaczającym go powietrzem.

Ponieważ oporność termiczna R_t wpływa na rozkład temperatur pomiędzy R_T i R_r (rys. 5), można mówić o pośrednim wpływie tej oporności



Rys. 5. Rozkład temperatur między złączem kolektora i powietrzem otaczającym

radiator dla różnych stosunków $\frac{R_T}{R_r}$ przy stałej wartości $(R_T + R_r)$, P , T_k i T_0

na wartość oporności termicznej radiatora. Tranzystor o oporności termicznej R'_T użyty zamiast tranzystora o oporności R_T , przy czym $R'_T < R_T$, pozwala na wzrost dopuszczalnej mocy strat w kolektorze, wynikający nie tylko ze zmniejszenia wartości R_T na R'_T , ale także z jednoczesnego zmniejszenia się wartości oporności termicznej radiatora.

Wpływ grubości płyty chłodzącej na wielkość oporności R_r może być znaczny w przypadku wykonania jej z cienkiej blachy o dużej oporności cieplnej, np. z blachy żelaznej. Różnica temperatur między punktem zamocowania tranzystora i brzegiem płyty może dochodzić w takim przypadku, zwłaszcza dla płyt o dużych rozmiarach, do kilkunastu stopni Celsjusza. Oczywiście oporność termiczna radiatora jest wtedy znacznie

większa w porównaniu z opornością płyty o takiej samej powierzchni, lecz wykonanej z grubszej blachy.

Dla płyt miedzianych i aluminiowych o wymiarach $150 \times 150 \times 1$ mm różnica temperatur między brzegami i środkiem płyty nie przekracza zazwyczaj $2-3^{\circ}\text{C}$ przy mocy strat w kolektorze rzędu 10 W, dlatego też zwiększenie grubości płyty w tym przypadku jest zbędne.

Należy jeszcze zwrócić uwagę na istotną zależność maksymalnej mocy strat w kolektorze od temperatury otoczenia. Jak widać z zależności (6a), wzrost temperatury otoczenia powoduje zmniejszenie dopuszczalnej mocy strat w kolektorze. Dla temperatury otoczenia $T_o = +85^{\circ}\text{C}$ dopuszczalna moc strat w kolektorze $P_{k\max}$ osiąga wartość równą zero, to znaczy, że przy tej temperaturze otoczenia tranzystor w ogóle nie może pracować, oczywiście w założeniu, że temperatura złącza kolektora nie powinna przekroczyć wartości $+85^{\circ}\text{C}$.

3. WPŁYW KONSTRUKCJI TRANZYSTORA NA WARTOŚĆ OPORNOŚCI CIEPLNEJ

Uzyskanie małej oporności termicznej między kolektorem i oprawką tranzystora wymaga, aby elementy, przez które przepływa strumień cieplny wytworzony w złączu kolektora, były możliwie krótkie, miały duży przekrój poprzeczny oraz by materiał użyty do ich konstrukcji miał jak największą przewodność cieplną. Zamieszczona poniżej tablica 1 podaje wartości współczynników przewodności cieplnej oraz współczynników rozszerzalności najważniejszych materiałów stosowanych w technologii tranzystorów mocy. Ze względu na dużą przewodność cieplną najbardziej przydatnymi materiałami do konstrukcji oprawek tranzystorów

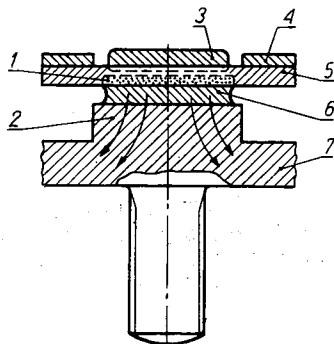
Tablica 1

Wartości współczynników przewodności cieplnej oraz współczynników rozszerzalności najważniejszych materiałów stosowanych w technologii tranzystorów mocy

Materiał	Przewodność cieplna $\text{cal}/^{\circ}\text{C cm sec}$	Współczynnik rozszerzalności $\alpha \cdot 10^6$	Materiał	Przewodność cieplna $\text{cal}/^{\circ}\text{C cm sec}$	Współczynnik rozszerzalności $\alpha \cdot 10^6$
Aluminium	0,48	23	Molibden	0,35	4,9
Cyna	0,16	24	Mosiądz	0,26	18,5
Cynk	0,29	14,1	Nikiel	0,14	13,1
German	0,125	6	Ołów	0,08	28,9
Ind	0,057	33	Srebro	1,0	18,7
Miedź	0,94	16,5	Stal niskowęglowa	0,14	11
Mika	0,0012	—	Złoto	0,7	14,3

byłyby srebro i miedź. Ponieważ jednak srebro jest metalem stosunkowo drogim, bywa ono stosowane jedynie w specjalnych rozwiązaniach konstrukcyjnych, usprawiedliwiających jego użycie, jak ma to miejsce w przypadku tranzystorów o mocy strat w kolektorze około 250 mW, przeznaczonych do pracy przy częstotliwościach rzędu kilkudziesięciu MHz.

W niniejszym rozdziale omówimy dwa zasadnicze sposoby odprowadzania ciepła od kolektora. Pierwszy z nich pokazany na rysunku 6 jest



Rys. 6. Odprowadzanie ciepła w typowym tranzystorze mocy. Linia przerywaną oznaczono złącza kolektora i emitera. Warstwa zrekrytalizowanego germanu znajdująca się między złączem kolektora i warstwą indy została oznaczona kropkami

1 — złącze kolektora, 2 — trzpień miedziany, 3 — emiter, 4 — pierścień doprowadzenia bazy, 5 — płytka germanowa, 6 — ind, 7 — podstawa oprawki tranzystora

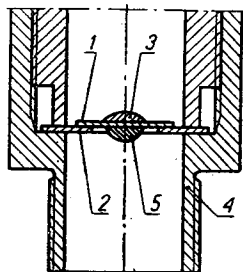
stosowany przede wszystkim w tranzystorach dużej i średniej mocy. Łatwo zauważyć, że całkowita oporność termiczna tranzystora germanowego wykonanego metodą stopową będzie sumą oporności cieplnych następujących elementów:

- a) warstwy zrekrytalizowanego germanu,
- b) warstwy indy,
- c) miedzianego trzpienia,
- d) podstawy oprawki tranzystora.

Ze względu na małą przewodność cieplną indy, wynoszącą zaledwie $0,057 \text{ cal/}^{\circ}\text{C cm}\cdot\text{sec}$, oporność termiczna, którą on wnosi, ma decydujący wpływ na wartość R_T . Zmniejszenie oporności cieplnej warstwy indy jest możliwe przez zmniejszenie jej grubości i zwiększenie przekroju poprzecznego. Średnice kolektora w produkowanych obecnie tranzystorach dużej mocy dochodzą do siedmiu i więcej milimetrów, co przy właściwej technologii tranzystora pozwala na zmniejszenie jego oporności cieplnej do wartości około $0,3^{\circ}\text{C/W}$.

Rysunek 7 przedstawia inny typ konstrukcji, stosowany jedynie w tranzystorach o mocy strat w kolektorze mniejszej od czterech watów.

Ciepło wydzielane w kolektorze jest odprowadzane do oprawki tranzystora przez płytkę germanową, a następnie przez metalowy pierścień, stanowiący jednocześnie elektryczne doprowadzenie bazy. Ponieważ grubość płytki germanowej zwykle nie przekracza $200\ \mu$, grubość zaś metalowego pierścienia — $500\ \mu$, wynikająca stąd oporność cieplna między złączem



Rys. 7. Przykład konstrukcji tranzystora stosowanej bardzo często w przypadku tranzystorów o mocy strat w kolektorze do 4 watów
1 — płytka germanowa, 2 — pierścień doprowadzenia bazy, 3 — kolektor, 4 — korpus oprawki, 5 — emiter

kolektora i oprawką zawiera się zazwyczaj w granicach od kilkunastu do kilkudziesięciu stopni Celsjusza na wat. Zaletą tej konstrukcji jest bardzo łatwa technologia montażu, co powoduje, że jest ona chętnie stosowana w przypadku tranzystorów o stosunkowo małej mocy strat w kolektorze.

Przykładem tranzystorów o tego rodzaju konstrukcji są tranzystory firmy Telefunken (OD-604), radzieckie (ПЗ-А, ПЗ-Б, ПЗ-В) oraz polskie (TM-1, TM-2, TM-3).

Często zachodzi potrzeba elektrycznego odizolowania oprawki tranzystora od płyty chłodzącej, na przykład, gdy płyta chłodząca i oprawka znajdują się na różnych potencjałach elektrycznych. Dokonuje się tego przy pomocy cienkich krążków mikowych, które wkłada się między oprawkę tranzystora i płytę chłodzącą, a dla lepszego styku cieplnego pokrywa się je dodatkowo smarem silikonowym. Oporność cieplna tranzystora powiększy się w takim przypadku o wartość oporności cieplnej warstwy miki, która zazwyczaj jest tego samego rzędu co oporność R_T .

4. WARUNEK STABILNOŚCI TERMICZNEJ TRANZYSTORA

Omawiane poprzednio zagadnienia dotyczyły wpływu temperatury złącza na wielkość takich parametrów elektrycznych, jak I_{k0} , U_{kmax} , r_k oraz zależności tej temperatury od mocy strat w kolektorze, konstrukcji tranzystora i płyty chłodzącej. W niniejszym rozdziale omówimy

po krótko wpływy wielkości temperatury złącza oraz oporności cieplnych R_T i R_r na stabilność termiczną tranzystora.

Niech temperatura tranzystora znajdującego się w stanie równowagi termicznej wynosi T [°C], napięcie zaś kolektora U_k [V]. Jeżeli z jakichkolwiek przyczyn temperatura złącza kolektora chwilowo zwiększy się o wartość dT , to odpowiadający temu przyrostowi temperatury przyrost prądu zerowego kolektora dI_{k_0} spowoduje zwiększenie się mocy strat w kolektorze o wartości dP , którą można wyrazić jako:

$$dP = U_k \cdot dI_{k_0}.$$

Jednocześnie wzrost temperatury złącza kolektora wywoła dodatkowy przepływ ciepła o wielkości:

$$dQ = \frac{dT}{R},$$

gdzie R jest opornością termiczną.

Tranzystor wróci do poprzedniego stanu równowagi, jeżeli zostanie spełniona zależność:

$$dP < dQ,$$

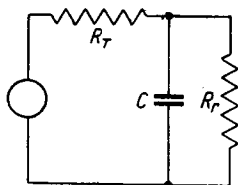
czyli:

$$U_k \cdot dI_{k_0} < \frac{dT}{R},$$

a stąd:

$$\frac{dI_{k_0}}{dT} < \frac{1}{U_k \cdot R}. \quad (7)$$

Jak wynika z powyższej zależności, warunek stabilności termicznej jest łatwiej spełniony przy niskim napięciu kolektora i małej oporności cieplnej R .



Rys. 8. Schemat zastępczy obwodu cieplnego tranzystora i płyty chłodzącej z uwzględnieniem pojemności cieplnej

Dla zadanych wartości U_k i R nierówność (7) określa krytyczną temperaturę złącza kolektora ze względu na warunek stabilności (wartość $\frac{dI_{k_0}}{dT}$ rośnie ze wzrostem temperatury).

Należy zaznaczyć, że oporność termiczna R występująca w wyrażeniu (7) jest właściwie rzecz biorąc impedancją, dlatego też w schemacie zastępczym obwodu cieplnego tranzystora musimy jeszcze uwzględnić pojemność cieplną oprawki tranzystora wraz z płytą chłodzącą. Pojemność ta jest oznaczona na rysunku 8 symbolem C .

W przypadku gdy czas trwania czynnika zakłócającego równowagę termiczną jest dużo mniejszy od iloczynu $R_T \cdot C$, na miejsce oporności R wstawiamy wartość oporności cieplnej tranzystora R_T .

Dla czynników zakłócających o czasie trwania większym od wartości $R_T \cdot C$ przyjmuje się, że oporność R jest równa sumie oporności cieplnej tranzystora i radiatora, a więc:

$$R = R_T \quad \text{dla} \quad t \ll R_T \cdot C$$

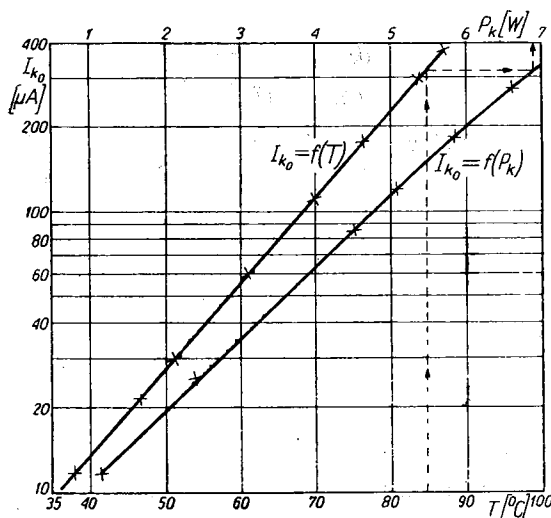
i

$$R = R_T + R_r \quad t \gg R_T \cdot C,$$

gdzie t jest czasem trwania czynnika zakłócającego.

5. POMIAR TEMPERATURY ZŁĄCZA KOLEKTORA I OPORNOŚCI TERMICZNYCH

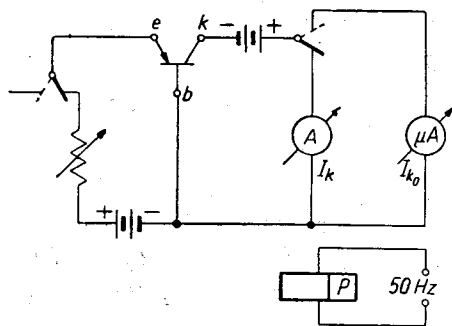
Prąd I_{k_0} , który płynie w kolektorze przy rozwartym obwodzie emitera, jest funkcją temperatury złącza kolektora. Własność tę wykorzystuje



Rys. 9. Przykład zależności I_{k_0} od temperatury i mocy traconej w kolektorze

się do pomiaru temperatury kolektora w funkcji traconej w nim mocy elektrycznej.

Porównując charakterystykę $I_{k_0} = f(P_k)$ z uprzednio zdjętą w termostacie charakterystyką I_{k_0} w funkcji temperatury (patrz rys. 9) możemy określić wielkość temperatury złącza odpowiadającą danej mocy strat i odwrotnie. Ze względu na wykładniczy charakter zależności I_{k_0} od temperatury, pomiar oparty na tej metodzie jest bardzo dokładny. Ponieważ



Rys. 10. Układ do pomiaru prądu I_{k_0} w funkcji mocy strat w kolektorze

prąd I_{k_0} jest sumą prądu termicznego złącza i prądu upływu, który zależy od przyłożonego napięcia, charakterystyka I_{k_0} w funkcji temperatury oraz charakterystyka I_{k_0} w funkcji mocy strat muszą być zdejmowane przy takim samym napięciu kolektora. Zależność prądu I_{k_0} od mocy strat w kolektorze bada się w ten sposób, że po przyłożeniu do kolektora znanej mocy prądu stałego i doprowadzeniu tranzystora do stanu równowagi termicznej przerywa się szybko dopływ mocy i mierzy wartość I_{k_0} .

Schemat układu umożliwiającego wykonanie takiego pomiaru podano na rysunku 10. Obwód zasilania tranzystora jest przerywany za pośrednictwem styków przekaźnika spolaryzowanego z częstotliwością 50 Hz. Przez pół okresu trwa zasilanie tranzystora, w drugiej zaś połowie okresu, po przerzuceniu kotwicy przekaźnika, odbywa się pomiar.

Przy założeniu prostokątnego kształtu impulsów i zaniedbaniu czasu przeskoiku kotwicy przekaźnika amplituda I_{k_0} będzie równa podwójnej wartości prądu wskazywanego przez mikroamperomierz.

W celu uniknięcia ewentualnych błędów mogących wynikać z niespełnienia tych założeń wygodniej jest porównywać ze sobą charakterystykę $I_{k_0} = f(P_k)$ z charakterystyką $I_{k_0} = f(T)$, zdjętą również w tym układzie. Jest rzeczą oczywistą, że w przypadku zdejmowania charakterystyki $I_{k_0} = f(T)$ obwód zasilania emitera musi być rozarty, aby tranzystor nie był dodatkowo podgrzewany.

Pomiar temperatury złącza kolektora wykonany przy mocy strat w kolektorze o wielkości P_k [W] oraz pomiar temperatury podstawy oprawki tranzystora pozwalają na wyznaczenie wartości oporności termicznej R_T .

z następującej zależności:

$$R_T = \frac{T_k - T_{opr}}{P_k} \quad [^{\circ}\text{C/W}], \quad (8)$$

gdzie T_{opr} — temperatura oprawki [$^{\circ}\text{C}$].

Wielkość R_T można również wyznaczyć w przypadku, gdy znana jest temperatura złącza kolektora dla dwóch określonych wartości mocy. Oporność termiczna tranzystora wyrazi się wtedy jako:

$$R_T = \frac{T_2 - T_1}{P_2 - P_1} \quad [^{\circ}\text{C/W}], \quad (9)$$

gdzie T_1 i T_2 są temperaturami złącza kolektora odpowiadającymi mocom strat P_1 i P_2 wyrażonym w watach.

Wzór (9) jest słuszny przy założeniu, że oporność radiatora R_r jest stała w przedziale mocy $[P_1, P_2]$, co jest w przybliżeniu słuszne dla niezbyt dużych różnic $P_2 - P_1$.

Korzystając z zależności (5) możemy określić sumę oporności cieplnej tranzystora i radiatora znając temperaturę złącza kolektora wywołaną przez moc P_k [W]:

$$R_T + R_r = \frac{T_k - T_o}{P_k} \quad [^{\circ}\text{C/W}]. \quad (10)$$

Pomiar temperatury radiatora przy znanej mocy strat w kolektorze umożliwia wyznaczenie jego oporności cieplnej ze wzoru:

$$R_r = \frac{T_r - T_o}{P_k} \quad [^{\circ}\text{C/W}], \quad (11)$$

gdzie: T_r i T_o są odpowiednio temperaturą radiatora i temperaturą otoczenia wyrażonymi w stopniach Celsjusza.

Jak wykazaliśmy uprzednio, oporność R_r jest funkcją różnicy tych temperatur, a więc przy stałej temperaturze otoczenia R_r jest także funkcją mocy strat w kolektorze.

WYKAZ LITERATURY

1. Falter: Probleme der Fertigung von Leistungstransistoren. — Radio und Fernsehen, 1956, Nr 16.
2. Fletcher N. H.: Some Aspects of the Design of Power Transistors. — PIRE, vol. 43, May 1955.
3. Gondry R.: Utilisation rationelle des transistors à différentes températures. — Electronique, mai—juin 1958.
4. Maloff I. G.: Heat Transfer in Power Transistors. — Electronic Industries, December 1957.

5. Mooers: Rozrobotka moszcznych połuprowodnikowych triodow. — Woprosy radiołokacyjnoy tiechniki, nr 6, 1955.
6. Połuprowodnikowyje pribory i ich primienienije. — Sowietiskoje Radio 1956.
7. Rosiński W.: Tranzystory i ich układy. PWT, Warszawa 1957 r.

Т. ЯНИЦКИ

ТЕРМИЧЕСКИЕ ВОПРОСЫ В МОЩНЫХ КРИСТАЛЛИЧЕСКИХ ТРИОДАХ

Резюме

В статье описано влияние термических сопротивлений выступающих между коллекторным $p-n$ переходом и окружающим радиатор воздухом на величину допускаемой рассеиваемой мощности в мощных германиевых триодах типа $p-n-p$. Приведена зависимость этих сопротивлений от конструкции транзистора, величины радиатора и обработки его поверхности.

В заключении статьи описан способ измерения температуры коллекторного $p-n$ перехода, термического сопротивления между коллекторным $p-n$ переходом и корпусом триода, а также термического сопротивления между радиатором и окружающим его воздухом.

T. JANICKI

THE THERMAL PROBLEM IN POWER TRANSISTORS

Summary

The paper discusses the influence of thermal resistances occurring between the collector junction and air surrounding the radiator upon the magnitude of the admissible power dissipation in germanium power transistors of the type $p-n-p$. The dependence of these resistances on construction of the transistor, dimensions of the radiator and its surface treatment, is presented.

In conclusion, the author discusses the method of measuring temperature of the collector junction, thermal resistance between the collector junction and the transistor container, and thermal resistance between the radiator and the surrounding air.

T. JANICKI

WÄRME-PROBLEME BEI LEISTUNGSTRANSISTOREN

Zusammenfassung

In diesem Aufsatz ist der Einfluss der zwischen dem Kollektor- pn -Übergang und der den Radiator umhüllenden Luft auftretenden Wärmewiderstände auf die Grösse der zulässigen Verlustleistung in Germanium-Leistungstransistoren von der ($p-n-p$)-Type besprochen worden.

Die Abhängigkeit dieser Widerstände von der Konstruktion des Transistors, von den Abmessungen des Radiators und von der Bearbeitung seiner Oberfläche ist angegeben worden.

Zur Beendigung des Aufsatzes ist ein Messverfahren der Temperatur des Kollektor- pn -Überganges, des Wärmewiderstandes zwischen dem Kollektor- pn -Übergang und dem Gehäuse des Transistors und des Wärmewiderstandes zwischen dem Radiator und der umhüllenden Luft besprochen worden.

WARUNKI PRENUMERATY CZASOPISMA

„Rozprawy Elektrotechniczne” — Kwartalnik

Cena w prenumeracie zł 100.— rocznie

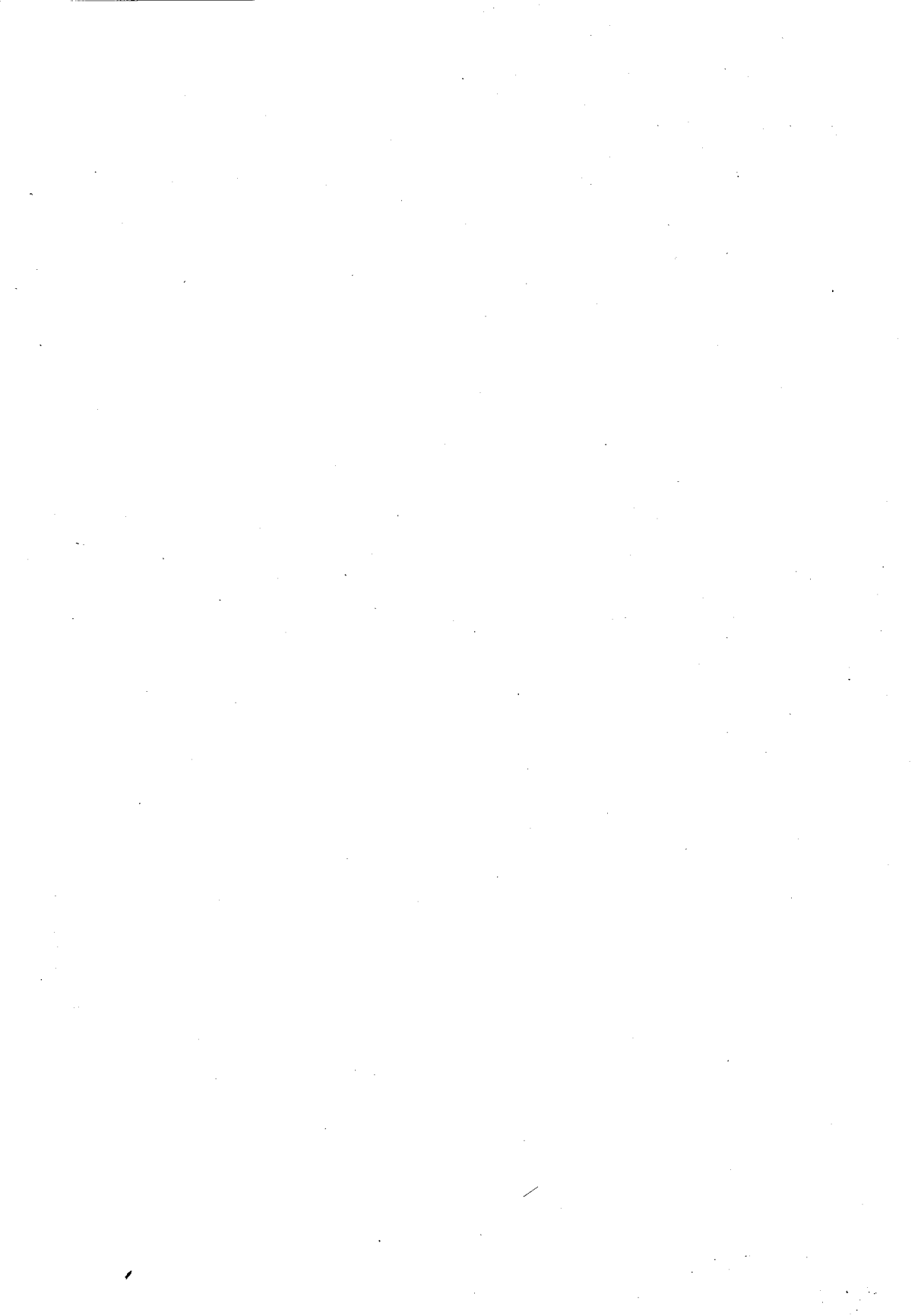
50.— półrocznie

Zamówienia i wpłaty przyjmują:

1. Centrala Kolportażu Prasy i Wydawnictw „Ruch”, Warszawa, ul. Srebrna 12, konto PKO Nr 1-6-100-020.
2. Urzędy pocztowe.

Prenumerata ze zleceniem wysyłki za granicę 40% drożej. Zamówienia dla zagranicy przyjmuje Przedsiębiorstwo Kolportażu Wydawnictw Zagranicznych „Ruch”, Warszawa, ul. Wilcza 46, konto PKO Nr 1-6-100-024.

Bieżące numery do nabycia w księgarniach naukowych „Dom. Książki” oraz w Ośrodku Rozpowszechniania Wydawnictw Naukowych Polskiej Akademii Nauk — Wzorcownia Wydawnictw Naukowych PAN — Ossolineum — PWN; Warszawa Pałac Kultury i Nauki (wysoki parter).



POLSKA AKADEMIA NAUK
INSTYTUT PODSTAWOWYCH PROBLEMÓW TECHNIKI

ROZPRAWY ELEKTROTECHNICZNE

TOM VI • ZESZYT 3

K W A R T A L N I K

PAŃSTWOWE WYDAWNICTWO NAUKOWE
WARSZAWA 1960

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 600 (479 + 121)	Oddano do składania 20.V.1960
Ark. wyd. 15,25, ark. druk. 15,25	Podpisano do druku 18.XI.1960
Pap. druk. sat. kl. III, 80 g, 70×100	Druk ukończono w grudniu 1960
Zamówienie nr 941/60	Cena zł 25.— C-23

Drukarnia im. Rewolucji Październikowej, Warszawa

621.3.01/02

JERZY OSIOWSKI

Układy elektryczne liniowe pobudzane okresowo

Rękopis dostarczono 29.5.1959

W pracy zostały omówione w sposób systematyczny zagadnienia analizy przypadku ogólnego sieci liniowej RLCM pobudzonej dowolnymi przebiegami okresowymi. W oparciu o rachunek macierzowy i operatorowy podano warunki istnienia i jednoznaczności rozwiązań okresowych i ustalonych oraz wyprowadzono dla tych rozwiązań odpowiednie wzory. Otrzymane wyniki zilustrowano przykładami. Na zakończenie rozpatrzono przypadek szczególnie sieci liniowej traktowanej jako czwórnik.

1. WSTĘP

We współczesnej elektrotechnice, a w szczególności w radiotechnice, technice impulsowej, telewizyjnej, radiolokacji, autoatmyce itp. występują bardzo liczne zagadnienia, których rozwiązanie sprowadza się do wyznaczenia przebiegów elektrycznych w układzie pobudzonym ze źródła prądowego lub napięciowego przebiegiem okresowym o określonym, aczkolwiek w zasadzie dowolnym, kształcie. Wymagania, jakie przy tym stawia praktyka odnośnie charakteru otrzymywanego rozwiązania, są oczywiście różne w rozmaitych konkretnych przypadkach, ogólnie jednak biorąc są one na ogół wysokie; chodzi zwykle o wyznaczenie rozwiązania dla stanu ustalonego i nieustalonego, możliwie o dużej dokładności, pozwalającego na badanie i dyskusję kształtu otrzymanych krzywych, rozmaitych współczynników zniekształceń itp.

Metody, jakie można przy tym stosować, mogą być oczywiście różne: numeryczne, graficzne, grafoanalityczne, analityczne, oparte na wykorzystaniu maszyn matematycznych itp. W dalszym ciągu pracy będzie mowa tylko o metodach analitycznych.

Trzeba od razu podkreślić oczywisty fakt, iż cała teoria prądów sinusoidalnie zmiennych jest przypadkiem szczególnym teorii układów pobudzanych okresowo, przypadkiem polegającym na badaniu w stanie

ustalonym układów pobudzanych przebiegami sinusoidalnymi, tj. przebiegami okresowymi typu:

$$f(t) = A \sin(\omega t + \varphi) \quad (1)$$

lub kombinacjami liniowymi takich przebiegów o postaci:

$$f(t) = \sum_{v=1}^n A_v \sin(v \omega t + \varphi_v). \quad (2)$$

Metody, jakie się wówczas stosuje przy wyznaczaniu rozwiązań, są powszechnie znane i stosowane (rachunek symboliczny dla stanów ustalonych, metody klasyczne lub operatorowe dla stanów nieustalonych) i wobec tego nie będą tutaj omawiane.

Z punktu widzenia zarówno teoretycznego, jak i praktycznego, inny problem wyłania się wówczas, gdy mamy do czynienia z przypadkiem ogólnym układu liniowego pobudzanego dowolnym przebiegiem okresowym. W tym przypadku, od dawna znaną i stosowaną, niejako klasyczną metodą rozwiązywania jest metoda szeregów Fouriera (metoda harmoniczných), pozwalająca na wyznaczenie rozwiązania okresowego dla stanu ustalonego w postaci szeregu trygonometrycznego. Obok bardzo dużych i powszechnie znanych zalet metoda ta w niektórych przypadkach jest uciążliwa lub nawet w ogóle nie nadaje się praktycznie do stosowania ze względu na to, iż rozwiązanie ma formę szeregu (lub w przypadku uproszczonym — wielomianu trygonometrycznego). Występuje to szczególnie w tych zagadnieniach, gdzie ważna jest analiza czasowa, a nie częstotliwościowa przebiegów, które w znacznym stopniu różnią się od przebiegu sinusoidalnego. Z tego względu, a także w związku z szybkim rozwojem i doskonaleniem metod operatorowych, powstał w okresie ostatnich kilkunastu lat szereg metod rozwiązywania układów pobudzanych dowolnym przebiegiem okresowym, pozwalających na wyznaczenie rozwiązania dla stanu ustalonego w postaci zamkniętej (tj. nie w postaci szeregu), a więc eliminujących podstawową niedogodność metody harmoniczných.

Zagadnienia układów liniowych pobudzanych okresowo mają już swoją historię. Pierwsze, wstępne prace z tego zakresu (Carnahan [1], Marchand [18], Hamlin [9] i in., a także Chin [3]) pojawiły się w okresie poprzedzającym ostatnią wojnę oraz w początkowych jej latach. Zagadnienie efektywnego wyznaczania przebiegów okresowych (ustalonych) w konkretnych układach elektrycznych zostało rozwiązane — i to kilkoma sposobami — w cyklu prac Waidelicha [25], [26], [28]. Niektóre z koncepcji Waidelicha zostały rozwinięte w książkach: Kontorowicza [11] oraz Carlsława, Jaegera [2]. Łurie [14] (a także [15]) sformułował problem ogólniej

wychodząc z układu równań różniczkowych rzędu pierwszego w postaci normalnej. Podał on także inną metodę wyznaczania rozwiązań okresowych oraz sformułował warunki istnienia i jednoznaczności takich rozwiązań. Warunki te były w ogólniejszej postaci sformułowane później przez Małkina [16] (a także [17]), który podał także warunki istnienia rozwiązań quasiokresowych. W późniejszych latach powstało szereg dalszych prac (Zaľmanzon [31], Gonorowski [8], Osiowski [20], Węgrzyn [29], Kożeśnik [12], Woronow [30] i in.) w rozmaitych ujęciach traktujących omawiany problem. Niedawno Kaczorek [10] podał uogólnienie znanego wzoru dla rozwiązania okresowego na przypadek biegunów wielokrotnych funkcji przenoszenia [por. wzór (100)].

Jak wynika choćby z tego niekompletnego wyliczenia, literatura dotycząca omawianego zagadnienia jest już bogata, poszczególne prace różnią się jednak dość znacznie, jeśli chodzi o sformułowanie samego zagadnienia (ujęcie elektryczne, mechaniczne, czysto matematyczne), zakres stosowalności wyników, przyjmowane założenia, metody rozumowania itp. W wielu pracach można z łatwością odnaleźć różnego rodzaju nieścisłości i niedomówienia.

W tej sytuacji autor niniejszej pracy postawił sobie za cel przedstawić możliwie ogólnie i ściśle zarys zagadnień elektrycznych układów liniowych pobudzanych okresowo oraz współczesnych metod ich rozwiązywania. Niezależnie od samego ujęcia niektóre partie tekstu mają charakter oryginalny i nie były dotąd publikowane.

2. OGÓLNA CHARAKTERYSTYKA ZAGADNIENIA

Rozpatrywać będziemy układ równań różniczkowo-całkowych liniowych o współczynnikach stałych o postaci:

$$\sum_{m=1}^r \left[\alpha_{lm} \frac{dy^{(m)}}{dt} + \beta_{lm} y^{(m)} + \gamma_{lm} \int_{-\infty}^t y^{(m)} dt \right] = f_l(t) \quad (3)$$

$$l = 1, 2, \dots, r$$

z warunkami początkowymi:

$$y^{(m)}|_{t=0} = A_m, \quad (4)$$

przy czym: $y^{(m)} \equiv y^{(m)}(t)$.

Wprowadzając oznaczenia:

$$\int_{-\infty}^0 y^{(m)} dt = B_m \quad (5)$$

można układ równań (3) przepisać w postaci:

$$\sum_{m=1}^r \left[a_{lm} \frac{dy^{(m)}}{dt} + \beta_{lm} y^{(m)} + \gamma_{lm} \left(\int_0^t y^{(m)} dt + B_m \right) \right] = f_l(t) \quad (6)$$

$$l = 1, 2, \dots, r,$$

w której stałe B_m mają charakter wartości początkowych całek funkcji niewiadomych $y^{(m)}$.

Układ równań (6) można zapisać w postaci jednego równania macierzowego. Przyjmując symboliczne oznaczenia:

$$D \equiv \frac{d}{dt}; \quad \frac{1}{D} \equiv \int_0^t \dots dt$$

oraz oznaczając:

$$Z_{lm}(D) = a_{lm} D + \beta_{lm} + \gamma_{lm} \frac{1}{D} \quad (7)$$

dostajemy ($i = 1, 2, \dots, r$):

$$\sum_{m=1}^r Z_{lm}(D) y^{(m)} + \sum_{m=1}^r \gamma_{lm} B_m = f_l(t). \quad (8)$$

Wprowadzając następujące macierze kwadratowe (wymiaru $r \times r$):

$$\mathfrak{A} = \|a_{lm}\|; \quad \mathfrak{B} = \|\beta_{lm}\|; \quad \mathfrak{G} = \|\gamma_{lm}\| \quad (9)$$

oraz:

$$\mathbf{Z}(D) = \|Z_{lm}(D)\| = \mathfrak{A} D + \mathfrak{B} + \mathfrak{G} \frac{1}{D}, \quad (10)$$

a także macierze-wektory (wymiaru $r \times 1$):

$$\mathbf{y}(t) = \|y^{(m)}\|; \quad \mathbf{f}(t) = \|f_l(t)\|; \quad \mathbf{A} = \|A_m\|; \quad \mathbf{B} = \|B_m\|, \quad (11)$$

możemy układ równań (6) zapisać w postaci równoważnego mu równania macierzowego:

$$\mathbf{Z}(D) \mathbf{y}(t) + \mathfrak{G} \mathbf{B} = \mathbf{f}(t) \quad (12)$$

z warunkiem początkowym:

$$\mathbf{y}(0) = \mathbf{A}. \quad (13)$$

Równanie (12), jak każde równanie różniczkowe, można rozpatrywać przy różnych warunkach początkowych (13), tj. przy różnych macierzach $\mathbf{A}$, przy czym, oczywiście zmiana macierzy $\mathbf{A}$ nie zmienia samego równania (12), podczas gdy

zmiana macierzy **B** formalnie zmienia równanie. Mimo tego, uwzględniając, że macierze **A** i **B** mają ten sam charakter i interpretację fizyczną (wartości początkowe funkcji niewiadomych i ich całek), będziemy w dalszym ciągu mówić o tym samym równaniu (12) nawet przy różnych macierzach **B** traktując zmianę macierzy **B** jako zmianę warunków początkowych, a nie zmianę samego równania.

Nadmieniamy jeszcze, że w zależności od macierzy **G** elementy macierzy **GB** występujące w równaniu (12) nie muszą zależeć od wszystkich elementów macierzy **B**. Jeśli np. macierz **G** ma ν -tą kolumnę zerową, to macierz **GB** nie będzie zależeć od elementu B_ν , który może pozostać nieokreślony.

(Przyjmujemy przy tym następujące założenia.

- a) Wielkości α_{lm} , β_{lm} , γ_{lm} , A_m , B_m są liczbami rzeczywistymi.
- b) Macierz **Z(s)**, gdzie s jest dowolną zmienną rzeczywistą lub zespoloną, jest nieosobliwa, tzn.

$$\Delta(s) \equiv \det \mathbf{Z}(s) \neq 0. \quad (14)$$

- c) Funkcje $f_l(t)$ są funkcjami rzeczywistymi, przedziałami regularnymi¹⁾; ponadto zakładamy, że są to funkcje albo stałe, albo okresowe o okresie $\frac{T}{k_l}$ (k_l — liczby naturalne), przy czym co najmniej jedna z nich jest okresowa o okresie T . W ten sposób macierz **f(t)** jest okresowa o okresie T , tzn.

$$\mathbf{f}(t + T) \equiv \mathbf{f}(t). \quad (15)$$

Jak wiadomo, funkcja **f(t)** okresowa i przedziałami regularna jest równa sumie swojego szeregu Fouriera²⁾:

$$f(t) = \frac{a_0}{2} + \sum_{k=1}^{\infty} (a_k \cos k\omega t + b_k \sin k\omega t) = \sum_{k=-\infty}^{+\infty} c_k e^{jk\omega t}, \quad (16)$$

gdzie $\omega = \frac{2\pi}{T}$ oraz:

$$a_k = \frac{2}{T} \int_0^T f(t) \cos k\omega t dt; \quad b_k = \frac{2}{T} \int_0^T f(t) \sin k\omega t dt \quad (17)$$

$$c_k = \frac{a_k - jb_k}{2} = \frac{1}{T} \int_0^T f(t) e^{-jk\omega t} dt.$$

¹⁾ O funkcji **f(t)** mówimy, iż jest ona przedziałami regularna, jeśli w każdym przedziale domkniętym jest ona prawie wszędzie różniczkowalna, przy czym **f(t)** i **f'(t)** mają w tym przedziale skończoną liczbę punktów nieciągłości tylko pierwszego rodzaju (skoków).

²⁾ Ściśle, równość (16) zachodzi w każdym punkcie, w którym funkcja **f(t)** jest ciągła, i może nie zachodzić w punktach nieciągłości. Utożsamiając funkcje różniące się tylko wartościami w punktach nieciągłości piszemy we wzorze (16) znak równości bez żadnych już zastrzeżeń.

Analogiczne rozwinięcie można więc napisać także dla macierzy $f(t)$:

$$f(t) = \frac{a_0}{2} + \sum_{k=1}^{\infty} (a_k \cos k\omega t + b_k \sin k\omega t) = \sum_{k=-\infty}^{+\infty} c_k e^{jk\omega t}, \quad (18)$$

gdzie: a_k , b_k , c_k są odpowiednimi macierzami-wektorami określonymi analogicznie jak współczynniki (17).

W równaniu macierzowym (12) występuje macierz stała $\mathcal{G}\mathbf{B}$. Przenosząc ją na prawą stronę równania otrzymujemy:

$$\mathbf{Z}(D)\mathbf{y}(t) = \mathbf{f}(t) - \mathcal{G}\mathbf{B}. \quad (19)$$

Zauważamy, że jeśli macierz $f(t)$ spełnia założenia p. c), to przy dowolnej macierzy $\mathbf{B}$ macierz $f(t) - \mathcal{G}\mathbf{B}$ też je spełnia¹⁾. Bez szkody dla ogólności rozważań można więc przyjąć, że równanie ma postać:

$$\mathbf{Z}(D)\mathbf{y}(t) = \mathbf{g}(t), \quad (20)$$

gdzie prawa strona:

$$\mathbf{g}(t) = \mathbf{f}(t) - \mathcal{G}\mathbf{B} \quad (21)$$

spełnia wszystkie poprzednie założenia.

W dalszych rozważaniach będziemy korzystać zarówno z postaci (12), jak i z postaci (20)²⁾.

Rozwiązaniem równania (12) lub (20) [a tym samym rozwiązaniem układu równań (6) lub (8)] będziemy nazywać każdą macierz:

$$\mathbf{y}(t) = \begin{bmatrix} y^{(1)}(t) \\ y^{(2)}(t) \\ \dots\dots\dots \\ y^{(r)}(t) \end{bmatrix}$$

spełniającą dla $t > 0$ równanie (12) przy danej macierzy $\mathbf{B}$.

Podstawowym zadaniem niniejszej pracy jest przedyskutowanie zagadnień istnienia i sposobów wyznaczania rozwiązań okresowych równania (12), przy czym w rozdz. 5 wyjaśnimy jeszcze szczegółowiej, co będziemy rozumieć pod pojęciem rozwiązania okresowego.

¹⁾ Dodanie składnika $-\mathcal{G}\mathbf{B}$ zmienia co najwyżej składowe stałe funkcji $f(t)$ (tj. macierz a_0) nie zmieniając charakteru tych funkcji.

²⁾ W postaci (20) przy z danej macierzy $f(t)$ prawa strona równania, tj. macierz $\mathbf{g}(t)$, nie jest jeszcze jednoznacznie określona, gdyż zależy także od macierzy $\mathbf{B}$. W wielu jednak zagadnieniach ta niejednoznaczność nie powoduje żadnych komplikacji i w przypadkach tych wygodnie jest posługiwać się równaniem w postaci (20).

3. INTERPRETACJA ELEKTRYCZNA

Rozpatrywany w poprzednim rozdziale układ równań różniczkowo-całkowych (6) ma swoją bezpośrednią i oczywistą interpretację elektryczną. Takim bowiem układem równań da się opisać dowolną sieć elektryczną liniową o parametrach skupionych R , C , L , M , w której działa pewna (skończona) liczba źródeł energii o przebiegach okresowych. Zgodnie z zasadą dualizmu wielkościom występującym w równaniach (6) można nadać jedno z dwóch możliwych znaczeń elektrycznych w zależności od tego, czy równania (6) zostały ułożone dla danej sieci metodą prądów obwodowych, czy metodą potencjałów węzłowych.

Dla wyjaśnienia znaczenia elektrycznego użytych symboli podajemy następującą tabelkę.

Wielkości	Metoda prądów obwodowych	Metoda potencjałów węzłowych
r	liczba obwodów niezależnych	liczba węzłów niezależnych
$y^{(m)}(t)$	m -ta funkcja niewiadoma prąd w m -tym obwodzie niezależnym	napięcie m -tego węzła względem węzła odniesienia
$f_m(t)$	m -ty przebieg wymuszający wypadkowa siła elektromotoryczna źródeł działających w m -tym obwodzie niezależnym	wypadkowa wydajność prądowa źródeł doprowadzających prąd do m -tego węzła
$y(t)$	macierz prądów obwodowych	macierz napięć węzłowych
$f(t)$	macierz sił elektromotorycznych obwodowych	macierz wydajności prądowych węzłowych
$\mathfrak{M}$	macierz indukcyjności obwodowych	macierz pojemności węzłowych
$\mathfrak{R}$	macierz oporności (rezystancji) obwodowych	macierz przewodności (konduktancji) węzłowych
$\mathfrak{G}$	macierz odwrotności pojemności obwodowych	macierz odwrotności indukcyjności węzłowych
$Z(j\omega)$	macierz impedancji obwodowych	macierz admitancji węzłowych

Wielkości	Metoda prądów obwodowych	Metoda potencjałów węzłowych
A	macierz wartości początkowych funkcji $y^{(m)}(t)$	
	macierz wartości początkowych prądów obwodowych określająca wartość energii magnetycznej zmagazynowanej w chwili $t = 0$ w sieci	macierz wartości początkowych napięć węzłowych określająca wartość energii elektrycznej zmagazynowanej w chwili $t = 0$ w sieci
B	macierz wartości początkowych całek funkcji $y^{(m)}(t)$	
	macierz wartości początkowych ładunków kondensatorów określająca wartość energii elektrycznej zmagazynowanej w chwili $t = 0$ w sieci	macierz wartości początkowych strumieni magnetycznych cewek określająca wartość energii magnetycznej zmagazynowanej w chwili $t = 0$ w sieci

Jak wiadomo, macierze $\mathcal{U}$, $\mathcal{B}$, $\mathcal{G}$, $\mathbf{Z}(j\omega)$ sieci realizowalnej nie mogą być dowolne, lecz muszą spełniać pewne warunki. Ponieważ warunki te nie będą potrzebne do dalszych rozważań, nie będziemy ich tu zakładać przyjmując za punkt wyjścia układ równań (6) opisujący zarówno wszystkie fizycznie realizowalne układy elektryczne o stałych skupionych (rzeczywiste i idealne), jak i inne typy układów liniowych, nie dających przedstawić się w postaci układów R , C , L , M .

4. ROZWIĄZANIE W POSTACI KLASYCZNEJ I OPERATOROWEJ

Jeśli $y_0(t)$ jest rozwiązaniem równania (12) z macierzą $\mathbf{B}_0$, tzn.

$$\mathbf{Z}(D)y_0(t) + \mathcal{G}\mathbf{B}_0 = \mathbf{f}(t), \quad (22)$$

to każde inne rozwiązanie $\mathbf{y}(t)$ tego równania z dowolną macierzą $\mathbf{B}$ da się przedstawić w postaci:

$$\mathbf{y}(t) = \mathbf{Y}(t) + \mathbf{y}_0(t), \quad (23)$$

gdzie $\mathbf{Y}(t)$ jest rozwiązaniem równania jednorodnego, tzn.¹⁾

$$\mathbf{Z}(D)\mathbf{Y}(t) + \mathcal{G}\mathbf{B}' = \mathbf{0}, \quad (24)$$

gdzie: $\mathbf{B}' = \mathbf{B} - \mathbf{B}_0$.

¹⁾ Symbolami: $\mathbf{1}$ i $\mathbf{0}$ będziemy oznaczać odpowiednio macierz jednostkową i macierz zerową (odpowiednich wymiarów).

Przystąpimy teraz do rozwiązania równania (12) stosując przy tym rachunek operatorowy. Oznaczając:

$$\mathcal{L}[y(t)] = \bar{y}(s); \quad \mathcal{L}[f(t)] = \bar{f}(s); \quad \mathcal{L}[g(t)] = \bar{g}(s)$$

otrzymamy po dokonaniu przekształcenia Laplace'a równania (20) i uwzględnieniu warunku (13):

$$\mathbf{Z}(s)\bar{y}(s) - \mathfrak{A}\mathbf{A} = \bar{g}(s), \quad (25)$$

gdzie:

$$\mathbf{Z}(s) = \mathfrak{A}s + \mathfrak{B} + \mathfrak{C}\frac{1}{s}. \quad (26)$$

Macierz $\mathbf{Z}(s)$ jest z założenia nieosobliwa, istnieje więc macierz odwrotna:

$$\mathbf{W}(s) = \|\mathbf{W}_{lm}(s)\| = \mathbf{Z}^{-1}(s), \quad (27)$$

której elementy $\mathbf{W}_{lm}(s)$ są funkcjami wymiernymi zmiennej s i mają postać:

$$\mathbf{W}_{lm}(s) = \frac{\Delta_{ml}(s)}{\Delta(s)}, \quad (28)$$

gdzie:

$$\Delta(s) \equiv \det \mathbf{Z}(s) \quad (29)$$

oraz $\Delta_{ml}(s)$ są odpowiednimi podwyznacznikami względnymi macierzy $\mathbf{Z}(s)$. Macierz $\mathbf{W}(s)$ będziemy nazywać macierzą funkcji przenoszenia.

Mnożąc równanie (25) lewostronnie przez macierz $\mathbf{W}(s)$ i uwzględniając, że:

$$\mathbf{W}(s)\mathbf{Z}(s) = \mathbf{Z}(s)\mathbf{W}(s) = \mathbf{1}, \quad (30)$$

dostajemy transformatę rozwiązania w postaci:

$$\bar{y}(s) = \mathbf{W}(s)\bar{f}(s) + \mathbf{W}(s)\mathfrak{A}\mathbf{A} - \frac{1}{s}\mathbf{W}(s)\mathfrak{C}\mathbf{B} \quad (31)$$

lub:

$$\bar{y}(s) = \mathbf{W}(s)\bar{g}(s) + \mathbf{W}(s)\mathfrak{A}\mathbf{A}. \quad (32)$$

Należy teraz rozróżnić dwa przypadki.

Przypadek 1: $|a| = \det \mathfrak{A} \neq 0$.

W tym przypadku macierz $\mathbf{W}(s)$ można przedstawić w postaci:

$$\mathbf{W}(s) = \frac{\mathbf{r}_1}{s} + \mathbf{W}_1(s), \quad (33)$$

gdzie macierz $\mathbf{r}_1$ jest stała oraz macierz $\mathbf{W}_1(s)$ jest macierzą funkcji wymiernych i ma własność

$$\lim_{s \rightarrow \infty} s \mathbf{W}_1(s) = \mathbf{0} \quad (34)$$

oraz:

$$\mathbf{r}_1 = \mathfrak{A}^{-1} \quad (35)$$

Elementy macierzy $\mathbf{W}(s)$ są zatem funkcjami wymiernymi, których liczniki są stopnia co najmniej o jeden niższego od stopnia mianownika (wspólnego dla wszystkich elementów).

Przytoczone własności wynikają z następującego prostego rozumowania¹⁾. Wobec związku (30) zachodzą oczywiście także związki:

$$\frac{1}{s} \mathbf{Z}(s) \cdot s \mathbf{W}(s) = s \mathbf{W}(s) \cdot \frac{1}{s} \mathbf{Z}(s) = \mathbf{1},$$

a więc i relacje graniczne:

$$\lim_{s \rightarrow \infty} \frac{1}{s} \mathbf{Z}(s) \cdot s \mathbf{W}(s) = \lim_{s \rightarrow \infty} s \mathbf{W}(s) \cdot \frac{1}{s} \mathbf{Z}(s) = \mathbf{1}. \quad (36)$$

Ponieważ ponadto, zgodnie z (26), istnieje granica:

$$\lim_{s \rightarrow \infty} \frac{1}{s} \mathbf{Z}(s) = \mathfrak{A} \quad (37)$$

i jest z założenia macierzą nieosobliwą, więc istnieje także granica:

$$\lim_{s \rightarrow \infty} s \mathbf{W}(s) = \mathbf{r}_1 \quad (38)$$

i przy tym:

$$\mathfrak{A} \mathbf{r}_1 = \mathbf{r}_1 \mathfrak{A} = \mathbf{1},$$

a więc dochodzimy do zależności (35), a także wzorów (33) i (34).

Oznaczając ($t \geq 0$):

$$\mathcal{L}^{-1}[\mathbf{W}(s)] = \mathbf{h}(t) \quad (39)$$

dostajemy z (32) po wykonaniu przekształcenia odwrotnego ($t \geq 0$):

$$\mathbf{y}(t) = \mathbf{h}(t) * \mathbf{g}(t) + \mathbf{h}(t) \mathfrak{A} \mathbf{A}. \quad (40)$$

Przy $t \rightarrow 0$ otrzymujemy:

$$\mathbf{y}(0) = \lim_{t \rightarrow 0+} \mathbf{y}(t) = \mathbf{h}(0) \mathfrak{A} \mathbf{A} = \mathbf{A}, \quad (41)$$

¹⁾ Do wniosków tych można też dojść wychodząc z ogólnych zależności podanych dla układu równań różniczkowych rzędu n -tego przez Mohra [19] (por. także [21]).

gdyż zgodnie z (38) i (39) mamy:

$$h(0) = \lim_{s \rightarrow \infty} s W(s) = r_1, \quad (42)$$

a więc rozwiązanie (40) istotnie spełnia zadane w postaci macierzy A warunki początkowe.

Przypadek 2: $|a| = \det \mathfrak{A} = 0$.

W tym przypadku macierz funkcji przenoszenia nie ma już, ogólnie biorąc, postaci (33). Przypadek ten będziemy rozpatrywać jedynie przy dodatkowym założeniu upraszczającym, że macierz funkcji przenoszenia ma postać:

$$W(s) = r_0 + W_0(s), \quad (43)$$

gdzie macierz r_0 jest stała oraz macierz $W_0(s)$ jest macierzą funkcji wymiernych mającą własność:

$$\lim_{s \rightarrow \infty} W_0(s) = 0. \quad (44)$$

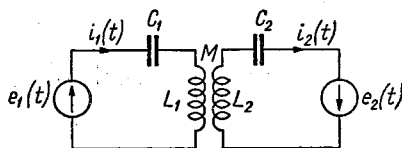
Macierz $W_0(s)$ wygodnie będzie czasami przedstawiać także w postaci analogicznej do (33), tzn.

$$W_0(s) = \frac{r_1}{s} + W_1(s), \quad (45)$$

gdzie r_1 jest macierzą stałą oraz $W_1(s)$ jest macierzą funkcji wymiernych o własności:

$$\lim_{s \rightarrow \infty} s W_1(s) = 0. \quad (46)$$

Ograniczamy się zatem do przypadku, gdy wszystkie funkcje przenoszenia występujące w układzie mają liczniki stopnia co najwyżej równego z mianownikiem.



Rys. 1. Dwa obwody LC sprzężone magnetycznie

Jak wiadomo (por. np. [22]), w dowolnym układzie liniowym o parametrach skupionych R , C , L , M dowolna funkcja przenoszenia jest zawsze funkcją wymierną, której licznik może być stopnia co najwyżej o jeden wyższego niż mianownik. Układy R , C , L , M , w których funkcja przenoszenia ma stopień licznika dokładnie o jeden większy od stopnia mianownika (np. układy idealnie różniczkujące), są jednak układami o bardzo dużym stopniu idealizacji w stosunku do układów rzeczywistych i nie mają praktycznego znaczenia. Do takich układów, obok znanych dwóch najprostszych przypadków (indukcyjność zasilana z idealnego źródła napięciowego i odpowiedni układ dualny — pojemność zasilana z idealnego źródła prądowego),

należy np. układ rozpatrywany przez Ghizetti'ego [7] i Doetscha [4] (por. także [6], str. 319) przedstawiony na rys. 1, dla którego, jak łatwo sprawdzić, macierz funkcji przenoszenia ma postać ($M^2 = L_1 L_2$):

$$W(s) = \frac{C_1 C_2 s}{1 + s^2 (L_1 C_1 + L_2 C_2)} \begin{vmatrix} \frac{1}{C_2} (1 + s^2 L_2 C_2), & -s^2 M \\ -s^2 M, & \frac{1}{C_1} (1 + s^2 L_1 C_1) \end{vmatrix}.$$

Ograniczenie się do układów o funkcjach przenoszenia typu (43) wydaje się więc uzasadnione i nie zmniejsza ono, w zasadzie, ogólności rozważań, jeśli chodzi o rzeczywiste układy elektryczne.

Ponieważ w dalszym ciągu jest słuszna zależność (30), więc zachodzi także równość:

$$\lim_{s \rightarrow \infty} W(s) Z(s) = \lim_{s \rightarrow \infty} Z(s) W(s) = 1. \quad (47)$$

Uwzględniając wzory (26) i (43) oraz (44) otrzymujemy:

$$\begin{aligned} \lim_{s \rightarrow \infty} \left[\mathfrak{A} \mathbf{r}_0 s + \mathfrak{B} \mathbf{r}_0 + \mathfrak{A} \mathbf{r}_1 + \mathfrak{G} \mathbf{r}_0 \frac{1}{s} + \mathfrak{A} \mathbf{W}_1(s) s + \left(\mathfrak{B} + \mathfrak{G} \frac{1}{s} \right) \mathbf{W}_0(s) \right] = \\ = \lim_{s \rightarrow \infty} \left[\mathbf{r}_0 \mathfrak{A} s + \mathbf{r}_0 \mathfrak{B} + \mathbf{r}_1 \mathfrak{A} + \mathbf{r}_0 \mathfrak{G} \frac{1}{s} + \mathbf{W}_1(s) \mathfrak{A} s + \mathbf{W}_0(s) \left(\mathfrak{B} + \mathfrak{G} \frac{1}{s} \right) \right] = 1. \end{aligned} \quad (48)$$

Z zależności tej uwzględniając (44) i (46) otrzymujemy związki między współczynnikami macierzy $Z(s)$ i $W(s)$ w postaci:

$$\mathfrak{A} \mathbf{r}_0 = \mathbf{r}_0 \mathfrak{A} = \mathbf{0}, \quad (49)$$

$$\mathfrak{B} \mathbf{r}_0 + \mathfrak{A} \mathbf{r}_1 = \mathbf{r}_0 \mathfrak{B} + \mathbf{r}_1 \mathfrak{A} = \mathbf{1}^1). \quad (50)$$

Transformatę rozwiązania można przedstawić teraz zgodnie z (32) i (49) w postaci:

$$\bar{\mathbf{y}}(s) = \mathbf{r}_0 \bar{\mathbf{g}}(s) + \mathbf{W}_0(s) \bar{\mathbf{g}}(s) + \mathbf{W}_0(s) \mathfrak{A} \mathbf{A}. \quad (51)$$

Oznaczając ($t \geq 0$):

$$\mathcal{L}^{-1}[\mathbf{W}_0(s)] = \mathbf{h}_0(t) \quad (52)$$

otrzymujemy po wykonaniu przekształcenia odwrotnego ($t \geq 0$):

$$\mathbf{y}(t) = \mathbf{r}_0 \mathbf{g}(t) + \mathbf{h}_0(t) * \mathbf{g}(t) + \mathbf{h}_0(t) \mathfrak{A} \mathbf{A}. \quad (53)$$

W tym przypadku rozwiązanie $\mathbf{y}(t)$ przybiera przy $t \rightarrow 0$ wartości:

$$\mathbf{y}(0) = \lim_{t \rightarrow 0+} \mathbf{y}(t) = \mathbf{r}_0 \mathbf{g}(0) + \mathbf{h}_0(0) \mathfrak{A} \mathbf{A} = \mathbf{r}_0 \mathbf{g}(0) + \mathbf{r}_1 \mathfrak{A} \mathbf{A}, \quad (54)$$

¹⁾ Z równań (49) i (50) wynika, że jeśli $|\alpha| = \det \mathfrak{A} = 0$ oraz $\mathfrak{A} \neq \mathbf{0}$, to: $\mathfrak{B} \neq \mathbf{0}$, a także jeśli $\mathfrak{A} = \mathbf{0}$, to: $|\beta| = \det \mathfrak{B} \neq 0$.

przy czym korzystając z (50) zależność (54) można przepisać w postaci:

$$\mathbf{y}(0) = \mathbf{A} + \mathbf{r}_0[\mathbf{g}(0) - \mathfrak{B}\mathbf{A}]. \quad (55)$$

Z zależności (55) wynika, że wartości osiągane przez rozwiązanie przy $t = 0+$ różnią się teraz, ogólnie biorąc, od wartości początkowych określonych macierzą $\mathbf{A}$. Wynika to stąd, że w przypadku, gdy wyznacznik macierzy współczynników przy najstarszych pochodnych w układzie równań (w naszym przypadku $|\alpha|$) jest równy zeru, wówczas nie zawsze istnieje rozwiązanie spełniające dowolnie zadane warunki początkowe. Sam układ równań narzuca wówczas pewne związki między warunkami początkowymi i wartościami w zerze prawych stron równań względnie ich pochodnych, tak że nie wszystkie z warunków początkowych (w naszym przypadku — nie wszystkie z elementów macierzy $\mathbf{A}$) można przyjąć w sposób niezależny ([19], [4]). Muszą być spełnione tzw. warunki zgodności, które w naszym przypadku wobec równości (55) przybierają postać:¹⁾

$$\mathbf{r}_0[\mathbf{g}(0) - \mathfrak{B}\mathbf{A}] = \mathbf{0} \quad (56)$$

lub wobec (21):

$$\mathbf{r}_0[\mathbf{f}(0) - \mathfrak{G}\mathbf{B} - \mathfrak{B}\mathbf{A}] = \mathbf{0}.$$

Jeśli jednak warunki te nie są spełnione i warunki początkowe (macierz $\mathbf{A}$) są zadane w sposób dowolny, to mimo wszystko otrzymuje się pewne „rozwiązanie” [w naszym przypadku (53)], które jednak przy $t = 0+$ przybiera inne, niż założono, wartości początkowe [wartości (55)]. Jak pokazał Doetsch [4], można w pewien sposób, uogólnić pojęcie rozwiązania układu równań dopuszczając nieciągłość tego rozwiązania w chwili początkowej $t = 0$. Z takim właśnie rozwiązaniem mamy do czynienia w naszym przypadku, jeśli nie jest spełniony warunek (56). Zadane z góry wartości początkowe (macierz $\mathbf{A}$) można traktować teraz jako istniejące w układzie przy $t = 0-$, a różnice $\mathbf{y}(0) - \mathbf{A}$ jako skoki nieciągłego przy $t = 0$ rozwiązania. Jak wiadomo, przy dokonywaniu przekształcenia Laplace’a układu równań [układu (20)] potrzeba jest znajomości wartości przyjmowanych przez rozwiązanie przy $t = 0+$, podczas gdy równanie (25) i wynikające z niego dalsze zależności otrzymaliśmy biorąc za te wartości elementy macierzy $\mathbf{A}$. Jak wykazał jednak Doetsch [4] (por. także [7], [19]), postępowanie takie prowadzi do poprawnych rezultatów, gdyż — jak łatwo sprawdzić — jeśli przy dokonywaniu prze-

¹⁾ Por. także: V. Doležal, O systémech lineárních integrodiferenciálních rovnic v technických problémech (Aplikace Matematiky, 1959, 4).

kształcenia Laplace'a równania (20) wziąć w miejsce macierzy $\mathbf{A}$ macierz $\mathbf{y}(0)$ określoną wzorem (55), tj. macierz rzeczywiście przyjmowanych przez rozwiązanie wartości przy $t = 0+$, otrzymuje się rezultat identyczny, a więc poprawny. Jest to niejako cecha samego przekształcenia Laplace'a, które mimo iż wychodzimy od wartości początkowych, które nie mogą być przyjęte przez rozwiązanie, prowadzi w pewnym sensie automatycznie do wartości rzeczywiście osiąganych.

Warto jeszcze zwrócić uwagę na to, że macierz funkcji przenoszenia o postaci (43) otrzymujemy również w przypadku, gdy $\mathfrak{A} = 0$ oraz $|\beta| = \det \mathfrak{B} \neq 0$ ¹⁾. Istotnie, w tym przypadku

$$\mathbf{Z}(s) = \mathfrak{B} + \mathfrak{C} \frac{1}{s} \quad (57)$$

oraz istnieje granica:

$$\lim_{s \rightarrow \infty} \mathbf{Z}(s) = \mathfrak{B}$$

i jest macierzą nieosobliwą. Z równości (47) wynika zatem, że istnieje wówczas również granica:

$$\lim_{s \rightarrow \infty} \mathbf{W}(s) = \mathbf{r}_0 \quad (58)$$

oraz że²⁾:

$$\mathbf{r}_0 = \mathfrak{B}^{-1}. \quad (59)$$

Wobec zależności (58) macierz funkcji przenoszenia $\mathbf{W}(s)$ daje się zatem przedstawić w postaci (43).

5. ROZWIĄZANIE OKRESOWE I USTALONE

Przy rozpatrywaniu sieci elektrycznych liniowych, pobudzanych okresowo i odpowiednich układów równań jednym z najistotniejszych zagadnień jest wyznaczenie rozwiązania okresowego, tj. takiego rozwiązania szczególnego, które jest okresowe i ma okres równy okresowi przebiegu pobudzającego. W przypadku sieci rzeczywistych zawierających rezystancje, których funkcje przenoszenia mają bieguny leżące tylko w lewej

¹⁾ Jeśli traktować rozważany układ równań jako układ równań sieci elektrycznej otrzymany metodą prądów obwodowych, to założenia te mają prostą interpretację. Warunek $\mathfrak{A} = 0$ oznacza, że w układzie nie występują indukcyjności (układ R, C), a warunek $\det \mathfrak{B} \neq 0$ świadczy o tym, że przy dowolnym wyborze obwodów niezależnych w każdym obwodzie znajduje się oporność rzeczywista (rezystancja).

²⁾ Zwracamy uwagę na to, że w przypadku, gdy $\mathfrak{A} \neq 0$ macierz $\mathbf{r}_0$ jest zawsze osobliwa, gdyż gdyby była nieosobliwa, z równości (49) otrzymalibyśmy: $\mathfrak{A} = 0$.

półpłaszczyźnie, przy pobudzaniu dowolnym przebiegiem okresowym istnieje zawsze stan ustalony i rozwiązanie okresowe jest jednocześnie tzw. rozwiązaniem ustalonym, tj. opisuje zachowanie się sieci w stanie ustalonym.

W przypadku jednak układów bezrezystancyjnych (bieguny funkcji przenoszenia leżą na osi urojonej) pobudzanych okresowo nie można już mówić, w zasadzie, o stanie ustalonym w poprzednim znaczeniu. Rozwiązań okresowych, jakie mogą w takich przypadkach istnieć, nie będziemy już zatem zaliczać do rozwiązań ustalonych.

Jeszcze szerszy zakres możliwości dostaniemy, jeśli będziemy przyjmować za punkt wyjścia nie realną sieć elektryczną, lecz układ równań (6), nie wnikając w możliwości realizacji sieci o równaniach (6). Otrzymuje się wówczas ogólny problem czysto matematyczny. Położenie biegunów funkcji przenoszenia może być teraz dowolne i ogólnie biorąc, nie istnieje stan ustalony oraz rozwiązanie ustalone. Rozwiązanie okresowe może natomiast istnieć tracąc jednak zupełnie ten sens fizyczny, jaki jest związany ze stanem ustalonym.

Wyda się więc celowe wyraźne rozróżnienie pomiędzy rozwiązaniem okresowym i ustalonym oraz ściśle zdefiniowanie tych pojęć.

Rozwiązanie okresowe określamy następująco:¹⁾

Rozwiązanie $y_0(t)$ równania (20) będziemy nazywać rozwiązaniem okresowym, jeśli wszystkie funkcje $y_0^{(m)}(t)$ będące współrzędnymi wektora

$$y_0(t) = \begin{bmatrix} y_0^{(1)}(t) \\ y_0^{(2)}(t) \\ \dots\dots\dots \\ y_0^{(x)}(t) \end{bmatrix}$$

są funkcjami okresowymi o okresach $\frac{T}{k_m}$ (k_m — liczby naturalne) lub stałymi, przy czym co najmniej jedna z nich nie jest stała.

Jeśli k oznacza największy wspólny dzielnik liczb k_m , to rozwiązanie $y_0(t)$ jest okresowe w okresie $\frac{T}{k}$, tzn.

$$y_0\left(t + \frac{T}{k}\right) \equiv y_0(t). \quad (60)$$

¹⁾ Określenia rozwiązania okresowego i ustalonego podajemy od razu w odniesieniu do rozpatrywanego układu równań (6) lub równania macierzowego (20), jest jednak oczywiste, iż analogiczne określenia można by wypowiedzieć i w przypadku ogólniejszym, na przykład układu równań różniczkowych rzędu n -tego.

Jeśli istnieje rozwiązanie okresowe $y_0(t)$, to dowolne inne, rozwiązanie $y(t)$ równania (20) można przedstawić w postaci (23), gdzie $Y(t)$ jest rozwiązaniem równania jednorodnego (24).

Podamy teraz określenie rozwiązania ustalonego.

Rozwiązanie okresowe $y_0(t)$ będziemy nazywać rozwiązaniem ustalonym równania (20), jeśli dla każdego rozwiązania tego równania zachodzi równość:

$$\lim_{t \rightarrow +\infty} [y(t) - y_0(t)] = 0. \quad (61)$$

Łatwo zauważyć, że wobec zależności (23) równość (61) zachodzi wtedy i tylko wtedy, gdy:

$$\lim_{+\infty} Y(t) = 0 \quad (62)$$

przy dowolnych warunkach początkowych, tj. przy dowolnej macierzy wartości początkowych funkcji $Y(t)$ i dowolnej macierzy B' w równaniu (24).

Widzimy więc, że rozwiązanie ustalone zostało określone jako szczególny przypadek rozwiązania okresowego. Rozwiązanie ustalone (jeśli istnieje) posiada tę własność, że każde inne rozwiązanie, tj. każde rozwiązanie „startujące” z dowolnych warunków początkowych, dąży asymptotycznie do rozwiązania ustalonego; wydaje się więc, iż przy takim właśnie formalnym zdefiniowaniu rozwiązania ustalonego, odpowiada ono fizycznemu pojmowaniu stanu ustalonego.

Podamy teraz kilka podstawowych własności rozwiązania ustalonego.

Równanie (20) może mieć przy danej macierzy $f(t)$ co najwyżej jedno rozwiązanie ustalone.

Istotnie, gdyby bowiem było przeciwnie, tj. jednocześnie $y_0(t)$ i $y_0^*(t)$ były rozwiązaniami ustalonymi równania (20), to wówczas dowolne rozwiązanie $y(t)$ spełniałoby, zgodnie z (61), jednocześnie zależności:

$$\lim_{t \rightarrow +\infty} [y(t) - y_0(t)] = 0; \quad \lim_{t \rightarrow +\infty} [y(t) - y_0^*(t)] = 0,$$

a więc byłoby także:

$$\lim_{t \rightarrow +\infty} [y_0(t) - y_0^*(t)] = 0. \quad (63)$$

Ponieważ zgodnie z podanym określeniem różnica w nawiasie kwadratowym może być tylko wektorem stałym lub okresowym, więc z (63) wynika, że¹⁾:

¹⁾ Przypominamy, iż zgodnie z uwagą podaną w przypisku na str. 131, utożsamiamy funkcje różniące się tylko wartościami w punktach nieciągłości, tj. funkcje równe prawie wszędzie.

$$y_0(t) \equiv y_0^*(t),$$

tzn. rozwiązanie ustalone jest jedyne.

Podamy teraz warunki istnienia rozwiązania ustalonego.

Do tego, aby istniało rozwiązanie ustalone równania (12) [lub (20)], trzeba i wystarczy, aby wszystkie bieguny macierzy $W(s)$ ¹⁾ leżały w lewej półpłaszczyźnie.

Nazkicujemy dowód tej własności. Przede wszystkim zauważymy, że do tego, aby istniało rozwiązanie ustalone, trzeba i wystarczy, aby: 1) istniało rozwiązanie okresowe oraz 2) zachodziła równość (62).

Założmy na razie, że wszystkie bieguny macierzy $W(s)$ leżą w lewej półpłaszczyźnie (dowód dostateczności). Wystarczy to, jak zobaczymy w rozdz. 7 (por. także [14]), do istnienia rozwiązania okresowego $y_0(t)$, a więc istnieją wówczas takie macierze A_0 i B_0 , że:

$$\bar{y}_0(s) = W(s) \bar{f}(s) - \frac{1}{s} W(s) \mathfrak{G} B_0 + W(s) \mathfrak{A} A_0 \quad (64)$$

jest transformata rozwiązania okresowego. Transformata dowolnego innego rozwiązania $\bar{y}(s)$ ma postać (31), a wobec tego transformata różnicy jest równa:

$$\bar{Y}(s) = \bar{y}(s) - \bar{y}_0(s) = -\frac{1}{s} W(s) \mathfrak{G} B' + W(s) \mathfrak{A} A', \quad (65)$$

gdzie:

$$A' = A - A_0; \quad B' = B - B_0$$

są macierzami dowolnymi. Ponieważ wszystkie bieguny macierzy $W(s)$ leżą w lewej półpłaszczyźnie, więc niezależnie od tego, czy zachodzi przypadek 1. czy 2. (rozdz. 4), mamy zawsze — por. (43) i (49):

$$\lim_{t \rightarrow +\infty} \mathcal{L}^{-1} [W(s) \mathfrak{A} A'] = 0 \quad (66)$$

przy dowolnej macierzy A' . Analogicznie przy dowolnej macierzy B' mamy:

$$\lim_{t \rightarrow +\infty} \mathcal{L}^{-1} \left[\frac{1}{s} W(s) \mathfrak{G} B' \right] = 0 \quad (67)$$

co wynika stąd, że wobec zależności:²⁾

$$\mathfrak{G} W(0) = W(0) \mathfrak{G} = 0 \quad (68)$$

¹⁾ Mówiąc o biegunach macierzy $W(s)$ rozumiemy zbiór biegunów wszystkich jej elementów.

²⁾ Zależność ta wynika stąd, że przy założeniu, że $s = 0$ nie jest wartością charakterystyczną macierzy $Z(s)$ [a więc nie jest biegunem $W(s)$], mamy zgodnie z (30):

$$\lim_{s \rightarrow 0} s Z(s) \cdot W(s) = \lim_{s \rightarrow 0} W(s) \cdot s Z(s) = 0,$$

co prowadzi wprost do zależności (68).

składnik stały odpowiadający biegunowi $s = 0$ jest przy dowolnej macierzy $\mathbf{B}'$ równy zeru:

$$\lim_{s \rightarrow 0} [e^{st} \mathbf{W}(s)] \mathcal{G} \mathbf{B}' = \mathbf{W}(0) \mathcal{G} \mathbf{B}' = 0 \quad (69)$$

a pozostałe składniki jako odpowiadające biegunom leżącym w lewej półpłaszczyźnie, jak poprzednio, dążą do zera. Mamy więc:

$$\lim_{t \rightarrow +\infty} \mathbf{Y}(t) = 0 \quad (70)$$

i dostateczność warunków jest wykazana.

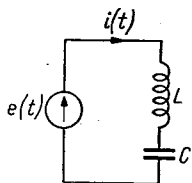
Konieczność podanych warunków wynika stąd, że gdyby macierz przenoszenia $\mathbf{W}(s)$ miała choć jeden biegun leżący na osi urojonej lub w prawej półpłaszczyźnie, to zawsze można by dobrać warunki początkowe tak, aby macierz $\mathbf{Y}(t)$ zawierała składnik stały lub nieograniczony przy $t \rightarrow +\infty$, a więc własność (70) nie zachodziłaby.

Z przytoczonej własności wynika również następujący prosty i ważny wniosek.

Jeśli istnieje rozwiązanie ustalone równania (12) przy danej macierzy $\mathbf{f}(t)$, to istnieje również rozwiązanie ustalone przy każdej innej macierzy okresowej $\mathbf{f}^(t)$.*

Przytoczone własności odróżniają w sposób istotny rozwiązania ustalone od ogólniejszego rozwiązania okresowego, które takich cech regularnościowych, ogólnie biorąc, nie posiada. Jak zobaczymy na prostych przykładach, równanie (12) może np. w zależności od prawej strony:

- a) mieć jedyne rozwiązanie okresowe nie będące rozwiązaniem ustalonym,
- b) mieć nieskończenie wiele rozwiązań okresowych,
- c) nie mieć rozwiązania okresowego.



Rys. 2. Dwójnik LC zasilany napięciowo

Przykład

Rozpatrzmy dwójnik LC (rys. 2) zasilany ze źródła napięciowego o sile elektromotorycznej:

$$e(t) = E \sin \omega t.$$

Równanie obwodu ma postać:

$$L \frac{di}{dt} + \frac{1}{C} \left(\int_0^t i dt + B \right) = e(t).$$

Przy oznaczeniu:

$$\omega_0^2 = \frac{1}{LC}$$

i warunku początkowym:

$$i(0) = A$$

równanie to ma rozwiązanie:

$$i(t) = \frac{\omega E}{L(\omega^2 - \omega_0^2)} (\cos \omega_0 t - \cos \omega t) + A \cos \omega_0 t - \omega_0 B \sin \omega_0 t \quad (71)$$

przy $\omega \neq \omega_0$ oraz przy $\omega = \omega_0$:

$$i(t) = \frac{E}{2L} t \sin \omega_0 t + A \cos \omega_0 t - \omega_0 B \sin \omega_0 t. \quad (72)$$

Mamy teraz następujące przypadki:

a. $\omega \neq \omega_0$; istnieje wówczas rozwiązanie okresowe:

$$i_0(t) = -\frac{\omega E}{L(\omega^2 - \omega_0^2)} \cos \omega t,$$

ale nie istnieje rozwiązanie ustalone, bowiem, ogólnie biorąc:¹⁾,

$$\lim_{t \rightarrow +\infty} [i(t) - i_0(t)] \neq 0.$$

Jeżeli przy tym:

$$\omega \neq \frac{\omega_0}{\nu}; \quad \nu = 2, 3, 4, \dots,$$

to rozwiązanie powyższe jest j e d y n y m rozwiązaniem okresowym. Jeżeli natomiast:

$$\omega = \frac{\omega_0}{\nu},$$

to k a ż d e rozwiązanie (71) przy dowolnych wartościach A i B jest rozwiązaniem okresowym.

b. $\omega = \omega_0$; jak wynika z (72), przy ż a d n y c h wartościach A i B rozwiązanie (72) nie jest okresowe, a więc rozwiązanie okresowe nie istnieje.

¹⁾ Jeśli przyjmiemy, że w obwodzie występuje pewna oporność rzeczywista R , to w przypadku tym będzie oczywiście istnieć rozwiązanie ustalone, które przy $R \rightarrow 0$ i $\omega \neq \omega_0$ przejdzie w granicę w rozwiązanie (72) obwodu LC. Mimo więc, iż rozwiązanie (72), zgodnie z przyjętą definicją, nie jest rozwiązaniem ustalonym, posiada ono jednak w pewnym zakresie cechy związane ze stanem ustalonym rozumianym w tym sensie, iż chodzi tu o przypadek graniczny przy $R \rightarrow 0$ stanu ustalonego w obwodzie RLC.

Uwaga ta dotyczy również innych układów bezrezystancyjnych.

6. WŁASNOŚCI TRANSFORMAT FUNKCJI OKRESOWYCH

Zanim przejdziemy do omówienia warunków istnienia i sposobów wyznaczania rozwiązań okresowych, podamy jeszcze kilka podstawowych własności transformat funkcji okresowych potrzebnych do dalszych rozważań.

Jak wiadomo (por. np. [15], [5]), transformatę dowolnej funkcji okresowej $f(t)$ o okresie T można przedstawić w postaci:

$$\bar{f}(s) = \frac{\bar{f}_T(s)}{1 - e^{-sT}}, \quad (73)$$

gdzie:

$$\bar{f}_T(s) = \int_0^T e^{-st} f(t) dt. \quad (74)$$

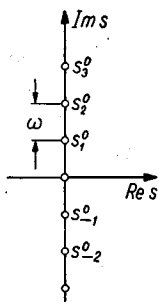
Funkcja $\bar{f}_T(s)$, często nazywana transformatą za okres funkcji okresowej $f(t)$, jest funkcją całkowitą zmiennej zespolonej s (por. np. [5], str. 145), a więc transformata $\bar{f}(s)$ jest funkcją meromorficzną, której biegunami mogą być tylko zera mianownika, tj. punkty spełniające równanie:

$$1 - e^{-sT} = 0, \quad (75)$$

czyli:

$$s = s_k^0 = jk \frac{2\pi}{T} = jk\omega, \quad (76)$$

gdzie k jest dowolną liczbą całkowitą. Wszystkie te punkty są pierwiast-



Rys. 3. Położenie biegunów transformaty funkcji okresowej o okresie $T = \frac{2\pi}{\omega}$

kami pojedynczymi równania (75) i leżą na osi urojonej w jednakowych odstępach, co ω (rys. 3).

Ponieważ:

$$\text{res}_{s=s_k^0} \bar{f}(s) = \frac{\bar{f}_T(s_k^0)}{\frac{d}{ds}(1 - e^{-sT})|_{s=s_k^0}} = \frac{1}{T} \bar{f}_T(s_k^0) = \frac{1}{T} \int_0^T e^{-jk\omega t} f(t) dt = c_k \quad (77)$$

oraz analogicznie¹⁾:

$$\text{res}_{s=s_{-k}^0} \bar{f}(s) = c_{-k} = \tilde{c}_k,$$

gdzie c_k jest odpowiednim współczynnikiem w zespolonym rozwinięciu Fouriera funkcji $f(t)$ — por. (16) i (17) —

$$f(t) = \sum_{k=-\infty}^{+\infty} c_k e^{jk\omega t} = \sum_{k=-\infty}^{+\infty} c_k e^{s_k^0 t} \quad (78)$$

Punkty $s_k^0 = jk\omega$ i $s_{-k}^0 = -jk\omega$ stanowią więc wtedy i tylko wtedy parę biegunów (pojedynczych) transformaty $\bar{f}(s)$, gdy $c_k \neq 0$, tzn. gdy w rozwinięciu Fourierowskim funkcji $f(t)$ występuje k -ta harmoniczna.

Jeśli natomiast $c_k = 0$, tzn. w rozwinięciu Fourierowskim nie występuje k -ta harmoniczna, wówczas punkty $s_{\pm k}^0$ nie są biegunami transformaty. To samo odnosi się do punktu $s^0 = 0$ i składowej stałej w rozkładzie Fourierowskim.

Podstawiając do (74) rozwinięcie (78) i całkując wyraz po wyrazie (jest to dopuszczalne zgodnie ze znaną własnością szeregów Fouriera — por. np. [24], str. 63) dostajemy uwzględniając wzór (76):

$$\begin{aligned} \bar{f}_T(s) &= \int_0^T e^{-st} \sum_{k=-\infty}^{+\infty} c_k e^{s_k^0 t} dt = \sum_{k=-\infty}^{+\infty} c_k \int_0^T e^{-(s-s_k^0)t} dt = \\ &= (1 - e^{-sT}) \sum_{k=-\infty}^{+\infty} \frac{c_k}{s - s_k^0}, \end{aligned} \quad (79)$$

a stąd zgodnie z (73) otrzymujemy:

$$\bar{f}(s) = \frac{\bar{f}_T(s)}{1 - e^{-sT}} = \sum_{k=-\infty}^{+\infty} \frac{c_k}{s - s_k^0}. \quad (80)$$

Zależność (80) przedstawia jednocześnie rozkład funkcji meromorficznej $\bar{f}(s)$ na ułamki proste (tzw. Mittag — Lefflera, por. np. [13]). Poszczególne wyrazy szeregu (80) są, zgodnie z (76) i (77), częściami głównymi szeregów Laurenta transformaty $f(s)$ w otoczeniach biegunów s_k .

¹⁾ Wężyk oznacza liczbę zespoloną sprzężoną.

Ponieważ:

$$\mathcal{L}\left[e^{\frac{s_k^0}{s}t}\right] = \frac{1}{s - s_k^0}, \quad (81)$$

widzimy zatem porównując zależności (78) i (80), iż można wykonywać przekształcenie Laplace'a szeregu (78) oraz odwrotne przekształcenie Laplace'a szeregu (80) wyraz po wyrazie.

Wynika stąd także, iż dla transformaty $\bar{f}(s)$ — zależność (73) — słuszne jest twierdzenie o rozkładzie¹⁾ ($t \geq 0$)

$$\mathcal{L}^{-1}[\bar{f}(s)] = \frac{1}{2\pi j} \int_L \bar{f}(s) e^{st} ds = \sum_{k=-\infty}^{+\infty} [\text{res } \bar{f}(s)] e^{\frac{s_k^0}{s}t} \quad (82)$$

7. WARUNKI ISTNIENIA ROZWIĄZANIA OKRESOWEGO

Zajmiemy się teraz zagadnieniem warunków, przy których istnieje rozwiązanie okresowe (por. [14], [26]).

Transformata dowolnego rozwiązania równania (12) wyrażona wzorem (31) jest (niezależnie od tego, czy zachodzi przypadek 1, czy 2 — por. rozdz. 4) funkcją meromorficzną zmiennej zespolonej s . Jak pokażemy, transformata ta spełnia założenia stosowalności twierdzenia o rozkładzie dla funkcji meromorficznych (por. np. [5], str. 276), a więc wynik przekształcenia odwrotnego może być zgodnie z tym twierdzeniem przedstawiony w postaci ($t \geq 0$):

$$y(t) = \mathcal{L}^{-1}[\bar{y}(s)] = \frac{1}{2\pi j} \int_L \bar{y}(s) e^{st} ds = \sum_{\substack{\text{wszystkie} \\ \text{bieguny}}} \text{res} [\bar{y}(s) e^{st}]. \quad (83)$$

Zauważymy przede wszystkim, iż aby dowieść słuszności wzoru (83) wystarczy wykazać, że można stosować twierdzenie o rozkładzie do każdej transformaty o postaci²⁾:

$$\bar{y}(s) = W(s) \bar{f}(s),$$

gdzie funkcja przenoszenia $W(s)$ ma stopień licznika niższy od stopnia mianownika,

¹⁾ Por. np. [11], [13] lub [5].

²⁾ Łatwo zauważyć, że wszystkie pozostałe składniki występujące w elementach macierzy $\bar{y}(s)$ określonej wzorem (31) lub (51) są albo funkcjami wymiernymi o stopniu licznika niższym od stopnia mianownika, albo transformatami funkcji okresowych przedziałami regularnych, a więc do wszystkich tych składników można twierdzenie o rozkładzie stosować bez zastrzeżeń.

$\bar{f}(s)$ jest transformatą dowolnej funkcji okresowej o okresie T , przedziałami regularnej (por. str. 131, tzn. że $t \geq 0$):

$$y(t) = \frac{1}{2\pi j} \int_L \bar{y}(s) e^{st} ds = \sum_{\substack{\text{wszystkie} \\ \text{bieguny}}} \text{res} [W(s) \bar{f}(s) e^{st}]. \quad (84)$$

Aby to wykazać, oprzemy się na udowodnionej w Dodatku własności, iż można tak dobrać ciąg okręgów $\{C_n\}$: $|s| = R_n$, $R_n \rightarrow \infty$, na których transformata $\bar{f}(s)$ dowolnej funkcji okresowej przedziałami regularnej jest ograniczona, tzn.

$$|\bar{f}(s)|_{C_n} \leq K_1.$$

Ponieważ stopień licznika funkcji przenoszenia $W(s)$ jest niższy od stopnia mianownika, więc: $W(s) \rightarrow 0$ przy $s \rightarrow \infty$ ($|s| > R$). Dobierając okręgi C_n tak, aby nie przechodziły przez bieguny s_ν funkcji $W(s)$ (zawsze jest to możliwe) dostajemy:

$$|s W(s)|_{C_n} \leq K_2,$$

a więc:

$$|s \bar{y}(s)|_{C_n} = |s W(s) \bar{f}(s)|_{C_n} \leq K_1 K_2 = K,$$

co wystarcza do stosowalności twierdzenia o rozkładzie i dowodzi słuszności zależności (84).

Rozpatrzmy teraz dwa przypadki wzajemnego położenia biegunów macierzy funkcji przenoszenia $W(s)$ i macierzy $\bar{f}(s)$.

Przypadek 1. Żaden z biegunów s_ν ($\nu = 1, 2, \dots, N$) macierzy funkcji przenoszenia $W(s)$ nie pokrywa się z żadnym z punktów s_k^0 określonych zależnością (76), w których macierz $\bar{f}(s)$ może mieć bieguny. Oznacza to, że dla wszystkich s_ν ($\nu = 1, 2, \dots, N$) mamy:

$$1 - e^{-s_\nu T} \neq 0. \quad (85)$$

Wykażemy, że w tym przypadku istnieje jedyne rozwiązanie okresowe $\bar{y}_0(t)$ określone wzorem:

$$y_0(t) = \sum_{k=-\infty}^{+\infty} \text{res}_{s=s_k^0} \left[W(s) \frac{\bar{f}_T(s)}{1 - e^{-sT}} e^{st} \right] = \sum_{k=-\infty}^{+\infty} W(s_k^0) c_k e^{s_k^0 t}, \quad (86)$$

gdzie — por. (77) — :

$$c_k = \text{res}_{s=s_k^0} \bar{f}(s) \quad (87)$$

i przy tym (str. 132):

$$f(t) = \sum_{k=-\infty}^{+\infty} c_k e^{j k \omega t} \quad (88)$$

Dowód można oprzeć na tym, iż jak łatwo stwierdzić wobec własności macierzy $W(s)$ i macierzy współczynników Fouriera c_k szereg (86) jest zbieżny i można go

różniczkować wyraz po wyrazie¹⁾. Podstawiając więc (86) do równania (12) dostajemy:

$$\begin{aligned} \mathbf{Z}(D) \mathbf{y}_0(t) + \mathfrak{G} \mathbf{B} &= \sum_{\substack{k=-\infty \\ k \neq 0}}^{+\infty} \left[\mathfrak{U} s_k^0 \mathbf{W}(s_k^0) + \mathfrak{B} \mathbf{W}(s_k^0) + \mathfrak{G} \frac{1}{s_k^0} \mathbf{W}(s_k^0) \right] \mathbf{c}_k e^{s_k^0 t} + \\ &- \sum_{\substack{k=-\infty \\ k \neq 0}}^{+\infty} \mathfrak{G} \frac{1}{s_k^0} \mathbf{W}(s_k^0) \mathbf{c}_k + \mathfrak{B} \mathbf{W}(0) \mathbf{c}_0 + \mathfrak{G} \mathbf{W}(0) \mathbf{c}_0 t + \mathfrak{G} \mathbf{B}. \end{aligned}$$

Z założenia punkt $s = 0$ nie jest biegunem macierzy $\mathbf{W}(s)$, przechodząc zatem w zależności (30) do granicy przy $s \rightarrow 0$ dostajemy — uwzględniając (68):

$$\mathfrak{B} \mathbf{W}(0) = 1 - \lim_{s \rightarrow 0} \frac{1}{s} \mathfrak{G} \mathbf{W}(s) = 1 - \mathfrak{G} \mathbf{W}'(0).$$

Wobec zależności (68) oraz (30) przy $s = s_k^0$ otrzymujemy więc:

$$\mathbf{Z}(D) \mathbf{y}_0(t) + \mathfrak{G} \mathbf{B} = \sum_{k=-\infty}^{+\infty} \mathbf{c}_k e^{s_k^0 t} + \mathfrak{G} \left[\mathbf{B} - \sum_{\substack{k=-\infty \\ k \neq 0}}^{+\infty} \frac{1}{s_k^0} \mathbf{W}(s_k^0) \mathbf{c}_k - \mathbf{W}'(0) \mathbf{c}_0 \right].$$

Dobierając teraz macierz $\mathbf{B}$ tak, aby:

$$\mathbf{B} = \sum_{\substack{k=-\infty \\ k \neq 0}}^{+\infty} \frac{1}{s_k^0} \mathbf{W}(s_k^0) \mathbf{c}_k + \mathbf{W}'(0) \mathbf{c}_0,$$

dostajemy wobec (88):

$$\mathbf{Z}(D) \mathbf{y}_0(t) + \mathfrak{G} \mathbf{B} = \mathbf{f}(t),$$

a więc $\mathbf{y}_0(t)$ jest istotnie rozwiązaniem równania (12).

Fakt, że rozwiązanie (86) jest jedynym rozwiązaniem okresowym, wynika stąd, że zgodnie z (23) każde inne rozwiązanie $\mathbf{y}(t)$ różni się od rozwiązania (86) o macierz $\mathbf{Y}(t)$, której transformata jest określona wzorem (66). Ponieważ jednak z założenia żaden z biegunów macierzy $\mathbf{W}(s)$ nie leży na osi urojonej oraz wobec zależności (70), macierz $\mathbf{Y}(t)$ przy żadnym doborze $\mathbf{A}'$ i $\mathbf{B}'$ nie jest ani stała, ani okresowa (z wyjątkiem przypadku $\mathbf{Y}(t) = 0$), a więc $\mathbf{y}_0(t)$ jest istotnie jedynym rozwiązaniem okresowym.

Rozwiązanie (86) ma prostą budowę. Jeśli wrócimy do zależności (65) i pominiemy w niej składniki zależne od warunków początkowych (np. przyjmiemy: $\mathbf{A} = \mathbf{B} = 0$), to zgodnie z (83) dostaniemy:

$$\mathbf{y}(t) = \frac{1}{2\pi j} \int_L \mathbf{W}(s) \bar{\mathbf{f}}(s) e^{st} ds = \sum_{\substack{\text{wszystkie} \\ \text{bieguny}}} \text{res}[\mathbf{W}(s) \bar{\mathbf{f}}(s) e^{st}]. \quad (89)$$

¹⁾ Szczegóły tego rozumowania opartego na teorii szeregów Fouriera (por. [24]) zostały tu pominięte.

Zbiór punktów mogących być biegunami macierzy $W(s)\bar{f}(s)$ jest przy dowolnej transformacie $\bar{f}(s)$ sumą dwóch rozłącznych zbiorów: zbioru biegunów s_v macierzy przenoszenia $W(s)$ i zbioru biegunów s_k^0 macierzy $\bar{f}(s)$. Wynika stąd, że rozwiązanie okresowe (86) otrzymuje się jako sumę (nieskończoną) residuów iloczynu $W(s)\bar{f}(s)e^{st}$ tylko w biegunach s_k^0 macierzy $\bar{f}(s)$ (por. [26], [8]).

Zależność (86) określa rozwiązanie okresowe $y_0(t)$ w postaci szeregu trygonometrycznego, który można by otrzymać również np. metodą harmonicznych (por. [26], [23]). W następnym paragrafie podamy już zależność dla $y_0(t)$ w postaci zamkniętej.

Przypadek 2. Istnieją takie bieguny s_v ($v = 1, 2, \dots, N_0 \leq N$) macierzy funkcji przenoszenia $W(s)$, dla których:

$$1 - e^{-s_v T} = 0, \quad (90)$$

tzn., które pokrywają się z odpowiednimi punktami s_k^0 , w których macierz $\bar{f}(s)$ może mieć bieguny.

W tym przypadku, w przeciwieństwie do przypadku poprzedniego, rozwiązanie okresowe może istnieć lub nie istnieć w zależności od macierzy $\bar{f}(s)$.

Jeśli punkt s_v jest biegunem m_v -krotnym macierzy przenoszenia $W(s)^1$, to w zależności od $\bar{f}(s)$ może on być biegunem macierzy $W(s)\bar{f}(s)$ rzędu co najwyżej $m_v + 1$. Wyróżniamy teraz dwie możliwości.

Przypadek 2a. Macierz $\bar{f}(s)$ jest taka, że co najmniej jeden z punktów s_v ($v = 1, 2, \dots, N_0$) jest biegunem macierzy $W(s)\bar{f}(s)$ rzędu wyższego niż w przypadku macierzy $W(s)$. Wówczas punkt s_v o wspomnianej własności jest istotnie biegunem macierzy $\bar{f}(s)^2$; tzn. przy $s_v = s^0$, mamy:

$$c_{k_v} = \operatorname{res}_{s=s_k^0} \bar{f}(s) = \frac{1}{T} \bar{f}_T(s_{k_v}^0) \neq 0. \quad (91)$$

W tym przypadku rozwiązanie okresowe nie istnieje.

Wynika to stąd, że jeśli punkt $s_v \neq 0$ jest biegunem rzędu m_v macierzy $W(s)$, to jest on również biegunem rzędu m_v macierzy $\frac{1}{s} W(s)$ oraz biegunem rzędu $m_v + 1$ macierzy $W(s)\bar{f}(s)$. W wyrażeniu (89) wystąpi więc składnik typu: $t^{m_v} e^{s_v t}$

¹⁾ Rozumiemy przez to, że co najmniej jeden element tej macierzy ma przy $s=s_v$ biegun rzędu m_v , żaden natomiast z elementów nie ma przy $s=s_v$ bieguna rzędu wyższego.

²⁾ Własność ta jest konieczna w przypadku 2a, jak zobaczymy jednak na przykładzie, nie jest ona wystarczająca do tego, aby przypadek ten zachodził.

którego nie będzie w wyrazach pochodzących od macierzy: $\mathbf{W}(s) \mathfrak{A} \mathbf{A}$ lub: $\frac{1}{s} \mathbf{W}(s) \mathfrak{G} \mathbf{B}^1$. Składnik ten wystąpi wobec tego zawsze w rozwiązaniu $\mathbf{y}(t)$, które zatem przy żadnych warunkach początkowych $\mathbf{A}$ i $\mathbf{B}$ nie jest okresowe.

Analogicznie jest w przypadku $s_v = 0$. Wprawdzie teraz punkt ten jest biegunem rzędu $m_v + 1$ macierzy $\frac{1}{s} \mathbf{W}(s)$, jednakże nie jest on biegunem rzędu $m_v + 1$ macierzy $\frac{1}{s} \mathbf{W}(s) \mathfrak{G} \mathbf{B}^2$; wobec czego, podobnie jak poprzednio, w wyrazach pochodzących od macierzy: $\mathbf{W}(s) \mathfrak{A} \mathbf{A}$ oraz $\frac{1}{s} \mathbf{W}(s) \mathfrak{G} \mathbf{B}$ nie wystąpi składnik typu t^{m_v} .

Przypadek 2b. Macierz $\bar{\mathbf{f}}(s)$ jest taka, że żaden z punktów s_v ($v = 1, 2, \dots, N_0$) nie jest biegunem macierzy $\mathbf{W}(s) \bar{\mathbf{f}}(s)$ rzędu wyższego niż w przypadku macierzy $\mathbf{W}(s)$. Przypadek ten zachodzi np. wówczas, gdy żaden z punktów s_v ($v = 1, 2, \dots, N_0$) nie jest biegunem macierzy $\mathbf{W}(s)$, tj. gdy

$$c_{k_v} = \operatorname{res}_{s=s_{k_v}^0} \bar{\mathbf{f}}(s) = \frac{1}{T} \bar{\mathbf{f}}_T(s_{k_v}^0) = 0, \quad (92)$$

gdzie $s_v = s_{k_v}^0$ — przy wszystkich k_v odpowiadających liczbom s_v ($v = 1, 2, \dots, N_0$)³.

W tym przypadku rozwiązań okresowych jest nieskończenie wiele.

Przypadkiem tym bliżej zajmować się nie będziemy. Dowód podanej własności (szkicujemy tu tylko samą jego zasadę) można oprzeć na następującym rozumowaniu. Jeśli wszystkie punkty s_v ($v = 1, 2, \dots, N_0$) są biegunami pojedynczymi macierzy $\mathbf{W}(s)$, to zależność (86) i w tym przypadku określa rozwiązanie okresowe. Nie jest ono jedyne. Można bowiem tak dobrać warunki początkowe $\mathbf{A}$ i $\mathbf{B}$, aby w wyrażeniu

¹⁾ Mogą tu wystąpić tylko składniki typu $t^\lambda e^{s_v t}$ przy $\lambda \leq m_v - 1$.

²⁾ Wystarczy wykazać, że

$$\lim_{s \rightarrow 0} s^{m_v+1} \cdot \frac{1}{s} \mathbf{W}(s) \mathfrak{G} = \lim_{s \rightarrow 0} s^{m_v} \mathbf{W}(s) \mathfrak{G} = 0. \quad (a)$$

Ponieważ punkt $s=0$ jest z założenia biegunem m_v -krotnym macierzy $\mathbf{W}(s)$, wobec tego istnieje granica:

$$\lim_{s \rightarrow 0} s^{m_v} \mathbf{W}(s) \neq 0;$$

mnożąc więc równość (30) stronami przez s^{m_v+1} dostajemy:

$$s^{m_v} \mathbf{W}(s) (\mathfrak{A} s^2 + \mathfrak{B} s + \mathfrak{G}) = s^{m_v+1} \cdot \mathbf{1},$$

a stąd po przejściu do granicy przy $s \rightarrow 0$ otrzymujemy od razu zależność (a).

³⁾ Własność ta jest wystarczająca do tego, aby zachodził przypadek 2b, jak zobaczymy jednak na przykładzie, nie jest ona w tym przypadku konieczna.

niu (89) wyrazy pochodzące od macierzy: $\mathbf{W}(s)\mathfrak{A}$ oraz $\frac{1}{s}\mathbf{W}(s)\mathfrak{B}$ redukowały się tylko do składników odpowiadających biegunom s_ν ($\nu = 1, 2, \dots, N_0$), a więc okresowych lub stałych. Składniki te będą zależne zawsze od pewnej liczby stałych dowolnych i po dodaniu do $y_0(t)$ określonego wzorem (86) stanowią całą rodzinę rozwiązań okresowych.

Podobnie można postępować w przypadku, gdy bieguny s_ν ($\nu = 1, 2, \dots, N_0$) są wielokrotne z tym, że wówczas rozwiązanie (86) nie jest już okresowe.

Zauważamy, że warunek (92) ma prostą interpretację fizyczną. Oznacza on, że w rozwinięciu w szereg Fouriera macierzy $\mathbf{f}(t)$ [por. (18)] nie występują harmoniczne rzędów k .

Przykłady.

1. Rozpatrzmy układ równań:

$$\mathbf{Z}(D)\mathbf{y}(t) = \mathbf{f}(t),$$

gdzie:

$$\mathbf{y}(t) = \begin{bmatrix} y_1(t) \\ y_2(t) \end{bmatrix}; \quad \mathbf{f}(t) = \begin{bmatrix} f_1(t) \\ f_2(t) \end{bmatrix}$$

oraz:

$$\mathbf{Z}(D) = \mathfrak{A}D + \mathfrak{B} = \begin{bmatrix} D + 1, & 1 \\ D, & D \end{bmatrix}$$

przy warunku początkowym:

$$\mathbf{y}(0) = \mathbf{A} = \begin{bmatrix} A_1 \\ A_2 \end{bmatrix}.$$

Transformata rozwiązania ma postać:

$$\bar{\mathbf{y}}(s) = \mathbf{W}(s)\bar{\mathbf{f}}(s) + \mathbf{W}(s)\mathfrak{A},$$

przy czym macierz funkcji przenoszenia jest równa:

$$\mathbf{W}(s) = \begin{bmatrix} \frac{1}{s}, & -\frac{1}{s^2} \\ -\frac{1}{s}, & \frac{s+1}{s^2} \end{bmatrix}.$$

Macierz ta posiada jedyny biegun (podwójny) $s = 0$, zachodzi więc przypadek 2.

Jeśli przyjmiemy:

$$\mathbf{f}(t) = \begin{bmatrix} 0 \\ 1 - \cos \omega t \end{bmatrix},$$

a więc:

$$\bar{f}(s) = \left\| \begin{array}{c} 0 \\ \omega^2 \\ s(s^2 + \omega^2) \end{array} \right\|,$$

wówczas dostajemy:

$$W(s) \bar{f}(s) = \left\| \begin{array}{c} -\frac{\omega^2}{s^2(s^2 + \omega^2)} \\ \frac{(s+1)\omega^2}{s^3(s^2 + \omega^2)} \end{array} \right\|; \quad W(s) \mathcal{A} = \left\| \begin{array}{cc} \frac{A_1}{s} - \frac{A_1 + A_2}{s^2} \\ \frac{A_2}{s} + \frac{A_1 + A_2}{s^2} \end{array} \right\|.$$

Macierz $W(s) \bar{f}(s)$ ma zatem przy $s = 0$ biegun rzędu trzeciego, a więc o wyższej krotności niż biegun macierzy $W(s)$. Rozwiązanie ogólne ma postać:

$$y(t) = \left\| \begin{array}{c} -\frac{t^2}{2} + \frac{1}{\omega^2}(1 - \cos \omega t) + A_1 - (A_1 + A_2)t \\ \frac{t^2}{2} + t - \frac{1}{\omega^2}(1 - \cos \omega t) - \frac{1}{\omega} \sin \omega t + A_2 - (A_1 + A_2)t \end{array} \right\|$$

i przy żadnym doborze stałych A_1 i A_2 macierz $y(t)$ nie jest okresowa. Rozwiązanie okresowe zatem, zgodnie z przypadkiem 2a, nie istnieje.

Jeśli natomiast przyjmiemy:

$$f(t) = \left\| \begin{array}{c} 1 - \cos \omega t \\ 0 \end{array} \right\|; \quad \bar{f}(s) = \left\| \begin{array}{c} \omega^2 \\ s(s^2 + \omega^2) \\ 0 \end{array} \right\|,$$

wówczas dostajemy:

$$W(s) \bar{f}(s) = \left\| \begin{array}{c} \frac{1}{s^2(s^2 + \omega^2)} \\ -\frac{1}{s^2(s^2 + \omega^2)} \end{array} \right\|.$$

Macierz $W(s) \bar{f}(s)$ ma teraz przy $s = 0$ biegun podwójny, tj. tego samego rzędu, co macierz $W(s)$ ¹⁾. Rozwiązanie ogólne ma postać:

¹⁾ Mimo, iż punkt $s = 0$ jest również biegunem macierzy $\bar{f}(s)$. Przykład ten świadczy o tym, że warunek (91) nie jest wystarczający dla przypadku 2a oraz że warunek (92) nie jest konieczny dla przypadku 2b.

$$y(t) = \left\| \begin{array}{l} t - \frac{1}{\omega} \sin \omega t + A_1 - (A_1 + A_2)t \\ -t + \frac{1}{\omega} \sin \omega t + A_2 + (A_1 + A_2)t \end{array} \right\|$$

Rozwiązanie to będzie okresowe, jeśli przyjąć: $A_2 = 1 - A_1$. Dostajemy więc, zgodnie z przypadkiem 2b, nieskończenie wiele rozwiązań okresowych o postaci:

$$y_0(t) = \left\| \begin{array}{l} -\frac{1}{\omega} \sin \omega t + A_1 \\ \frac{1}{\omega} \sin \omega t + 1 - A_1 \end{array} \right\|$$

zależnych od jednej stałej dowolnej A_1 .

2. W podanym w rozdz. 5 (str. 28) przykładzie przy: $\omega \neq \omega_0$ oraz $\omega \neq \frac{\omega_0}{\nu}$ ($\nu = 2, 3, \dots$) mamy przypadek 1, przy: $\omega = \frac{\omega_0}{\nu}$ przypadek 2b i przy: $\omega = \omega_0$ przypadek 2a.

8. WYZNACZANIE ROZWIĄZANIA OKRESOWEGO

Zajmiemy się teraz zadaniem podstawowym — wyznaczeniem rozwiązania okresowego $y_0(t)$ w postaci zamkniętej. Zakładamy przy tym, że zachodzi przypadek 1 (rozdz. 7), tzn. żaden z biegunów s_ν ($\nu = 1, 2, \dots, N$) macierzy przenoszenia $\mathbf{W}(s)$ nie pokrywa się z żadnym z punktów s_k^0 określonych zależnością (76), w których ogólnie biorąc, macierz $\bar{\mathbf{f}}(s)$ ma bieguny. Zachodzi więc nierówność (85) i, jak już zaznaczyliśmy w rozdz. 7, istnieje jedyne rozwiązanie okresowe $y_0(t)$, które można wyrazić w postaci szeregu (86). Naszym obecnym celem jest więc wyznaczenie tego rozwiązania w postaci zamkniętej.

Jeśli przez $y(t)$ oznaczymy rozwiązanie przy warunkach początkowych zerowych, tj. przy: $\mathbf{A} = \mathbf{B} = \mathbf{0}$, to można je przedstawić w postaci (89).

Porównując (89) z zależnością (86) oraz uwzględniając, jakie punkty mogą być biegunami macierzy $\mathbf{W}(s)$ $\bar{\mathbf{f}}(s)$ (rozdz. 7), możemy napisać rozwiązanie okresowe (86) w postaci ($t \geq 0$):

$$\begin{aligned} y_0(t) = \sum_k \operatorname{res}_{s=s_k^0} [\mathbf{W}(s) \bar{\mathbf{f}}(s) e^{st}] &= \sum_{\substack{\text{wszystkie} \\ \text{bieguny}}} \operatorname{res} [\mathbf{W}(s) \bar{\mathbf{f}}(s) e^{st}] + \\ &- \sum_{\nu=1}^N \operatorname{res}_{s=s_\nu} [\mathbf{W}(s) \bar{\mathbf{f}}(s) e^{st}] = \frac{1}{2\pi j} \int_L \mathbf{W}(s) \frac{\bar{\mathbf{f}}_T(s)}{1 - e^{-sT}} e^{st} ds + \\ &- \sum_{\nu=1}^N \operatorname{res}_{s=s_\nu} \left[\mathbf{W}(s) \frac{\bar{\mathbf{f}}_T(s)}{1 - e^{-sT}} e^{st} \right]. \end{aligned}$$

Ponieważ w przedziale $0 \leq t \leq T$ rozwiązanie $y(t)$ równania (12) przy prawej stronie $f(t)$ jest identyczne z rozwiązaniem tego równania przy prawej stronie równej:

$$f_T(t) = \begin{cases} f(t); & 0 \leq t \leq T; \\ 0; & T < t, \end{cases} \quad (93)$$

wobec tego dla przedziału $0 \leq t \leq T$ mamy:

$$y_0(t) = \frac{1}{2\pi j} \int_L W(s) \bar{f}_T(s) e^{st} ds - \sum_{\nu=1}^N \operatorname{res}_{s=s_\nu} \left[W(s) \frac{\bar{f}_T(s)}{1 - e^{-sT}} e^{st} \right]. \quad (94)$$

Zakładając teraz, że:

$$W(s) = r_0 + W_0(s), \quad (95)$$

przy czym macierz $W_0(s)$ spełnia warunek (44)¹⁾, otrzymujemy (94) w postaci ($0 \leq t \leq T$):

$$y_0(t) = r_0 f(t) + \frac{1}{2\pi j} \int_L W_0(s) \bar{f}_T(s) e^{st} ds - \sum_{\nu=1}^N \operatorname{res}_{s=s_\nu} \left[W_0(s) \frac{\bar{f}_T(s)}{1 - e^{-sT}} e^{st} \right]. \quad (96)$$

Oznaczając:

$$h_0(t) = \mathcal{L}^{-1}[W_0(s)] = \sum_{\nu=1}^N \operatorname{res}_{s=s_\nu} [W_0(s) e^{st}] \quad (97)$$

otrzymujemy ($t \geq 0$):

$$\mathcal{L}^{-1}[W_0(s) \bar{f}_T(s)] = \frac{1}{2\pi j} \int_L W_0(s) \bar{f}_T(s) e^{st} ds = h_0(t) * f_T(t). \quad (98)$$

Wyrażenie to możemy przy $0 \leq t \leq T$ przekształcić następująco:

$$\begin{aligned} h_0(t) * f_T(t) &= \int_0^t h_0(t-\tau) f(\tau) d\tau = \int_0^t \left\{ \sum_{\nu=1}^N \operatorname{res}_{s=s_\nu} [W_0(s) e^{s(t-\tau)}] \right\} f(\tau) d\tau = \\ &= \sum_{\nu=1}^N \int_0^t \left\{ \operatorname{res}_{s=s_\nu} [W_0(s) e^{s(t-\tau)}] \right\} f(\tau) d\tau = \sum_{\nu=1}^N \operatorname{res}_{s=s_\nu} \left[W_0(s) e^{st} \int_0^t e^{-s\tau} f(\tau) d\tau \right]^2 \end{aligned} \quad (99)$$

¹⁾ Zakładamy więc, że zachodzi przypadek 2 dyskutowany w rozdz. 4. Wzór (95) obejmuje jednak także i przypadek 1 (rozdz. 4), gdyż wówczas po prostu: $r_0 = 0$.

²⁾ Poprawność czynionych tu przekształceń wynika stąd, że w wyrażeniach typu:

$$\int_0^t \lim_{s \rightarrow s_\nu} \frac{d^{k-1}}{ds^{k-1}} [W(s) (s-s_\nu)^k e^{s(t-\tau)} f(\tau)] d\tau,$$

jakie spotykamy przy dokonywaniu tych przekształceń, funkcja w nawiasie kwa-

Wstawiając (99) do (96) otrzymujemy poszukiwany wzór dla rozwiązania okresowego $y_0(t)$ w przedziale: $0 \leq t \leq T$:

$$y_0(t) = r_0 f(t) + \sum_{\nu=1}^N \operatorname{res}_{s=s_\nu} \left\{ W_0(s) e^{st} \left[\int_0^t e^{-s\tau} f(\tau) d\tau - \frac{\int_0^T e^{-s\tau} f(\tau) d\tau}{1 - e^{-sT}} \right] \right\} \quad (100)$$

lub po prostych przekształceniach ($0 \leq t \leq T$):

$$y_0(t) = r_0 f(t) + \sum_{\nu=1}^N \operatorname{res}_{s=s_\nu} \frac{W_0(s) e^{st} \int_t^{t+T} e^{-s\tau} f(\tau) d\tau}{1 - e^{-sT}}. \quad (101)$$

Prosty przypadek szczególny wzoru (100) lub (101) dostaniemy, jeśli założymy, że wszystkie bieguny s_ν macierzy $W(s)$ są pojedyncze. Wówczas

$$\begin{aligned} \operatorname{res}_{s=s_\nu} \left\{ W_0(s) e^{st} \left[\int_0^t e^{-s\tau} f(\tau) d\tau - \frac{\int_0^T e^{-s\tau} f(\tau) d\tau}{1 - e^{-sT}} \right] \right\} = \\ = \operatorname{res}_{s=s_\nu} W_0(s) \cdot e^{s_\nu t} \left[\int_0^t e^{-s_\nu \tau} f(\tau) d\tau - \frac{\int_0^T e^{-s_\nu \tau} f(\tau) d\tau}{1 - e^{-s_\nu T}} \right]. \end{aligned}$$

Oznaczając więc:

$$d_\nu = \operatorname{res}_{s=s_\nu} W_0(s) = \lim_{s \rightarrow s_\nu} W_0(s) (s - s_\nu) \quad (102)$$

dostajemy dla przypadku biegunów pojedynczych przy $0 \leq t \leq T$:

$$y_0(t) = r_0 f(t) + \sum_{\nu=1}^N d_\nu e^{s_\nu t} \left[\int_0^t e^{-s_\nu \tau} f(\tau) d\tau - \frac{\int_0^T e^{-s_\nu \tau} f(\tau) d\tau}{1 - e^{-s_\nu T}} \right] \quad (103)$$

lub ($0 \leq t \leq T$):

$$y_0(t) = r_0 f(t) + \sum_{\nu=1}^N \frac{d_\nu e^{s_\nu t} \int_t^{t+T} e^{-s_\nu \tau} f(\tau) d\tau}{1 - e^{-s_\nu T}}. \quad (104)$$

dratowym po skróceniu czynnika $(s - s_\nu)^k$ z takim samym czynnikiem w mianowniku $W(s)$ jest funkcją meromorficzną względem s w punkcie $s = s_\nu$ i przedziałami ciągłą względem t w przedziale $[0, t]$, wobec czego można zmienić kolejność całkowania i brania granicy oraz całkowania i różniczkowania.

Podane wzory stanowią rozwiązanie ogólne postawionego zagadnienia przy przyjętych założeniach.

9. PRZYPADEK SIECI TRAKTOWANEJ JAKO CZWÓRNIK

Dotychczas traktowaliśmy układ równań (12) i sieć elektryczną opisywaną nim jako całość. W sieci działała dowolna liczba źródeł energii i rozwiązaniem była macierz kolumnowa $y(t)$ o r składowych (elementach). Często w praktyce zagadnienie nie jest stawiane tak ogólnie. Jeśli ograniczymy się do rozpatrzenia działania tylko jednego źródła oraz interesuje nas nie całe rozwiązanie sieci w postaci macierzy, np. prądów, lecz tylko jedna jego składowa, np. prąd w określonej gałęzi, zagadnienie znacznie się upraszcza, gdyż całą sieć można traktować jako czwórnik liniowy bierny¹⁾. Zajmiemy się teraz po krótkce takim właśnie przypadkiem szczególnym.

Założmy, iż macierze $f(t)$ i $\bar{f}(s)$ mają postać:

$$f(t) = \begin{pmatrix} f(t) \\ 0 \\ \dots\dots\dots \\ 0 \end{pmatrix}; \quad \bar{f}(s) = \begin{pmatrix} \bar{f}(s) \\ 0 \\ \dots\dots\dots \\ 0 \end{pmatrix}, \quad (105)$$

gdzie $f(t)$ jest funkcją okresową o okresie T i przedziałami regularną, oraz rozważmy rozwiązanie $y(t)$ równania (12) przy warunkach: $A = B = 0$. Zgodnie z (31) transformata tego rozwiązania jest równa:

$$\bar{y}(s) = \begin{pmatrix} \bar{y}^{(1)}(s) \\ \bar{y}^{(2)}(s) \\ \dots\dots\dots \\ \bar{y}^{(r)}(s) \end{pmatrix} = W(s) \bar{f}(s) = \begin{pmatrix} W_{11}(s) \\ W_{12}(s) \\ \dots\dots\dots \\ W_{1r}(s) \end{pmatrix} \bar{f}(s) = P(s) \bar{f}(s), \quad (106)$$

gdzie:

$$P(s) = \begin{pmatrix} P_1(s) \\ P_2(s) \\ \dots\dots\dots \\ P_r(s) \end{pmatrix}; \quad \begin{matrix} P_m(s) = W_{1m}(s) \\ \\ \\ m = 1, 2, \dots, r \end{matrix}$$

¹⁾ Większość prac dotyczących rozwiązań okresowych (np. [8], [10], [26] i in.) odnosi się właśnie do tego przypadku.

Jeśli przyjmiemy, że chodzi nam o wyznaczenie tylko jednej dowolnie wybranej składowej rozwiązania $y(t)$, np. $y^{(m)}(t)$, i pozostałe składowe nas nie interesują, dostaniemy:

$$\bar{y}^{(m)}(s) = P_m(s) \bar{f}(s)$$

lub po opuszczeniu zbędnego już indeksu:

$$\bar{y}(s) = P(s) \bar{f}(s). \quad (107)$$

Do przypadku tego można zastosować natychmiast wyniki poprzednich naszych rozważań. Jeśli np. wszystkie bieguny s_ν ($\nu = 1, 2, \dots, N$) macierzy $P(s)$ spełniają warunek

$$1 - e^{-s_\nu T} \neq 0, \quad (108)$$

to w rozważanym układzie istnieje jedyne rozwiązanie okresowe $y_0(t)$. Jego m -ta składowa $y_0^{(m)}(t)$, tj. właściwe rozwiązanie, o które nam chodzi, możemy zatem wyrazić zgodnie ze wzorem (100) przy: $0 \leq t \leq T$ w postaci (opuszczamy zbędny indeks):

$$y_0(t) = \varrho_0 f(t) + \sum_{\nu=1}^N \operatorname{res}_{s=s_\nu} \left\{ P_0(s) e^{st} \left[\int_0^t e^{-s_\nu \tau} f(\tau) d\tau - \frac{\int_0^T e^{-s_\nu \tau} f(\tau) d\tau}{1 - e^{-s_\nu T}} \right] \right\} \quad (109)$$

lub w postaci ($0 \leq t \leq T$):

$$y_0(t) = \varrho_0 f(t) + \sum_{\nu=1}^N \operatorname{res}_{s=s_\nu} \frac{P_0(s) e^{st} \int_t^{t+T} e^{-s_\nu \tau} f(\tau) d\tau}{1 - e^{-s_\nu T}}, \quad (110)$$

gdzie:

$$P(s) = \varrho_0 + P_0(s) \quad (111)$$

i przy tym:

$$\lim_{s \rightarrow \infty} P_0(s) = 0 \quad (112)$$

Kilka słów wyjaśnienia wymaga jeszcze sprawa okresowości rozwiązania $y_0(t)$ określonego wzorem (109) lub (110). Spełnienie warunku (108) gwarantuje istnienie rozwiązania okresowego $y_0(t)$, jednakże nie wszystkie jego składowe muszą być funkcjami okresowymi (por. określenie rozwiązania okresowego podane na str. 141). Rozwiązanie (109) może więc być

przy spełnieniu warunku (108) albo okresowe (o okresie T lub $\frac{T}{k}$, gdzie k — liczba naturalna), albo stałe, a więc jest okresowe w sensie rozszerzonym.

Zajmiemy się teraz sprawą warunków dodatkowych, przy których rozwiązanie (109) jest istotnie okresowe. Zauważymy przede wszystkim, że gdy $\varrho_0 \neq 0$, to funkcja $y_0(t)$ jest istotnie okresowa. Pozostaje więc rozpatrzyć bliżej przypadek $\varrho_0 = 0$.

Jeśli funkcję $f(t)$ rozwiniemy w szereg Fouriera o postaci (16), to zgodnie ze wzorem (86) rozwiązanie $y_0(t)$ można przedstawić jako:

$$y_0(t) = \sum_{k=-\infty}^{+\infty} P(s_k^0) c_k e^{s_k^0 t},$$

gdzie s_k^0 są punktami określonymi wzorem (76). Rozwiązanie to jest stałe wtedy i tylko wtedy, gdy

$$\sum_{k=-\infty}^{+\infty} P(s_k^0) c_k e^{s_k^0 t} = \text{const},$$

tzn. gdy dla wszystkich $s_k^0 \neq 0$ jest:

$$P(s_k^0) c_k = 0. \quad (113)$$

Jeśli przez s_k' oznaczymy wszystkie te punkty $s_k^0 \neq 0$, przy których $c_k \neq 0$ tzn. te spośród wszystkich punktów $s_k^0 \neq 0$, które są istotnie biegunami transformaty $\bar{f}(s)$, to warunek (113) jest równoważny warunkowi, aby:

$$P(s_k') = 0 \quad (114)$$

dla wszystkich punktów s_k' .

Warunek (114) oznacza, że funkcja przenoszenia $P(s)$ ma zera wszędzie tam, gdzie transformata $\bar{f}(s)$ ma bieguny. Ponieważ zbiór biegunów $\bar{f}(s)$ określa widmo harmoniczych funkcji $f(t)$, wobec tego warunek (114) może być spełniony tylko wtedy, gdy funkcja $f(t)$ ma skończoną liczbę harmoniczych.

Wykazaliśmy zatem, że rozwiązanie $y_0(t)$ określone wzorem (109) jest stałe wtedy i tylko wtedy, gdy stopień licznika funkcji przenoszenia $P(s)$ jest niższy od stopnia jej mianownika (tzn. $\varrho_0 = 0$) oraz jednocześnie gdy funkcja przenoszenia $P(s)$ ma zera we wszystkich różnych od zera biegunach transformaty $\bar{f}(s)$.

W każdym innym przypadku rozwiązanie $y_0(t)$ jest okresowe.

Przykład

Rozpatrzmy obwód przedstawiony na rys. 4. Wybierając jako wielkość niewiadomą napięcie $u(t)$ na oporności R dostaniemy przy warunkach początkowych zerowych:

$$\bar{u}(s) = P(s) \bar{e}(s),$$

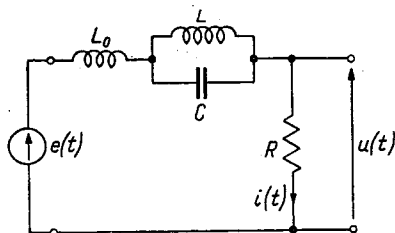
gdzie:

$$P(s) = a_2 \frac{s^2 + \omega_0^2}{s^3 + a_2 s^2 + a_1 s + a_0},$$

$$\bar{e}(s) = \mathcal{L}[e(t)]$$

oraz:

$$\omega_0^2 = \frac{1}{LC}; \quad a_0 = \frac{R}{L_0 LC}; \quad a_1 = \frac{L + L_0}{L_0 LC}; \quad a_2 = \frac{R}{L_0}.$$



Rys. 4. Czwórnik RLC zasilany napięciowo

Wszystkie pierwiastki mianownika funkcji $P(s)$ leżą w lewej półpłaszczyźnie.

Przyjmijmy:

$$e(t) = E \sin \omega t; \quad \bar{e}(s) = E \frac{\omega}{s^2 + \omega^2}.$$

Wówczas:

$$\bar{u}(s) = k \frac{s^2 + \omega_0^2}{(s^2 + \omega^2)(s^3 + a_2 s^2 + a_1 s + a_0)},$$

gdzie: $k = Ea_2 \omega$.

Rozwiązanie okresowe $u_o(t)$ w sensie rozszerzonym oczywiście istnieje (i jest ono rozwiązaniem ustalonym).

Przy: $\omega \neq \omega_0$ jest ono istotnie rozwiązaniem okresowym i wynosi:

$$u_o(t) = 2 \operatorname{Re} \left[\frac{k (s^2 + \omega_0^2) e^{st}}{(s + j\omega)(s^3 + a_2 s^2 + a_1 s + a_0)} \right]_{s = j\omega},$$

a przy $\omega = \omega_0$ jest ono stałe i równe:

$$u_o(t) \equiv 0.$$

W przypadku tym mamy oczywiście do czynienia z działaniem filtrującym obwodu równoległego LC.

DODATEK. OSZACOWANIE TRANSFORMATY FUNKCJI OKRESOWEJ NA CIĄGU OKRĘGÓW

Udowodnimy teraz następujące twierdzenie, na które powoływaliśmy się w tekście.

Można tak dobrać ciąg okręgów $\{C_n\}$: $|s| = R_n$, $R_n \rightarrow \infty$, aby transformata

$$\bar{f}(s) = \frac{\bar{f}_T(s)}{1 - e^{-st}} = \sum_{k=-\infty}^{+\infty} \frac{c_k}{s - s_k^0} \quad (D1)$$

dowolnej funkcji okresowej o okresie T przedziałami regularnej była na tych okręgach ograniczona, tj.

$$|\bar{f}(s)|_{C_n} \leq K. \quad (D2)$$

Weźmy ciąg okręgów $\{C_n\}$, $|s| = R_n$, gdzie:

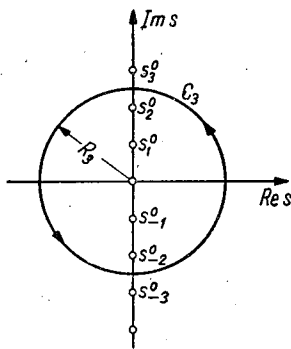
$$R_n = \left(n + \frac{1}{2}\right) \frac{2\pi}{T} = \left(n + \frac{1}{2}\right) \omega$$

$n = 1, 2, \dots$. W kole: $|s| < R_n$ znajduje się zatem (rys. 5) $2n + 1$ biegunów s_k^0 — wzór (76) — o indeksach: $k = -n, \dots, 0, \dots, n$.

Zgodnie z (D1) mamy:

$$\begin{aligned} |\bar{f}(s)|_{C_n} &= \left| \sum_{k=-\infty}^{+\infty} \frac{c_k}{s - s_k^0} \right|_{C_n} \leq \sum_{k=-\infty}^{+\infty} \frac{|c_k|}{|s - s_k^0|_{C_n}} = \\ &= \frac{|c_0|}{R_n} + \sum_{k=1}^{\infty} |c_k| \left[\frac{1}{|s - s_k^0|} + \frac{1}{|s - s_{-k}^0|} \right]_{C_n}, \end{aligned}$$

bowiem: $c_{-k} = \tilde{c}_k$ a więc $|c_{-k}| = |c_k|$.



Rys. 5. Położenie okręgu C_n względem biegunów transformaty funkcji okresowej

Słuszne są następujące oszacowania wynikające z położenia punktów s_k^0 względem okręgu C_n (por. rys. 5):

$$\begin{aligned} |s - s_k^0|_{C_n} &\geq \left| j \left(n + \frac{1}{2}\right) \omega - j k \omega \right| = \left| n + \frac{1}{2} - k \right| \omega \\ |s - s_{-k}^0|_{C_n} &\geq \left| -j \left(n + \frac{1}{2}\right) \omega + j k \omega \right| = \left| n + \frac{1}{2} - k \right| \omega, \end{aligned}$$

czyli:

$$\left[\frac{1}{|s - s_k^0|} + \frac{1}{|s - s_{-k}^0|} \right]_{C_n} \leq \frac{2}{\omega \left| n + \frac{1}{2} - k \right|},$$

a zatem także:

$$|\bar{f}(s)|_{C_n} \leq \frac{|c_0|}{R_n} + \frac{2}{\omega} \sum_{k=1}^{\infty} \frac{c_k}{\left| n + \frac{1}{2} - k \right|}.$$

Współczynniki szeregu Fouriera funkcji okresowej przedziałami regularnej spełniają nierówność:¹⁾

$$k |c_k| \leq M$$

przy dowolnym k , mamy więc:

$$\begin{aligned} |\bar{f}(s)|_{C_n} &\leq \frac{|c_0|}{R_n} + \frac{2M}{\omega} \sum_{k=1}^{\infty} \frac{1}{k \left| n + \frac{1}{2} - k \right|} = \\ &= \frac{|c_0|}{R_n} + \frac{2M}{\omega} \sum_{k=1}^n \frac{1}{k \left(n + \frac{1}{2} - k \right)} + \frac{2M}{\omega} \sum_{k=n+1}^{\infty} \frac{1}{k \left(k - n - \frac{1}{2} \right)}. \end{aligned}$$

Rozpatrzmy poszczególne składniki tej zależności. Ponieważ przy $1 \leq k \leq n$ wobec nierówności $(n-k)(k-1/2) \geq 0$ mamy także:

$$k \left(n + \frac{1}{2} - k \right) \geq \frac{n}{2}; \quad \frac{1}{k \left(n + \frac{1}{2} - k \right)} \leq \frac{2}{n},$$

a więc:

$$\sum_{k=1}^n \frac{1}{k \left(n + \frac{1}{2} - k \right)} \leq n \cdot \frac{2}{n} = 2 \quad (D4)$$

przy dowolnym n .

Druga suma występująca po prawej stronie wzoru (D3) po zamianie indeksów sumowania $k = n + \nu$ przybiera postać:

$$\sum_{k=n+1}^{\infty} \frac{1}{k \left(k - n - \frac{1}{2} \right)} = \sum_{\nu=1}^{\infty} \frac{1}{(n+\nu) \left(\nu - \frac{1}{2} \right)} = \sum_{\nu=1}^{\infty} \frac{1}{\nu^2 + n\nu - \frac{1}{2}(n+\nu)}.$$

Ponieważ przy $n \geq 1$ oraz $\nu \geq 1$ mamy:

$$n\nu - \frac{1}{2}(n+\nu) = \left(n - \frac{1}{2} \right) \nu - \frac{1}{2}n \geq n - \frac{1}{2} - \frac{1}{2}n = \frac{1}{2}(n-1) \geq 0,$$

¹⁾ Oszacowanie to wynika z prostego uogólnienia znanego z teorii szeregów Fouriera twierdzenia o rzędzie małości współczynników Fouriera funkcji różniczkowalnej — por. np. [24], str. 133.

a więc:

$$\frac{1}{v^2 + n v - \frac{1}{2}(n + v)} \leq \frac{1}{v^2}.$$

Otrzymujemy zatem:

$$\sum_{k=n+1}^{\infty} \frac{1}{k \left(k - n - \frac{1}{2} \right)} = \sum_{v=1}^{\infty} \frac{1}{v^2 + n v - \frac{1}{2}(n + v)} \leq \sum_{v=1}^{\infty} \frac{1}{v^2} = \frac{\pi^2}{6}. \quad (D5)$$

Uwzględniając oszacowania (D4) i (D5) dostajemy z nierówności (D3):

$$|\bar{f}(s)|_{C_n} \leq \frac{|c_0|}{R_n} + \frac{4M}{\omega} + \frac{2M}{\omega} \cdot \frac{\pi^2}{6}. \quad (D6)$$

Ponieważ przy $n \rightarrow \infty$ także $R_n \rightarrow \infty$, mamy więc:

$$\frac{|c_0|}{R_n} \rightarrow 0,$$

wobec czego z nierówności (D6) wynika wprost, iż transformata $\bar{f}(s)$ jest na ciągu okręgów $\{C_n\}$ ograniczona, a więc nierówność (D2).

*Instytut Podstawowych Problemów Techniki PAN
Zakład Teorii Łączności*

WYKAZ LITERATURY

1. Carnahan C. W.: The Steady State Response of a Network to a Periodic Driving Force of Arbitrary Shape and Applications to Television Circuits. — Proc. I. R. E., 23 (1935), nr 11.
2. Carslaw H. S., Jaeger J. C.: Operational Methods in Applied Mathematics. Oxford University Press, 1949 (II wyd).
3. Chin P. T.: Circuit Response to Non-Sinusoidal Wave Forms. — Electronics, October 1944.
4. Doetsch G.: Das Anfangswertproblem für Systeme linearer Differentialgleichungen unter unzulässigen Anfangsbedingungen. — Ann. Mat. pura appl., 39 (1955), nr 4.
5. Doetsch G.: Handbuch der Laplace-Transformation. Band 1. Birkhäuser Verlag 1950.
6. Doetsch G.: Handbuch der Laplace-Transformation. Band 2. Birkhäuser Verlag 1955.
7. Ghizzetti A.: Sull'uso della trasformazione di Laplace nello studio dei circuiti elettrici. — Rend. Circ. mat. Palermo, 61 (1937).
8. Gonorowski J. S. Wozdziejstwiye složnych pieriodiczeskich elektrodwizuszczich sił na liniejnyje sistemy. — Radiotekhnika, 8 (1953), nr 1.
9. Hamlin E. W.: The Response of a Network to an Arbitrary Periodic Driving Force. — Journ. Frankl. Inst., 233 (1942), nr 3.
10. Kaczorek T.: Obliczanie stanów nieustalonych w obwodach liniowych zasilanych periodycznymi siłami elektromotorycznymi. — Arch. Elektrotechniki, VI (1957), nr 4.
11. Kontorowicz M. J.: Rachunek operatorowy i stany nieustalone w obwodach elektrycznych. PWT, Warszawa 1956.

12. Kožešnik J.: Periodický stav v systému o stálých parametrech. — Rozprawy ČSAV, 65 (1955), nr TV 5.
13. Ławrentiew M. A., Szabat B. W.: Metody teorii funkcji kompleksnowo piereminnowo. Moskwa 1958.
14. Łurie A. I.: O pieriodiczeskom reszenii sistemy liniejnych urawnienij s postojannymi koefficientami. Prikl. mat. i mech., XII (1948).
15. Łurie A. I.: Opieracjonnoje isczislenije i priłożenija jewo k zadaczam miechaniki. Moskwa 1950.
16. Małkin I. G.: K teorii kolebanij kwaziliniyjnych sistem so mnogimi stiepieniami swobody. Prikl. mat. i mech., XIV (1950), nr 4.
17. Małkin I. G.: Niekotoryje zadacz i teorii nieliniejnych kolebanij. Moskwa 1956.
18. Marchand N.: The Response Electrical Networks to Nonsinusoidal Periodic Waves. — Proc. I.R.E. 29 (1941), nr 6.
19. Mohr E.: Integration von gewöhnlichen Differentialgleichungen mit konstanten Koeffizienten mittels Operatorenrechnung. — Math. Nachr., 10 (1953), nr 1—2.
20. Osowski J.: Stany nieustalone w układach pobudzanych ciągiem impulsów. — Arch. Elektrotechniki, II (1953), nr 1—2.
21. Pipes L. A.: Matrix Generalization of Heaviside's Expansion Theorem. — Journ. Frankl. Inst., 230 (1940), nr 4.
22. Pol B. van der, Bremmer H.: Operational Calculus Based on the Two-Sided Laplace Integral. Cambridge University Press, 1950.
23. Stiepanow W. W.: Równania różniczkowe. PWN Warszawa 1956.
24. Tołstow D. L.: Szeregi Fouriera. PWN, Warszawa 1954.
25. Waidelich D. L.: Steady State Currents of Electrical Networks. — J. Appl. Ph., 13, November 1942.
26. Waidelich D. L.: The Steady-State Response of Circuits. — Communications, 23 (1943), nr 4.
27. Waidelich D. L.: Steady-State Testing with Saw-Tooth Waves. — Proc. I.R.E., 32 (1944), nr 6.
28. Waidelich D. L.: The Steady-State Operational Calculus. — Proc. I.R.E., 34 (1946), nr 2.
29. Węgrzyn S.: Stany nieustalone w układach pobudzanych periodycznie. — Arch. Elektrotechniki III (1954), nr 4.
30. Woronow R. A.: Rasczot pieriodiczeskich tokow i napriaženij pri niesinusoidalnoj formie e.d.s. — Elektrizestwo, 1956, nr 8.
31. Załmanzon W. B.: Formuła razłożenija dla pieriodiczeskich funkcij — Radiotekhnika, 6 (1951), nr 6.

Е. ОСЕВСКИ

ЛИНЕЙНЫЕ ЭЛЕКТРИЧЕСКИЕ СХЕМЫ С ПЕРИОДИЧЕСКИМ ВОЗБУЖДЕНИЕМ

Резюме

В работе обсуждаются систематическим образом вопросы существования, а также определение в замкнутом виде периодических решений системы диф-

ференциально-интегральных уравнений в виде (6) или (12), при чём правые стороны этих уравнений являются периодическими функциями, регулярными в интервалах. В особенности, уравнениями (6) или (12) описывается любая линейная электрическая цепь построена из сосредоточенных элементов $RLCM$, возбуждаемая периодически. Общее операторное решение матричного уравнения (12) определяется зависимостью (31), где $\mathbf{W}(s)$ обозначает матрицу переходных функций определяемую формулой (27), $\mathbf{A}$ и $\mathbf{B}$ являются матрицами соответственных начальных значений. При выполнении обратного преобразования различаются два случая: 1) при условии $|\alpha| = \det \mathfrak{A} \neq 0$ получается решение $y(t)$ в виде (40), а также 2) при условии $|\alpha| = \det \mathfrak{A} = 0$ и добавочных предположениях (43) и (44) получается $y(t)$ в виде (53).

В главе 5 приведены определения периодического решения $y_0(t)$ и его частного случая — стационарного решения [формулы (61), (62)], доказана одно-значности от взаимного расположения полюсов s_v матрицы переходной функции условия его существования.

В главе 7 приведены условия существования периодического решения в зависимости от взаимного расположения полюсов s_v матрицы переходных функций $\mathbf{W}(s)$ и полюсов s_k^0 матрицы $\bar{\mathbf{f}}(s)$. Эти условия выражены в виде следующих теорем:

1) Если никакой из полюсов s_v матрицы $\mathbf{W}(s)$ не покрывается с какой-либо из точек s_k^0 определённых зависимостью (76), то есть происходит неравенство (85), то при данной матрице $\bar{\mathbf{f}}(s)$ существует всегда единственное периодическое решение $y_0(t)$, которое можно представить в виде (86).

2) Если существует по крайней мере один такой полюс s_v матрицы $\mathbf{W}(s)$, для которого происходит неравенство (90), то периодическое решение может существовать, или же не существовать, в зависимости от матрицы $\bar{\mathbf{f}}(s)$.

В частности доказано в дальнейшем:

2а) Если матрица $\bar{\mathbf{f}}(s)$ является такой, что по крайней мере одна из точек s_v есть полюсом матрицы $\mathbf{W}(s) \cdot \bar{\mathbf{f}}(s)$ высшего порядка, чем в случае матрицы $\mathbf{W}(s)$, то периодическое решение не существует.

2б) Если матрица $\bar{\mathbf{f}}(s)$ является такой, что никакая из точек s_v не есть полюсом матрицы $\mathbf{W}(s) \cdot \bar{\mathbf{f}}(s)$ высшего порядка, чем в случае матрицы $\mathbf{W}(s)$, то периодических решений бесконечно много.

В главе 8 на основании формулы Риманна-Меллина выведена формула (100), определяющая периодическое решение $y_0(t)$ в замкнутом виде в пределах $0 \leq t \leq T$.

В главе 9 использованы полученные результаты для решения частного случая цепи в виде четырёхполюсника, а также приведены условия (стр. 160), при которых периодическое решение вырождается, в случае четырёхполюсника, в постоянное решение.

В Приложении приведено доказательство следующей теоремы:

Можно так подобрать последовательность окружностей (C_n) ; $|s| = R_n$, $R_n \rightarrow \infty$, чтобы трансформата любой периодической функции, регулярной в интервалах, была на этих окружностях ограниченной, то есть чтобы было соблюдено неравенство (D₂).

J. OSIOWSKI

LINEAR ELECTRIC NETWORKS EXCITED PERIODICALLY

Summary

The present paper discusses in a systematic way the problems of existence and determination in a closed form of periodic solutions for a set of integro-differential equations with the form (6) or (12), the right-hand sides of these equations being periodic functions regular in intervals. In particular by equations (6) or (12) it is possible to describe any linear electric network made up of lumped RLCM elements and excited periodically. The general operational solution of the matrix equation (12) is determined by the relation (31), where $\mathbf{W}(s)$ denotes the matrix of the transfer function determined by formula (27), and $\mathbf{A}$ and $\mathbf{B}$ are matrices of the corresponding initial values. In performing the inverse transformation, two cases are distinguished: 1) under the condition $|\mathbf{A}| = \det \mathbf{A} \neq 0$ we obtain the solution $\mathbf{y}(t)$ in the form of (40), and 2) under the condition $|\mathbf{A}| = \det \mathbf{A} = 0$ and additional assumptions (43) and (44) we obtain the solution $\mathbf{y}(t)$ in the form of (53).

In Chapter 5, there is a definition of the periodic solution $\mathbf{y}_0(t)$ and its particular case, i.e. the steady state solution (formulae 61 and 62); moreover, there is a proof for the uniqueness of the steady state solution and determination of the necessary and sufficient conditions of its existence.

Chapter 7 deals with the conditions of the existence of the periodic solution, depending on the mutual position of the poles s_p , the matrix of the transfer function $\mathbf{W}(s)$ and the poles s_k^0 of the matrix $\bar{\mathbf{f}}(s)$. These conditions are formulated in the form of the following theorems:

1) If no one of the poles s_p of the matrix $\mathbf{W}(s)$ covers any of the points s_k^0 determined by the relation (76), i.e. there holds the inequality (85), then with a given matrix $\bar{\mathbf{f}}(s)$ there always exists only one periodic solution $\mathbf{y}_0(t)$ which may be represented in the form of (86).

2) If there exists at least one such a pole s^* of the matrix $\mathbf{W}(s)$, for which there holds the equality (90), then the periodic solution may exist or not depending on the matrix $\bar{\mathbf{f}}(s)$.

Other proofs are given too and in particular the following ones:

2a) If the matrix $\bar{\mathbf{f}}(s)$ is such that at least one of the points s_p is a pole of the matrix $\mathbf{W}(s) \cdot \bar{\mathbf{f}}(s)$ of a higher order than in the case of the matrix $\mathbf{W}(s)$, then there is no periodic solution.

2b) If the matrix $\bar{\mathbf{f}}(s)$ is such that there are no points s_p being a pole of the matrix $\mathbf{W}(s) \cdot \bar{\mathbf{f}}(s)$ of a higher order than in the case of the matrix $\mathbf{W}(s)$, then there is an infinite number of periodic solutions.

In Chapter 8, on the basis of the Riemann-Mellin formula, there is presented formula (100) determining the periodic solution $\mathbf{y}_0(t)$ in a closed form in the interval $0 \leq t \leq T$.

In Chapter 9, the obtained results are applied to a particular case of a network treated as a four-pole, and conditions are given (p. 160) under which the periodic solution reduces in the case of a four-pole to a Constant state solution.

In the Annex, there is the following theorem:

It is possible to choose the sequence of circles $\{C_n\} : |s| = R_n, R_n \rightarrow \infty$, so that the transform of an arbitrary periodic function regular in intervals be limited to these circles, that is to make the inequality (D2) hold.

STANISŁAW BELLERT

Rachunek operatorowy w przestrzeniach liniowych

Rękopis dostarczono 11.12.1959

W pracy przedstawiona jest metoda stanowiąca uogólnienie klasycznych metod rachunku operatorowego. Podstawy tej metody zostały opublikowane przez autora w roku 1957 w pracy [1], a artykuł niniejszy stanowi rozszerzenie metody zarówno pod względem założeń teoretycznych, jak i zastosowań praktycznych. Celem pracy jest jednolite ujęcie metod operatorowych dostosowanych do różnych typów zagadnień, na przykład do rozwiązywania równań różniczkowych i różnicowych o stałych współczynnikach, równań typu Eulera, równań różniczkowo-różnicowych, równań Bernoulli'ego i innych. Omawiana metoda oparta jest na podstawowych pojęciach analizy funkcjonalnej. Liczne przykłady zamieszczone w pracy mają na celu zapoznanie pracowników technicznych ze sposobem praktycznego wykorzystania metody.

1. WSTĘP

Rachunek operatorowy jest metodą matematyczną posiadającą szerokie pole zastosowań w technice, a w szczególności w elektrotechnice. Wartość rachunku operatorowego dla elektrotechniki wypływa może nie tyle z faktu, że upraszcza on rozwiązywanie równań różniczkowych, lecz raczej stąd, że umożliwia badanie szeregu zagadnień w ogóle bez konieczności posługiwania się równaniami różniczkowymi. Rachunek operatorowy odgrywa szczególną rolę przy badaniu stanów nieustalonych w układach elektrycznych. Jego znaczenie w tej dziedzinie wynika również stąd, że „tłumaczy” on skomplikowane relacje matematyczne opisujące zjawiska stanów nieustalonych na prosty i przejrzysty „język algebry”. Zagadnienia stanów nieustalonych nie stanowią jednak wyłącznej domeny tego rachunku. Metody operatorowe zyskują prawo obywatelstwa w coraz to nowych dziedzinach elektrotechniki teoretycznej: w zagadnieniach syntezy, w zagadnieniach stabilności układów, w zagadnieniach procesów stochastycznych, a nawet w teorii układów nieliniowych.

W okresie ostatnich lat pojawiło się wiele, nieraz bardzo interesujących prac, omawiających podstawy teoretyczne i zastosowania rachunku

operatorowego. W obecnej chwili istnieje kilka odrębnych interpretacji rachunku operatorowego, najbardziej jednak powszechną jest metoda oparta na przekształceniach całkowych, a mianowicie na tzw. przekształceniach Laplace'a.

Mimo że rachunek operatorowy jest metodą stosowaną już od kilkudziesięciu lat w technice, stanowi on jeszcze dziedzinę stosunkowo mało zbadaną. Jest więc rzeczą możliwą, że kryje on w sobie możliwości nowych, dotąd niewykorzystanych zastosowań.

Dla uzyskania możliwie największej ogólności rachunku operatorowego ważne jest, by w uzasadnieniu rachunku korzystać jedynie z istotnie niezbędnych założeń ograniczających. Metody operatorowe oparte na przekształceniach Laplace'a (i również na innych, ogólniejszych przekształceniach całkowych) nie spełniają tego postulatu, gdyż narzucają ograniczenia na klasę funkcji transformowanych, a ograniczenia te, jak się okazuje, nie są niezbędne dla uzasadnienia rachunku.

Praca niniejsza jest próbą oparcia rachunku operatorowego na takich podstawach teoretycznych, które umożliwiłyby stosowanie go w jak najszerszym zakresie.

W pracy — oprócz podstaw teoretycznych — rozważono kilka operacji szczególnych. Za pomocą operacji tych można rozwiązywać zagadnienia praktyczne. Na szczególną uwagę zasługują — zdaniem autora — metody efektywnego rozwiązywania pewnych specjalnych równań różnicowych o zmiennych współczynnikach i równań różnicowo-różniczkowych. W pracy podana jest również próba stosowania rachunku operatorowego w zagadnieniach równań nieliniowych.

Podstawy omawianej w niniejszej pracy metody zostały przez autora opublikowane w roku 1957 w Biuletynie PAN [1]. Artykuł niniejszy stanowi rozszerzenie metody i jest przeznaczony przede wszystkim dla pracowników technicznych. Z tego względu autor podał liczne przykłady praktycznych zastosowań omówionej metody.

Przedstawiona w pracy metoda uogólnionego rachunku operatorowego może stanowić podstawę do szerzej pojętych metod analizy i syntezy układów elektrycznych, zbudowanych z elementów, które stanowią fizyczną realizację operatorów ogólnych $F(p)$.

2. PODSTAWY TEORETYCZNE

2.1. Algebra operatorów

2.1.1. Określenie operacji podstawowej

Założmy, że w pewnej przestrzeni liniowej X na ciele liczb zespolonych Z określona została operacja liniowa T spełniająca następujące warunki:

$$T(\mathbf{X}) \subset \mathbf{X}, \quad (1)$$

$$T[\alpha x_1 + \beta x_2] = \alpha T x_1 + \beta T x_2, \quad (2)$$

gdzie α i β są liczbami zespolonymi, a x_1 i x_2 są dwoma dowolnymi elementami przestrzeni $\mathbf{X}$. Operację taką nazywamy endomorfizmem.

Dzięki warunkowi (1) operację T możemy w przestrzeni $\mathbf{X}$ wykonywać wielokrotnie. Operację n -krotną $T^n[x]$ określamy wzorem rekurencyjnym

$$T^n[x] = T[T^{n-1}[x]]; \quad n = 1, 2, \dots \quad (3)$$

Umawiamy się przy tym konsekwentnie symbolem T^0 oznaczać operator tożsamościowy posiadający własność następującą

$$T^0[x] = x. \quad (4)$$

Dzięki tej własności operator T^0 będziemy również zapisywali jako liczbę 1:

$$T^0 = 1. \quad (5)$$

Z określenia (3) wynika od razu przemienność operatorów T^m i T^n , to znaczy, że mamy spełnione następujące równanie

$$T^m[T^n x] = T^n[T^m x] = T^{m+n} x. \quad (6)$$

Zakładamy, że zbiór wyników operacji potęgowych $x_i = T^i x$ stanowi układ elementów liniowo niezależnych, a więc że spełniony jest również następujący warunek:

$$\sum_{i=0}^n \alpha_i T^i x = 0 \equiv \alpha_0 = \alpha_1 = \dots = \alpha_n = 0 \quad \text{lub} \quad x = 0, \quad (7)$$

gdzie n jest dowolną liczbą naturalną.

2.1.2. Operatory wielomianowe

Rozważmy obecnie zbiór $\mathbf{P}$ operatorów wielomianowych

$$W(T) = \alpha_0 + \alpha_1 T + \dots + \alpha_n T^n; \quad \alpha_i \in \mathbf{Z}. \quad (8)$$

Dla operatorów wielomianowych określamy równość, sumę i iloczyn w sposób analogiczny jak dla wielomianów algebraicznych. Tak na przykład dwa operatory wielomianowe

$$W_1(T) = \sum_i \alpha_i T^i \quad \text{i} \quad W_2(T) = \sum_i \beta_i T^i$$

nazywamy równymi i piszemy

$$W_1(T) = W_2(T),$$

wtedy i tylko wtedy gdy $\alpha_i = \beta_i$; $i = 0, 1, 2, \dots$

Operator wielomianowy (8), dla którego $\alpha_0 = \alpha_1 = \dots = \alpha_n = 0$ nazywamy operatorem zerowym i zapisujemy go po prostu jako liczbę zero.

Wynik operacji $W(T)x$ określonej za pomocą operatora wielomianowego określimy wzorem

$$W(T)x = \sum_{i=0}^n \alpha_i T^i x. \quad (9)$$

Można zauważyć, że zachodzą następujące związki słuszne dla dwóch dowolnych operatorów wielomianowych W_1 i W_2

$$\left. \begin{aligned} W_1 x + W_2 x &= (W_1 + W_2)x, \\ W_1(W_2 x) &= (W_1 W_2)x. \end{aligned} \right\} \quad (10)$$

Uzasadnimy obecnie następujące

Twierdzenie 1. *Zbiór $\mathbf{P}$ operatorów wielomianowych tworzy pierścień przemienny bez dzielników zera.*

Dowód tego twierdzenia jest niemal natychmiastowy. Ponieważ bowiem na zbiorze operatorów wielomianowych określiliśmy działania analogicznie jak dla wielomianów algebraicznych, więc zbiór $\mathbf{P}$ ze względu na dodawanie tworzy grupę abelową. Ponadto mnożenie operatorów wielomianowych jest przemienne oraz również łączne i rozdzielne względem dodawania. Zbiór $\mathbf{P}$ tworzy więc pierścień przemienny. Pierścień ten nie ma dzielników zera ze względu na warunek (7), mamy bowiem

$$Wx = 0 \equiv W = 0 \quad \text{lub} \quad x = 0$$

oraz

$$W_1 W_2 = 0 \equiv W_1 = 0 \quad \text{lub} \quad W_2 = 0. \quad (11)$$

Z twierdzenia 1 wynika od razu następujący

Wniosek. *Wyniki dwóch operacji $W_1 x$ i $W_2 x$, gdzie $x \in \mathbf{X}$, $x \neq 0$, są równe wtedy i tylko wtedy, gdy*

$$W_1 = W_2.$$

Istotnie, jeśli założymy

$$W_1 x = W_2 x,$$

to dzięki pierwszemu ze związków (10)

$$(W_1 - W_2)x = 0; \quad x \neq 0,$$

ale na mocy (11) równanie to jest spełnione wtedy i tylko wtedy, gdy

$$W_1 = W_2.$$

2.1.3. Operatory wymierne

Ponieważ pierścień operatorów wielomianowych $W(T)$ nie posiada dzielników zera, więc można go w zwykły sposób uzupełnić do ciała ilorazowego wprowadzając operatory wymierne postaci

$$\frac{P(T)}{Q(T)} = \frac{a_0 + a_1 T + \dots + a_n T^n}{\beta_0 + \beta_1 T + \dots + \beta_m T^m}; \quad Q \neq 0, \quad (12)$$

gdzie a_i oraz β_i są liczbami zespolonymi.

Wynikiem operacji

$$\frac{P}{Q} x \quad (13)$$

określonej przez operator wymierny nazywać będziemy element $y \in X$ spełniający równanie

$$\boxed{Px = Qy}. \quad (14)$$

Wykażemy, że operacja (13) jest jednoznaczna, tzn. że dla danych dwóch operatorów wielomianowych P i Q i dla danego elementu $x \in X$ istnieje jeden i tylko jeden element $y \in X$ spełniający równanie (14).

Istotnie, jeśli założymy, że istnieją dwa elementy y_1 i y_2 spełniające (14)

$$P(x) = Q(y_1) \quad \text{i} \quad P(x) = Q(y_2),$$

to dzięki addytywności i jednorodności operacji potęgowych otrzymamy

$$P(x - x) = Q(y_1 - y_2)$$

i w konsekwencji

$$Q(y_1 - y_2) = 0; \quad Q \neq 0.$$

Ale na mocy twierdzenia 1 równanie powyższe jest spełnione wtedy i tylko wtedy, gdy

$$y_1 - y_2 = 0,$$

a więc wówczas, gdy elementy y_1 i y_2 są identyczne.

Warto podkreślić, że operacja (13) określona za pomocą operatora wymiernego nie zawsze jest wykonalna na rozważanym zbiorze liniowym $\mathbf{X}$. Przypadek niewykonalności tej operacji ma miejsce wówczas, gdy nie jest na zbiorze $\mathbf{X}$ rozwiązalne równanie operatorowe (14).

Jeśli na przykład $\mathbf{X}$ jest zbiorem funkcji całkowalnych według Lebesque określonych na półosi rzeczywistej $[0, \infty)$, to równanie

$$Tx(t) = 1; \quad t \geq 0, \quad x(t) \in \mathbf{X}$$

przy określeniu operacji $Tx(t)$ przez wzór

$$Tx(t) = \int_0^t x(\tau) d\tau$$

nie jest rozwiązalne w zbiorze $\mathbf{X}$, gdyż żądamy, by prawa strona równania była równa 1 dla $t \geq 0$, a to nie jest możliwe, gdyż wynik operacji

$$Tx(t) = \int_0^t x(\tau) d\tau$$

zawsze przedstawia funkcję, która jest równa zeru dla $t = 0$. Równanie to byłoby natomiast rozwiązalne w przypadku, gdy $\mathbf{X}$ jest zbiorem dystrybucji.

Można zauważyć, że warunkiem wystarczającym (*ale nie koniecznym*) istnienia operacji określonej przez operator wymierny (12) jest warunek $\beta_0 \neq 0$.

Łatwo stwierdzić, że przy założeniu istnienia odpowiednich wyników mają miejsce następujące związki:

$$\frac{P}{Q} [ax_1 + \beta x_2] = a \frac{P}{Q} x_1 + \beta \frac{P}{Q} x_2, \quad (15)$$

$$\frac{P_1}{Q_1} x + \frac{P_2}{Q_2} x = \frac{P_1 Q_2 + P_2 Q_1}{Q_1 \cdot Q_2} x, \quad (16)$$

$$\frac{P_1}{Q_1} \left[\frac{P_2}{Q_2} x \right] = \frac{P_1 \cdot P_2}{Q_1 \cdot Q_2} x. \quad (17)$$

Związki te oznaczają, że operator wymierny $\frac{P}{Q}$ jest operatorem liniowym oraz że operatory wymierne możemy dodawać i mnożyć w sposób formalny tak jak ułamki algebraiczne. Dla operatorów wymiernych słuszne są więc wszelkie prawa działań na ułamkach, a między innymi prawa przemienności, łączności i rozdzielności¹⁾. **Wyrażamy to mówiąc, że ciało**

¹⁾ Jako wynik operacji $\frac{P}{Q} x$ możemy również rozumieć parę uporządkowaną $\left(\frac{P}{Q}, x \right)$, dla której jako aksjomaty przyjmujemy wzory (15, ... 17). W ten sposób równanie (14) zawsze będzie posiadało jednoznaczne rozwiązanie na zbiorze par $\left(\frac{P}{Q}, x \right)$. Takie uogólnienie nie jest jednak konieczne dla rachunku operatorów.

operatorów wymiernych jest izomorficzne z ciałem funkcji wymiernych. Dzięki istnieniu powyższego izomorfizmu operator wymierny (12) ma jednoznaczny rozkład na ułamki proste postaci

$$\frac{1}{(T - \lambda)^k}, \quad (18)$$

co ma podstawowe znaczenie w zastosowaniach rachunku operatorowego.

Szczególnie ważnym operatorem wymiernym jest operator $\frac{1}{T^i}$. Operator ten będziemy oznaczali również symbolem T^{-i} pisząc

$$\frac{1}{T^i} = T^{-i}. \quad (19)$$

Z uwagi na powyższe oznaczenie mamy

$$T^i T^{-i} = T^{-i} T^i = 1.$$

Ogólnie operatory P i Q , których iloczyn jest równy jedności (a więc również i operatory T^i , T^{-i}), nazywamy operatorami odwrotnymi. Operator $\frac{1}{T}$ będziemy również oznaczali literą p

$$p = \frac{1}{T}. \quad (20)$$

2.1.4. Operatory uogólnione

Rozważmy z kolei operator wielomianowy ogólniejszy

$$W(T) = A_0 + A_1 T + \dots + A_n T^n, \quad (21)$$

gdzie $A_0, A_1, \dots, A_n$ są zadanymi endomorfizmami przemiennymi z endomorfizmem T .

Taki operator A_i , który sprowadza dowolny element y zbioru X do elementu zerowego, będziemy nazywali operatorem zerowym, pisząc $A_i = 0$. Jeśli więc

$$A_i y = 0 \text{ przy } y \neq 0, \text{ to } A_i = 0. \quad (22)$$

Operator wielomianowy (21), dla którego $A_1 = A_2 = \dots = A_n = 0$, będziemy również nazywali operatorem zerowym, pisząc $W = 0$.

Żałómy, że wyniki operacji $x_i = A_i T^i x$ stanowią układ elementów liniowo niezależnych w tym sensie, że

$$\sum_{i=0}^n A_i T^i x = 0 \equiv A_i = 0 \text{ lub } x = 0; \quad i = 0, 1, \dots, n. > \quad (23)$$

Przy określeniu działań arytmetycznych na operatorach (21) tak jak dla wielomianów algebraicznych, zbiór tych operatorów będzie oczywiście stanowił pierścień przemienny nie posiadający — dzięki założeniu (23) — dzielników zera. Pierścień ten można zatem uogólnić do ciała ilorazowego w sposób analogiczny jak poprzednio otrzymując operatory wymierne $\frac{P}{Q}$. Dla operatorów postaci (21) pozostają oczywiście w mocy wzory (15, ... 17), dzięki czemu na uogólnionych operatorach wymiernych również możemy przeprowadzać działania algebraiczne w sposób formalny tak jak na ułamkach algebraicznych. W szczególności może mieć miejsce rozkład operatora uogólnionego $\frac{P}{Q}$ na ułamki proste postaci

$$\frac{B}{(1 - ET)^k}, \quad (24)$$

gdzie B i E — endomorfizmy, zaś k — liczba naturalna.

Operatory uogólnione mogą między innymi być stosowane w zagadnieniach rozwiązywania równań różniczkowych postaci

$$x^{(n)}(t) + A_1 x^{(n-1)}(t) + \dots + A_n x(t) = a(t), \quad (25)$$

gdzie $A_1, \dots, A_n$ są ciągłymi endomorfizmami, a funkcje $x(t)$, $a(t)$ są ciągłymi funkcjami zmiennej rzeczywistej t , o wartościach z liniowej przestrzeni topologicznej X . Równania takie były np. rozważane w pracy [6].

2.2. Analiza operatorów

2.2.1. Operatory rzeczywiste

Założmy obecnie, że X jest liniową przestrzenią topologiczną np. typu L^* Frécheta. Zbieżność ciągu x_n w przestrzeni Frécheta określamy za pomocą następujących aksjomatów:

1°. Dla niektórych ciągów x_n utworzonych z elementów przestrzeni X przyporządkowany jest element $x \in X$, który nazywamy granicą ciągu x_n , pisząc

$$x = \lim_n x_n \text{ lub } x_n \rightarrow x.$$

Ciągi posiadające granicę nazywamy zbieżnymi.

2°. Każdy ciąg x_n posiada co najwyżej jedną granicę

3°. Jeśli $x_n = x$ dla $n = 1, 2, \dots$, to $\lim_n x_n = x$

4°. Podciąg ciągu zbieżnego do x jest też zbieżny do x , tzn. jeśli $x = \lim_n x_n$

oraz $m_1 < m_2 < \dots$, to również $x = \lim_n x_{m_n}$

5°. Jeśli każdy podciąg x_{m_n} ciągu x_n zawiera podciąg $x_{m_{k_n}}$ zbieżny do x , to $x = \lim_n x_n$.

Oprócz powyższych aksjomatów przyjmujemy również następujące warunki ciągłości:

$$\left. \begin{aligned} x_n \rightarrow x, \quad y_n \rightarrow y &\supset x_n + y_n \rightarrow x + y, \\ x_n \rightarrow x, \quad a_n \rightarrow a &\supset a_n x_n \rightarrow ax. \end{aligned} \right\} \quad (26)$$

Przestrzeń topologiczną X nazywamy metryzowalną, jeśli istnieje możliwość zdefiniowania w zbiorze X metryki $\rho(x, y)$ w taki sposób, aby zbieżność, określona w przestrzeni X za pomocą powyższych aksjomatów, była identyczna ze zbieżnością implikowaną przez metrykę. Oczywiście nie wszystkie przestrzenie topologiczne są przestrzeniami metryzowalnymi.

Pojęcie szeregu utworzonego z elementów przestrzeni topologicznej określamy analogicznie jak w analizie klasycznej. Sumą szeregu nazywamy granicę ciągu sum częściowych. Szereg nazywamy zbieżnym, jeśli ciąg jego sum częściowych posiada granicę na zbiorze X .

Określmy obecnie w liniowej przestrzeni topologicznej X endomorfizm T , a więc operację spełniającą warunki (1) i (2) i założmy ponadto, że endomorfizm ten spełnia następujące warunki dodatkowe:

$$\lim_n T x_n = T \lim_n x_n, \quad (27)$$

$$\sum_{i=0}^{\infty} a_i T^i x = 0 \equiv a_0 = a_1 = \dots = a_i = 0 \dots = 0 \text{ lub } x = 0, \quad (28)$$

gdzie $a_i \in \mathbb{Z}$ oraz $x \in X$,

a więc warunek ciągłości i liniowej niezależności.

Weźmy teraz pod uwagę ciąg operatorów wymiernych

$$\frac{P_n}{Q_n}(T) = \frac{\sum_i a_{in} T^i}{\sum_i \beta_{in} T^i}. \quad (29)$$

Ciąg $\frac{P_n}{Q_n}(T)$ będziemy nazywali zbieżnym do operatora wymiernego

$\frac{P}{Q}(T)$ pisząc

$$\lim_n \frac{P_n}{Q_n}(T) = \frac{P}{Q}(T)$$

lub

$$\frac{P_n}{Q_n}(T) \rightarrow \frac{P}{Q}(T),$$

wtedy i tylko wtedy, gdy ciąg funkcyjny $\frac{P_n}{Q_n}(z)$ zmiennej zespolonej z jest niemal jednostajnie zbieżny do funkcji wymiernej $\frac{P}{Q}(z)$

$$\frac{P_n}{Q_n}(z) \Rightarrow \frac{P}{Q}(z),$$

Można zauważyć, że tak określona zbieżność spełnia aksjomaty 1° ... 5° zbieżności w przestrzeni topologicznej oraz spełnia również warunki ciągłości (26).

Ciąg operatorów $\frac{P_n}{Q_n}(T)$ będziemy z kolei nazywali ciągiem podstawowym, jeśli ciąg funkcyjny $\frac{P_n}{Q_n}(z)$ jest niemal jednostajnie zbieżny do funkcji $F(z)$ niekoniecznie wymiernej

$$\frac{P_n}{Q_n}(z) \Rightarrow F(z).$$

Wprowadzimy obecnie następujące

Określenie. Klasę ciągów podstawowych operatorów wymiernych $u_n = \frac{P_n}{Q_n}(T)$ wyznaczających jednoznacznie funkcję zmiennej zespolonej

$F(z) = \lim_{n \rightarrow \infty} \frac{P_n}{Q_n}(z)$, przy określeniu równości sumy i iloczynu analogicznie jak dla funkcji $F(z)$, nazywamy operatorem rzeczywistym oznaczając go symbolem $F(T)$. Mamy zatem

$$\frac{P_n}{Q_n}(T) \sim F(T) \equiv \frac{P_n}{Q_n}(z) \Rightarrow F(z). \quad (30)$$

Zbiór operatorów rzeczywistych będziemy oznaczali przez R_0 .

Odejmowanie i dzielenie operatorów rzeczywistych określamy jako działania odwrotne względem dodawania i mnożenia, a zbieżność ciągu operatorów $F_n(T)$ rozumiemy w sensie zbieżności ciągów podstawowych $F_n(z)$. Z przyjętych powyżej założeń wynika natychmiast wniosek, że zbiór operatorów rzeczywistych R_0 tworzy przestrzeń zupełną izomorficzną z pewnym podzbiorem funkcji zmiennej zespolonej. Przestrzeń R_0 jest poza tym liniową przestrzenią topologiczną typu L^* Frécheta.

Wynikiem operacji $F(T)x$:

$$y = F(T)x, \quad (31)$$

określonej za pomocą operatora rzeczywistego $F(T) \sim \frac{P_n}{Q_n}(T)$ nazywać

będziemy element y liniowej przestrzeni topologicznej X , spełniający równanie

$$y = \lim_n y_n = \lim_n \frac{P_n}{Q_n} x; \quad x, y, y_n \in X. \quad (32)$$

przy założeniu, że $\frac{P_n}{Q_n} \rightarrow F \supset \frac{P_n}{Q_n} x \rightarrow Fx \in X$.

Oczywiście operacja (31) nie zawsze jest wykonalna na zbiorze X . W związku z powyższym określimy warunki wystarczające istnienia tej operacji.

Dla sformułowania powyższych warunków rozważymy przypadek szczególny, gdy X jest przestrzenią lokalnie wypukłą z topologią określoną za pomocą ciągu pseudonorm $\|x\|_k$

tj. funkcjonałów nieujemnych spełniających warunki

- (1.1) $(i_1) \|x\|_k = 0$, gdy $x = 0$;
 (i_2) z warunku $\|x\|_k = 0$ dla $k = 1, 2, \dots$ wynika, że $x = 0$;
 $(i_3) \|x + y\|_k \leq \|x\|_k + \|y\|_k$; $x, y \in X$;
 $(i_4) \|ax\|_k = |a| \cdot \|x\|_k$; $a \in \mathbb{Z}$, $x \in X$.

Przy metryce

$$\varrho(x, y) = \sum_{k=1}^{\infty} \frac{1}{2^k} \min(1, \|x - y\|_k) \quad (33)$$

przestrzeń X jest liniową przestrzenią topologiczną.

Przestrzeń taka jest szczególnie ważna w praktycznych zastosowaniach omawianej metody.

Pseudonormą operatora T nazwiemy wyrażenie

$$\|T\|_k = \sup_{\|x\|_k \leq 1} \|Tx\|_k. \quad (34)$$

Operator T będziemy nazywali

1° — słabo ograniczonym, jeśli

$$\lim_n \sqrt[n]{\|T^n\|_k} < M; \quad k = 1, 2, \dots, \quad (35)$$

gdzie M jest dowolną dodatnią liczbą rzeczywistą (skończoną);

2° — istotnie ograniczony, jeśli

$$\lim_n \sqrt[n]{\|T^n\|_k} = 0; \quad k = 1, 2, \dots \quad (36)$$

Uzasadnimy obecnie następujące

Twierdzenie 2. Jeżeli endomorfizm T jest operatorem istotnie ograniczonym oraz istnieje $x_0 \in X$ taki, że $\frac{P_n}{Q_n} x = \frac{P_n^*}{Q_n^*} x_0$, a $\lim_{n \rightarrow \infty} \frac{P_n^*}{Q_n^*}(z)$ jest

funkcją holomorficzną w otoczeniu punktu $z = 0$, to operacja $F(T)x$ określona za pomocą operatora rzeczywistego $F(T) \sim \frac{P_n}{Q_n}(T)$ jest wykonalna na zbiorze $\mathbf{X}$.

Dowód. Granicą ciągu $y_n = \frac{P_n^*}{Q_n^*} x_0$ jest suma szeregu potęgowego od operatora T . Mamy zaś

$$\left\| \sum_{i=0}^{\infty} \gamma_i T^i x_0 \right\|_k \leq \sum_{i=0}^{\infty} \left\| \gamma_i T^i x_0 \right\|_k \leq \sum_{i=0}^{\infty} |\gamma_i| \left\| T^i \right\|_k \left\| x_0 \right\|_k < M_k \sum_{i=0}^{\infty} \frac{1}{r_1^i} \left\| T^i \right\|_k ;$$

$$k = 1, 2, \dots; r_1 < r,$$

gdzie r jest promieniem zbieżności szeregu $\sum_{i=0}^{\infty} \gamma_i z^i$. Jeśli więc operator T spełnia warunek (36), to ciąg y_n posiada granicę na zbiorze $\mathbf{X}$, skąd wynika słuszność tezy twierdzenia.

Można zauważyć, że teza twierdzenia 2 jest również spełniona, jeśli

$$\lim_i \sqrt[i]{\left\| T^i \right\|_k} < r_1 < r; k = 1, 2, \dots, \quad (37)$$

a więc również dla niektórych operatorów słabo ograniczonych. W przypadku operatorów słabo ograniczonych zbieżność ciągu $W_n(T)$ nie implikuje jednak zbieżności $y_n = W_n(T)x$, co zawsze ma miejsce dla operatorów istotnie ograniczonych i co jest warunkiem istnienia operacji określonej za pomocą operatora rzeczywistego (30). (W a r u n e k 32).

Istotnie ograniczonym jest np. operator Heaviside'a, dla którego $\mathbf{X}$ jest zbiorem funkcji $x(t)$ ciągłych (lub całkownych) w każdym przedziale $[0, t_0]$ półośi rzeczywistej oraz

$$Tx(t) = \int_0^t x(\tau) d\tau. \quad (38)$$

Pseudonormę w tym przypadku określamy wzorem

$$\|x\|_k = \sup \left| x(t) \right|; k = 1, 2, \dots$$

$$\left(\frac{1}{k}, k \right) \quad (39)$$

Istotna ograniczoność operatora Heaviside'a T natychmiast wynika z całki Cauchy'go

$$T^i x(t) = \frac{1}{(i-1)!} \int_0^t (t-\tau)^{i-1} x(\tau) d\tau.$$

Zbieżność implikowana przez ciąg pseudonorm (39) jest identyczna ze zbieżnością niemal jednostajną. Liniowa niezależność w sensie (28) dla operatorów Heaviside'a wynika z twierdzenia Titchmarsh'a.

Ze wzoru (28) wynika, że zbiór operatorów rzeczywistych nie ma dzielników zera, tzn.

$$F(T)x = 0 \equiv F'(T) = 0 \text{ lub } x = 0. \quad (40)$$

Łatwo również stwierdzić słuszność związków analogicznych do (10).

Dla określania operatorów rzeczywistych bardzo przydatna jest metoda aproksymacji w sensie Padé.

2.2.2. Obliczanie wyników operacji metodą iteracyjną

Rozważmy równanie operatorowe

$$y - E_1 T y = 1, \quad (41)$$

gdzie E_1 i T są swoma zadanymi endomorfizmami.

Zgodnie z określeniem operatora wymiernego, równanie to określa wynik następującej operacji

$$y = \frac{1}{1 - E_1 T} (1). \quad (42)$$

Ponieważ równanie (41) nie zawsze ma rozwiązanie w dziedzinie funkcji elementarnych, dlatego dogodnym sposobem wyznaczania wyniku operacji (42) jest stosowanie procesu iteracyjnego.

Weźmy mianowicie dowolny element $y_0 \in X$ i znając operację $E_1 T$, z równania

$$y_1 - E_1 T y_0 = 1$$

obliczmy element y_1 . Kontynuując ten proces otrzymamy

$$\begin{aligned} y_2 &= E_1 T y_1 + 1, \\ y_3 &= E_1 T y_2 + 1, \\ &\dots\dots\dots \\ y_n &= E_1 T y_{n-1} + 1, \\ &\dots\dots\dots \end{aligned} \quad (43)$$

Wobec tego element y_n można wyrazić przez element początkowy y_0 następująco:

$$y_n = E_1^n T^n y_0 + E_1^{n-1} T^{n-1} + \dots + E_1 T + 1.$$

Jeżeli powyższy proces iteracyjny jest zbieżny w sensie przyjętej dla przestrzeni X topologii, to wówczas zakładając dla prostoty $y_0 = 1$ otrzymamy szereg potęgowy od operatora $E_1 T$ zbieżny do y . W tym przypadku słuszne jest więc następujące rozwinięcie operatora potęgowego (42)

$$\frac{1}{1 - E_1 T} = \sum_{i=0}^{\infty} E_1^i T^i. \quad (44)$$

Z rozwinięcia (44) wynika natychmiast związek ogólniejszy

$$\frac{1}{(1 - E_1 T)^k} = \sum_{i=0}^{\infty} \binom{k+i-1}{i} E_1^i T^i, \quad (45)$$

który otrzymujemy podnosząc obie strony wzoru (44) do potęgi k . Przybliżoną wartość operacji określonej za pomocą operatora wymiernego (44) można więc wyznaczyć obliczając skończoną liczbę wyrazów szeregu, a mianowicie

$$\frac{1}{(1 - E_1 T)^k} x \cong \sum_{i=0}^n \binom{k+i-1}{i} E_1^i T^i x. \quad (46)$$

Postępowanie takie jest równoważne obliczeniu skończonej liczby iteracji. Wzory (45) i (46) pozostają oczywiście słuszne dla przypadku, gdy endomorfizm E_1 zastąpimy liczbą zespoloną α .

Rozważając w analogiczny sposób jak poprzednio równanie ogólniejsze

$$y - F(T)y = 1$$

można zauważyć, że w przypadku zbieżności odpowiedniego procesu iteracyjnego istnieje również następujący związek

$$\frac{1}{[1 - F(T)]^k} = \sum_{i=0}^{\infty} \binom{k+i-1}{i} F^i(T). \quad (46a)$$

2.2.3. Operacje w przestrzeniach operatorów

Ponieważ zbiór $\mathbf{R}_0$ operatorów rzeczywistych tworzy przestrzeń liniową, więc na zbiorze tym możemy określić endomorfizm T_1

$$T_1(\mathbf{R}_0) \subset \mathbf{R}_0, \quad (47)$$

$$T_1[\alpha F_1 + \beta F_2] = \alpha T_1 F_1 + \beta T_1 F_2, \quad (48)$$

$$F_1, F_2 \in \mathbf{R}_0; \alpha, \beta \in \mathbf{Z}, \quad (49)$$

spełniający warunki ustalone ogólnie w poprzednich rozdziałach pracy. Jeśli więc na endomorfizm T_1 nałożymy warunki zapewniające istnienie wyników operacji $F^{(1)}(T_1)$ $[\mathbf{R}_0]$, to moglibyśmy skonstruować nową przestrzeń $\mathbf{R}_1$ operatorów $F^{(1)}(T_1)$. Oczywiście postępowanie to możemy w tym przypadku kontynuować dalej dowolną liczbę razy otrzymując nowe przestrzenie $\mathbf{R}_2, \mathbf{R}_3 \dots, \mathbf{R}_n, \dots$

Na zbiorze $\mathbf{R}_0$ można określić operację różniczkową odwzorowującą przestrzeń $\mathbf{R}_0$ w siebie i spełniającą warunki

$$\begin{aligned} D(F_1 F_2) &= F_1 D(F_2) + F_2 D(F_1), \\ D(\alpha F_1 + \beta F_2) &= \alpha D(F_1) + \beta D(F_2). \end{aligned} \quad (50)$$

Przykładem operacji tego typu może być operacja różnicowa określona na zbiorze ciągów F_n operatorów za pomocą wzoru

$$\Delta^k F_n = \sum_{\nu=0}^k (-1)^\nu \frac{k!}{\nu! (k-\nu)!} F_{n+k-\nu}. \quad (51)$$

Określając endomorfizm T_1 wzorem

$$T_1 F_n = \sum_{m=0}^{n-1} F_m, \quad (52)$$

możemy operatorowo rozwiązywać równania różnicowe definiowane w przestrzeni operatorów.

Innym przykładem operacji różniczkowej może być operacja różniczkowania funkcji $F(\lambda)$ zmiennej rzeczywistej λ , o wartościach z przestrzeni R_0 . Pochodną $F'(\lambda)$ można przy tym określić przez izomorfizm z pochodną $F'_\lambda(z, \lambda)$ funkcji zmiennej zespolonej zależnej od parametru rzeczywistego λ . W ten sposób na pochodną $F'(\lambda)$ przenoszą się własności pochodnej $F'_\lambda(z, \lambda)$. Można zatem zdefiniować równanie różniczkowe dla przestrzeni funkcji $F(\lambda)$, co umożliwia operatorowe rozwiązywanie równań różniczkowych cząstkowych.

Przykład. Równanie

$$u_{t\lambda} - a u = 0, \quad (53)$$

z warunkami brzegowymi $u(x, 0) = 0$, $u(0, t) = f(t)$ rozwiązujemy określając endomorfizm T wzorem (38).

Jeśli oznaczymy $F(\lambda) f = u$, to równanie cząstkowe (53) przekształca się w równanie zwyczajne w przestrzeni $R_0(\lambda)$

$$(p F'_\lambda - a F) f = 0; \quad F(0) = I = 1,$$

które prowadzi do wyniku

$$F(\lambda) = e^{\frac{a}{p} \lambda} = e^{a\lambda T},$$

oraz

$$u = e^{a\lambda T} f.$$

Wynik tej operacji otrzymujemy prosto przez rozłożenie $e^{a\lambda T}$ na szereg potęgowy od operatora T .

3. ZASTOSOWANIA

3.1. Operatory Heaviside'a

Założmy, że zbiór X jest zbiorem funkcji $f(t)$ całkowalnych w każdym przedziale $[0, t_0]$, i określmy na tym zbiorze operację $T f(t)$ następującą:

$$T f(t) = \int_0^t f(\tau) d\tau. \quad (54)$$

Z definicji (54) natychmiast wynika, że operacja $T f(t)$ jest addytywna i jednorodna, a więc, że spełnione są związki (2). Można również zauważyć, że operacja ta spełnia związek (7), a wobec tego zbiór operatorów wielomianowych $W(T)$ tworzy pierścień bez dzielników zera.

Podstawiając we wzorze (54) na miejsce funkcji $f(t)$ jej pochodną, otrzymamy

$$T f'(t) = f(t) - f(0)$$

oraz

$$f'(t) = \frac{1}{T} f(t) - \frac{1}{T} f(0)$$

lub wprowadzając oznaczenie $\frac{1}{T} = p$

$$f'(t) = p f(t) - p f(0). \quad (55)$$

Wzór powyższy można drogą prostej indukcji uogólnić na pochodne wyższych rzędów, a wówczas otrzymamy

$$\begin{aligned} f''(t) &= p^2 f(t) - p^2 f(0) - p f'(0), \\ f^{(v)}(t) &= p^v f(t) - \sum_{n=0}^{v-1} p^{v-n} f^{(n)}(0); \quad v = 1, 2, \dots, \end{aligned} \quad (56)$$

przy czym należy oczywiście założyć istnienie pochodnych funkcji $f(t)$ aż do rzędu v włącznie.

Operator $p = \frac{1}{T}$ nazywamy różniczkowym, a $p^{-1} = T$ — całkowym operatorem Heaviside'a.

Korzystając ze związków (56) możemy rozwiązywać liniowe równania różniczkowe o stałych współczynnikach w sposób analogiczny jak w przypadku metody opartej np. na przekształceniach Laplace'a.

W zastosowaniach dogodnie jest przy tym korzystać z dobrze znanego w rachunku operatorowym

Twierdzenia. Jeżeli $h(t)$ jest wynikiem operacji $\frac{P}{Q}$ nad funkcją jednostkową

$$h(t) = \frac{P}{Q}(1),$$

to wynik operacji nad dowolną funkcją całkowalną $f(t)$ można wyrazić za pomocą całki

$$\frac{P}{Q}[f(t)] = \frac{d}{dt} \int_0^t h(t-\tau) f(\tau) d\tau. \quad (57)$$

Dowód powyższego twierdzenia można znaleźć w literaturze.

Całka (57) bywa zwykle nazywana „całką Duhamela”. Należy podkreślić, że przedstawiona w niniejszym rozdziale metoda jest ogólniejsza od metody przekształcenia Laplace’a w tym sensie, że nie narzuca ograniczeń „na szybkość wzrostu” funkcji $f(t)$.

Wyniki niektórych operacji można obliczać wprost ze wzoru definicyjnego. Zakładając np. we wzorze (54) $f(t) = e^{at}$, znajdziemy:

$$T e^{at} = \int_0^t e^{a\tau} d\tau = \frac{1}{a}(e^{at} - 1)$$

oraz

$$(1 - aT)[e^{at}] = 1$$

i ostatecznie po podstawieniu $p = T^{-1}$:

$$\boxed{\frac{p}{p-a}(1) = e^{at}.} \quad (58)$$

Dogodnie jest umówić się, by w przypadku operacji wykonywanych nad funkcją jednostkową zamiast $\frac{P(p)}{Q(p)}(1)$ pisać po prostu $\frac{P(p)}{Q(p)}$. Przy takiej umowie zależność (58) zapisalibyśmy wzorem

$$\frac{p}{p-a} = e^{at}. \quad (58a)$$

Nic nie stoi na przeszkodzie, by podaną w niniejszym rozdziale interpretację rachunku Heaviside’a stosować również w zagadnieniach równań różniczkowych cząstkowych. W tym celu należało by wprowadzić kilka operatorów niewymiernych (nieregularnych), a mianowicie np. następujące operatory

$$e^{-\lambda p}; \quad e^{-\lambda \sqrt{p}}; \quad \sqrt{p}; \quad \frac{1}{\sqrt{p}}. \quad (59)$$

Stosowanie tej metody w zagadnieniach równań cząstkowych również miałoby zalety w stosunku do metody przekształcenia Laplace’a,

gdyż umożliwiałoby np. przeprowadzanie dowodów jednoznaczności rozwiązań tych równań.

3.2. Operatory równań Eulera

Załóżmy, że operacja $Tf(t)$ określona jest na zbiorze X za pomocą wzoru następującego

$$Tf(t) = \begin{cases} \int_1^t \frac{f(\tau)}{\tau} d\tau; & t \geq 1, \\ 0; & t \leq 1. \end{cases} \quad (60)$$

Można wykazać, że w ten sposób zdefiniowana operacja spełnia warunki postawione w rozdziale 2.

Jeśli we wzorze definicyjnym (60) na miejsce $f(t)$ podstawimy funkcję $t f'(t)$, to otrzymamy następujący związek

$$T t f'(t) = f(t) - f(1),$$

a więc, po podstawieniu $T^{-1} = p$:

$$t f'(t) = p f(t) - p f(1). \quad (61)$$

Związek ten można również uogólnić na pochodne wyższych rzędów. Jeżeli mianowicie oznaczmy $f(t) = t f_1(t)$, to znajdziemy

$$t^2 f_1''(t) + t f_1'(t) = p t f_1'(t) - p f_1'(1),$$

a stąd, po wykorzystaniu (61):

$$t^2 f_1''(t) + p f_1(t) - p f_1(1) = p^2 f_1(t) - p^2 f_1(1) - p f_1'(1),$$

oraz po uporządkowaniu i podstawieniu na miejsce $f_1(t)$ ponownie funkcji $f(t)$:

$$t^2 f''(t) = p(p-1)f(t) - p(p-1)f(1) - p f'(1). \quad (62)$$

Uogólniając następnie uzyskaną zależność na pochodne wyższych rzędów otrzymamy następujący wzór¹⁾:

$$t^{n+1} f^{(n+1)}(t) = p(p-1) \dots (p-n) f(t) - p(p-1) \dots (p-n) f(1) + \dots - p(p-n) f^{(n-1)}(1) - p f^{(n)}(1), \quad (63)$$

gdzie n jest liczbą całkowitą, nieujemną.

¹⁾ Powyższe wzory można również wyprowadzić w oparciu o przekształcenie Mellina; wówczas jednak, dzięki ograniczeniom nałożonym na klasę funkcji $f(t)$ uzyskuje się mniejszą ogólność rozważań.

W przypadku szczególnym, gdy

$$f(1) = f'(1) = \dots = f^{(n)}(1) = 0,$$

wzór (41) upraszcza się i otrzymujemy

$$t^{n+1} \cdot f^{(n+1)}(t) = p(p-1) \dots (p-n) f(t). \quad (63a)$$

Obliczmy wyniki niektórych prostszych operacji.

1. Załóżmy we wzorze definicyjnym (60) $f(t) = 1$; wówczas

$$T(1) = p^{-1}(1) = \int_1^t \frac{d\tau}{\tau} = \ln t. \quad (64)$$

Wykonując operację $T(1)$ wielokrotnie, otrzymamy ogólnie

$$p^{-\nu}(1) = \frac{\ln^\nu t}{\nu!}, \quad (64a)$$

gdzie ν jest liczbą naturalną i $\ln^\nu t$ oznacza uproszczony zapis $[\ln t]^\nu$.

2. Załóżmy we wzorze (60) $f(t) = t^a$; wówczas

$$T(t^a) = \int_1^t \tau^{a-1} d\tau = \frac{\tau^a}{a} \Big|_1^t = \frac{1}{a} (t^a - 1),$$

oraz po wykonaniu elementarnych przekształceń i po podstawieniu $p = T^{-1}$

$$\boxed{t^a = \frac{p}{p-a}(1).} \quad (65)$$

Korzystając zaś z oczywistej tożsamości

$$\frac{p}{p-a}(1) = \frac{1}{a} \cdot \frac{1}{p-a}(1) - \frac{1}{a}$$

mamy również

$$\frac{1}{a} (t^a - 1) = \frac{1}{p-a}(1). \quad (65a)$$

Za pomocą przedstawionej w niniejszym rozdziale operacji można zrećcznie rozwiązywać równania różniczkowe Eulera. Równaniem Eulera nazywamy — jak wiadomo — równanie postaci

$$a_0 t^n f^{(n)}(t) + \dots + a_{n-1} t f'(t) + a_n f(t) = \varphi(t),$$

gdzie $a_0, \dots, a_{n-1}, a_n$ są stałymi dowolnymi oraz $\varphi(t)$ jest daną funkcją zmiennej rzeczywistej t .

Poniższy przykład ilustruje sposób rozwiązywania tych równań metodą omawianej operacji.

Przykład. Wyznaczyć funkcję $f(t)$ spełniającą równanie

$$t^2 f''(t) - f(t) = \ln t$$

i warunki $f(1) = 0$, $f'(1) = 1$.

Rozwiązanie. Dzięki zależności (62) otrzymamy następujące równanie operatorowe

$$p(p-1)f(t) - f(t) = p1 + \ln t.$$

Zatem

$$f(t) = \frac{p}{p(p-1)-1}(1) + \frac{p}{p(p-1)-1}(\ln t).$$

Dalej, z uwagi na wzory (64) i (65)

$$f(t) = \frac{a_1^2 - 1}{a_1^2(a_1 - a_2)}(t^{a_1} - 1) + \frac{a_2^2 - 1}{a_2^2(a_2 - a_1)}(t^{a_2} - 1) - \frac{\ln t}{a_1 a_2},$$

gdzie $a_1, a_2 = \frac{1}{2} \pm \frac{\sqrt{5}}{2}$.

3.3. Operatory równań różnicowych

3.3.1. Równania różnicowe o stałych współczynnikach

W bardzo prosty sposób można metodą operacji $Tf(t)$ rozwiązywać liniowe równania różnicowe o stałych współczynnikach.

Z równania definicyjnego

gdzie $T \Delta_{\lambda} f(t) = f(t) - f(0),$ (66)

$$\Delta_{\lambda} f(t) = f(t + \lambda) - f(t),$$

po podstawieniu $T^{-1} = p$, otrzymamy

$$\Delta_{\lambda} f(t) = p f(t) - p f(0). \quad (67)$$

Uogólniając następnie powyższy wzór, znajdziemy

$$\begin{aligned} \Delta_{\lambda}^2 f(t) &= p^2 f(t) - p^2 f(0) - p \Delta_{\lambda} f(0), \\ \Delta_{\lambda}^n f(t) &= p^n f(t) - \sum_{\nu=0}^{n-1} p^{n-\nu} \Delta_{\lambda}^{\nu} f(0), \end{aligned} \quad (68)$$

gdzie oznaczono $\Delta_{\lambda}^0 f(0) = f(0)$.

Trzeba zauważyć, że operacja powyższa — w odróżnieniu od wszystkich poprzednio rozważanych operacji — nie jest operacją jednoznaczną. Można jednak uzasadnić następującą

Własność. Jeżeli dwie funkcje $f_1(t)$ i $f_2(t)$ spełniają równocześnie równanie $T \left[\Delta_{\lambda} f(t) \right] = x(t)$:

$$T \left(\Delta_{\lambda} f_1 \right) = x \quad \text{ i } \quad T \left(\Delta_{\lambda} f_2 \right) = x, \quad (69)$$

to różnica funkcji $f_1(t)$ i $f_2(t)$ jest funkcją okresową o okresie λ .

Istotnie, jeżeli założymy, że funkcje f_1 i f_2 spełniają równania (69), to z uwagi na liniowość operacji $T(f)$:

$$T \left(\Delta_{\lambda} f_1 - \Delta_{\lambda} f_2 \right) = 0;$$

ale ze wzoru definicyjnego (66) wynika, że równość powyższa pociąga za sobą

$$\Delta_{\lambda} f_1 - \Delta_{\lambda} f_2 = 0,$$

czyli

$$\Delta_{\lambda} (f_1 - f_2) = \Delta_{\lambda} N(t) = N(t + \lambda) - N(t) = 0,$$

gdzie

$$N(t) = f_1(t) - f_2(t).$$

Funkcja $N(t)$ jest więc funkcją okresową o okresie λ .

Wobec powyższej własności, wyniki rozważanej w niniejszym rozdziale operacji są „jednoznaczne” z dokładnością do funkcji okresowych o okresie λ . Należy podkreślić, że uzasadniona wyżej własność operacji $T f(t)$ jest naturalną cechą wszelkich operacji różnicowych definiowanych w zbiorze funkcji ciągłych. Jednoznaczność operacji $T f(t)$ można by np. uzyskać ograniczając dziedzinę operacji $T f(t)$ do zbioru funkcji schodkowych lub zakładając, że $t = 0, 1, 2, \dots$

Obliczmy obecnie wyniki niektórych prostszych operacji.

1. Załóżmy we wzorze definicyjnym (66) $f(t) = (a + 1)^{\frac{t}{\lambda}}$, wówczas

$$\Delta_{\lambda} f(t) = (a + 1)^{\frac{t + \lambda}{\lambda}} - (a + 1)^{\frac{t}{\lambda}} = a(a + 1)^{\frac{t}{\lambda}}$$

oraz

$$a(a + 1)^{\frac{t}{\lambda}} = p(a + 1)^{\frac{t}{\lambda}} - p \cdot 1$$

i ostatecznie

$$\boxed{(a+1) \frac{t}{\lambda} = \frac{p}{p-a} (1).} \quad (70)$$

2. Załóżmy we wzorze (66) $f(t) = t$, wówczas

$$\Delta_{\lambda} f(t) = t + \lambda - t = \lambda$$

oraz

$$p^{-1}(1) = \frac{t}{\lambda}.$$

Uogólniając uzyskaną zależność łatwo znajdziemy następujący wzór

$$p^{-v}(1) = \frac{1}{v!} \left(\frac{t}{\lambda} \right)^{(v)}, \quad (71)$$

gdzie wyrażenie $\left(\frac{t}{\lambda} \right)^{(v)}$ oznacza ważną w rachunku różnicowym tzw. „uogólnioną potęgę”, określoną przez wzór

$$\left(\frac{t}{\lambda} \right)^{(v)} = \frac{t}{\lambda} \left(\frac{t}{\lambda} - 1 \right) \left(\frac{t}{\lambda} - 2 \right) \dots \left(\frac{t}{\lambda} - v + 1 \right). \quad (72)$$

Można zauważyć, że dla uogólnionej potęgi słuszny jest następujący wzór

$$\Delta_{\lambda} \left(\frac{t}{\lambda} \right)^{(v)} = v \left(\frac{t}{\lambda} \right)^{(v-1)}, \quad (73)$$

przypominający różniczkowanie „zwykłej” potęgi $\left(\frac{t}{\lambda} \right)^v$.

Dla wyników operacji różnicowych można oczywiście ułożyć tablice analogiczne do tablic stosowanych w metodzie przekształcenia Laplace'a.

Zastosowanie omawianej operacji do rozwiązywania równań różnicowych ilustruje następujący

Przykład. Znaleźć rozwiązanie równania różnicowego

$$\Delta_{\lambda}^2 f(t) - 3 \Delta_{\lambda} f(t) + 2 f(t) = 0,$$

spełniające warunki początkowe

$$f(0) = 0 \quad \text{ i } \quad \Delta_{\lambda} f(0) = 1.$$

Rozwiązanie. Wobec wzorów (68) zamiast równania różnicowego rozwiązujemy następujące równanie operatorowe

$$p^2 f(t) - 3 p f(t) + 2 f(t) = p \cdot 1.$$

Zatem

$$f(t) = \frac{p}{p-2} (1) - \frac{p}{p-1} (1)$$

oraz, z uwagi na wzór (69):

$$f(t) = 3 \frac{t}{\lambda} - 2 \frac{t}{\lambda}.$$

3.3.2. Równanie różnicowe o zmiennych współczynnikach

Rozważmy obecnie jeszcze jedną operację różnicową. Operacji tej nie można wyprowadzić opierając się na przekształceniu Laplace'a. Weźmy mianowicie pod uwagę równanie

$$x \Delta y(x) = \varphi(x).$$

Równanie to przy założeniu warunku początkowego wyznacza jednoznacznie funkcję $y(x)$. Wobec tego wzór

$$T \varphi(x) = y(x) - y(1),$$

czyli

$$T x \Delta y(x) = y(x) - y(1), \quad (74)$$

określa endomorfizm T . W ten sposób otrzymujemy metodę, za pomocą której można rozwiązywać interesujące równania różnicowe o zmiennych współczynnikach, a mianowicie równania postaci

$$a_0(x+n\lambda) \dots (x+\lambda) x \Delta_{\lambda}^{n+1} y + \dots + a_{n-1}(x+\lambda) x \Delta_{\lambda}^2 y + a_n x \Delta_{\lambda} y + a_{n+1} y = f(x), \quad (75)$$

gdzie λ jest zadaną liczbą rzeczywistą i $a_0, \dots, a_n, a_{n+1}$ są stałymi dowolnymi.

Podstawiając we wzorze definicyjnym (74) $T^{-1} = p$ otrzymamy następujący związek

$$x \Delta_{\lambda} y = py - py(\lambda). \quad (76)$$

Okazuje się, że związek powyższy można uogólnić na różnice wyższych rzędów. Kładąc mianowicie

$$y = x \Delta_{\lambda} y_1 \quad (77)$$

mamy

$$\Delta_{\lambda} y = \Delta_{\lambda} [x \Delta_{\lambda} y_1] = (x + \lambda) \Delta_{\lambda}^2 y_1 + \lambda \Delta_{\lambda} y_1$$

oraz

$$x \Delta_{\lambda} y = (x + \lambda) x \Delta_{\lambda}^2 y_1 + \lambda x y_1. \quad (78)$$

Uwzględniając teraz wzory (77) i (78) w zależności (76) znajdziemy

$$(x + \lambda) x \Delta_{\lambda}^2 y_1 = p x \Delta_{\lambda} y_1 - \lambda p \Delta_{\lambda} y_1(1) - \lambda x \Delta_{\lambda} y_1 = (p - \lambda) x \Delta_{\lambda} y_1 - \lambda p \Delta_{\lambda} y_1(1),$$

oraz, po ponownym wykorzystaniu (76) i po podstawieniu na miejsce y^1 funkcji y :

$$(x + \lambda) x \Delta_{\lambda}^2 y = p(p - \lambda) y - p(p - \lambda) y(\lambda) - \lambda p \Delta_{\lambda} y(\lambda). \quad (79)$$

Drogą indukcji można wyprowadzić również wzór ogólny wyrażający funkcję $\prod_{i=0}^n (x + i\lambda) \Delta y^{n+1}(x)$ jako wynik operacji wykonanej nad funkcją $y(x)$. Otrzymamy mianowicie wzór następujący

$$\prod_{i=0}^n (x + i\lambda) \Delta^{n+1} y(x) = \prod_{i=0}^n (p - i\lambda) y - \sum_{v=0}^n \prod_{i=v+1}^n p(p - i\lambda) v! \lambda^v \Delta^v y(\lambda), \quad (80)$$

gdzie założono, że dla $i > n$

$$\prod_{i=v+1}^n p(p - i\lambda) = p.$$

Obliczymy obecnie wynik operacji

$$\frac{p}{p - a} (1). \quad (81)$$

Otóż zauważmy, że wynikiem operacji (81) jest funkcja spełniająca równanie

$$x \Delta_{\lambda} y(x) - a y(x) = 0 \quad (82)$$

i warunek początkowy $y(\lambda) = 1$.

Równanie powyższe można rozwiązać dla „dyskretnych” wartości zmiennej niezależnej x :

$$1\lambda, 2\lambda, 3\lambda, \dots, k\lambda, \dots$$

Przepisując równanie (82) w następującej postaci równoważnej

$$y(x + \lambda) = \left(1 + \frac{a}{x}\right) y(x),$$

otrzymamy następujące równości

$$y(2\lambda) = \left(1 + \frac{a}{x}\right) y(\lambda),$$

$$y(3\lambda) = \left(1 + \frac{a}{x}\right) y(2\lambda),$$

$$\dots \dots \dots$$

$$y(k\lambda) = \left(1 + \frac{a}{x}\right) y(k\lambda - \lambda).$$

W wyniku pomnożenia powyższych równości stronami, po skróceniu przez iloczyn

$$y(2\lambda) \cdot y(3\lambda) \dots y[(k-1)\lambda],$$

otrzymamy poszukiwane rozwiązanie w postaci wzoru

$$y(x) = y(\lambda) \prod_{t=1}^{x-1} \left(1 + \frac{\alpha}{\lambda t}\right).$$

Zatem, po uwzględnieniu warunku początkowego $y(\lambda) = 1$, otrzymamy:

$$\boxed{\frac{p}{p-\alpha}(1) = \prod_{t=1}^{x-1} \left(1 + \frac{\alpha}{\lambda t}\right).} \quad (83)$$

Wynik powyższej operacji szczególnie prosto wyraża się w przypadku, gdy $\alpha = \lambda$, mamy bowiem

$$\frac{p}{p-\lambda}(1) = x, \quad (83a)$$

Można zauważyć, że wynikiem operacji

$$p^{-1}(1)$$

jest funkcja spełniająca równanie

$$\Delta_{\lambda} y(x) = \frac{1}{x}; \quad y(\lambda) = 0.$$

Warto nadmienić, że równanie to nie posiada rozwiązania w zbiorze funkcji elementarnych (Gelfond [7], str. 311).

Sposób praktycznego wykorzystania rozważanej operacji ilustruje następujący

Przykład. Rozwiązać równanie różnicowe

$$(x+1)x \Delta_{\frac{1}{2}} y(x) - 2x \Delta_1 y(x) + 2y(x) = 0; \quad x \geq 1,$$

z warunkami początkowymi $y(1) = 0$, $\Delta y(1) = 1$.

Rozwiązanie. Korzystając ze związku (76) i (79) zamiast równania różnicowego rozwiązujemy następujące równanie operatorowe

$$p(p-1)y - 2py + 2y = p1.$$

Zatem

$$y = \frac{p}{p(p-1) - 2p + 2}(1) = \frac{p}{p-2}(1) - \frac{p}{p-1}(1).$$

i wobec równości (83) i (83a):

$$y(x) = \prod_{t=1}^{x-1} \left(1 + \frac{2}{t}\right) - x; \quad x \geq 1.$$

3.4. Operatory równań różnicowo - różniczkowych

Weźmy pod uwagę równanie różnicowo-różniczkowe o postaci

$$a_0 y^{(n)}(x+n) + \dots + a_{n-2} y'(x+1) + a_n y(x) = 0, \quad (84)$$

gdzie $a_0, \dots, a_{n-1}, a_n$ są stałymi dowolnymi.

Równania takie są zaliczane do klasy równań różniczkowych nieskończonego rzędu ze stałymi współczynnikami:

$$\sum_{v=0}^{\infty} \frac{A_v}{v!} y^{(v)}(x) = 0, \quad (85)$$

których funkcje charakterystyczne

$$\varphi(t) = \sum_{v=0}^{\infty} \frac{A_v}{v!} t^v$$

są funkcjami całkowitymi pierwszego rzędu. Równoważność równań (84) i (85) wynika ze wzoru Taylora.

Zbadajmy, jakie należy przyjąć założenia, by uzyskać jednoznaczność rozwiązań równań typu (84). Przede wszystkim umówmy się, że rozwiązaniem równania (84) będziemy nazywali każdą funkcję $y(x)$, która sprowadza do zera lewą stronę równania dla prawie wszystkich wartości x na półosi rzeczywistej $(0, \infty)$. Zbadajmy teraz, jak należałoby formułować warunki początkowe, by równanie (84) posiadało jedno i tylko jedno rozwiązanie.

Pisząc równanie (84) w postaci następującej

$$y^{(n)}(x+n) = \Phi[y(x), y'(x+1), \dots, y^{(n-1)}(x+n-1)],$$

łatwo można zauważyć, że jeśli odpowiednio w przedziałach $[0,1]$, $[1,2]$, $\dots$, $[n-2, n-1]$ założymy funkcję ciągłą $y(x)$ i jej pochodne $y'(x)$, $\dots$, $y^{(n-1)}(x)$, to tym samym na całej półosi rzeczywistej $[0, \infty)$ będzie jednoznacznie określona funkcja ciągła $y(x)$, spełniająca równanie (84). Wobec powyższego, dla uzyskania jednoznaczności rozwiązania równania różnicowo-różniczkowego nie wystarcza założyć wartości funkcji $y(x)$ i pochodnych $y^{(v)}(x)$ w punkcie, lecz należy z góry narzucić przebieg $y(x)$ i $y^{(v)}(x)$ w przedziałach jednostkowych $[0,1]$, $\dots$, $[n-2, n-1]$.

Poniżej rozważona jest możliwość efektywnego całkowania równań różnicowo-różniczkowych za pomocą metody opartej na jednoznacznej operacji funkcjonalnej $T y(x)$.

Założmy, że X jest zbiorem funkcji $y(x)$, ciągłych i różniczkowalnych na całej osi rzeczywistej x . Na zbiorze tym definiujemy operację liniową, następującą

$$T y'(x+1) = y(x) - y(0,1), \quad (86)$$

gdzie $y(0,1)$ oznacza funkcję ciągłą określoną przez wzór

$$y(0,1) = \begin{cases} y(x); & 0 \leq x < 1 \\ \text{const}; & x < 0; x \geq 1. \end{cases}$$

Podstawiając we wzorze definicyjnym (86) $p = T^{-1}$, otrzymamy

$$y'(x+1) = p[y(x) - y(0,1)] = py(x) - py(0,1). \quad (86a)$$

Dla wykazania jednoznaczności operacji (86a) uzasadnimy następującą

Własność. Równanie operatorowe

$$p[y(x) - y(0,1)] = 0; \quad x \geq 0 \quad (87)$$

przy zadanym przebiegu funkcji $y(x)$ w przedziale $x \in [0, 1]$ ma dla $x \geq 0$ jedno i tylko jedno „trywialne” rozwiązanie na zbiorze funkcji ciągłych:

$$y(x) = y(0,1).$$

Dowód. Ze wzoru definicyjnego (86a) wynika, że równanie (87) pociąga za sobą

$$y'(x+1) = 0; \quad 0 \leq x < \infty.$$

Mamy więc

$$y'(x) = 0; \quad 1 \leq x < \infty,$$

czyli

$$y(x) = \text{const}; \quad 1 \leq x < \infty.$$

Zatem, zgodnie z określeniem funkcji $y(0, 1)$

$$y(x) = y(0,1).$$

Funkcja $y(0, 1)$ jest więc jedynym rozwiązaniem równania (87) w zbiorze funkcji ciągłych.

Dzięki powyższej własności operacja $T y(x)$, zdefiniowana wzorem (86), jest operacją jednoznaczną¹⁾.

¹⁾Liniową niezależność układu funkcji $T^n y(x)$ można uzasadnić wyrażając wynik operacji $T y(x)$ za pomocą całki i korzystając następnie z całki Cauchy'ego.

Uogólniając wzór (86a) na pochodne wyższych rzędów, łatwo znajdziemy następujące związki

$$\begin{aligned} y''(x+2) &= p^2 y(x) - p^2 y(0,1) - p y'(1,2), \\ &\dots\dots\dots \\ y^{(v)}(x+v) &= p^v y(x) - \sum_{n=0}^{v-1} p^{v-n} y^{(n)}(n, n+1), \end{aligned} \quad (88)$$

gdzie $y^{(n)}(n, n+1)$ oznacza funkcję ciągłą określoną wzorem

$$y^{(n)}(n, n+1) = \begin{cases} y^{(n)}(x); & n \leq x < n+1 \\ \text{const}; & x < n; x \geq n+1. \end{cases} \quad (89)$$

Obliczymy obecnie wyniki niektórych prostszych operacji.

1. Weźmy pod uwagę operację

$$\frac{p}{p-1} (1). \quad (90)$$

Można zauważyć, że wynikiem powyższej operacji jest funkcja ciągła spełniająca równanie

$$y'(x+1) = y(x), \quad (91)$$

z warunkiem początkowym $y(x) = 1; x \in [0, 1]$.

Istotnie, stosując wzór (86a) do równania (91) otrzymamy, w założeniu $y(0, 1) = 1$:

$$p y(x) - y(x) = p \cdot 1.$$

Zatem

$$y(x) = \frac{p}{p-1} (1).$$

Łatwo można zauważyć, że rozwiązaniem równania (91) spełniającym warunek $y(x) = 1$ dla $x \in [0, 1]$ jest funkcja określona przez następujący ciąg

$$y(x) = \frac{x^n}{n!}; \quad n \leq x < n+1. \quad (92)$$

Jeśli (jak to jest ogólnie przyjęte) symbolem $[x]$ oznaczmy tzw. „całość z x ”, to funkcję (92) można będzie prosto zapisać wzorem następującym

$$y(x) = \frac{x^{[x]}}{[x]!}, \quad (93)$$

Nietrudno sprawdzić, że funkcja powyższa spełnia związek

$$\left(\frac{(x+1)^{[x+1]}}{[x+1]!} \right)' = \frac{x^{[x]}}{[x]!},$$

a więc faktycznie jest rozwiązaniem równania (91), spełniającym ponadto warunek

$$y(x) = 1; \quad x \in [0, 1].$$

Funkcja (93) jest więc wynikiem operacji (90):

$$\frac{p}{p-1}(1) = \frac{x^{[x]}}{[x]!}. \quad (94)$$

2. Obliczymy wynik operacji ogólniejszej

$$\frac{p}{p-a}(1), \quad (95)$$

gdzie a jest dowolną liczbą zespoloną.

Nietrudno zauważyć, że wynikiem operacji (95) jest funkcją $y(x)$ spełniająca równanie

$$y'(x+1) = ay(x), \quad (96)$$

z warunkiem początkowym $y(x) = 1; \quad x \in [0, 1]$.

Udowodnimy następujące

Twierdzenie. *Funkcję ciągłą spełniającą równanie (96) i warunek początkowy $y(x) = 1; \quad x \in [0, 1]$ jest funkcja określona wzorem*

$$P(a, x) = a^{[x]} \frac{x^{[x]}}{[x]!} - \sum_{v=0}^{[x-1]} a^v (a-1) \frac{x^v}{v!}. \quad (97)$$

Dowód. Ponieważ

$$P(a, x+1) = a^{[x+1]} \frac{(x+1)^{[x+1]}}{[x+1]!} - \sum_{v=0}^{[x]} a^v (a-1) \frac{x^v}{v!}$$

czyli

$$P'(a, x+1) = a a^{[x]} \frac{x^{[x]}}{[x]!} - a \sum_{v=0}^{[x-1]} a^v (a-1) \frac{x^v}{v!},$$

więc

$$P'(a, x+1) = a P(a, x),$$

zatem funkcja $P(a, x)$ spełnia równanie (96).

Dla wykazania, że $P(a, x)$ jest funkcją ciągłą, wystarczy uzasadnić, że różnica

$$P(a, n+0) - P(a, n-0)$$

jest równa zero dla każdego naturalnego n . Otóż ponieważ

$$P(a, n+0) = a^n \frac{n^n}{n!} - \sum_{\nu=0}^{n-1} a^\nu (a-1) \frac{n^\nu}{\nu!}$$

oraz

$$P(a, n-0) = a^{n-1} \frac{n^{n-1}}{(n-1)!} - \sum_{\nu=0}^{n-2} a^\nu (a-1) \frac{n^\nu}{\nu!},$$

więc

$$\begin{aligned} P(a, n+0) - P(a, n-0) &= a^n \frac{n^n}{n!} - a^{n-1} \frac{n^{n-1}}{(n-1)!} - a^{n-1} \frac{n^{n-1}}{(n-1)!} (a-1) = \\ &= a^{n-1} \frac{n^{n-1}}{(n-1)!} (a-1) - a^{n-1} \frac{n^{n-1}}{(n-1)!} \cdot (a-1) = 0. \end{aligned}$$

Zatem $P(a, x)$ jest funkcją ciągłą na całej półosi rzeczywistej $(0, \infty)$.

Funkcja $P(a, x)$ jest więc wynikiem operacji (95):

$$\boxed{\frac{p}{p-a}(1) = P(a, x).} \quad (98)$$

Ze wzoru (97) dla przypadku szczególnego $a = 1$ otrzymujemy poprzednio badaną funkcję $\frac{x^{[x]}}{[x]!}$:

$$P(1, x) = \frac{x^{[x]}}{[x]!},$$

która jest wynikiem operacji (94).

3. Obliczymy z kolei wyniki następujących operacji

$$a \frac{p}{p^2 + a^2}(1) \quad \text{ i } \quad \frac{p^2}{p^2 + a^2}(1). \quad (99)$$

Otóż, jeśli wprowadzimy oznaczenia $\sin r(a, x)$ i $\cos r(a, x)$ dla następujących funkcji ciągłych

$$\begin{aligned}\sin r(a, x) &= \frac{P(ja, x) - P(-ja, x)}{2j}, \\ \cos r(a, x) &= \frac{P(ja, x) + P(-ja, x)}{2},\end{aligned}\quad (100)$$

to rozbijając operacje złożone (99) na operacje prostsze typu (98), znajdziemy od razu

$$\begin{aligned}a \frac{p}{p^2 + a^2} (1) &= \sin r(a, x), \\ \frac{p^2}{p^2 + a^2} (1) &= \cos r(a, x).\end{aligned}\quad (101)$$

Funkcje $\sin r(a, x)$ i $\cos r(a, x)$ możemy wyznaczyć wprost ze wzoru (97). Po wykonaniu całkiem elementarnych przekształceń otrzymamy

$$\begin{aligned}\sin r(a, x) &= \sin \frac{\pi}{2} [x] a^{[x]} \frac{x^{[x]}}{[x]!} - \sum_{\nu=0}^{[x-1]} \varrho_a \sin \left(\nu \frac{\pi}{2} + \psi_a \right) a^\nu \frac{x^\nu}{\nu!}, \\ \cos r(a, x) &= \cos \frac{\pi}{2} [x] a^{[x]} \frac{x^{[x]}}{[x]!} - \sum_{\nu=0}^{[x-1]} \varrho_a \cos \left(\nu \frac{\pi}{2} + \psi_a \right) a^\nu \frac{x^\nu}{\nu!},\end{aligned}\quad (102)$$

gdzie

$$\operatorname{tg} \psi_a = -a; \quad \varrho_a = \sqrt{1 + a^2}.$$

Funkcje $\sin r(a, x)$ i $\cos r(a, x)$ są szczególnie ważne w zagadnieniach równań różnicowo-różniczkowych. Można zauważyć, że funkcje powyższe spełniają związki następujące:

$$\begin{aligned}\sin' r(a, x+1) &= a \cos r(a, x), \\ \cos' r(a, x+1) &= -a \sin r(a, x),\end{aligned}\quad (103)$$

przypominające różniczkowanie funkcji trygonometrycznych $\sin a x$ i $\cos a x$.

Ciągłość funkcji $\sin r(a, x)$ i $\cos r(a, x)$ wynika wprost z definicji (100); słuszne są więc związki

$$\begin{aligned}\sin r(a, n+0) &= \sin r(a, n-0), \\ \cos r(a, n+0) &= \cos r(a, n-0),\end{aligned}\quad (104)$$

gdzie n jest liczbą naturalną.

Ze wzoru (102) wynikają również następujące zależności:

$$\sin r(\alpha, 0) = 0; \quad \cos r(\alpha, 0) = 1. \quad (105)$$

Sposób praktycznego wykorzystania wyników podanych w niniejszym rozdziale ilustruje następujący przykład.

Przykład. Rozwiązać równanie

$$y'''(x+3) - 2y''(x+2) + 9y'(x+1) - 18y(x) = 0$$

z warunkami początkowymi

$$y(0, 1) = y'(1, 2) = 0; \quad y''(2, 3) = 1.$$

Rozwiązanie. Korzystając z zależności (88), otrzymamy

$$p^3 y - 2p^2 y + 9p y - 18y = p1.$$

Zatem

$$y = \frac{p}{p^3 - 2p^2 + 9p - 18} (1) = \frac{1}{13} \cdot \frac{p}{p-2} (1) - \frac{1}{13} \frac{p^2}{p^2 + 3^2} (1) - \frac{2}{13} \cdot \frac{p}{p^2 + 3^2} (1).$$

i na mocy wzorów (98), (101):

$$y(x) = \frac{1}{13} P(2, x) - \frac{1}{13} \cos r(3, x) - \frac{2}{13} \sin r(3, x).$$

3.5. Operatory pseudonieliniowe

Niech X będzie zbiorem funkcji ciągłych $y(x)$ określonych na półosi rzeczywistej $[0, \infty)$. Określmy na tym zbiorze następującą operację

$$Tf(y) = \int_0^x f(y) g'(\xi) d\xi; \quad y = y(x), \quad (106)$$

gdzie $g(x)$ (funkcja wagi) jest z góry zadaną funkcją ciągłą i różniczkowalną oraz spełniającą warunek $g(0) = 0$. Zakładamy przy tym, że całka (106) istnieje.

Operacja powyższa jest oczywiście operacją liniową, a więc spełnia związki (2). Można wykazać, że spełnia ona również inne warunki postawione w 2 rozdziale niniejszej pracy. Operator T określony wzorem (106) nazywać będziemy operatorem pseudonieliniowym.

Jeśli założymy, że funkcja $y(x)$ jest funkcją różniczkowalną, to na podstawie wzoru (106) otrzymamy

$$\frac{1}{g'(x)} f'_y y'_x = p f(y) - p f_0,$$

gdzie $f_0 = f[y(0)] = f(y)|_{x=0}$