

POLSKA AKADEMIA NAUK
INSTYTUT PODSTAWOWYCH PROBLEMÓW TECHNIKI

ROZPRAWY ELEKTROTECHNICZNE

TOM VII • ZESZYT 1

KWARTALNIK

PAŃSTWOWE WYDAWNICTWO NAUKOWE
WARSZAWA 1961

RADA REDAKCYJNA

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

KOMITET REDAKCYJNY

Redaktor Naczelny

PROF. WITOLD NOWICKI

Sekretarz

JULIUSZ MIERZEJEWSKI

ADRES REDAKCJI

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

PAŃSTWOWE WYDAWNICTWO NAUKOWE
Warszawa, Miodowa 10

Nakład 640 (501 + 139)	Oddano do składania 1.XI.1960
Ark. wyd. 11,75, ark. druk. 10,5	Podpisano do druku w marcu 1961
Papier druk. sat. kl. III, 80 g, 70×100	Druk. ukończono w marcu 1961
Zamówienie nr 2318/60	Cena zł 25.— S-18

Drukarnia im. Rewolucji Październikowej, Warszawa

621.317.374:621.319.4

ANDRZEJ JELLONEK

Możliwości realizacji wzorców tangensa kąta strat kondensatorów

Rękopis dostarczono 10.7.1959

Tangens kąta strat kondensatorów jest wielkością, którą mierzy się tak często i tak różnorodnymi przyrządami, że wygodne byłoby rozporządzanie odpowiednimi wzorcami tej wielkości umożliwiającymi uzgodnienie wyników pomiarów pochodzących z różnych ośrodków. Przyrządy takie można zrealizować jako wzorce stałej lub regulowanej wartości tangensa kąta strat, o stałej lub o regulowanej pojemności. We wszystkich tych przypadkach przyrząd składa się z możliwie bezstratnych kondensatorów i bezpojemnościowych, i bezindukcyjnych oporników połączonych z sobą w odpowiedni sposób. Możliwe jest również zastąpienie oporności skupionych elektrodami oporowymi albo też użycie kondensatora z dielektrykiem o stałych i znanych stratach. Skonstruowane w ten sposób wzorce pracują w sposób zadowalający w obrębie średnich wartości znamionowych $\tan \delta$, oraz w pasmie małych częstotliwości. Granice pracy co do częstotliwości i co do dokładności zakresłone są wpływem działania współczynników resztkowych użytych podzespołów. Dotychczasowe wykonania zarówno o stałej, jak i o regulowanej wartości $\tan \delta$ odtwarzają pojemności od 100 pF do 1000 pF i $\tan \delta$ od $10 \cdot 10^{-4}$ do $10^3 \cdot 10^{-4}$, przy czym dokładność pojemności jest około 0,5%, a $\tan \delta$ nie gorsza od kilku procentów.

1. WSTĘP

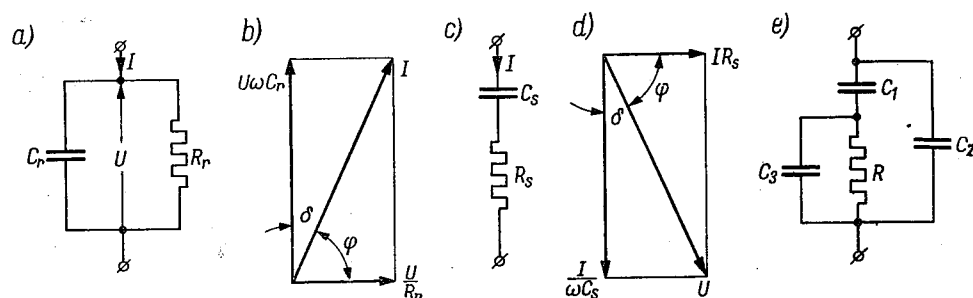
Tangens kąta strat ($\tan \delta$) jest wielkością, która jest mierzona zawsze w przypadku badania właściwości materiałów dielektrycznych i podzespołów, których działanie oparte jest na właściwościach dielektryków, jak kable, kondensatory itp. Wartość $\tan \delta$ mierzona jest również często w czasie sprawdzania właściwości niektórych urządzeń elektrycznych, jak transformatory, generatory, motory itp.; wielkość kąta strat i jego zmiany stanowią bowiem czuły wskaźnik zjawisk zachodzących w każdym urządzeniu zawierającym materiały dielektryczne.

Dlatego pomiary $\tan \delta$ wykonywane są masowo; stosuje się przy tym rozmaite metody i przyrządy: układy mostkowe i obwody rezonansowe przy małej częstotliwości i wielkiej częstotliwości, linie długie i wnęki w pasmie mikrofalowym, metody kalorymetryczne w pomiarach bezwzględnych o charakterze laboratoryjnym itp. Ta różnorodność stosowanych metod i używanych przyrządów jak również duża zależność wartości $\tan \delta$

od warunków pomiaru i otoczenia, przede wszystkim od temperatury i częstotliwości, w końcu masowość wykonywanych pomiarów i związane z tym posługiwanie się personelem przyuczonym, powodują występowanie znacznych rozbieżności w wynikach otrzymywanych dla tych samych materiałów przez różnych eksperymentatorów i w rozmaitych instytucjach.

Z tych względów odczuwa się potrzebę wprowadzenia odpowiednio skonstruowanych wzorców $\tan \delta$, które mogłyby być wskaźnikiem poprawności zastosowanej metody i przyrządu oraz wykonanego pomiaru.

Koncepcja wzorca $\tan \delta$ wynika z definicji tej wielkości i obrazującego ją układu zastępczego. Kąt δ podaje przesunięcie fazowe pomiędzy napię-



Rys. 1. Możliwości przedstawiania własności kondensatora

a) Układ zastępczy równoległy, b) Wykres wektorowy odpowiadający układowi a), c) Układ zastępczy szeregowy, d) Wykres wektorowy odpowiadający układowi c), e) Układ zastępczy szeregowy, z uwzględnieniem pojemności resztkowych podzespołów

ciem lub prądem w bezstratnym, idealnym kondensatorze i spadkiem napięcia — lub prądem — na kondensatorze rzeczywistym, obciążonym stratami (rysunki 1b i 1d).

2. WZORCE O STAŁEJ WARTOŚCI $\tan \delta$ I STAŁEJ POJEMNOŚCI

Najprostszy wzorec tangensa kąta strat składa się z opornika odwzorowującego straty i z możliwie bezstratnego kondensatora obrazującego zachowanie się idealnego dielektryka. Podzespoły te mogą być połączone równoległe lub szeregowo (rysunki 1a i 1c); oba układy są równorzędne ze względu na zasadę działania wzorca; natomiast ten lub inny układ może poprawniej odtwarzać zjawiska zachodzące w danych warunkach w dielektryku; jeden z układów może też być wygodniejszy w przypadku konstruowania wzorca przy założeniu konkretnych warunków jego pracy, np. dla określonej wartości pojemności itp. Zamiast oddzielnego opornika można również zastosować elektrody o dużej oporności. Tego rodzaju elektrody uzyskuje się np. przez napylenie cieplne w próżni cienkich warstw Ni-chromu lub manganinu. Działanie takiego kondensatora można zobrazować układem zastępczym o rozłożonej równomiernie pojemności i oporności.

Wartości tangensa kąta strat wynikające z przytoczonych na rys. 1 układów zastępczych mają postać:
układ równoległy

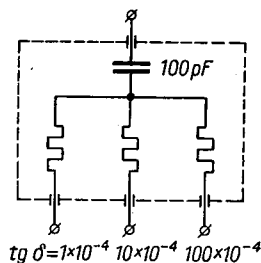
$$\operatorname{tg} \delta = \frac{1}{R_r C_r \omega},$$

układ szeregowy

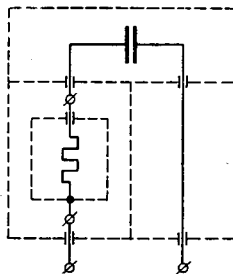
$$\operatorname{tg} \delta = R_s \cdot C_s \cdot \omega,$$

przy czym ω oznacza pulsację prądu lub napięcia.

Regulację wartości $\operatorname{tg} \delta$ przy stałej pulsacji i nieziennej pojemności uzyskuje się w takim wykonaniu przez przełączanie lub wymianę oporników. W wykonaniu zastosowanym np. we wzorcu firmy angielskiej¹⁾ po-



Rys. 2. Wzorec $\operatorname{tg} \delta$ z przełączanymi oporami, pochodzenia angielskiego, o danych nominalnych $C = 100$ pF, $\operatorname{tg} \delta = (1, 10, 100) \cdot 10^{-4}$, $U_{\max} = 10$ kV, $\omega = 2\pi \cdot 50$ Hz



Rys. 3. Wzorec $\operatorname{tg} \delta$ z wymienianymi oporami: C i $\operatorname{tg} \delta$ dowolne, $U_{\max} = 300$ V (konstr. P.T.B. wzgl. K.M.E.)

szczególnym zakresom $\operatorname{tg} \delta$ odpowiadają różne oporniki (rys. 2) oraz oddzielne zaciski wejściowe wzorca. W konstrukcji takiej trudno uniknąć zmian pojemności wzorca występujących przy przełączaniu zakresów. Korzystniejsze wydaje się rozdzielenie wzorca na części zawierające kondensator i oporniki i wymiana tej drugiej części przy zmianie zakresów $\operatorname{tg} \delta$ (rys. 3). Dodatkowe zastosowanie bezpojemnościowegołączenia obu części [1] zapewnia stałość pojemności układu. Wykonania takie zaproponowane i skonstruowane zostały niezależnie przez K.M.E.²⁾ i Phys. techn. Bundesanstalt [2], [3]. Ostatnia praca zawiera również szczegółową analizę pożądaných warunków pracy i dokładności takiego wzorca.

W rozumowaniu dotychczasowym pomijano straty własne kondensatora oraz działanie pojemności (C_3) bocznikującej opornik i pojemności (C_2) zacisków i podzespołów względem otoczenia (rys. 1e). Przy uwzględnieniu pojemności C_2 i C_3 występujących zawsze w wykonaniach

¹⁾ British Electricity Laboratories — Power factor unit instrument type 657.

²⁾ Kat. Miernictwa Elektronowego Politechniki Wrocławskiej.

rzeczywistych tangens kąta strat czynny na zaciskach przyrządu ma wartość [3]:

$$\operatorname{tg} \delta = \frac{\omega R C_1^2}{C_2 [1 + \omega^2 R^2 (C_1 + C_3)^2] + C_1 [1 + \omega^2 R^2 C_3 (C_1 + C_3)]}.$$

Wpływ pojemności resztkowej opornika można pominąć [3], jeżeli:

$$\omega^2 R^2 (C_1 + C_3)^2 \ll 1 \quad \text{oraz} \quad \omega^2 R^2 C_3 (C_1 + C_3) \ll 1.$$

Decydujący wpływ ma zwykle warunek pierwszy. Można go jednak prawie zawsze spełnić w konstrukcji wzorców przeznaczonych do pracy w obrębie małej częstotliwości, przy której wzorce $\operatorname{tg} \delta$ są najczęściej używane.

Przykład: Dla wzorca o wartościach znamionowych $\operatorname{tg} \delta = 1000 \cdot 10^{-4}$, $C_1 = 1000$ pF, $f = 10$ kHz należy zastosować opornik $R = 1595 \Omega$, którego pojemność własna $C_3 \leq 10$ pF, a $C_2 \approx 20$ pF. Poprawki wartości $\operatorname{tg} \delta$ względem układu idealnego wynoszą:

$$\omega^2 R^2 (C_1 + C_3)^2 \approx 1,04 \cdot 10^{-2} \ll 1; \quad \omega^2 R^2 C_3 (C_1 + C_3) \approx 0,8 \cdot 10^{-4} \ll 1.$$

Przybliżona wartość kąta strat po pominięciu wpływu pojemności C_3 jest:

$$\operatorname{tg} \delta \approx \omega R C_1 \frac{1}{1 + \frac{C_2}{C_1}}.$$

Stratność własna kondensatora wzorcowego ma najczęściej również niewielki wpływ na wartość wzorca $\operatorname{tg} \delta$ stosowanego przy m. cz. Stratność tę można przedstawić w postaci dodatkowej oporności (ΔR) załączonej w szereg z kondensatorem. Wartość $\operatorname{tg} \delta$ występująca na końcówkach takiego zespołu jest:

$$\operatorname{tg} \delta = (R + \Delta R) \omega C = R \omega C + \Delta R \omega C = \operatorname{tg} \delta' + \operatorname{tg} \delta''.$$

Pierwszy człon tego wyrażenia przedstawia straty układu kondensatora bezstratnego załączonego w szereg z opornikiem R ; człon drugi reprezentuje straty kondensatora. Wpływ oporności ΔR należy natomiast uwzględnić jako poprawkę przy obliczaniu wartości oporności dodatkowej:

$$R = \frac{\operatorname{tg} \delta - \operatorname{tg} \delta''}{\omega C} \quad \text{zamiast wartości} \quad R = \frac{\operatorname{tg} \delta}{\omega C}.$$

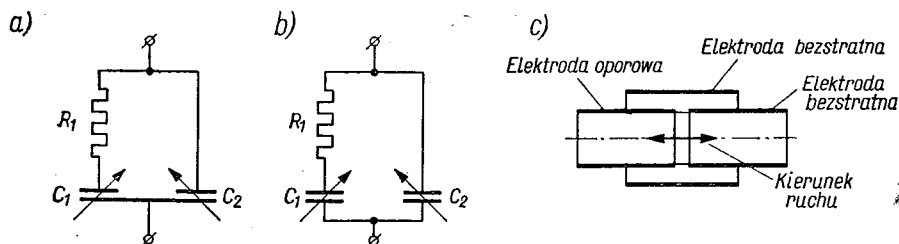
Jeżeli kondensator wzorcowy jest odpowiednio skonstruowany, a dielektrykiem jego jest powietrze, wówczas można stosunkowo nietrudno uzyskać jego stratność $\operatorname{tg} \delta \leq 1 \cdot 10^{-4}$. Uwzględnianie stratności własnej kondensatora jest zatem konieczne w przypadku wykonywania wzorców o wartościach $\operatorname{tg} \delta \leq 10 \cdot 10^{-4}$. Trudność konstruowania wzorców o małych wartościach stratności polega zatem na uzyskaniu kondensatora o odpowiednio małych i stałych stratach. Trudny do spełnienia jest zwłaszcza drugi warunek, ponieważ stratność kondensatora zależy i od jego budowy

i od trudno uchwytnych czynników zewnętrznych, jak starzenie się dielektryka użytego jako izolatory, wilgotność otoczenia, przyczepność do elektrod kurzu i pary wodnej itp. Ze względu na te trudności niechętnie wykonuje się wzorce $\operatorname{tg} \delta$ o wartościach znamionowych mniejszych od $10 \cdot 10^{-4}$; również dokładność wykonanych wzorców maleje szybko wraz ze zmniejszaniem się wartości znamionowej.

3. WZORCE O WARTOŚCI $\operatorname{tg} \delta$ REGULOWANEJ W SPOSÓB CIĄGŁY I O STAŁEJ POJEMNOŚCI

Wzorzec o wartości $\operatorname{tg} \delta$ regulowanej w sposób ciągły i o stałej pojemności można uzyskać najłatwiej przy pomocy dwu kondensatorów o pojemnościach zmieniających się odwrotnie, z których jeden jest praktycznie bezstratny, drugi posiada odpowiednio duży $\operatorname{tg} \delta$. Układ taki można zrealizować albo przy pomocy kondensatora różnicowego, albo też dwoma oddzielnymi kondensatorami połączonymi równolegle, o rotorach przedstawionych o 180° , sprzężonych na wspólnej osi. W obu przypadkach opornik reprezentujący straty załączony jest w szereg z częścią regulowanej pojemności (rys. 4).

Kondensatory wzorcowe wykonane są w ten sposób, że ich pojemności C_1 lub też C_2 zmieniają się w czasie regulowania proporcjonalnie do



Rys. 4. Wzorzec $\operatorname{tg} \delta$ o regulowanej w sposób ciągły wartości $\operatorname{tg} \delta$ i o stałej pojemności

a) Wykonanie z kondensatorem różnicowym, b) Wykonanie z dwoma kondensatorami o rotorach sprzężonych mechanicznie, przestawionych o 180° , c) Wykonanie z elektrodą oporową

kąta obrotu (φ), przy czym zmiana pojemności przypadająca na 1° kąta obrotu rotora wynosi α wzgl. β pF/działkę. Wartość aktualnej pojemności każdej z części można określić wzorem:

$$C_1 = C_{1\min} + \alpha \cdot \varphi \quad \text{oraz} \quad C_2 = C_{2\max} - \beta \varphi.$$

Jeżeli oba kondensatory wykonane są możliwie identycznie, to można założyć, że $\alpha = \beta$; wówczas pojemność na zaciskach wzorca ma wartość: $C = C_1 + C_2 = C_{2\max} + C_{1\min} = K = \text{const.}$ Wypadkowy tangens kąta strat można w tych warunkach obliczyć ze wzoru:

$$\operatorname{tg} \delta = \frac{C_1 \operatorname{tg} \delta_1 + C_2 \operatorname{tg} \delta_2^1)}{C_1 + C_2}.$$

¹⁾ Przy czym $\operatorname{tg} \delta_1$ i $\operatorname{tg} \delta_2$ przedstawiają zastępcze wartości tg kąta strat każdej z równoległych gałęzi.

Jeżeli dodatkowo oba kondensatory są praktycznie bezstratne, tj: $\operatorname{tg} \delta_2 \approx 0$ oraz $\operatorname{tg} \delta_1 = R_1 C_1 \omega$, wówczas przybliżona wartość wypadkowej stratności jest:

$$\begin{aligned} \operatorname{tg} \delta &\approx \frac{C_1}{C_1 + C_2} \operatorname{tg} \delta_1 = \frac{C_1^2}{C_1 + C_2} R \cdot \omega = \frac{(C_{1\min} + a\varphi)^2}{K} R \cdot \omega = \\ &= (C_{1\min}^2 + 2 \cdot C_{1\min} \cdot a \cdot \varphi + a^2 \cdot \varphi^2) \left(\frac{R \omega}{K} \right). \end{aligned}$$

Zależność tangensa kąta strat i kąta obrotu rotora ma zatem różny charakter na poszczególnych swych odcinkach. I tak, jeżeli:

$$C_{1\min} \ll a \cdot \varphi,$$

wówczas:

$$\operatorname{tg} \delta \approx \frac{a^2 \cdot R \cdot \omega}{K} \varphi^2 = \operatorname{const} \cdot \varphi^2.$$

Natomiast w przypadku małych zmian pojemności i dużej wartości pojemności początkowej, tj. dla: $C_{1\min} \gg a \cdot \varphi$,

$$\operatorname{tg} \delta \approx \frac{C_{1\min}}{K} R \cdot \omega = \operatorname{const}.$$

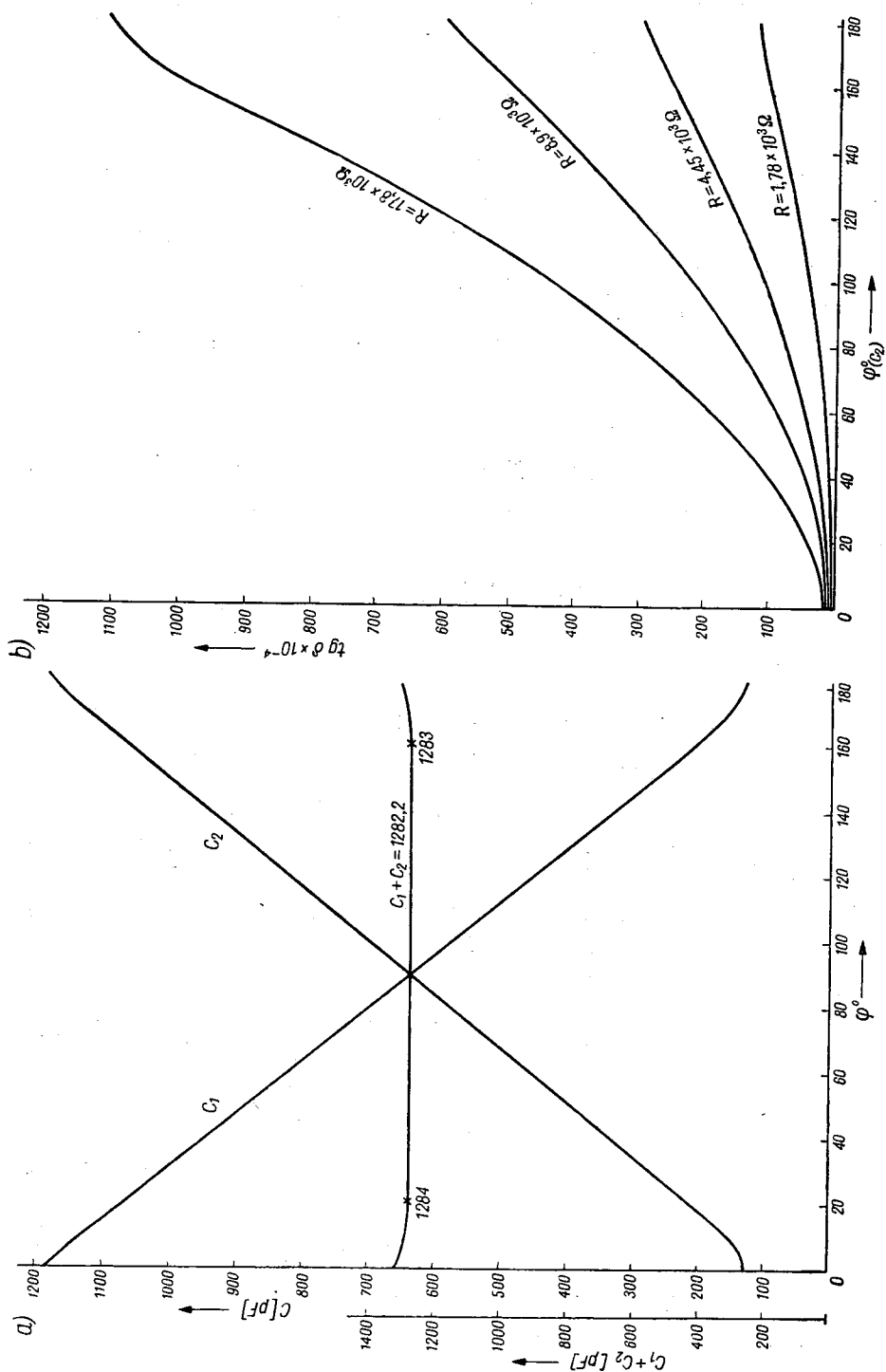
W końcu w zakresie przejściowym, w którym:

$$2 \cdot C_{1\min} a \cdot \varphi \gg C_{1\min}^2 \quad \text{oraz} \quad 2 \cdot C_{1\min} a \cdot \varphi \gg a^2 \varphi^2, \quad \text{tj.} \quad \frac{C_1}{2} < a \cdot \varphi < 2 C_1$$

$$\operatorname{tg} \delta \approx 2 \cdot C_{1\min} a \cdot \varphi \frac{R \omega}{K} = \operatorname{const} \cdot \varphi.$$

Wartość tangensa kąta strat wzrasta zatem przeważnie proporcjonalnie do kwadratu kąta obrotu rotora kondensatora załączonego w gałęzi zawierającej opór dodatkowy; jedynie w obrębie niewielkich, początkowych zmian pojemności zależność ta ma teoretycznie inny charakter; w rzeczywistych wykonaniach tak początek, jak i koniec krzywej przedstawiającej powyższą zależność mają przebiegi różniące się znacznie od wartości teoretycznych. Zjawisko to spowodowane jest wpływem pojemności szcztątkowych rotora do obudowy i niestałością wartości pojemności $C = C_1 + C_2$ (rys. 5).

W rozumowaniu dotychczasowym pominięto wpływ strat własnych obu kondensatorów wzorcowych oraz pojemności własnej (C_3) opornika. Jak wykazano w dyskusji działania wzorca stałego, wpływ pojemności resztkowej opornika jest prawie zawsze do pominięcia w przypadku pracy wzorca w pasmie m. cz. Nieco większą poprawkę względem wartości przybliżonej mogą wprowadzić natomiast straty obu kondensatorów wzor-



Rys. 5. Krzywe skalowania wzorca $\tan \delta$ o ciągłej zmianie $\tan \delta$ i o stałej pojemności
a) Wartości pojemności, b) Wartości $\tan \delta$

cowych. Składają się one ze strat w dielektryku, z którego utworzone są izolatory oddzielające elektrody kondensatora, oraz ze strat na ciepło Joula powstających w częściach metalowych, tj. w elektrodach i doprowadzeniach. W przypadku pracy wzorca w pasmie m. cz. straty pochodzące z drugiego źródła można praktycznie pominąć [2]. Każdy z kondensatorów układu można zatem rozpatrywać jako złożony z dwu równolegle połączonych części: stałej, obciążonej stratami w dielektryku (C_{min}), i regulowanej ($\alpha \cdot \varphi$), praktycznie bezstratnej, bo zawierającej powietrze jako dielektryk. Można zatem napisać:

$$\operatorname{tg} \delta_2 = \frac{C_{2min} \operatorname{tg} \delta_{2izol} + \alpha \cdot \varphi \operatorname{tg} \delta_{2pow}}{C_2} \approx \frac{C_{2min}}{C_2} \operatorname{tg} \delta_{2izol},$$

$$\operatorname{tg} \delta_1 = (R_1 + R_{1izol}) \omega \cdot C_1 = R_1 \cdot \omega \cdot C_1 + R_{1izol} \cdot \omega \cdot C_1 = \operatorname{tg} \delta'_1 + \operatorname{tg} \delta''_1,$$

przy czym $\operatorname{tg} \delta'_1$ oznacza stratność układu kondensatora idealnego, bezstratnego o pojemności C_1 załączonej w szereg z opornikiem dodatkowym R_1 , natomiast $\operatorname{tg} \delta''_1$ jest stratnością własną kondensatora C_1 . Ale

$$\operatorname{tg} \delta''_1 = \frac{C_{1min} \operatorname{tg} \delta_{1izol} + \alpha \cdot \varphi \operatorname{tg} \delta_{1pow}}{C_1} \approx \frac{C_{1min}}{C_1} \operatorname{tg} \delta_{1izol};$$

podstawiając:

$$C_1 = C_{1min} + \alpha \cdot \varphi; \quad C_2 = C_{2max} - \beta \cdot \varphi; \quad \alpha \approx \beta; \quad C_1 + C_2 = K$$

otrzymuje się ostatecznie:

$$\operatorname{tg} \delta = \frac{C_1 \operatorname{tg} \delta_1 + C_2 \operatorname{tg} \delta_2}{C_1 + C_2} = \frac{C_{1min}}{K} \operatorname{tg} \delta_{1izol} + \frac{C_{2min}}{K} \operatorname{tg} \delta_{2izol} + \frac{(C_{1min} + \alpha \cdot \varphi)^2}{K} R_1 \omega.$$

Wzór ten upraszcza się do postaci analogicznej jak podana w przykładzie liczenia uproszczonego, jeżeli założyć się:

$$\operatorname{tg} \delta_{izol1} = \operatorname{tg} \delta_{izol2} = 0 \quad \text{albo:} \quad C_{1min} \ll K, \quad C_{2min} \ll K.$$

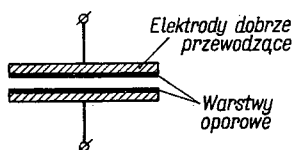
Wzorce tego rodzaju wykonywane są albo jako kondensatory różnicowe, walcowe [3], albo też w postaci kondensatora obrotowego [4]. W Katedrze Miernictwa Elektronowego wykonywane są również wzorce walcowe, w których zastąpiono opór skupiony — oporowymi elektrodami (rys. 9).

Jak się wydaje, wykonanie walcowe zapewnia lepszą dokładność i stałość pojemności; może być ono jednak zrealizowane tylko w przypadku stosunkowo niewielkich pojemności znamionowych, natomiast konstrukcja drugiego rodzaju jest łatwiejsza do wykonania również przy znacznych wartościach pojemności. Krzywą skalowania wzorca tego rodzaju¹⁾ podaje rys. 5.

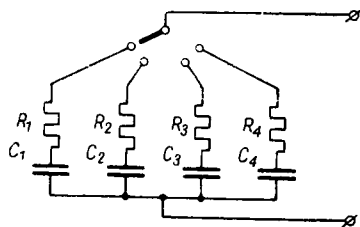
¹⁾ Zawierającego kondensator różnicowy.

4. WZORCE O REGULOWANEJ W SPOSÓB CIĄGŁY POJEMNOŚCI I STAŁYM $\operatorname{tg} \delta$

Wzorzec o stałym $\operatorname{tg} \delta$ i regulowanej pojemności można uzyskać zmieniając opór R w czasie regulacji pojemności C w ten sposób, aby iloczyn $R \cdot C = A$ pozostawał stały, a tym samym stały był $\operatorname{tg} \delta = RC\omega$. Układ tego rodzaju można zrealizować m. in. przez zastosowanie elektrod kondensatora wzorcowego o dużej oporności właściwej [4]. Możliwe są dwa rodzaje konstrukcji tego rodzaju. W pierwszej z nich oporowa warstwa elektrod ułożona jest na dobrze przewodzącym



Rys. 6. Zasada konstrukcji wzorca o stałym $\operatorname{tg} \delta$ i regulowanej pojemności



Rys. 7. Układ zastępczy wzorca o stałym $\operatorname{tg} \delta$ i regulowanej pojemności

podłożu (rys. 6). Prąd elektryczny dopływający przez przewody doprowadzające rozplywa się równomiernie na całą powierzchnię czynną. Układ zastępczy (rys. 7) kondensatora można sobie wyobrazić jako złożony z dużej ilości kondensatorów elementarnych połączonych równolegle, których powierzchnie, pojemności i opory szeregowe mają identyczne wartości. Tangens kąta strat takiego układu ma wartość:

$$\operatorname{tg} \delta = \frac{C_1 \operatorname{tg} \delta_1 + C_2 \operatorname{tg} \delta_2 + \dots C_n \operatorname{tg} \delta_n}{C_1 + C_2 + \dots C_n}.$$

Jeżeli $C_1 = C_2 = \dots C_n$ oraz $R_1 = R_2 = \dots R_n$,

a więc i $\operatorname{tg} \delta_1 = \operatorname{tg} \delta_2 = \dots \operatorname{tg} \delta_n = \operatorname{tg} \delta$,

wówczas $\operatorname{tg} \delta = \frac{n \operatorname{tg} \delta_n}{n} = \operatorname{tg} \delta_n = \operatorname{tg} \delta$.

Konkretna wartość $\operatorname{tg} \delta$ wzorca tego rodzaju jest $\operatorname{tg} \delta = RC\omega$, przy czym C oznacza pojemność przypadającą na $S \text{ cm}^2$ czynnej powierzchni elektrod, a

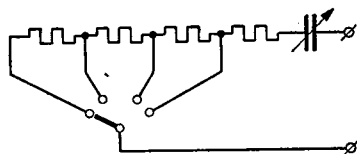
$$R/\Omega = \frac{\rho/\Omega \text{ cm}^2}{\text{cm}} \cdot g/\text{cm} \cdot \frac{1}{S/\text{cm}^2},$$

gdzie g jest grubością warstwy oporowej.

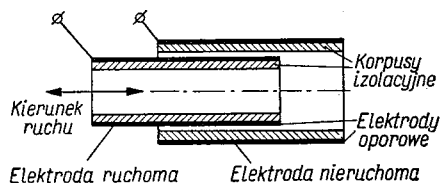
Wartość ta jest niezależna od ilości czynnych kondensatorów cząstkowych, a więc i od regulacji wartości pojemności wypadkowej. Trudność przy konstrukcji wzorca $\operatorname{tg} \delta$ o działaniu opartym na powyższej zasadzie

stanowi uzyskanie jednolitej, nie zmieniającej swych własności w czasie warstwy oporowej, o dostatecznie dużej oporności właściwej. Poza tym wartość $\operatorname{tg} \delta$ wzorca zależy od częstotliwości, jak w każdym układzie z szeregowo połączonymi: opornością i pojemnością.

W wykonaniu drugim przełączane są oporności szeregowe, skupione równocześnie z regulacją wartości pojemności i to w taki sposób, aby iloczyn RC pozostawał stały (rys. 8). Jeżeli powierzchnia czynna kondensatora jest S , to zmiany pojemności ΔC są proporcjonalne do zmian tej powierzchni ($\Delta C = B \cdot \Delta S$); wobec tego zmiany oporności odwrotnie proporcjonalne do zmian powierzchni czynnej ($\Delta R = D \frac{1}{\Delta S}$) pozwalają na uzyskanie stałej wartości $\operatorname{tg} \delta$ przy regulowanej pojemności. Układ o takich własnościach uzyskuje się np. w kondensatorze walcowym, w którym



Rys. 8. Druga odmiana zasady konstrukcji wzorca stałego $\operatorname{tg} \delta$ i regulowanej pojemności



Rys. 9. Układ wzorca stałego $\operatorname{tg} \delta$ i regulowanej pojemności — schemat konstrukcji

elektrody wykonane są w sposób specjalny z materiału o dużej oporności właściwej (rys. 9).

Wszystkie omówione dotychczas metody uzyskiwania stałej pojemności i regulowanego $\operatorname{tg} \delta$ lub stałego $\operatorname{tg} \delta$ przy regulowanej pojemności mają wspólną wadę: zależność wartości $\operatorname{tg} \delta$ od częstotliwości. Wadę tę można usunąć budując wzorce z kondensatorów o stałych $\operatorname{tg} \delta$, niezależnych od częstotliwości i od kątów obrotu rotorów. I tak: Wzorzec o stałej pojemności i regulowanym $\operatorname{tg} \delta$ uzyskuje się łącząc równolegle dwa kondensatory: bezstratny ($\operatorname{tg} \delta_2 \approx 0$) i stratny ($\operatorname{tg} \delta_1 = \text{const}$). Wówczas:

$$\operatorname{tg} \delta = \frac{C_1 \operatorname{tg} \delta_1 + C_2 \operatorname{tg} \delta_2}{C_1 + C_2} \approx \frac{C_1}{C_1 + C_2} \operatorname{tg} \delta_1 = \frac{1}{1 + \frac{C_2}{C_1}} \operatorname{tg} \delta_1.$$

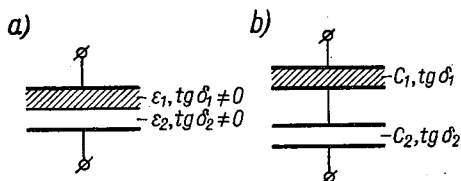
Jeżeli układ wykona się różnicowo, tak że $C_1 + C_2 = \text{const} = K$, a wartość $\operatorname{tg} \delta$ jest praktycznie niezależna od kąta obrotu rotora, wówczas:

$$\operatorname{tg} \delta = \text{const} \cdot C_1 = K(C_{\min} + \alpha \varphi) = \text{const} + \text{const} \cdot \varphi.$$

Stała regulacji $\operatorname{tg} \delta$ ma charakter liniowy, wartość skalowania jest niezależna od częstotliwości.

Konstrukcyjnie wzorzec taki można zrealizować z dwu kondensatorów regulowanych, o identycznych układach elektrod, z których jeden ma jako dielektryk powietrze, drugi dielektryk stratny. Ten ostatni kondensator najłatwiej uzyskać zanurzając odpowiednio mniejszą ilość płytek elektrodowych w cieczy o dostatecznie dużym $\operatorname{tg} \delta$. Trudność stanowi znalezienie cieczy dostatecznie stratnej, o $\operatorname{tg} \delta$ niezależnym od pulsacji. Dlatego wygodniejsze jest użycie dielektryka stałego, np. odpowiedniej sztucznej żywicy, który wypełnia częściowo lub możliwie całkowicie przestrzeń międzyelektrodową.

W przypadku kondensatorów rzeczywistych, o ruchomym jednym z systemów elektrod, dielektryk nie może nigdy wypełniać całkowicie przestrzeni międzyelektrodowej. W dodatku szczeliny powietrzne pomiędzy dielektrykiem i każdą z elektrod nie mogą być za cienkie ze względu na występujące wówczas zmniejszenie wytrzymałości napięciowej. Dla



Rys. 10. Gałąź stratna wzorca stałego $\operatorname{tg} \delta$ i regulowanej pojemności

a) Wzorzec złożony z dwu kondensatorów połączonych szeregowo, b) Wzorzec uzyskany przy pomocy poprzecznego uwarstwienia dielektryka

celów obliczeniowych układ taki można rozpatrywać jako dwa szeregowo połączone kondensatory, z których jeden ma jako dielektryk powietrze, drugi dielektryk jest stały (rys. 10a i 10b). Pojemność występująca na zaciskach takiego zespołu ma wartość

$$C = \frac{C_{pow} \cdot C_{diel}}{C_{pow} + C_{diel}} = \frac{C_{min\ pow} + \alpha_p \cdot \varphi}{(C_{min\ pow} + C_{min\ diel}) + (\alpha_p \cdot \varphi + \alpha_{diel} \cdot \varphi)} (C_{min\ diel} + \alpha_{diel} \cdot \varphi)$$

W środkowej części charakterystyki, w której:

$$C_{min\ pow} \ll C_p = \alpha_p \cdot \varphi \quad \text{oraz} \quad C_{min\ diel} \ll C_d = \alpha_d \cdot \varphi,$$

pojemność wypadkowa zmienia się w przybliżeniu proporcjonalnie do kąta obrotu rotora:

$$C \approx \frac{C_{pow} \cdot C_{diel}}{C_{pow} + C_{diel}} = \frac{\alpha_p \cdot \varphi \cdot \alpha_d \cdot \varphi}{\alpha_p \cdot \varphi + \alpha_d \cdot \varphi} = \frac{\alpha_p \cdot \alpha_d}{\alpha_p + \alpha_d} \varphi,$$

natomiast tangens kąta strat takiej gałęzi jest:

$$\operatorname{tg} \delta = \frac{C_{pow} \cdot \operatorname{tg} \delta_{diel} + C_{diel} \cdot \operatorname{tg} \delta_{pow}}{C_{pow} + C_{diel}} \approx \frac{C_{pow}}{C_{pow} + C_{diel}} \operatorname{tg} \delta_{diel} = \text{const.}$$

Wzorzec regulowanego $\operatorname{tg} \delta$ i stałej pojemności realizuje się przy pomocy takiego układu przez dołączenie równoległe kondensatora C_2 o znikomych stratach i o pojemności zmieniającej się od-

wrotnie w zależności od kąta obrotu rotora, zupełnie podobnie jak w wykonaniu wzorca z zewnętrzną opornością dodatkową. Wartość zmian tangensa kąta strat jest w tym przypadku mniej więcej proporcjonalna do kąta obrotu rotora; natomiast największa osiągalna wartość $\operatorname{tg} \delta$ jest znacznie mniejsza niż $\operatorname{tg} \delta_d$ dielektryka stałego.

$$\operatorname{tg} \delta = \frac{C_1 \operatorname{tg} \delta_2 + C_2 \operatorname{tg} \delta_1}{C_1 + C_2} \approx \frac{C_2}{C_1 + C_2} \operatorname{tg} \delta_1,$$

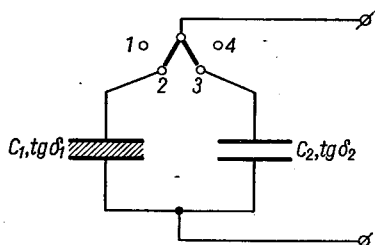
jeżeli

$$\operatorname{tg} \delta_1 \approx \frac{C_{1\text{pow}}}{C_{1\text{pow}} + C_{1\text{diel}}} \operatorname{tg} \delta_{\text{diel}} \quad \text{oraz} \quad C_1 + C_2 = K$$

oraz po przyjęciu, że $C_{1\text{min}} \ll a_1 \varphi$, $C_{2\text{min}} \ll a_2 \varphi$, otrzymuje się:

$$\operatorname{tg} \delta \approx \frac{a_2 \varphi \cdot a_{1\text{pow}} \cdot \varphi}{a_{1\text{pow}} \cdot \varphi + a_{1\text{diel}} \cdot \varphi} \cdot \frac{1}{K} \operatorname{tg} \delta_{\text{diel}} = \frac{a_2}{K} \cdot \frac{a_{1\text{pow}}}{a_{1\text{pow}} + a_{1\text{diel}}} \operatorname{tg} \delta_{\text{diel}} \cdot \varphi.$$

Dodatkową zaletą wzorców złożonych z kondensatorów: bezstratnego i wykazującego stały $\operatorname{tg} \delta$ jest możliwość użycia poszczególnych części z osobna, jako oddzielnych wzorców o różnych własnościach (rys. 11).



Rys. 11. Wzorec $\operatorname{tg} \delta$ przełączany

Położenie przełącznika 1-2, wzorec stałego $\operatorname{tg} \delta$ regulowanej pojemności; położenie przełącznika 2-3, wzorec regul. $\operatorname{tg} \delta$ i stałej pojemności; położenie przełącznika 3-4, wzorec zerowego $\operatorname{tg} \delta$ i regulowanej pojemności

I tak w zależności od sposobu załączenia poszczególnych gałęzi otrzymuje się:

1. Załączenie równoległe — wzorec o stałej pojemności i o $\operatorname{tg} \delta$ regulowanym proporcjonalnie do kąta obrotu rotora.
2. Załączenie na końcówki wyjściowe tylko kondensatora stratnego — wzorec o stałym $\operatorname{tg} \delta$ i o pojemności regulowanej.
3. Załączenie na końcówki tylko kondensatora bezstratnego — wzorec o regulowanej pojemności i $\operatorname{tg} \delta \approx 0$.

WYKAZ LITERATURY

1. Giebe E., Zickner G.: Über die Kapazitätsnormale der PTR. — Zt. f. Instrumentenkunde 53, 1933, 1, 49, 97.
2. Jellonek A.: Wzorce stratności dielektrycznej w zastosowaniu do sprawdzania mostków prądu zmiennego. — Zesz. Naukowe Politechniki Wrocław. Elektryka 5, PWN 1954/55.
3. Hoyer H.: Verlustwinkelnormale und Verlustwinkelvariatoren. — Archiv für Elektrotechnik 41, 1954, 347.
4. Jellonek A.: Standards of the Loss Angle of Capacitors. — Posiedzenie Komitetu Nr 1 CIGRE, Paryż 1958, Mediolan 1959.
5. Siciński Zb.: Application of $\tan \delta$ Measurements for the Evaluation of Insulating Oils Quality. — Posiedzenie komitetu CIGRE, Paryż 1958, Mediolan 1959.

A. ЕЛЛЕНЕК

ВОЗМОЖНОСТЬ РЕАЛИЗАЦИИ ЭТАЛОНОВ ТАНГЕНСА УГЛА
ДИЭЛЕКТРИЧЕСКИХ ПОТЕРЬ КОНДЕНСАТОРОВ

Резюме

Тангенс угла диэлектрических потерь конденсаторов является величиной измеряемой так часто и при таком разнообразии методов и измерительных приборов, что весьма удобным было бы иметь соответствующие эталоны этой величины, которые давали бы возможность сверки результатов измерений по различным источникам.

Такие приборы можно построить в виде эталонов постоянной или регулируемой величины тангенса угла диэлектрических потерь с постоянной или регулируемой ёмкостью. В каждом случае прибор составлен из соответственно соединённых конденсаторов по возможности без потерь и реостатов или сопротивлений не имеющих собственной ёмкости или индуктивности. Возможна также замена сосредоточенных сопротивлений электродами сопротивления равно как и применение конденсатора с диэлектриком, имеющим постоянные и известные потери. Построенные таким образом эталоны работают удовлетворительно в диапазоне средних номинальных значений $\tan \delta$ и низких частот. Ограничение диапазона по частоте и точности определяется влиянием остаточных коэффициентов параметров, составляющих элементов. Исполненные до настоящего времени экземпляры эталонов постоянного и регулируемого значения $\tan \delta$ отображают ёмкости от 100 pF до 1000 pF и $\tan \delta$ от $10 \cdot 10^{-4}$ до $10^3 \cdot 10^{-4}$, причем точность по ёмкости около 0,5%, а по $\tan \delta$ не менее нескольких ‰.

A. JELLONEK

FEASIBILITY OF STANDARDS FOR LOSS-ANGLE TANGENT OF CONDENSERS

Summary

The measurement of loss-angle tangent being a frequent operation carried out with a variety of instruments calls for the adequate standards which would enable to concert the records compiled from various sources. Actually the instruments may be designed as the standards for the constant or controlled quantity of loss-angle tangent with the constant or controlled capacitance.

Such an instrument consists of possibly loss-free condensers and noncapacitive and noninductive resistors interconnected correspondingly. It is equally possible to substitute the separate resistances by the resistant electrodes or to make use of a condenser with the dielectric of constant and known losses.

The standards designed in such a way have proved they can operate satisfactorily within the mean rating valued of $\operatorname{tg} \delta$ and the low frequency band.

The influence exerted by the residual coefficients of the elements used delimits their operation range with relation to the higher frequencies and the better accuracy.

The instruments so far designed with the constant as well as with the controlled value of $\operatorname{tg} \delta$, can reproduce the capacitances from 100 pF to 1000 pF and $\operatorname{tg} \delta$ from $10 \cdot 10^{-4}$ to $10^3 \cdot 10^{-4}$ at the accuracy of capacitance of about 0,5 percent and that of $\operatorname{tg} \delta$ by no amount worse than a few percent.

Optymalne serwomechanizmy przekaźnikowe drugiego rzędu

Rękopis dostarczono 21.3.1960

W pracy przedstawiono i zastosowano elementy usystematyzowanej teorii optymalnych układów regulacji automatycznej — układów charakteryzujących się najkrótszym czasem przebiegu przejściowego — do syntezy i analizy podstawowych typów serwomechanizmów przekaźnikowych drugiego rzędu: z silnikiem wykonawczym sterowanym prądowo i napięciowo. Rozważania teoretyczne oparto na kilku założeniach upraszczających, idealizujących rozpatrywane serwomechanizmy, nie zniekształcających jednakże w sposób jakościowy rozpatrywanych zjawisk. Analizy układów dokonano w oparciu o aparat matematyczny płaszczyzny fazowej, zwracając jednocześnie szczególną uwagę na zagadnienia słabo dotychczas opracowane w literaturze (np. zagadnienie stanu ustalonego) jak również na interpretację fizyczną opisu matematycznego.

1. WSTĘP

W ostatnim okresie rozwoju teorii i praktyki regulacji automatycznej metoda stabilizacji i poprawy jakości układów sterowania za pomocą celowo wprowadzanych elementów nieliniowych znajduje coraz szersze zastosowanie. Człony korekcyjne oparte na elementach nieliniowych (próżniowych, półprzewodnikowych, magnetycznych itp.) umożliwiają w szeregu przypadków uzyskanie znacznie większego wzmocnienia w układzie bez straty stabilności, niż byłoby to możliwe przy zastosowaniu liniowych członów korekcyjnych. Korekcja wprowadza do układu parametry zależne od jego współrzędnych dynamicznych, a więc również i od sygnałów wejściowych, co ma zasadnicze znaczenie przy realizacji przebiegów o właściwej jakości. Wreszcie właściwe wprowadzenie korekcji nieliniowej umożliwia realizację przebiegów przejściowych, których czas trwania jest — przy danych ograniczeniach — najkrótszy z możliwych, tzw. przebiegów optymalnych.

Zagadnienia związane z korekcją nieliniową zajmują od kilku lat poczesne miejsce w literaturze światowej. Znaczna część tych prac poświęcona jest przebiegom i układom optymalnym w sensie wyżej wspomnianym. Do najważniejszych, podstawowych prac w tej dziedzinie należy zaliczyć prace A.A. Feldbauma [3], [4], [5], A. M. Hopkina [8], A. J. LERNIERA [15], L. M. Silvy [18], [19], L. F. Kazdy [11] i T. M. Stouta [24].

O rozwoju tej dziedziny może świadczyć fakt zainteresowania się nią światowej sławy matematyka L. S. Pontriagina oraz to, że jedna z najnowszych książek z teorii regulacji automatycznej [20] w obszernej części poświęconej syntezie układów nieliniowych omawia wyłącznie układy optymalne. Należy dodać, że teoria układów optymalnych, jakkolwiek jeszcze nie ukształtowana, stanowi ważny fragment teorii układów adaptacyjnych i samooptymalizujących — najnowszych osiągnięć regulacji automatycznej [3].

W niniejszej pracy przedstawiono i zastosowano elementy usystematyzowanej teorii układów optymalnych do syntezy i analizy podstawowych typów serwomechanizmów przekąźnikowych drugiego rzędu — z silnikiem wykonawczym sterowanym prądowo i napięciowo — ze szczególnym podkreśleniem zagadnień słabo dotychczas opracowanych w literaturze (np. zagadnienie stanu ustalonego). Rozważania teoretyczne oparto na kilku założeniach upraszczających, idealizujących rozpatrywane serwomechanizmy. Tak więc założono idealną i bezinercyjną charakterystykę przekąźnika i idealne elementy różniczkujące. Założono również, że w układach nie występują ograniczenia prędkości, z wyjątkiem ograniczeń naturalnych w silniku sterowanym napięciowo. Przyjęto ponadto, że opory mechaniczne w części wykonawczej serwomechanizmu są pomijalne oraz że w przekładni mechanicznej nie występuje luz. Praca dotyczy głównie serwomechanizmów małych mocy z silnikami prądu stałego, sterowanymi prądem lub napięciem twornika, szereg zaś rozważań, w szczególności związanych z serwomechanizmem z silnikiem sterowanym prądowo może mieć zastosowanie w przypadkach bardziej ogólnych. Założenia upraszczające wydają się być przy tym umotywowane, tym bardziej, że serwomechanizmy optymalne jest celowe traktować jako swojego rodzaju układy „wzorcowe”.

Za sygnał wejściowy serwomechanizmu przyjęto sygnał paraboliczny, który w szeregu przypadków może skutecznie aproksymować sygnały rzeczywiste.

Syntezy i analizy układów dokonano w oparciu o aparat matematyczny płaszczyzny fazowej, starając się najpełniej, o ile to było możliwe, zinterpretować opis matematyczny układów ich działaniem fizycznym.

Wszystkich obliczeń przeprowadzonych w niniejszej pracy dokonano na jednostkach względnych, aby umożliwić ewentualne bardziej skuteczne porównanie obu rozpatrzonych typów serwomechanizmów.

2. PRZEBIEGI OPTYMALNE W UKŁADACH REGULACJI AUTOMATYCZNEJ

2.1. Pojęcia zasadnicze

Zagadnienie syntezy UAR (Układu Automatycznej Regulacji), który by w pewnym określonym sensie spełniał swoje zadanie w sposób opty-

malny, najkorzystniejszy, jest w chwili obecnej zagadnieniem centralnym, do którego sprowadza się większość prac z dziedziny nowoczesnej teorii i praktyki automatycznej regulacji.

Zadanie większości układów regulacji automatycznej stanowi realizacja równości¹⁾

$$D(p)x_0(t) = x(t), \quad (1)$$

gdzie:

$x_0(t)$ — wielkość prowadząca (lub wejściowa),

$x(t)$ — wielkość regulowana (lub wyjściowa),

$D(p)$ — liniowy operator różniczkowy.

W przypadku szczególnym [do którego da się również sprowadzić przypadek ogólny przez podstawienie $x_{01}(t) = D(p)x_0(t)$], gdy $D(p) = 1$, otrzymamy

$$x_0(t) = x(t). \quad (2)$$

Związek (2) pozostanie obowiązujący w stosunku do dalszych rozważań.

W „klasycznej” teorii regulacji optymalna jakość przebiegu w UAR oparta jest najczęściej o kryteria całkowe [3], [12] typu:

$$I_l = \int_0^{\infty} w(t) |x_0(t) - x(t)|^l dt = \int_0^{\infty} w(t) |\varepsilon(t)|^l dt, \quad (l = 1, 2 \dots), \quad (3)$$

gdzie:

$w(t)$ — funkcja wagowa, formująca kryterium całkowe,

$\varepsilon(t) = x_0(t) - x(t)$ — uchyb regulacji²⁾,

lub o uogólnione kryterium całkowe [3]

$$I_v = \int_0^{\infty} V dt, \quad (4)$$

gdzie V — forma kwadratowa współrzędnych układu³⁾, bądź też dla procesów stochastycznych o kryterium uchybu średniokwadratowego [3], [27]

$$\bar{I} = (\overline{\varepsilon^2}) = \lim_{T \rightarrow \infty} \frac{1}{2T} \int_{-T}^{+T} [x_0(t) - x(t)]^2 dt = \lim_{T \rightarrow \infty} \frac{1}{2T} \int_{-T}^{+T} [\varepsilon(t)]^2 dt, \quad (5)$$

¹⁾ Autor nie rozważa zagadnień związanych z układami ekstremalnymi i samo-optymalizującymi, których zadanie różni się od (1).

²⁾ W zagadnieniach konkretnych korzystnie jest operować składową przejściową uchybu regulacji — $\varepsilon_p(t)$, nie zawsze równą uchybowi całkowitemu. W przypadku ogólnym, a taki jest właśnie rozważany przez autora, jest to zbędne.

³⁾ W najprostszym przypadku funkcja V może mieć postać

$$V = \varepsilon^2 + a \dot{\varepsilon}^2 + b \ddot{\varepsilon}^2 + \dots,$$

gdzie $a, b, \dots$ — współczynniki stałe.

gdzie x_0, x, ε — odpowiednio sygnały reprezentujące stochastyczny proces stacjonarny.

Coraz częściej optymalną jakość przebiegu opiera się o kryterium maksymalnego uchybu w przedziale $t_1 \leq t \leq t_2$

$$L = \varepsilon_{\max} = \max_{(t_1, t_2)} |x_0(t) - x(t)| = \max_{(t_1, t_2)} |\varepsilon(t)|. \quad (6)$$

W szeregu zagadnień optymalną jakość przebiegu przejściowego korzystnie jest łączyć z czasem regulacji¹⁾, tzn. z najkrótszym czasem T_0 , po upływie którego będzie spełniona nierówność

$$|x_0(t) - x(t)| \leq \Delta \quad (7)$$

lub

$$\left. \begin{array}{l} |\varepsilon(t)| \leq \Delta \\ t \geq T_0 \end{array} \right\} \quad \text{przy} \quad (8)$$

gdzie Δ — wartość dana, zwana często odchyleniem dopuszczalnym. Reprezentantem takiego kryterium jest funkcjonal

$$T_0 = \int_{\varepsilon(0)}^{\Delta} \frac{d\varepsilon}{\dot{\varepsilon}} \quad (9)$$

z uwzględnieniem warunku (8) przy przebiegach niemonotonicznych.

W przypadku szczególnym można założyć $\Delta = 0$. Założenie to eliminuje w pewnym sensie subiektywizm definicji „czas regulacji”. Wówczas otrzymamy

$$x_0(t) = x(t) \quad \text{przy} \quad t > T_0 \quad (10)$$

lub

$$\varepsilon(t) = 0 \quad \text{przy} \quad t > T_0. \quad (11)$$

Funkcjonał (9) przybiera zatem postać

$$T_0 = \int_{\varepsilon(0)}^0 \frac{d\varepsilon}{\dot{\varepsilon}}. \quad (12)$$

Dobór odpowiedniego kryterium uwarunkowany jest oczywiście charakterem pracy danego układu.

Charakterystyki układu, przy których funkcjonały (3), (4), (5), (9) lub (12) bądź wyrażenie (6) osiągają minimum, noszą nazwę *charakterystyk optymalnych*, a przebiegi $x(t)$ [lub $\varepsilon(t)$] uzyskane w takim układzie — *przebiegów optymalnych*.

Wyznaczenie charakterystyk i przebiegów optymalnych sprowadza się z reguły do rozwiązania pewnych zagadnień wariacyjnych [12], [27]. Mi-

¹⁾ „Czas regulacji” można również nazwać „czasem przebiegu przejściowego”.

nimalizacja rozpatrzonych funkcjonałów bez wprowadzenia jakichkolwiek założeń dodatkowych musi prowadzić oczywiście do trywialnej tożsamości

$$x_0(t) = x(t) \quad (13)$$

przy $t > 0$ dla dowolnych warunków początkowych i dowolnej klasy sygnałów x_0 .

Ze związku (13) wynika bezpośrednio

$$K(s) = 1, \quad (14)$$

gdzie $K(s) = \frac{L[x(t)]}{L[x_0(t)]}$ — przepustowość operatorowa UAR.

Fizycznie odpowiada to warunkom nieograniczonej mocy źródła, nieograniczonej wytrzymałości (termicznej, mechanicznej itp.) części wykonawczej układu (tj. obiektu regulacji, organu wykonawczego i ewentualnie ostatniego stopnia organu wzmacniającego) oraz braku zakłóceń na wejściu układu. Wówczas, istotnie, przy odpowiednio „inteligentnej” części sterującej równość (13) może być osiągnięta w czasie nieskończenie krótkim, co oczywiście pociąga za sobą $I_l = \bar{I} = T_0 = 0$ [16].

2.2. Ograniczenia w UAR. Przebieg optymalny w sensie Feldbauma

W rzeczywistości związki (13) i (14) nie są realizowalne technicznie, bowiem na część wykonawczą układu nałożone są ograniczenia wynikające bądź z fizycznej zasady działania poszczególnych organów i ograniczonej mocy źródeł zasilających (ograniczenia modułu, nasycenia), bądź też z celowo wprowadzonych elementów zabezpieczających wytrzymałość układu (zawory bezpieczeństwa, ograniczniki prądu itp.). Istotny, ale odrębny charakter mają ograniczenia wynikające ze skończonego pasma częstotliwości przenoszonych przez część wykonawczą układu oraz z istnienia skończonego poziomu zakłóceń na wejściu.

Wymienione ograniczenia są reprezentowane analitycznie przez pewne warunki dodatkowe, w wyniku których minimalizacja rozpatrzonych funkcjonałów prowadzić będzie do wyniku różnego od tożsamości (13).

Ogólnie warunki te mają postać ograniczenia współrzędnych części wykonawczej układu [4], [5], [15]

$$x_{i1} \leq x_i \leq x_{i2}, \quad (i = 1, 2, \dots), \quad (15)$$

gdzie:

x_i — i -ta współrzędna układu,

x_{i1}, x_{i2} — wartości graniczne i -tej współrzędnej.

Warunki te nie dotyczą ograniczeń wynikających ze skończonego pasma częstotliwości przenoszonych przez część wykonawczą układu i ze skończonego poziomu zakłóceń na wejściu. Zagadnienia związane z tym

rodzajem ograniczeń, jako odrębne, nie będą rozpatrywane w niniejszej pracy [3], [12], [14], [27].

W przypadku szczególnym warunki (15) mogą mieć postać ograniczenia modułu kombinacji liniowej pochodnych wielkości wyjściowej

$$\left| a_n \frac{d^n x}{dt^n} + a_{n-1} \frac{d^{n-1} x}{dt^{n-1}} + \dots + a_1 \frac{dx}{dt} + a_0 x \right| \leq M, \quad (16)$$

gdzie:

M — wartość graniczna modułu,

a_i — współczynniki stałe, nie wszystkie równe zeru.

Zagadnienia rozpatrywane w niniejszej pracy oparte są na ograniczeniach typu (16) i założeniu, że część wykonawcza UAR jest układem liniowym o parametrach stałych.

Sygnały $x(t)$, które spełniają dane ograniczenie nazywamy sygnałami dopuszczalnymi. Oczywiście, jeżeli ma być $x_0(t) = x(t)$, to na sygnał wejściowy $x_0(t)$ musi być nałożone to samo ograniczenie, co na sygnał $x(t)$. Zatem funkcja $(x_0 t)$ musi spełniać nierówność

$$\left| a_n \frac{d^n x_0}{dt^n} + a_{n-1} \frac{d^{n-1} x_0}{dt^{n-1}} + \dots + a_1 \frac{dx_0}{dt} + a_0 x_0 \right| \leq M. \quad (17)$$

Sygnały $x_0(t)$ spełniające nierówność (17) nazywamy sygnałami odzwierciedlanymi.

Łatwo wykazać, że ograniczenia (16) nałożone na wielkość wyjściową warunkują w przypadku ogólnym skończony czas przebiegu przejściowego, a więc i nietrywialną minimalizację funkcjonałów np. (3) lub (12).

Istotnie, spełnienie równości

$$x_0(t) = x(t) \quad (18)$$

lub

$$\varepsilon(t) = 0, \quad (19)$$

a zatem również równości wszystkich pochodnych

$$\frac{d^i x_0}{dt^i} = \frac{d^i x}{dt^i}, \quad (i = 1, 2, \dots) \quad (20)$$

lub

$$\frac{d^i \varepsilon}{dt^i} = 0, \quad (i = 1, 2, \dots) \quad (21)$$

przy $t > 0$, przy występowaniu ograniczenia (16) jest teoretycznie możliwe tylko wówczas, jeżeli jest spełniona równość

$$\left. \frac{d^i x_0}{dt^i} \right|_{t=0} = \left. \frac{d^i x}{dt^i} \right|_{t=0}, \quad (i = 0, 1, \dots, n-1) \quad (22)$$

lub

$$\left. \frac{d^i \varepsilon}{d t^i} \right|_{t=0} = 0, \quad (i = 0, 1, \dots, n-1), \quad (23)$$

gdzie n — najwyższy rząd pochodnej występującej w ograniczeniu (16).

Jeżeli dla jakiegokolwiek $i = \mu$, gdzie $0 \leq \mu \leq n-1$, równość (18) lub (19) nie byłaby spełniona, to oznaczałoby to utratę ciągłości μ -tej pochodnej $\frac{d^\mu x}{d t^\mu}$ w chwili $t = 0$ [w założeniu, że spełnione są równości (18)÷(21)] i nieskończenie wielką wartość wszystkich wyższych pochodnych $\frac{d^i x}{d t^i}$, ($i > \mu$) w chwili $t = 0$, a więc niespełnienie ograniczenia (16) [2], [4].

Zatem w wyniku niezgodności wszystkich n warunków początkowych na $x_0(t)$ i $x(t)$ (22) musi nastąpić przebieg przejściowy o czasie trwania T_0 , po upływie którego może nastąpić realizacja równości (18) i (20) lub (19) i (21).

A więc — rozumując analogicznie — jeżeli

$$x_0(t) = x(t) \quad \text{przy} \quad t > T_0 \quad (24)$$

lub

$$\varepsilon(t) = 0 \quad \text{przy} \quad t > T_0, \quad (25)$$

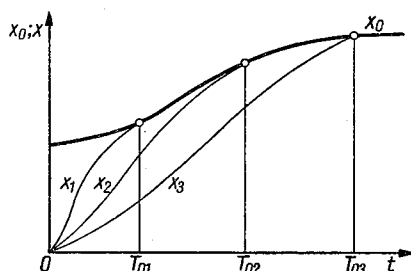
to musi być spełniona równość

$$\lim_{t \rightarrow T_0} \frac{d^i x}{d t^i} = \left. \frac{d^i x_0}{d t^i} \right|_{t=T_0}, \quad (i = 0, 1, \dots, n-1) \quad (26)$$

lub

$$\lim_{t \rightarrow T_0} \frac{d^i \varepsilon}{d t^i} = 0, \quad (i = 0, 1, \dots, n-1). \quad (27)$$

Geometrycznie jest to realizacja styczności $(n-1)$ -go rzędu krzywych $x_0(t)$ i $x(t)$. Po czasie $t > T_0$ obie krzywe schodzą się w jedną. Na przykład



Rys. 1. Przebiegi przejściowe typu (24)

dla układu z ograniczeniem drugiej pochodnej ($n = 2$) w chwili $t = T_0$ muszą być równe nie tylko położenia $x_0(T_0) = x(T_0)$, ale i prędkości

$\dot{x}_0(T_0) = \dot{x}(T_0)$. Przebiegów typu (24), tj. przebiegów, w których $\varepsilon(t) = 0$ przy $t > T_0$, może być oczywiście nieograniczenie wiele (rys. 1).

Przebieg $x(t)$, którego czas T_0 określony ze związków (24) lub (25) z uwzględnieniem ograniczenia (16) jest minimalny dla określonego sygnału $x_0(t)$ z danej klasy sygnałów odtwarzalnych i określonych warunków początkowych [a więc będący konsekwencją minimalizacji funkcjonatu (12) uzupełnionego warunkiem (16)], będziemy nazywać idealnym przebiegiem optymalnym w sensie Feldbauma¹⁾, a układ, w którym taki przebieg zostaje realizowany — układem optymalnym w sensie Feldbauma.

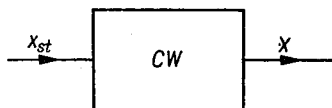
W dalszych rozważaniach przebiegi i układy optymalne będą rozpatrywane zgodnie z powyższą definicją — chyba, że będzie to zaznaczone inaczej.

2.3. Realizacja przebiegu optymalnego

Jeżeli równanie

$$a_n s^n + a_{n-1} s^{n-1} + \dots + a_1 s + a_0 = 0 \quad (28)$$

ma pierwiastki rzeczywiste i niedodatnie²⁾ (co z reguły występuje w zastosowaniach praktycznych), to można udowodnić [4], [5], [26], że przebieg optymalny uzyskuje się w co najwyżej n , niekoniecznie równych, przedziałach czasowych, w których ograniczona w module liniowa kombinacja



Rys. 2. Część wykonawcza UAR

pochodnych (16) przybiera na przemian wartości graniczne $+M$ i $-M$. Zatem przebieg optymalny należy do klasy sygnałów C_{n-1} i jest opisany równaniem

$$a_n \frac{d^n x}{dt^n} + a_{n-1} \frac{d^{n-1} x}{dt^{n-1}} + \dots + a_1 \frac{dx}{dt} + a_0 x = y(t), \quad (29)$$

gdzie

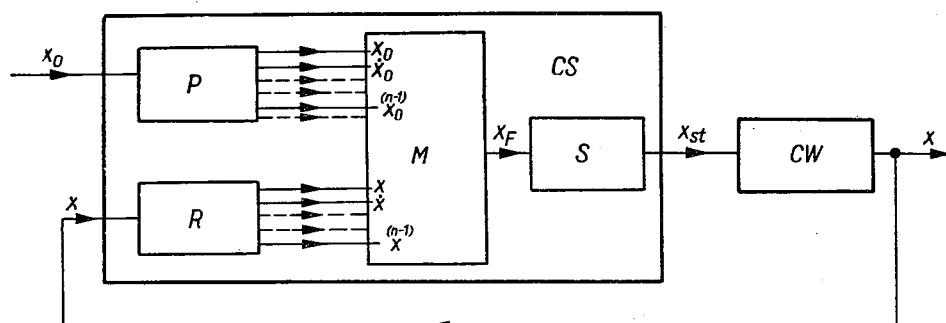
$$y(t) = M \sigma(t) \quad \text{i} \quad \sigma(t) = \begin{cases} +1 \\ -1 \end{cases}, \quad (30)$$

przy czym znaki σ są przemienne w kolejnych przedziałach. Równanie (29) w określonym i -tym przedziale czasowym (w którym $\sigma = +1$ lub $\sigma = -1$)

¹⁾ Określenie „w sensie Feldbauma” jakkolwiek nie jest ogólnie przyjęte w literaturze światowej, znajduje w Polsce zwolenników z uwagi na poważną rolę A. A. Feldbauma w rozwoju teorii optymalnych UAR.

²⁾ W przypadku istnienia pierwiastków zespolonych równania (28) zagadnienie wyznaczenia przebiegu optymalnego jest bardziej złożone [4].

jest oczywiście równaniem liniowym. Nie można go jednak traktować jako równania liniowego w przedziale $0 \leq t < \infty$, bowiem σ nie jest jawną funkcją czasu, lecz jest funkcją współrzędnej x i sygnału x_0 oraz ich pochodnych. W przypadku ogólnym w układzie optymalnym sygnał sterujący x_{st} na wejściu części wykonawczej UAR (rys. 2) różni się od wielkości y , bowiem zależność między x i x_{st} opisana jest równaniem części wykonawczej, które może być różne od równania (29) wynikającego



Rys. 3. Zasadniczy schemat strukturalny układu optymalnego

z ograniczenia (16). Zadaniem części sterującej UAR jest zatem uformowanie takiego sygnału $x_{st}(t)$, aby było spełnione równanie (29). Na podstawie powyższych rozważań wynika prawie bezpośrednio struktura układu optymalnego (rys. 3).

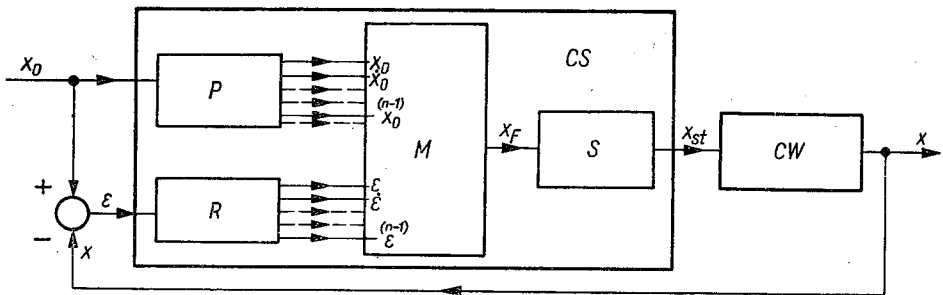
Zgodnie z definicją układ powinien realizować przebieg optymalny przy każdym sygnale x_0 należącym do określonej klasy sygnałów odtwarzalnych. Wobec tego w przypadku ogólnym realizacja przebiegu optymalnego wymaga zdolności przewidywania przebiegu x_0 . Zatem sygnał x_0 zostaje skierowany na wejście tzw. organu prognozy, oznaczonego na rysunku 3 literą P. Jeżeli układ ma odtwarzać sygnał należący do pewnej klasy sygnałów regularnych, to zdolność przewidywania jego przebiegu może być oparta na pomiarze wartości chwilowych sygnału i jego pochodnych. Zatem sygnałami wyjściowymi organu prognozy będą:

$x_0(t), \dot{x}_0(t), \dots, x_0^{(n-1)}(t), \dots$, przy czym liczba pochodnych x_0 i rząd najwyższej pochodnej zależą od określenia klasy sygnałów wejściowych. Im bardziej ograniczona będzie klasa sygnałów wejściowych, tym mniej będzie potrzeba danych świadczących o charakterze przebiegu x_0 . Oczywiście pomiar wartości chwilowych sygnału i jego pochodnych nie jest jedyną metodą przewidywania przebiegu x_0 ; schemat organu prognozy może być oparty również na założeniach statystycznych [20], [27].

Istota przebiegu optymalnego wymaga oczywiście pomiaru sygnału wyjściowego x i jego pochodnych do $(n-1)$ -ej włącznie. Zadanie to w układzie optymalnym spełnia organ rejestrujący R, którego zasada działania jest analogiczna do zasady działania organu prognozy.

Dane z organów P i R zostają wprowadzone do organu operacyjnego M . Wynikiem operacji matematycznych dokonanych w organie operacyjnym jest taki sygnał x_F skierowany na wejście organu sterującego S , że sygnał sterujący x_{st} zapewnia spełnienie równania (29), a zatem realizację przebiegu optymalnego x .

W szeregu przypadków na wejście organu rejestrującego jest doprowadzany zamiast sygnału x sygnał $\varepsilon = x_0 - x$ (rys. 4). Z punktu widzenia



Rys. 4. Struktura układu optymalnego z detektorem uchybu

istoty układów optymalnych nie ma to oczywiście żadnego znaczenia. Z definicji optymalnego przebiegu przejściowego i z rozważań nad strukturą układu jest oczywiste, że

$$x_F = F_x(x, \dot{x}, \dots, x^{(n-1)}, x_0, \dot{x}_0, \dots, x_0^{(n-1)}, \dots) \quad (31)$$

lub

$$x_F = F_\varepsilon(\varepsilon, \dot{\varepsilon}, \dots, \varepsilon^{(n-1)}, x_0, \dot{x}_0, \dots, x_0^{(n-1)}, \dots). \quad (32)$$

Organy: prognozy, rejestrujący, operacyjny i sterujący stanowią część sterującą CS optymalnego UAR.

Należy zaznaczyć, że struktura układów optymalnych nie zawsze da się sprowadzić ściśle do schematów przedstawionych na rysunkach 3 lub 4. W niektórych przypadkach na przykład organ prognozy może okazać się zbędny, w innych zaś konieczne jest wprowadzenie pewnych urządzeń dodatkowych.

Z uwagi na metodę uzyskania przebiegu optymalnego jest oczywiste, że w każdym układzie optymalnym musi występować w jakiegokolwiek postaci działanie przekąźnikowe [(29) i (30)], co determinuje nieliniowość układów optymalnych. Również sama definicja przebiegu optymalnego, w którym $\varepsilon(t) = 0$ przy $t > T_0$ (tj. po czasie skończonym, w odróżnieniu od układów liniowych), wskazuje na występowanie istotnej nieliniowości w układzie, w którym taki przebieg ma zostać zrealizowany.

Należy zaznaczyć, że istnieje możliwość realizacji przebiegu, w którym $\varepsilon(t) = 0$ przy $t > T_0$, gdzie T_0 jest minimalnym czasem regulacji, również za pomocą metod liniowych, gdzie działanie przekąźnikowe uży-

skane jest w sposób „ciągły” [21]. Na przykład układ o przepustowości pętli otwartej

$$K(s) = \frac{\frac{4}{T_0^2} \left(1 - 2e^{-\frac{T_0}{2}s} + e^{-T_0 s} \right)}{s^2 - \frac{4}{T_0^2} \left(1 - 2e^{-\frac{T_0}{2}s} + e^{-T_0 s} \right)}$$

umożliwia realizację takiego przebiegu przy skokowym sygnale wejściowym. Łatwo spostrzec, że $K(s)$ jest funkcją czasu — parametru T_0 , który z kolei zależy zarówno od wartości sygnału wejściowego, jak i od warunków początkowych (w danym przypadku — zerowych). Zatem uzyskanie przebiegu o skończonym i minimalnym czasie regulacji T_0 w układzie liniowym (nb. tylko o parametrach rozłożonych) jest możliwe jedynie dla jednego i tylko jednego sygnału wejściowego i konkretnych warunków początkowych. Tego rodzaju przebieg nie jest przebiegiem optymalnym w sensie Feldbauma.

Dlatego też układy optymalne muszą zawierać niezależnie od elementów przekąźnikowych pewne dodatkowe, lecz zasadnicze elementy nieliniowe, umożliwiające realizację przebiegu optymalnego przy dowolnym sygnale $x_0(t)$ z danej klasy sygnałów odtwarzalnych i dla dowolnych warunków początkowych.

Zatem układy optymalne są układami istotnie nieliniowymi, w których nieliniowość wynika z samej istoty ich działania, tj. w których parametry zależą od sygnału wejściowego.

Z fizycznego punktu widzenia optymalizacja układu sprowadza się do najkorzystniejszego wykorzystania energetycznego części wykonawczej układu [6].

2.4. Przebiegi optymalne w przestrzeni fazowej

Podstawową metodą analizy i syntezy układów optymalnych jest metoda przestrzeni fazowej współrzędnych układu [5], [8], [16], [25]. Korzystniej jest rozpatrywać nie całkowitą przestrzeń fazową, której współrzędne jednoznacznie określają stan układu, lecz jedynie pewną podprzestrzeń o tych współrzędnych¹⁾, których prędkości są ograniczone od góry, jako wynik ograniczeń nałożonych na część wykonawczą układu [5]. Jeżeli na przykład w układzie występuje ograniczenie typu (16) (przy czym rząd równania opisującego część wykonawczą może być większy od n), to

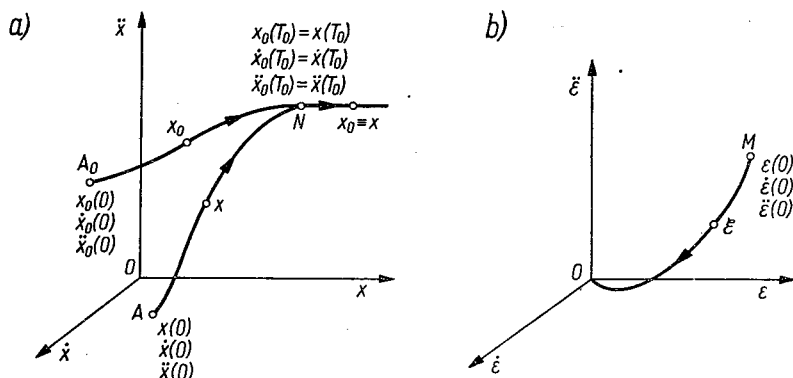
¹⁾ W rzeczywistości w każdym układzie fizycznym wszystkie jego współrzędne muszą mieć prędkości ograniczone od góry. W praktyce jednak szereg z tych ograniczeń nie jest istotny, dzięki czemu można operować przestrzenią fazową niższego rzędu.

celowe jest korzystać z n -wymiarowej przestrzeni fazowej o współrzędnych $x, \dot{x}, \dots, x^{(n-1)}$ [n jest rzędem najwyższej pochodnej występującej w ograniczeniu (16)]. W tego rodzaju przestrzeni fazowej punkt opisujący porusza się w sposób ciągły ze skończoną prędkością. Zamiast operować przestrzenią fazową o współrzędnych $x, \dot{x}, \dots, x^{(n-1)}$, celowe jest w szeregu przypadków korzystać z przestrzeni fazowej o współrzędnych

$$\left. \begin{aligned} \varepsilon &= x_0 - x \\ \dot{\varepsilon} &= \dot{x}_0 - \dot{x} \\ &\dots \dots \dots \\ \varepsilon^{(n-1)} &= x_0^{(n-1)} - x^{(n-1)} \end{aligned} \right\} \quad (33)$$

Z teoretycznego punktu widzenia jest nieistotne, którą z przestrzeni fazowych stosujemy do analizy i syntezy układów optymalnych. Ze względu na pewne korzyści praktyczne w dalszych rozważaniach będziemy stosować przestrzeń fazową o współrzędnych (33). Na rysunku 5a, b przedstawione są oba typy trójwymiarowych przestrzeni fazowych z przykładowymi torami przebiegów optymalnych.

Jeżeli w chwili $t = 0$ punkt opisujący X znajduje się w położeniu $A[x(0), \dot{x}(0), \ddot{x}(0)]$, a punkt opisujący X_0 w położeniu $A_0[x_0(0), \dot{x}_0(0), \ddot{x}_0(0)]$



Rys. 5. Tor przebiegu optymalnego w trójwymiarowej przestrzeni fazowej

(rys. 5a) [X_0 jest punktem opisującym, który odpowiada współrzędnym układu, zgodnym z sygnałem wejściowym $x_0(t)$], to zadanie układu optymalnego polega na spotkaniu się punktów x i x_0 (punkt N na rys. 5a) w chwili $t = T_0$, gdzie T_0 jest czasem minimalnym przy danym ograniczeniu. Torem fazowym przebiegu optymalnego jest zatem tor AN . Dla $t > T_0$ powinna nastąpić współbieżność obu punktów opisujących $x_0 \equiv x$.

Analogicznie w przestrzeni fazowej o współrzędnych $\varepsilon, \dot{\varepsilon}, \ddot{\varepsilon}$ punkt opisujący $\mathcal{E}$ startuje z położenia początkowego M [$\varepsilon(0), \dot{\varepsilon}(0), \ddot{\varepsilon}(0)$], odpowiadającego różnicy położenia punktów A_0 i A , i biegnie do początku układu współrzędnych O , odpowiadającego punktowi N , bowiem po zakończeniu przebiegu przejściowego uchyb ε i jego pochodne $\dot{\varepsilon}$ i $\ddot{\varepsilon}$ stają się równe zeru. Tor fazowy przebiegu optymalnego MO w przestrzeni fazowej o współrzędnych $\varepsilon, \dot{\varepsilon}, \ddot{\varepsilon}$ odpowiada torowi AN w przestrzeni fazowej o współrzędnych $x, \dot{x}, \ddot{x}$.

Jak zaznaczono wyżej, w dalszych rozważaniach mówiąc o przestrzeni fazowej będziemy mieli na myśli przestrzeń fazową o współrzędnych $\varepsilon, \dot{\varepsilon}, \ddot{\varepsilon}, \dots, \varepsilon^{(n-1)}$. Zagadnienie sprowadza się zatem do wyznaczania toru fazowego łączącego punkt określony przez warunki początkowe $\varepsilon(0), \dot{\varepsilon}(0), \dots, \varepsilon^{(n-1)}(0)$ z początkiem układu współrzędnych w czasie minimalnym. Jest to zagadnienie wariacyjne, lecz z uwagi na jego specyfikę można je rozwiązać niekoniecznie bezpośrednimi metodami wariacyjnymi.

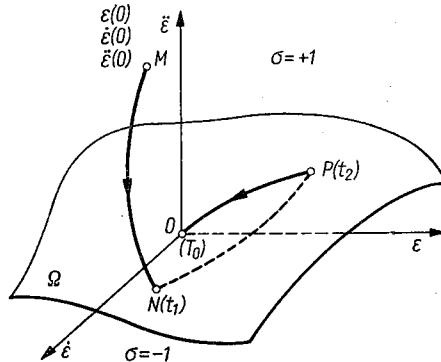
Zgodnie z zasadniczym twierdzeniem o realizacji przebiegu optymalnego [3] jest oczywiste, że tor przebiegu optymalnego składa się co najwyżej z n odcinków, którym odpowiadają kolejno wartości $y = +M$ i $y = -M$ lub $\sigma = +1$ i $\sigma = -1$. Zatem każdemu punktowi przestrzeni fazowej traktowanemu jako początkowemu odpowiada wartość $\sigma = +1$ lub $\sigma = -1$, z czego wynika, że n -wymiarową przestrzeń fazową można podzielić na dwa obszary, odpowiadające wartościom $\sigma = +1$ i $\sigma = -1$. Granicą dzielącą oba obszary będzie oczywiście $(n-1)$ -wymiarowa, jednorodna hiperpowierzchnia Ω , przechodząca przez początek układu współrzędnych¹⁾. Na samej hiperpowierzchni Ω wartości σ bądź nie określa się, bądź też dla ujednoznacznienia rozważań zakłada się równą zeru. Można w sposób ścisły udowodnić [5], że być może z wyjątkiem pierwszego (licząc od chwili początkowej) odcinka toru fazowego przebiegu optymalnego wszystkie następne odcinki do ostatniego włącznie leżą właśnie na tej hiperpowierzchni²⁾ [ściślej — nieskończenie blisko od obu stron hiperpowierzchni Ω , tj. na przemian w obszarach $\sigma = +1$ i $\sigma = -1$ (rys. 6)]. Hiperpowierzchnia Ω , zwana *hiperpowierzchnią przełączania* ma zasadnicze znaczenie w teorii układów optymalnych.

Przykładowy „mechanizm” powstawania przebiegu optymalnego w trójwymiarowej przestrzeni fazowej przedstawiono na rysunku 6; w chwili $t = 0$ punkt opisujący znajduje się w położeniu M , w obszarze $\sigma = +1$. Pierwszy odcinek toru fazowego odpowiada zatem wartości $y = +M$. W chwili $t = t_1$ tor fazowy przebija powierzchnię Ω w punkcie N , w któ-

¹⁾ W początku układu współrzędnych wartość σ nie jest określona.

²⁾ Jeżeli w chwili początkowej punkt opisujący leży na hiperpowierzchni Ω , to liczba odcinków toru fazowego, a więc liczba przedziałów czasowych będzie mniejsza od n .

rym następuje przełączenie na $y = -M$. Następnie tor fazowy biegnie w obszarze $\sigma = -1$, nieskończenie blisko powierzchni Ω . W chwili $t = t_2$ tor fazowy przebiega ponownie powierzchnię Ω w punkcie P , co odpowiada kolejnemu przełączeniu na wartość $y = +M$. Ostatni odcinek toru fazowego biegnie nieskończenie blisko powierzchni Ω , lecz w obszarze $\sigma = +1$



Rys. 6. Dwuwymiarowa powierzchnia przełączania i tor przebiegu optymalnego

i w chwili $t = T_0$ osiąga nieskończenie małe otoczenie początku układu współrzędnych. W przypadku ogólnym równanie hiperpowierzchni przełączania może zależeć od czasu. Taką ruchomą hiperpowierzchnię można określić jako *niestacjonarną*. Jeżeli hiperpowierzchnia nie zależy od czasu, a jedynie od parametrów sygnału wejściowego x_0 , to taka hiperpowierzchnia nazywa się *kwazistacjonarną*. Jeżeli zaś hiperpowierzchnia nie zależy również od parametrów sygnału, to taką hiperpowierzchnię będziemy nazywać hiperpowierzchnią *stacjonarną*¹⁾.

Tak więc równanie $(n-1)$ -wymiarowej hiperpowierzchni niestacjonarnej lub kwazistacjonarnej ma postać

$$F(\varepsilon, \dot{\varepsilon}, \dots, \varepsilon, x_0, \dot{x}_0, \dots, x_0^{(n-1)}) = 0, \quad (34)$$

przy czym niestacjonarność dotyczy przypadku, gdy w argumentach funkcji F występują składniki zawierające czas w postaci jawnej [np. $x_0(t)$]. Równanie hiperpowierzchni stacjonarnej ma postać prostszą

$$F(\varepsilon, \dot{\varepsilon}, \dots, \varepsilon^{(n-1)}) = 0. \quad (35)$$

2.5. Podstawowa metoda syntezy i analizy układów optymalnych

W przypadku ogólnym zagadnienie syntezy części sterującej układu optymalnego można podzielić (jako zagadnienie zasadnicze) na syntezę

¹⁾ Klasyfikacja typów hiperpowierzchni przełączania wiąże się ściśle z zagadnieniem klasy sygnałów odtwarzalnych [5].

organu operacyjnego M i syntezę organu sterującego S (rys. 4). Zgodnie z dotychczasowymi rozważaniami wydaje się oczywiste, że synteza organu operacyjnego sprowadza się do wyznaczenia hiperpowierzchni przełączania $F = 0$ i następnie do określenia takiej funkcji, która by po obu stronach hiperpowierzchni miała różny znak. Istnieje nieskończony zbiór takich funkcji i dlatego zagadnienie nie jest jednoznaczne. Najprościej jest operować po prostu funkcją F (32), o ile nie posiada ona niedopuszczalnych osłabień w całej przestrzeni fazowej¹). Zadaniem organu operacyjnego jest zatem generacja sygnału

$$x_F(t) = F[\varepsilon(t), \dot{\varepsilon}(t), \dots, \varepsilon^{(n-1)}(t), x_0(t), \dot{x}_0(t), \dots, x_0^{(n-1)}(t) \dots]. \quad (36)$$

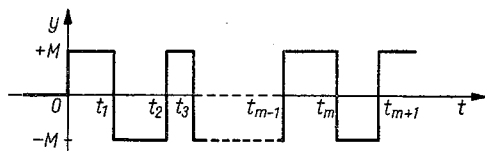
Jeżeli znak funkcji F jest zgodny ze znakiem σ , tj.

$$\sigma = \operatorname{sgn} F, \quad (37)$$

to zadaniem organu sterującego S jest wytworzenie takiego sygnału x_{st} , aby zgodnie z (30) była spełniona równość:

$$y(t) = M \sigma(t) = M \operatorname{sgn} x_F(t). \quad (38)$$

Równanie hiperpowierzchni przełączania można wyznaczyć stosunkowo prosto, korzystając z zależności (29). Wielkość y (30), która, jak zaznaczono, może być różna od sygnału sterującego x_{st} i którą można



Rys. 7. „Sygnał” $y(t)$ jako ciąg impulsów prostokątnych

wówczas traktować jako pewien sygnał fikcyjny, stanowi ciąg impulsów prostokątnych o przemiennych znakach i wysokości M (rys. 7).

Sygnał tego typu w przedziale $0 \leq t < t_{m+1}$ można przedstawić w postaci następującej:

$$y(t) = \sigma_{01} M \left[1(t) + 2 \sum_{k=1}^m (-1)^k 1(t - t_k) \right], \quad (39)$$

gdzie:

$1(t)$ — funkcja jednostkowa,

σ_{01} — wartość σ w przedziale $0 \leq t < t_1$ (tj. $\sigma_{01} = +1$ lub $\sigma_{01} = -1$),

t_k — kolejne chwile przełączania, przy czym $t_0 = 0$.

Oznaczmy rozwiązanie równania

$$a_n \frac{d^n x}{dt^n} + a_{n-1} \frac{d^{n-1} x}{dt^{n-1}} + \dots + a_1 \frac{dx}{dt} + a_0 x = 1(t) \quad (40)$$

¹) Niekiedy jest korzystniej, z punktu widzenia realizowalności technicznej układu, stosować funkcję opartą na funkcji F , lecz różną od niej.

przy zerowych warunkach początkowych przez $h(t)$. Wówczas, zgodnie z równaniem (29), przebieg optymalny w przedziale $0 \leq t \leq t_{m+1}$ można przedstawić w postaci¹⁾

$$x(t) = \sigma_{01} M \left[h(t) + 2 \sum_{k=1}^m (-1)^k h(t - t_k) \right] + x^* [x(0), \dot{x}(0), \dots, x^{(n-1)}(0), t], \quad (41)$$

gdzie $x^*[x(0), \dot{x}(0), \dots, x^{(n-1)}(0), t]$ jest rozwiązaniem równania jednorodnego

$$a_n \frac{d^n x^*}{dt^n} + a_{n-1} \frac{d^{n-1} x^*}{dt^{n-1}} + \dots + a_1 \frac{dx^*}{dt} + a_0 x^* = 0 \quad (42)$$

przy warunkach początkowych

$$x^{*(i)}(0) = x^{(i)}(0), \quad (i = 0, 1, \dots, n-1). \quad (43)$$

Należy zaznaczyć, że w wyrażeniu (39), a więc również i w (40) chwile t_k nie są na razie wiadome i dlatego wyrażenia (40) nie można na razie traktować jako rozwiązania równania (29). Należy również zastrzec, że w przypadku ogólnym podobnie jak $y(t)$ może być różne od sygnału sterującego, $h(t)$ może nie być odpowiedzią jednostkową części wykonawczej układu.

Jeżeli w wyrażeniach

$$\varepsilon^{(i)}(t) = x_0^{(i)}(t) - x^{(i)}(t), \quad (i = 0, 1, \dots, n-1), \quad (44)$$

gdzie $x_0^{(i)}(t)$ jest dane, bowiem dana jest klasa sygnałów wejściowych, a $x(t)$ jest dane z wyrażenia (41), podstawić $t = t_n = T_0$, to zgodnie z określeniem przebiegu optymalnego otrzymamy n równości

$$\varepsilon^{(i)}(t_n) = x_0^{(i)}(t_n) - x^{(i)}(t_n) = 0, \quad (i = 0, 1, \dots, n-1). \quad (45)$$

Wobec tego, że przy danych określonych warunkach chwila drugiego przełączenia t_2 zależy od chwili poprzedniej t_1 , chwila t_3 od chwili t_2 i — ogólnie — chwila t_k od chwili t_{k-1} , to jest oczywiste, że chwila n -tego przełączenia t_n zależy od wszystkich poprzednich chwil przełączenia $t_1, t_2, \dots, t_{n-1}$, przy czym chwila t_1 determinuje wszystkie następne, co można zapisać

$$t_n = f\{t_{n-1}[t_{n-2}(\dots(t_1)\dots)]\}. \quad (46)$$

Wartość $x^{(i)}(t_n)$ również zatem zależy od wszystkich poprzednich chwil przełączania²⁾, tj.

$$x^{(i)}(t_n) \equiv x^{(i)}(t_n, t_{n-1}, \dots, t_1). \quad (47)$$

¹⁾ W literaturze (np. [2], [17] zależność typu (41) obowiązuje w przedziale $t_m < t < t_{m+1}$. Łatwo zauważyć, że jest ona słuszna również dla $0 \leq t \leq t_m$, bowiem $h(t - t_k) = 0$ przy $t \leq t_k$.

²⁾ Wynika to zresztą bezpośrednio z wyrażenia (41).

Rozwiązując układ równań (45), otrzymamy chwile przełączania $t_1, t_2, \dots, t_n$ jako funkcje parametrów części wykonawczej, sygnału wejściowego (lub jego parametrów) oraz warunków początkowych. Rozwiązanie układu (45) i wyrażenie (41) określają zatem jednoznacznie przebieg optymalny, tzn. stanowią rozwiązanie równania (29). Znak σ_{01} jest określony jednoznacznie dowolnym warunkiem typu

$$t_k > t_{k-1}. \quad (48)$$

W chwili t_1 punkt opisujący znajduje się na hiperpowierzchni przełączania, podsawiając więc do związków (44) $t = t_1$

$$\varepsilon^{(i)}(t_1) = x_0^{(i)}(t_1) - x^{(i)}(t_1), \quad (i = 0, 1, \dots, n-1) \quad (49)$$

otrzymamy parametryczne równania hiperpowierzchni przełączania.

Eliminując z równań (49) czas t_1 i warunki początkowe $x(0), \dot{x}(0), \dots, x^{(n-1)}(0)$ [korzystając z rozwiązania układu równań (45)] uzyskujemy ostateczne równanie hiperpowierzchni przełączania o postaci (34) lub (35).

Zagadnienie syntezy organu sterującego na razie pominiemy.

Jeżeli równanie hiperpowierzchni przełączania jest dane, a zatem jest dana funkcja x_F , to wyznaczenie przebiegu optymalnego sprowadza się do rozwiązania równania

$$W(p)x(t) = M \operatorname{sgn} F[\varepsilon(t), \dot{\varepsilon}(t), \dots, \varepsilon^{(n-1)}(t), x_0(t), \dot{x}_0(t), \dots, x_0^{(n-1)}(t), \dots], \quad (50)$$

gdzie $W(p)$ — liniowy operator różniczkowy n -tego rzędu. Rozwiązanie tego równania ma oczywiście postać (41), przy czym

$$\sigma_{01} = \operatorname{sgn} x_F(0) = \operatorname{sgn} F[\varepsilon(0), \dot{\varepsilon}(0), \dots, \varepsilon^{(n-1)}(0), x_0(0), \dot{x}_0(0), \dots, x_0^{(n-1)}(0), \dots]. \quad (51)$$

Algorytm wyznaczenia chwil przełączania sprowadza się do kolejnego rozwiązania

$$F[\varepsilon(t_k), \dot{\varepsilon}(t_k), \dots, \varepsilon^{(n-1)}(t_k), x_0(t_k), \dot{x}_0(t_k), \dots, x_0^{(n-1)}(t_k), \dots] = 0, \quad (k = 1, 2, \dots, n). \quad (52)$$

Należy podkreślić, że zarówno w syntezie, jak i w analizie układów optymalnych nie występuje w zasadzie zagadnienie stabilności, bowiem układy optymalne są stabilne z założenia.

2.6. Układy optymalne z ograniczonym sygnałem sterującym

W przypadku szczególnym, będącym podstawą niniejszej pracy, jeżeli część wykonawcza UAR jest liniowa i jej przepustowość operatorowa ma postać

$$K_0(s) = \frac{b_0}{a_n s^n + a_{n-1} s^{n-1} + \dots + a_1 s + a_0} \quad (53)$$

odpowiadającą równaniu

$$a_n \frac{d^n x}{dt^n} + a_{n-1} \frac{d^{n-1} x}{dt^{n-1}} + \dots + a_1 \frac{dx}{dt} + a_0 x = b_0 x_{st}(t), \quad (54)$$

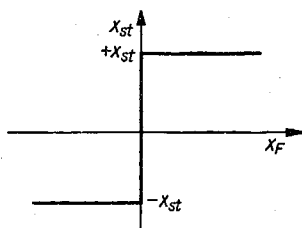
gdzie b_0, a_i — współczynniki stałe, nie wszystkie równe zero, i jeżeli sygnał sterujący jest ograniczony w module

$$|x_{st}(t)| \leq x_{st}, \quad (55)$$

to ograniczenie nałożone na część wykonawczą ma postać

$$\frac{1}{b_0} \left| a_n \frac{d^n x}{dt^n} + a_{n-1} \frac{d^{n-1} x}{dt^{n-1}} + \dots + a_1 \frac{dx}{dt} + a_0 x \right| \leq x_{st}. \quad (56)$$

Z porównania związków (16) i (56) oraz (29) i (54) wynika zatem, że



Rys. 8. Charakterystyka przekaźnika idealnego

sygnał $y(t)$ jest w tym przypadku sygnałem sterującym

$$y(t) = x_{st}(t), \quad (57)$$

gdzie

$$x_{st}(t) = x_{st} \sigma(t), \quad (58)$$

a $h(t)$ we wzorze (41) jest odpowiedzią jednostkową części wykonawczej.

Tak więc w tego rodzaju przypadku szczególnym zagadnienie syntezy organu sterującego w ogóle nie występuje i optymalizacja UAR sprowadza się do syntezy organu operacyjnego i sterowania części wykonawczej przekaźnikiem idealnym (rys. 8). Schemat strukturalny układu optymalnego z przekaźnikiem idealnym jako organem sterującym jest przedstawiony na rys. 9.

Charakterystyka przekaźnika idealnego ma postać

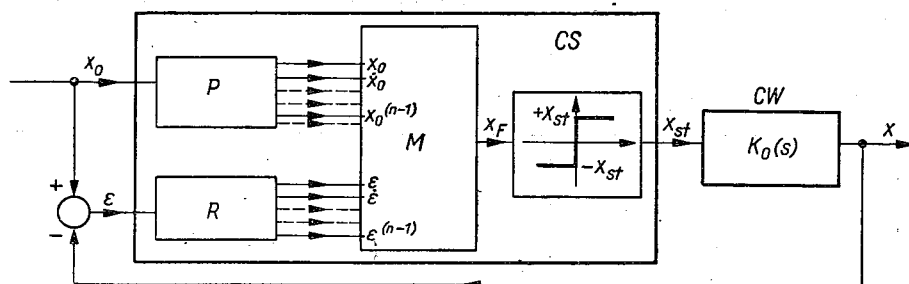
$$x_{st} = \Phi(x_F) x_{st} \operatorname{sgn} x_F, \quad (59)$$

a czas zadziałania jest równy zero.

Z przytoczonych rozważań wynika, że w układach z liniową częścią wykonawczą i ograniczonym sygnałem sterującym, w szczególności w układach przekaźnikowych, optymalizacja przebiegu przejściowego jest prostsza niż w pozostałych przypadkach [4], [16], [25]. Oczywiście sygnał ste-

rujący jest ograniczony w każdym rzeczywistym UAR; wspomniana prostota jest związana jedynie z charakterem ograniczenia (56).

W układach rzeczywistych charakterystyka organu sterującego różni się rzecz prosta od charakterystyki przedstawionej na rys. 8. W tablicy 1 przedstawione są typowe, rzeczywiste charakterystyki statyczne organów sterujących z ograniczonym sygnałem wyjściowym.



Rys. 9. Schemat strukturalny układu optymalnego z przekąźnikiem idealnym

Odstępstwa charakterystyki organu sterującego od charakterystyki idealnej (jak również odstępstwa rzeczywistego organu operacyjnego od wyznaczonego teoretycznie) powodują oczywiście odchylenia przebiegu rzeczywistego od idealnego przebiegu optymalnego [10], [22]. Dlatego też w praktyce możliwa jest jedynie realizacja przebiegów zbliżonych do optymalnych, przy czym przybliżenie przebiegu rzeczywistego, w sensie np. całkowym, do przebiegu idealnego może stanowić wskaźnik jakości układu.

Zagadnienia rozpatrywane w niniejszej pracy oparte są na idealnej charakterystyce przekąźnikowej i charakterystykach 1 i 3 z tablicy 1 (charakterystyka 1 jest szczególnym przypadkiem charakterystyki 3, gdy $\lambda = 0$), a więc dotyczą układów przekąźnikowych, które z istoty swej mogą być najszybciej działającymi układami automatycznej regulacji.

Należy zaznaczyć, że układy zbliżone do optymalnych, oparte na charakterystykach 4 i 5 przy małych odchyleniach mogą być traktowane jako układy liniowe, co w szeregu przypadków jest szczególnie korzystne [5], [9], [16].

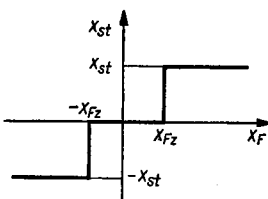
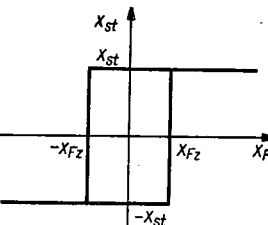
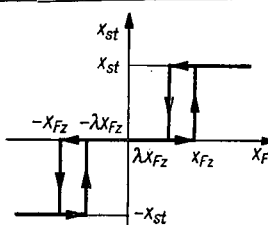
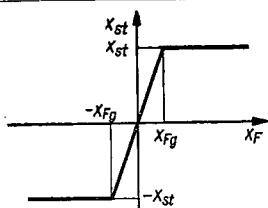
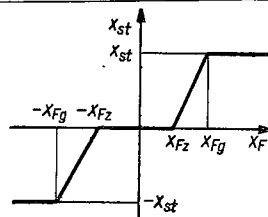
3. OPTIMALNE SERWOMECHANIZMY PRZEKĄŹNIKOWE Z PRZEKĄŹNIKIEM IDEALNYM

3.1. Część wykonawcza serwomechanizmu

Tematem naszych rozważań będzie optymalny serwomechanizm przekąźnikowy małej mocy, w którym częścią wykonawczą jest obcowzbudny silnik prądu stałego z przekładnią zębatą bezluzową. Przepustowość tak

Tablica 1

Charakterystyki statyczne organów sterujących

Nr	Charakterystyka	$x_{st} = \Phi(x_F)$
1	 Przekaznik ze strefą martwą	$x_{st} = \frac{1}{2} X_{st} [\operatorname{sgn}(x_F - x_{Fz}) + \operatorname{sgn}(x_F + x_{Fz})]$
2	 Przekaznik z histerezą	$x_{st} = \begin{cases} X_{st} \operatorname{sgn}(x_F - x_{Fz}) & \text{przy } \dot{x}_F > 0 \\ X_{st} \operatorname{sgn}(x_F + x_{Fz}) & \text{przy } \dot{x}_F < 0 \end{cases}$
3	 Przekaznik ze strefą martwą i histerezą, $0 \leq \lambda \leq 1$	$x_{st} = \begin{cases} \frac{1}{2} X_{st} [\operatorname{sgn}(x_F - x_{Fz}) + \operatorname{sgn}(x_F + \lambda x_{Fz})] & \text{przy } \dot{x}_F > 0 \\ \frac{1}{2} X_{st} [\operatorname{sgn}(x_F + x_{Fz}) + \operatorname{sgn}(x_F - \lambda x_{Fz})] & \text{przy } \dot{x}_F < 0 \end{cases}$
4	 Wzmacniacz z nasyceniem	$x_{st} = \begin{cases} \frac{X_{st}}{x_{Fg}} x_F & \text{przy } x_F \leq x_{Fg} \\ X_{st} \operatorname{sgn} x_F & \text{przy } x_F \geq x_{Fg} \end{cases}$
5	 Wzmacniacz z nasyceniem i strefą martwą	$x_{st} = \begin{cases} \frac{1}{2} \cdot \frac{X_{st}}{x_{Fg} - x_{Fz}} (x_F - x_{Fz}) [\operatorname{sgn}(x_F - x_{Fz}) + \operatorname{sgn}(x_F + x_{Fz})] & \text{przy } x_F \leq x_{Fg} \\ X_{st} \operatorname{sgn} x_F & \text{przy } x_F \geq x_{Fg} \end{cases}$

określonej części wykonawczej ma postać:

$$K_0(s) = \frac{\mathcal{L}[\theta(t)]}{\mathcal{L}[U(t)]} = \frac{k_0}{s(T_{em}s + 1)}, \quad (60)$$

gdzie:

$\theta(t)$ — przesunięcie kątowe wału wyjściowego [rd],

$U(t)$ napięcie wyjściowe nieobciążonego organu sterującego [V],

k_0 — współczynnik wzmacnienia części wykonawczej $\left[\frac{1}{Vs} \right]$.

T_{em} — elektromechaniczna stała czasowa części wykonawczej [s].

Wyrażenie (60) jest równoważne równaniu różniczkowemu drugiego rzędu o postaci

$$T_{em} \frac{d^2 \theta}{dt^2} + \frac{d\theta}{dt} = k_0 U(t) \quad (61)$$

wyprowadzonemu w założeniu, że statyczny moment obciążenia M_{obc} , a więc i moment tarcia suchego M_t , moment mechanicznego¹⁾ tarcia lepkiego M_f oraz elektryczna stała czasowa silnika T_e , a więc także indukcyjność twornika L_f równają się zeru. Założenia $M_{obc} = 0$, $M_t = 0$ i $T_e = 0$ są oparte na często spotykanych warunkach praktycznych. Obciążenie statyczne serwomechanizmu jak i tarcie lepkie jest niejednokrotnie pomijalne w porównaniu z obciążeniem dynamicznym, a elektryczna stała czasowa jest z reguły o rząd wartości mniejsza od stałej czasowej elektromechanicznej [3].

Parametry części wykonawczej są związane z parametrami silnika i przekładni następująco:

$$\left. \begin{aligned} T_{em} &= \frac{JR}{c_M c_E}, \\ k_0 &= \frac{1}{c_E i}, \end{aligned} \right\} \quad (62)$$

gdzie:

J — moment bezwładności mas wirujących sprowadzony do wału silnika [$G \text{ cms}^2$],

$$J = J_t + \frac{1}{i^2} J_{obc}, \quad (63)$$

J_t — moment bezwładności twornika,

J_{obc} — moment bezwładności obciążenia,

i — przełożenie przekładni; $i > 1$,

R — oporność czynna obwodu twornika [Ω],

¹⁾ Elektromagnetyczne „tarcie lepkie” jest uwzględnione w równaniu (61).

$$R = R_t + R_z, \quad (64)$$

R_t — oporność twornika,

R_z — oporność wyjściowa organu sterującego,

c_M — stała silnika [Gcm/A],

$$c_M = \frac{M_n}{I_n}, \quad (65)$$

M_n — moment znamionowy silnika [Gcm],

I_n — prąd znamionowy silnika [A],

c_E — stała silnika [Vs],

$$c_E = \frac{P_n}{\omega_n I_n}, \quad (66)$$

P_n — moc znamionowa silnika [W],

ω_n — prędkość obrotowa znamionowa silnika $\left[\frac{\text{rd}}{\text{s}} \right]$.

Związek (66) wynika bezpośrednio z zależności

$$i \quad \left. \begin{aligned} E &= c_E \omega \\ P &= E I, \end{aligned} \right\} \quad (67)$$

gdzie:

E — siła przeciwelektromotoryczna twornika,

ω — prędkość obrotowa,

P — moc oddawana przez silnik.

W dalszych rozważaniach korzystnie będzie operować jednostkami względnymi. Oznaczmy:

$$\left. \begin{aligned} x &= \frac{\theta}{\theta_N}, \\ x_{st} &= \frac{U}{U_{max}}, \\ \tau &= \frac{t}{T_{em}}, \end{aligned} \right\} \quad (68)$$

gdzie:

θ_N — wartość odniesienia kąta wyjściowego,

U_{max} — maksymalna wartość napięcia sterującego.

Podstawiając wyrażenia (68) do równania (61) i wyprowadzając względny współczynnik wzmocnienia

$$\xi_0 = \frac{k_0 T_{em} U_{max}}{\theta_N} \quad (69)$$

otrzymamy równanie części wykonawczej w jednostkach względnych

$$\frac{d^2 x}{d \tau^2} + \frac{d x}{d \tau} = \xi_0 x_{st}(\tau). \quad (70)$$

Przepustowość części wykonawczej w jednostkach względnych ma postać

$$K_0(q) = \frac{\xi_0}{q(q+1)}, \quad (71)$$

gdzie $q = T_{em}s$.

Jeżeli twornik silnika jest zasilany ze źródła o sterowanej wydajności prądowej $I(t)$, to przepustowość silnika upraszcza się:

$$K'_0(s) = \frac{\mathcal{L}[\theta(t)]}{\mathcal{L}[I(t)]} = \frac{k'_0}{s^2}, \quad (72)$$

gdzie

$$k'_0 = \frac{c_M}{J i}. \quad (73)$$

Równanie (61) przybiera postać

$$\frac{d^2 \theta}{d t^2} = k'_0 I(t), \quad (74)$$

a równanie (70) i przepustowość (71) odpowiednio:

$$\frac{d^2 x}{d \tau^2} = \xi_0 x_{st}(\tau), \quad (75)$$

$$K_0(q) = \frac{\xi_0}{q^2}, \quad (76)$$

gdzie

$$x_{st} = \frac{I}{I_{max}}, \quad (77)$$

a I_{max} — maksymalna wartość prądu sterującego. Względny współczynnik wzmocnienia ξ_0 zachowuje postać (69), z oczywistym zastrzeżeniem, że wielkości k_0 , T_m , $U_{max} = RI_{max}$ odnoszą się do silnika sterowanego napięciowo.

3.2. Serwomechanizm optymalny z częścią

wykonawczą o przepustowości $K_0(q) = \frac{\xi_0}{q^2}$

Synteza najprostszego, optymalnego serwomechanizmu przekąźnikowego o przepustowości części wykonawczej $K_0(q) = \frac{\xi_0}{q^2}$ opiera się na

naturalnym ograniczeniu, wynikającym z równania (75):

$$\frac{1}{\xi_0} \left| \frac{d^2 x}{d\tau^2} \right| \leq X_{st}, \quad (78)$$

gdzie X_{st} — maksymalny moduł sygnału sterującego x_{st} [por. (56)].

Nierówność (78) reprezentuje zatem ograniczenie przyspieszenia sygnału wyjściowego, a więc ograniczenie momentu obrotowego silnika. Z uwagi na to, że

$$X_{st} = 1, \quad (79)$$

nierówność (78) przybiera postać

$$\frac{1}{\xi_0} \left| \frac{d^2 x}{d\tau^2} \right| \leq 1. \quad (80)$$

Ponieważ ograniczenie (80) jest rzędu drugiego, zatem zgodnie z ogólnymi rozważaniami w rozdziale 2.4 synteza i analiza układu optymalnego będzie oparta na płaszczyźnie fazowej o współrzędnych $(\varepsilon, \dot{\varepsilon})$, a hiperpowierzchnia przełączania $F = 0$ sprowadzi się do krzywej płaskiej, zwanej w dalszym ciągu krzywą przełączania.

Założmy, że serwomechanizm pracuje na sygnał paraboliczny

$$x_0(\tau) = (A_0 + A_1\tau + A_2\tau^2)\mathbf{1}(\tau), \quad (81)$$

gdzie A_0, A_1, A_2 — stałe wyrażone w jednostkach względnych wynikających ze związków (68).

Odpowiedź jednostkowa części wykonawczej ma postać:

$$h(\tau) = \frac{1}{2} \xi_0 \tau^2, \quad (82)$$

a rozwiązanie równania jednorodnego

$$x^*(\tau) = x(0) + \dot{x}(0)\tau. \quad (83)$$

Zgodnie z (41) przebieg optymalny $x(\tau)$ wyraża się następująco:

$$\left. \begin{array}{l} \text{w przedziale} \quad 0 \leq \tau \leq \tau_1 \\ x(\tau) = \frac{1}{2} x(0) + \dot{x}(0)\tau + \frac{1}{2} \sigma_{01} \xi_0 \tau^2, \\ \text{w przedziale} \quad \tau_1 \leq \tau \leq \tau_2, \quad \tau_2 = T_0 \\ x(\tau) = x(0) + \dot{x}(0)\tau + \sigma_{01} \left[\frac{1}{2} \xi_0 \tau^2 - \xi_0 (\tau - \tau_1) \right]^2. \end{array} \right\} \quad (84)$$

A więc

$$\left. \begin{array}{l} \varepsilon(\tau) = x_0(\tau) - x(\tau) = A_0 + A_1\tau + A_2\tau^2 - x(0) - \dot{x}(0)\tau - \frac{1}{2} \sigma_{01} \xi_0 \tau^2; \\ 0 \leq \tau \leq \tau_1; \\ \varepsilon(\tau) = x_0(\tau) - x(\tau) = A_0 + A_1\tau + A_2\tau^2 - x(0) + \\ - \dot{x}(0)\tau - \frac{1}{2} \sigma_{01} \xi_0 \tau^2 + \sigma_{01} \xi_0 (\tau - \tau_1)^2; \quad \tau_1 \leq \tau \leq \tau_2. \end{array} \right\} \quad (85)$$

Wyrażenia (85) można napisać w postaci

$$\left. \begin{aligned} \varepsilon(\tau) &= [A_0 - x(0)] + [A_1 - \dot{x}(0)]\tau + A_2\tau^2 - \frac{1}{2}\sigma_{01}\xi_0\tau^2; \quad 0 \leq \tau \leq \tau_1; \\ \varepsilon(\tau) &= [A_0 - x(0)] + [A_1 - \dot{x}(0)]\tau + A_2\tau^2 - \frac{1}{2}\sigma_{01}\xi_0\tau^2 + \\ &\quad + \sigma_{01}\xi_0(\tau - \tau_1)^2; \quad \tau_1 \leq \tau \leq \tau_2. \end{aligned} \right\} \quad (86)$$

Oznaczmy

$$\left. \begin{aligned} \varepsilon(0) &= x_0(0) - x(0) = A_0 - x(0), \\ \dot{\varepsilon}(0) &= \dot{x}_0(0) - \dot{x}(0) = A_1 - \dot{x}(0). \end{aligned} \right\} \quad (87)$$

Uwzględniając (87), z wyrażeń (86) otrzymamy

$$\left. \begin{aligned} \varepsilon(\tau) &= \varepsilon(0) + \dot{\varepsilon}(0)\tau + \left(A_2 - \frac{1}{2}\sigma_{01}\xi_0\right)\tau^2; \quad 0 \leq \tau \leq \tau_1, \\ \dot{\varepsilon}(\tau) &= \dot{\varepsilon}(0) + (2A_2 - \sigma_{01}\xi_0)\tau, \quad 0 \leq \tau \leq \tau_1 \end{aligned} \right\} \quad (88)$$

oraz

$$\left. \begin{aligned} \varepsilon(\tau) &= \varepsilon(0) + \dot{\varepsilon}(0)\tau + \left(A_2 - \frac{1}{2}\sigma_{01}\xi_0\right)\tau^2 + \sigma_{01}(\tau - \tau_1)^2, \\ \dot{\varepsilon}(\tau) &= \dot{\varepsilon}(0) + (2A_2 - \sigma_{01}\xi_0)\tau + 2\sigma_{01}\xi_0(\tau - \tau_1), \end{aligned} \right\} \quad \tau_1 \leq \tau \leq \tau_2 \quad (89)$$

Równania wyjściowe (49) i (45) mają zatem postać

$$\left. \begin{aligned} \varepsilon(\tau_1) &= \varepsilon(0) + \dot{\varepsilon}(0)\tau_1 + \left(A_2 - \frac{1}{2}\sigma_{01}\xi_0\right)\tau_1^2 \\ \dot{\varepsilon}(\tau_1) &= \dot{\varepsilon}(0) + (2A_2 - \sigma_{01}\xi_0)\tau_1 \end{aligned} \right\} \quad (90)$$

oraz

$$\left. \begin{aligned} \varepsilon(\tau_2) &= \varepsilon(0) + \dot{\varepsilon}(0)\tau_2 + \left(A_2 - \frac{1}{2}\sigma_{01}\xi_0\right)\tau_2^2 + \sigma_{01}\xi_0(\tau_2 - \tau_1)^2 = 0, \\ \dot{\varepsilon}(\tau_2) &= \dot{\varepsilon}(0) + (2A_2 - \sigma_{01}\xi_0)\tau_2 + 2\sigma_{01}\xi_0(\tau_2 - \tau_1) = 0. \end{aligned} \right\} \quad (91)$$

Rozwiązując układ równań (91) względem niewiadomych τ_1 i τ_2 i uwzględniając, że

$$\sigma_{01}^2 = 1, \quad (92)$$

otrzymujemy

$$\tau_1 = -\frac{\dot{\varepsilon}(0)}{2A_2 - \sigma_{01}\xi_0} \pm \frac{2A_2 + \sigma_{01}\xi_0}{2\xi_0(2A_2 - \sigma_{01}\xi_0)} \sqrt{\frac{2\sigma_{01}\xi_0[\dot{\varepsilon}^2(0) - 2\varepsilon(0)(2A_2 - \sigma_{01}\xi_0)]}{2A_2 + \sigma_{01}\xi_0}}, \quad (93)$$

$$\tau_2 = -\frac{\dot{\varepsilon}(0)}{2A_2 - \sigma_{01}\xi_0} \mp \frac{\sigma_{01}}{2A_2 - \sigma_{01}\xi_0} \sqrt{\frac{2\sigma_{01}\xi_0[\dot{\varepsilon}^2(0) - 2\varepsilon(0)(2A_2 - \sigma_{01}\xi_0)]}{2A_2 + \sigma_{01}\xi_0}} \quad (94)$$

Warunek odtwarzalności sygnału wejściowego [por. (17)] ma postać

$$\frac{1}{\xi_0} \left| \frac{d^2 x_0}{dt^2} \right| < 1. \quad (95)$$

Warunek ten jest ostrzejszy od (17) i eliminuje sygnały określające „granicę odtwarzalności” jako mniej typowe.

Zgodnie z (95) przyspieszenie sygnału wejściowego $\ddot{x} = 2A_2$ powinno zatem spełniać nierówność:

$$|2A_2| < \xi_0 \quad (96)$$

i wobec tego w wyrażeniach (93) i (94) obowiązuje znak „-”, co jest uzasadnione¹⁾ oczywistym warunkiem

$$\tau_2 > \tau_1 > 0. \quad (97)$$

Wyrażenie (93) można przedstawić w postaci

$$\varepsilon(0) = -\frac{1}{2(2A_2 + \sigma_{01}\xi_0)} [2\sigma_{01}\xi_0(2A_2 - \sigma_{01}\xi_0)\tau_1^2 + 4\sigma_{01}\xi_0\dot{\varepsilon}(0)\tau_1 - \dot{\varepsilon}^2(0)] \quad (98)$$

Związek (98) wiąże parametr sygnału wejściowego A_2 i warunki początkowe $\varepsilon(0)$ i $\dot{\varepsilon}(0)$ z chwilą pierwszego przełączenia τ_1 . Podstawiając (98) do równań parametrycznych (90) i następnie eliminując z nich czas τ_1 uzyskamy związek między uchybem ε i jego pochodną $\dot{\varepsilon}$ w chwili pierwszego przełączenia

$$\varepsilon - \frac{\dot{\varepsilon}^2}{2(2A_2 + \sigma_{01}\xi_0)} = 0. \quad (99)$$

Związek (99) jest równaniem krzywej przełączania na płaszczyźnie fazowej $(\varepsilon, \dot{\varepsilon})$ [por. (34)].

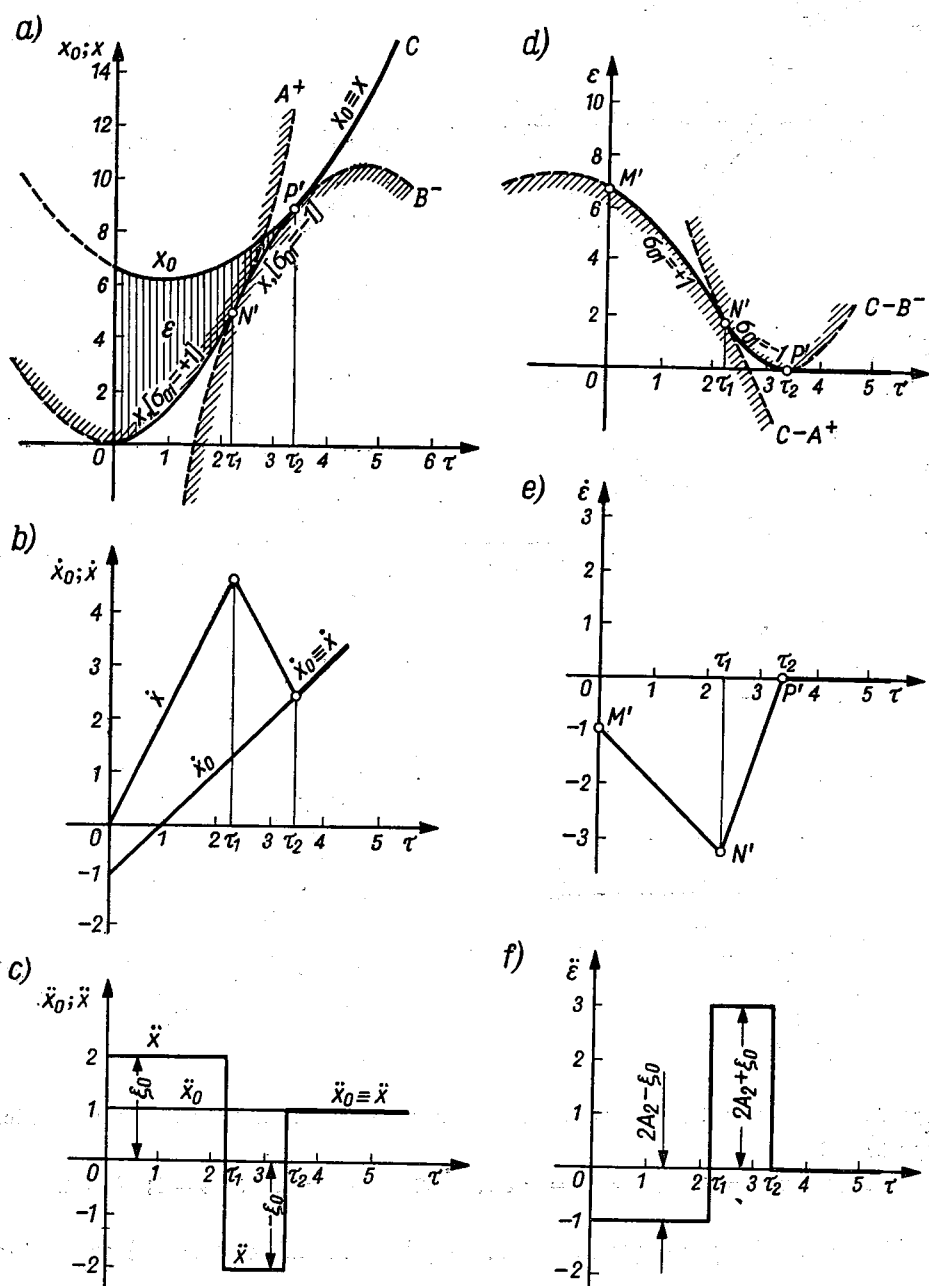
Odejmując drugie z równań (91) od pierwszego z równań (90) po elementarnych przekształceniach otrzymujemy

$$\tau_2 - \tau_1 = -\frac{\dot{\varepsilon}(\tau_1)}{2A_2 + \sigma_{01}\xi_0}. \quad (100)$$

Z uwagi na nierówność (96) oraz nierówność (97) wynika, że

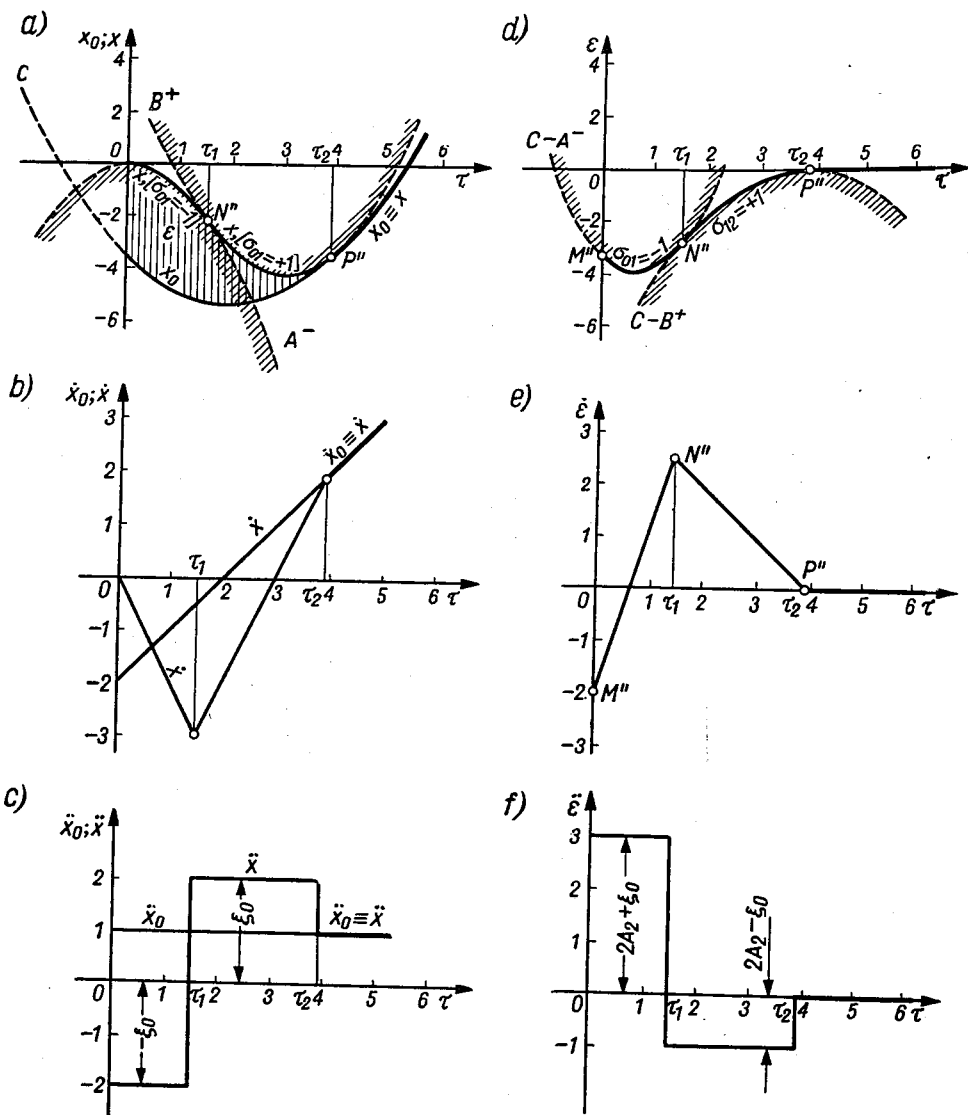
$$\sigma_{01} = -\operatorname{sgn} \dot{\varepsilon}(\tau_1). \quad (101)$$

¹⁾ Czynniki przed pierwiastkiem w wyrażeniu (93) jest mniejszy od czynnika w wyrażeniu (94).



Rys. 10. Przebiegi optymalne w serwo mechanizmie przekąźnikowym z ograniczeniem typu $|\ddot{x}| \leq \xi_0$

$$x(0) = \dot{x}(0) = 0; \quad \epsilon(0) = A_0 = 6,8; \quad \dot{\epsilon}(0) = A_1 = -1; \quad 2A_2 = 1; \quad \xi_0 = 2$$



Rys. 11. Przebiegi optymalne w serwo mechanizmie przekąźnikowym z ograniczeniem typu $|\ddot{x}| \leq \xi_0$

$$x(0) = \dot{x}(0) = 0; \quad \varepsilon(0) = A_0 = -3,4; \quad \dot{\varepsilon}(0) = A_1 = -2; \quad 2A_2 = 1; \quad \xi_0 = 2$$

Równanie krzywej przełączania (99) można zatem napisać w postaci

$$F(\varepsilon, \dot{\varepsilon}, \ddot{x}_0) = \varepsilon - \frac{\dot{\varepsilon}^2}{2(2A_2 - \xi_0 \operatorname{sgn} \dot{\varepsilon})} = 0. \quad (102)$$

Jest to krzywa przełączania typu kwazistacjonarnego, ponieważ kształt jej zależy od parametru sygnału wejściowego — przyspieszenia $\ddot{x}_0 = 2A_2$.

Geometrycznie układ równań (91) prowadzi do wyznaczenia dwóch parabol A i B , stycznych do siebie w punkcie τ_1 , z których jedna ma być styczna do trzeciej paraboli C w punkcie $\tau_2 > \tau_1 > 0$. Drugie pochodne parabol A i B są dane i jednakowe co do wartości bezwzględnej, zaś druga pochodna paraboli C , również dana, jest mniejsza bezwzględnie od drugiej pochodnej paraboli A lub B . Dane jest również położenie paraboli C względem paraboli A (lub B).

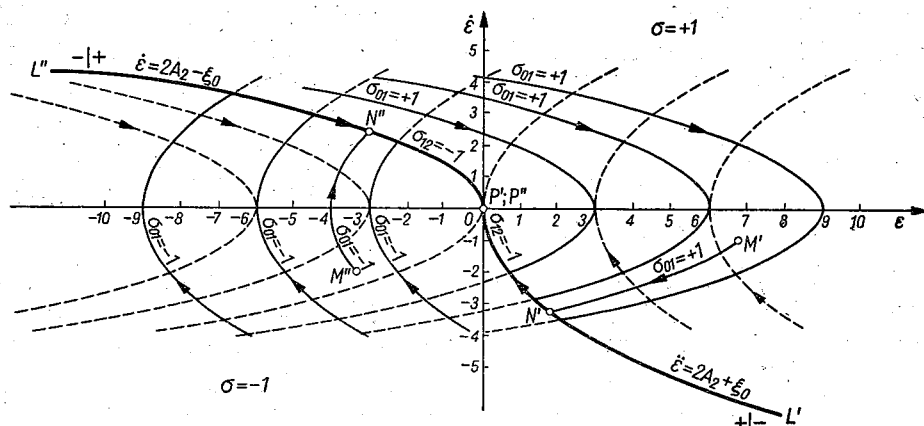
Dwa charakterystyczne przypadki przebiegów optymalnych opisanych równaniami (84) z sygnałem wejściowym typu (81) są przedstawione na rysunkach 10 i 11. W obu przypadkach założono zerowe warunki początkowe $x(0) = \dot{x}(0) = 0$, co nie zwięża jednak ogólności rozważań. W pierwszym przypadku (rys. 10) $\sigma_{01} = +1$, a więc przyspieszenie przebiegu $x(\tau)$ w pierwszym przedziale $0 \leq \tau < \tau_1$ (tj. druga pochodna pierwszej paraboli A^+) $\ddot{x} = \xi_0$ jest dodatnie, natomiast w drugim przedziale $\tau_1 < \tau < \tau_2$ (parabola B^-) $\sigma_{12} = -1$ i przyspieszenie przebiegu $\ddot{x} = -\xi_0$ jest ujemne (rys. 10c); przełączenie w chwili $\tau = \tau_1$ będziemy oznaczać symbolicznie $+|-$ (punkt N_1 na rys. 10a). W drugim przypadku (rys. 11) $\sigma_{01} = -1$ i kolejność jest przeciwna, w chwili zaś $\tau = \tau_1$ następuje przełączenie typu $-|+$. Przebiegi uchybu $\varepsilon(\tau)$ opisane wyrażeniami (89), jako różnice rzędnych parabol C i A oraz C i B są przedstawione na rys. 10d i 11d. Jednakże przyspieszenia sygnału uchybu $\varepsilon(\tau)$ nie są równe co do wartości bezwzględnej w obu przedziałach czasowych. W pierwszym przypadku w przedziale $0 < \tau < \tau_1$ przebieg uchybu (parabola $C - A^+$) ma przyspieszenie ujemne $\ddot{\varepsilon}(\tau) = 2A_2 - \xi_0$ (rys. 10f), a w przedziale $\tau_1 < \tau < \tau_2$ przebieg uchybu (parabola $C - B^-$) ma przyspieszenie dodatnie $\ddot{\varepsilon} = 2A_2 + \xi_0$. W drugim przypadku — przeciwnie. Zatem w jednym z przedziałów przyspieszenie uchybu jest różnicą przyspieszeń sygnału wejściowego $\ddot{x}_0 = 2A_2$ i sygnału wyjściowego $\ddot{x}$, a w drugim — sumą, przy czym znak przyspieszenia uchybu w danym przedziale jest przeciwny do znaku przyspieszenia sygnału wyjściowego.

Dla czasów $\tau \geq \tau_2$, $x_0(\tau) = x(\tau)$ i $\varepsilon(\tau) = 0$ zachowanie się optymalnych serwomechanizmów przekładnikowych w stanie ustalonym zostanie rozpatrzone w rozdz. 4.

Rozpatrzmy przebiegi optymalne na płaszczyźnie fazowej $(\varepsilon, \dot{\varepsilon})$ (rys. 12). Równanie krzywej przełączania (102) można napisać w postaci równań jej dwóch gałęzi:

$$\varepsilon - \frac{\dot{\varepsilon}^2}{2(2A_2 + \xi_0)} = 0 \quad \text{przy } \operatorname{sgn} \dot{\varepsilon} = -1; \quad (\sigma_{01} = +1), \quad (103)$$

$$\varepsilon - \frac{\dot{\varepsilon}^2}{2(2A_2 - \xi_0)} = 0 \quad \text{przy } \operatorname{sgn} \dot{\varepsilon} = +1; \quad (\sigma_{01} = -1). \quad (104)$$



Rys. 12. Tory przebiegów optymalnych na płaszczyźnie fazowej: $\xi_0 = 2$, $2A_2 = 1$

Gałąź (103) — L' — leży w 4 kwadrancie płaszczyzny $(\varepsilon, \dot{\varepsilon})$; gałąź (104) — L'' — w 2 kwadrancie. Takie położenie krzywej przełączania jest właściwe dla wszystkich układów przekąźnikowych drugiego rzędu.

Z równań (88) i (89) uwzględniając (90) łatwo uzyskać ogólną postać równania toru fazowego:

$$\varepsilon - \frac{\dot{\varepsilon}^2}{2(2A_2 - \sigma_{01}\xi_0)} - \left[\varepsilon(0) - \frac{\dot{\varepsilon}^2(0)}{2(2A_2 - \sigma_{01}\xi_0)} \right] = 0, \quad 0 \leq \tau \leq \tau_1, \quad (105)$$

$$\varepsilon - \frac{\dot{\varepsilon}^2}{2(2A_2 - \sigma_{12}\xi_0)} - \left[\varepsilon(\tau_1) - \frac{\dot{\varepsilon}^2(\tau_1)}{2(2A_2 - \sigma_{01}\xi_0)} \right] = 0, \quad \tau_1 \leq \tau < \tau_2,$$

gdzie $\sigma_{12} = -\sigma_{01}$ — wartość σ w przedziale $\tau_1 \leq \tau \leq \tau_2$.

Zgodnie ze (101) tory (105) w przedziale $0 \leq \tau \leq \tau_1$ prowadzące do gałęzi L' charakteryzują się wartością $\sigma_{01} = +1$, a więc punkty leżące na gałęzi L' reprezentują przełączanie typu $+|-$. Tory prowadzące do gałęzi L'' charakteryzują się wartością $\sigma_{01} = -1$ i punkty gałęzi L'' są typu $-|+$. Zatem (zob. rozdz. 2) krzywa przełączania dzieli płaszczyznę fazową na dwa obszary: obszar nad krzywą przełączania, określony wartością $\sigma = +1$ i obszar pod krzywą przełączania, określony wartością $\sigma = -1$. W wyniku tego tory (105) w przedziale $0 \leq \tau \leq \tau_1$, oznaczone na rysunku 12 linią przerywaną nie mogą mieć miejsca wobec niezgodności znaków σ i σ_{01} ¹⁾.

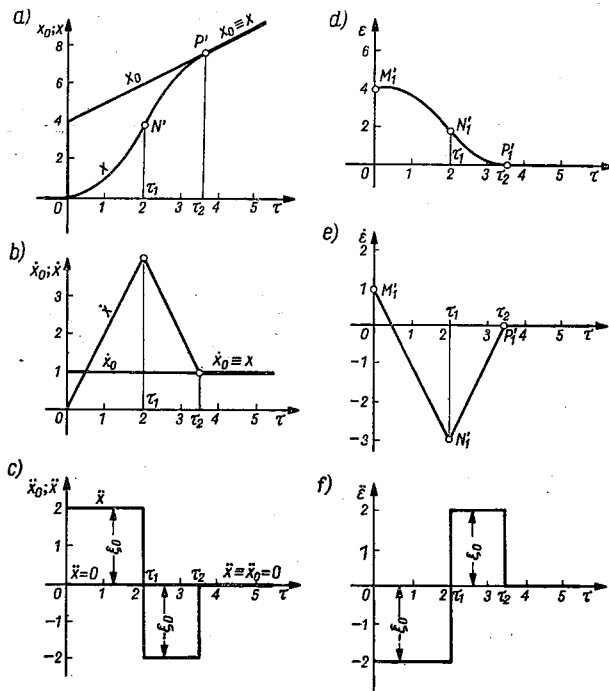
Łatwo zauważyć, że drugim odcinkiem toru fazowego, tj. toru w przedziale $\tau_1 \leq \tau \leq \tau_2$ może być tylko odcinek odpowiedniej gałęzi krzywej

¹⁾ Należy pamiętać, że σ_{01} jest oznaczeniem wielkości σ w pierwszym przedziale czasowym. Oznaczenie „ σ ” ma charakter ogólny.

przełączania¹⁾. Istotnie składnik w nawiasach kwadratowych w drugim z wyrażeń (105) jest, zgodnie z wyrażeniem (99), równy zero — w chwili τ_1 punkt opisujący znajduje się na krzywej przełączania. Zatem drugie z równań (105) jest identyczne z równaniem (103) lub (104), a więc odcinek toru fazowego w przedziale $\tau_1 \leq \tau \leq \tau_2$ leży nieskończenie blisko gałęzi L' w obszarze $\sigma = -1$ lub nieskończenie blisko gałęzi L'' w obszarze $\sigma = +1$ (ponieważ na samej krzywej przełączania wartość σ jest równa zero).

Tak więc, w przypadku ogólnym, tor fazowy optymalnego przebiegu przejściowego składa się z dwóch łuków parabolicznych, opisanych równaniami (105). W chwili $\tau = \tau_2$ $\varepsilon(\tau_2) = \dot{\varepsilon}(\tau_2) = 0$, punkt opisujący osiąga początek układu współrzędnych i przebieg przejściowy zostaje zakończony.

Tory fazowe $M'N'P'$ i $M''N''P''$ omawianych powyżej przebiegów optymalnych są przedstawione na rys. 12 zgodnie z oznaczeniami odpowied-



Rys. 13. Przebiegi optymalne przy sygnale liniowym

$$x(0) = \dot{x}(0) = 0; \quad \varepsilon(0) = A_0 = 4; \quad \dot{\varepsilon}(0) = A_1 = 1; \quad \ddot{\varepsilon}(0) = 2$$

nich wykresów czasowych na rysunkach 10d i 10e oraz 11d i 11e. Zatem pierwszymi odcinkami torów fazowych są odcinki $M'N'$ i $M''N''$, drugi-

¹⁾ Może to w zasadniczy sposób uprościć wyznaczenie krzywej przełączania w optymalnych UAR drugiego rzędu. W niniejszej pracy własność ta celowo nie została wykorzystana z uwagi na ogólność rozważań.

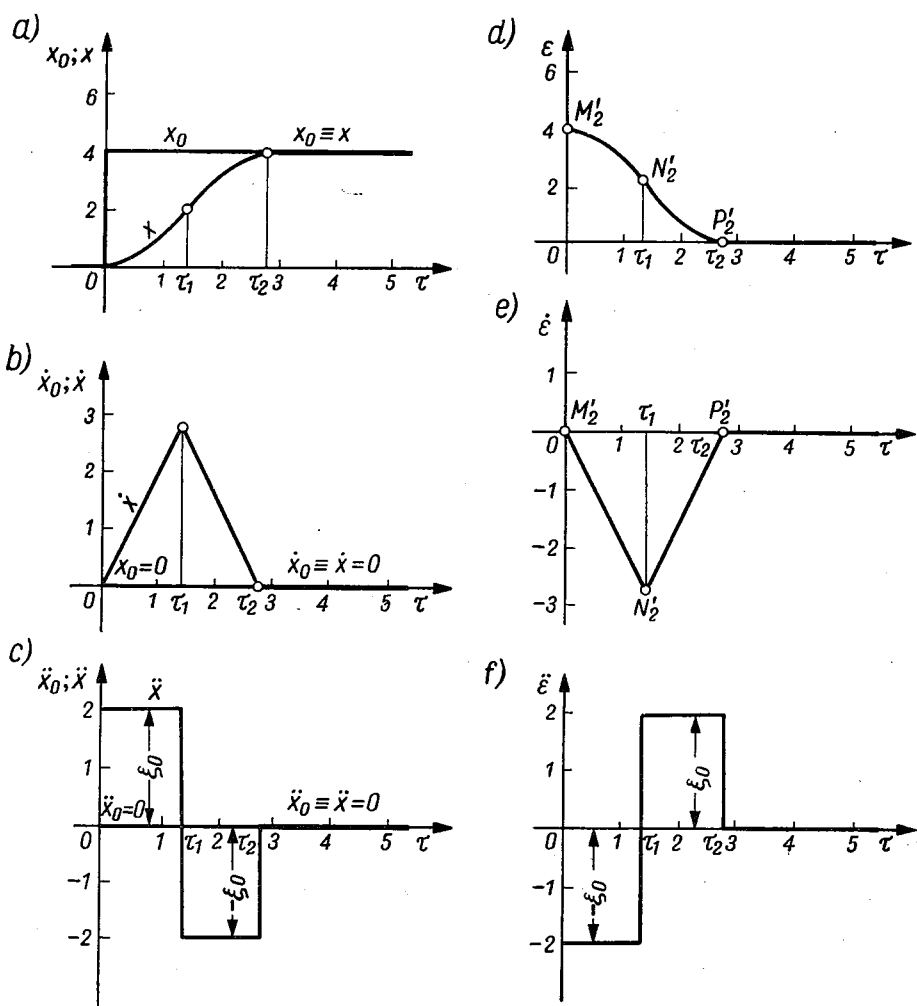
mi — $N'O$ i $N''O$. Punkt N' jest typu $+|-$, punkt N'' — typu $-|+$. Odcinki $N'O$ (a więc gałąź L') i $M''N''$ odpowiadają przyspieszeniu dodatniemu, równemu sumie przyspieszeń sygnału wyjściowego i wejściowego, odcinki $N''O$ (a więc gałąź L'') i $M'N'$ — przyspieszeniu ujemnemu, równemu różnicy tych przyspieszeń.

Jeżeli przyspieszenie sygnału wejściowego $2A_2 = 0$ i serwomechanizm ma odtwarzać sygnał liniowo narastający

$$x_0(\tau) = (A_0 + A_1 \tau) \mathbf{1}(\tau) \quad (106)$$

lub sygnał skokowy

$$x_0(\tau) = A_0 \mathbf{1}(\tau), \quad (107)$$



Rys. 14. Przebiegi optymalne przy sygnale skokowym

$$x(0) = \dot{x}(0) = 0; \quad \epsilon(0) = A_0 = 4; \quad \dot{\epsilon}(0) = 0; \quad \ddot{\epsilon}(0) = -2$$

to równanie krzywej przełączania (102) upraszcza się:

$$F(\varepsilon, \dot{\varepsilon}) = \varepsilon - \frac{\dot{\varepsilon}^2}{2\xi_0 \operatorname{sgn} \dot{\varepsilon}} = 0. \quad (108)$$

Funkcja $F(\varepsilon, \dot{\varepsilon}, \ddot{x}_0)$, gdzie $\ddot{x}_0 = 2A_2$ jest traktowane jako parametr spełniający warunek (96), jest analityczna w całej płaszczyźnie $(\varepsilon, \dot{\varepsilon})$, z wyjątkiem punktów $\dot{\varepsilon} = 0$, w których występuje nieciągłość pierwszego rodzaju. Łatwo spostrzec, że znak funkcji F jest różny po obu stronach krzywej przełączania i zgodny ze znakiem σ . Zatem [por. (37) i (38)]

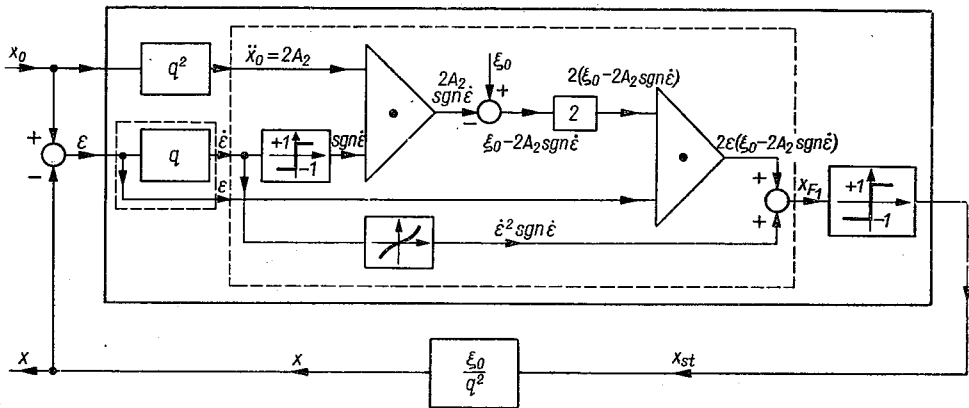
$$\operatorname{sgn} F(\varepsilon, \dot{\varepsilon}, \ddot{x}_0) = \sigma$$

lub

$$\operatorname{sgn} x_F(\tau) = \sigma(\tau).$$

(111)

Realizacja techniczna sygnału (110) wymaga zastosowania członów wykonujących operację dzielenia, co nastręcza trudności praktyczne. Korzystniej jest zatem operować funkcją F_1 , której znak jest zgodny ze zna-



Rys. 16. Schemat blokowy optymalnego serwomechanizmu przekąźnikowego z ograniczeniem typu $|x| \leq \xi_0$, pracującego na sygnał wejściowy o postaci

$$x_0(\tau) = (A_0 + A_1 \tau + A_2 \tau^2) \cdot 1(\tau)$$

kiem funkcji F , lecz w której nie występuje dzielenie. Zagadnienie syntezy jest zatem niejednoznaczne (por. rozdz. 2). Taka funkcja F_1 może mieć na przykład postać:

$$F_1(\varepsilon, \dot{\varepsilon}, \ddot{x}_0) = -2(2A_2 - \xi_0 \operatorname{sgn} \dot{\varepsilon})F = 2\varepsilon(\xi_0 - 2A_2 \operatorname{sgn} \dot{\varepsilon}) + \dot{\varepsilon}^2 \operatorname{sgn} \dot{\varepsilon}, \quad (112)$$

z której wynika, że

$$\operatorname{sgn} F_1 = \operatorname{sgn} F. \quad (113)$$

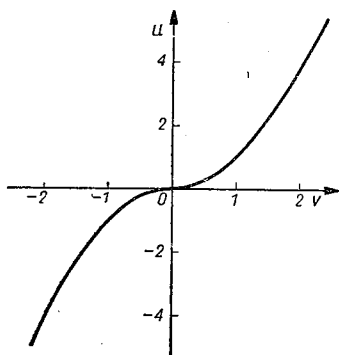
Zatem strukturę układu warunkuje sygnał operacyjny o postaci:

$$x_{F_1}(\tau) = 2\varepsilon(\tau)[\xi_0 - 2A_2 \operatorname{sgn} \dot{\varepsilon}(\tau)] + \dot{\varepsilon}^2(\tau) \operatorname{sgn} \dot{\varepsilon}(\tau). \quad (114)$$

Warto dodać, że z racji ciągłości funkcji F_1 w otoczeniu krzywej przełączania w przedziale $\tau_1 < \tau < \tau_2$ przekąźnik idealny pracuje na progu wzbudzenia, tj. przy wartościach granicznych $x_{F_1} = 0+$ lub $x_{F_1} = 0-$.

Schemat blokowy układu jest przedstawiony na rys. 16. Człon podwójnie różniczkujący q^2 , generujący drugą pochodną sygnału wejściowego x_0 ,

spełnia w danym przypadku zadanie organu prognozy¹⁾ (por. rys. 3 i 9). Rolę organu rejestrującego, dostarczającego danych o sygnale uchybu ε i jego pochodnej $\dot{\varepsilon}$ spełnia człon różniczkujący (q), poprzedzony węzłem



Rys. 17. Charakterystyka przetwornika nieliniowego $u = v|v| = v^2 \operatorname{sgn} v$

zaczepowym. Dane z organów prognozy i rejestrującego zostają wprowadzone do organu operacyjnego, zaznaczonego podobnie jak organ rejestrujący linią przerywaną. Sygnał $\ddot{x}_0 = 2A_2$ zostaje skierowany na jedno z wejść członu mnożącego, oznaczonego trójkątem. Na drugie wejście członu mnożącego jest doprowadzony sygnał $\operatorname{sgn} \dot{\varepsilon}$ wytworzony w członie sygnatury. Sygnał wyjściowy członu mnożącego — iloczyn $2A_2 \operatorname{sgn} \dot{\varepsilon}$ — zostaje w węźle sumującym dodany ze znakiem ujemnym do wprowadzonego z zewnątrz stałego składnika ξ_0 . Suma $\xi_0 - 2A_2 \operatorname{sgn} \dot{\varepsilon}$, pomnożona przez 2, tj. $2(\xi_0 - 2A_2 \operatorname{sgn} \dot{\varepsilon})$, oraz sygnał uchybu ε są skierowane na oba wejścia drugiego członu mnożącego. Iloczyn $2\varepsilon(\xi_0 - 2A_2 \operatorname{sgn} \dot{\varepsilon})$ zostaje wprowadzony jako jeden ze składników do węzła sumującego. Drugi składnik, a mianowicie $\dot{\varepsilon}^2 \operatorname{sgn} \dot{\varepsilon}$, zostaje wytworzony po przejściu sygnału $\dot{\varepsilon}$ przez bezinercyjny przetwornik nieliniowy, o charakterystyce przedstawionej na rysunku 17. Ogólnie wielkość wyjściowa u tego rodzaju przetwornika jest związana z wielkością wejściową v zależnością

$$u = v|v| = v^2 \operatorname{sgn} v. \quad (115)$$

Suma obu składników stanowi sygnał operacyjny x_{F_1} o postaci (114), działający na przekąźnik idealny, który z kolei steruje (sygnał x_{st}) silnikiem o przepustowości $\frac{\xi_0}{q^2}$. Część sterująca i część wykonawcza (silnik) serwomechanizmu są wyodrębnione na schemacie blokowym tłustymi liniami. Pozostałe elementy schematu nie wymagają omówienia.

¹⁾ W danym przypadku druga pochodna sygnału wejściowego wystarczająco określa jego przebieg (zob. rozdz. 2). Wynika to z ograniczenia się do stosunkowo wąskiej klasy sygnałów parabolicznych, jakie serwomechanizm ma odtwarzać.

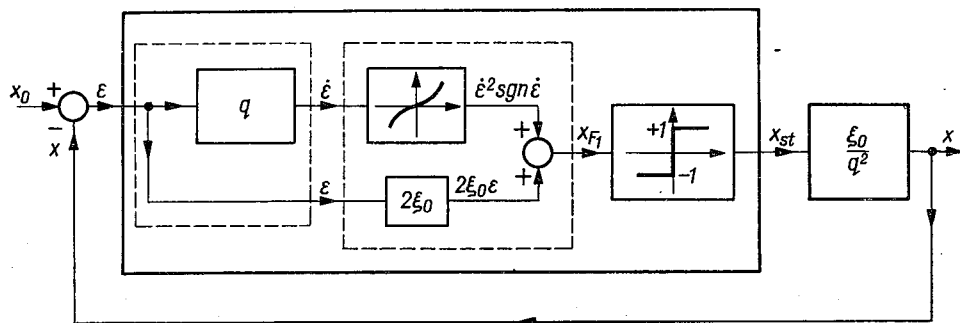
Łatwo zauważyć, że nawet sygnał operacyjny x_{F_1} nie określa jednoznacznie struktury układu. Wystarczy na przykład w wyrażeniu (114) wykonać działanie mnożenia i pozbyć się nawiasów kwadratowych, a struktura organu operacyjnego ulegnie zmianie.

Nie wnikając w szczegóły schematowe należy zaznaczyć, że człony różniczkujące, mnożące jak i przetwornik nieliniowy, zawarte w części sterującej serwomechanizmu są analogiczne do odpowiednich elementów operacyjnych maszyn matematycznych o działaniu ciągłym.

Jeżeli klasa sygnałów wejściowych jest węższa i ogranicza się do sygnałów typu (106) lub — w przypadku szczególnym — typu (107), wówczas struktura serwomechanizmu opiera się na wyrażeniu (108) i sygnał operacyjny (114) upraszcza się:

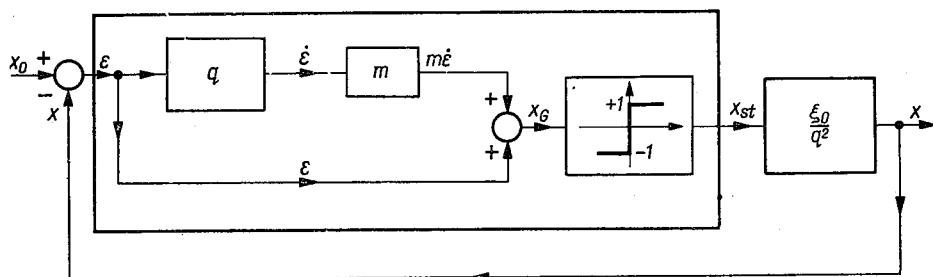
$$x_{F_1}(\tau) = 2\xi_0 \varepsilon(\tau) + \dot{\varepsilon}^2(\tau) \operatorname{sgn} \dot{\varepsilon}(\tau). \quad (116)$$

Odpowiedni schemat blokowy jest przedstawiony na rysunku 18. Jest on znacznie prostszy od schematu poprzedniego. Z racji stacjonarności krzywej przełączania nie zawiera on organu prognozy, tj. członu podwójnie różniczkującego, co ma zasadnicze znaczenie ze względu na występowanie zakłóceń i konieczność ich filtracji. Nie zawiera on również członów mnożących, których rozwiązanie schematowe i konstrukcyjne zawsze nastrocza trudności. Jedyną operacją nieliniową w organie operacyjnym



Rys. 18. Optymalny serwomechanizm przekąźnikowy z ograniczeniem typu $|x| \leq \xi_0$, pracujący na sygnał wejściowy o postaci

$$x_0(\tau) = (A_0 + A_1 \tau) \cdot \mathbf{1}(\tau)$$



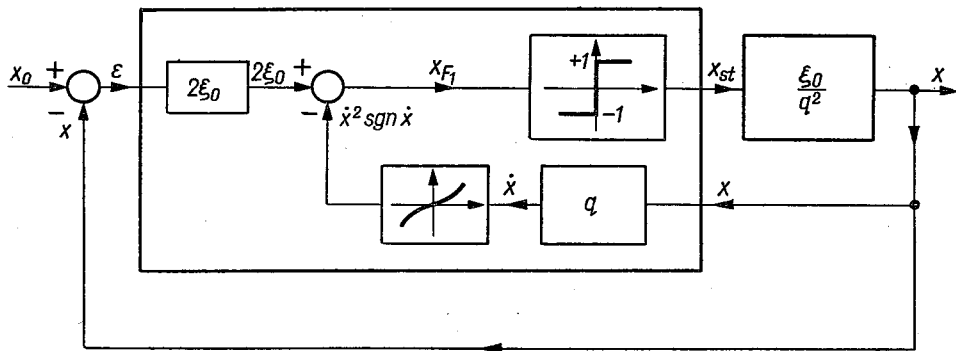
Rys. 19. Schemat blokowy serwomechanizmu przekąźnikowego z wprowadzoną pochodną sygnału uchybu

tego uproszczonego serwomechanizmu jest przekształcenie pochodnej sygnału uchybu ε w przetworniku nieliniowym o charakterystyce przedstawionej na rys. 18. Serwomechanizm ten był rozpatrywany w pracach Cypkina [2], Feldbauma [4], Mac Neale'a [16] i innych. Jest on zbliżony do konwencjonalnego serwomechanizmu przekąźnikowego (rys. 19) z wprowadzoną pochodną uchybu jako sygnałem forsującym [3]. W takim konwencjonalnym serwomechanizmie przełączania dokonuje się nie w punkcie gałęzi parabol (110), lecz w punktach prostej

$$\sigma_1(\varepsilon, \dot{\varepsilon}) = \varepsilon + m \dot{\varepsilon} = 0, \quad (117)$$

gdzie $m = \text{const}$.

Optymalizacja sprowadza się zatem w tym przypadku do bardziej intensywnego forsowania przez wprowadzenie nie pochodnej uchybu, lecz jej kwadratu z odpowiednim znakiem.



Rys. 20. Serwomechanizm optymalny z wewnętrznym sprzężeniem zwrotnym
 $x_0(\tau) = A_0 \cdot 1(\tau)$

Jeżeli sygnał wejściowy jest typu (107), tj. $x_0(\tau) = A_0 = \text{const}$, przy $\tau > 0$, to schemat na rysunku 18 można zastąpić równoważnym mu schematem, przedstawionym na rysunku 20. Serwomechanizm na rys. 20 zawiera wewnętrzne elastyczne sprzężenie zwrotne o sygnale wyjściowym proporcjonalnym do kwadratu pochodnej sygnału wejściowego [15]. Równoważność układów przedstawionych na rysunkach 18 i 20 jest oczywista z uwagi na to, że $\varepsilon = -x$. Zastosowanie wewnętrznego, tachometrycznego sprzężenia zwrotnego jest w danym przypadku korzystniejsze niż korekcja szeregową, która wymaga wprowadzenia kłopotliwych elementów różniczkujących.

3.3. Serwomechanizm optymalny z częścią wykonawczą o przepustowości $K_0(q) = \frac{\xi_0}{q(q+1)}$

Metoda syntezy optymalnego serwomechanizmu przekąźnikowego z silnikiem sterowanym napięciowo i analiza jego własności jest analogicz-

na jak w przypadku serwomechanizmu z ograniczonym momentem obrotowym. Dlatego też obecnie zostaną omówione tylko te własności, które są szczególnie charakterystyczne dla serwomechanizmu o przepustowości części wykonawczej $K_0(q) = \frac{\xi_0}{q(q+1)}$ i różnią go w sposób istotny od układu rozpatrzonego poprzednio.

Ograniczenie nałożone na część wykonawczą ma postać wynikającą z równania (70):

$$\frac{1}{\xi_0} \left| \frac{d^2 x}{d\tau^2} + \frac{dx}{d\tau} \right| \leq X_{st}, \quad (118)$$

gdzie X_{st} — maksymalny moduł sygnału sterującego x_{st} .

Ponieważ $X_{st} = 1$ [zob. (68)], więc

$$\frac{1}{\xi_0} \left| \frac{d^2 x}{d\tau^2} + \frac{dx}{d\tau} \right| \leq 1. \quad (119)$$

Warunek odtwarzalności sygnału wejściowego ma zatem postać:

$$\frac{1}{\xi_0} \left| \frac{d^2 x_0}{d\tau^2} + \frac{dx_0}{d\tau} \right| < 1. \quad (120)$$

A więc z interesujących nas sygnałów parabolicznych wykluczone są sygnały drugiego i wyższych stopni [5]. Wobec tego klasa odtwarzalnych sygnałów parabolicznych ogranicza się do przebiegów typu:

$$x_0(\tau) = (A_0 + A_1 \tau) \mathbf{1}(\tau), \quad (121)$$

w których

$$|A_1| < \xi_0. \quad (122)$$

Zatem serwomechanizm przekaźnikowy z silnikiem sterowanym napięciowo może odtwarzać optymalnie sygnały o stałej prędkości $x_0 = A_1$, lecz nie może odtwarzać sygnałów o stałym przyspieszeniu, w przeciwieństwie do serwomechanizmu z ograniczonym momentem [por. (95)]. Jest to oczywiste z racji ograniczonej prędkości obrotowej silnika w stanie ustalonym.

Odpowiedź jednostkowa części wykonawczej ma postać:

$$h(\tau) = \xi_0(\tau + e^{-\tau} - 1), \quad (123)$$

a rozwiązanie równania jednorodnego

$$x^*(\tau) = x(0) + \dot{x}(0)(1 - e^{-\tau}). \quad (124)$$

Przebieg optymalny $x(\tau)$ można przedstawić zatem w postaci

$$\left. \begin{aligned} x(\tau) &= x(0) + \dot{x}(0)(1 - e^{-\tau}) + \sigma_{01} \xi_0(\tau + e^{-\tau} - 1); & 0 \leq \tau \leq \tau_1, \\ x(\tau) &= x(0) + \dot{x}(0)(1 - e^{-\tau}) + \sigma_{01} [\xi_0(\tau + e^{-\tau} - 1) - 2\xi_0[(\tau - \tau_1) + \\ &\quad + e^{-(\tau - \tau_1)} - 1]]; & \tau_1 \leq \tau \leq \tau_2. \end{aligned} \right\} \quad (125)$$

A więc, uwzględniając (121),

$$\left. \begin{aligned} \varepsilon(0) &= x_0(\tau) - x(\tau) = \varepsilon(0) + \dot{\varepsilon}(0)(1 - e^{-\tau}) + \\ &+ (A_1 - \sigma_{01}\xi_0)(\tau + e^{-\tau} - 1); \quad 0 \leq \tau \leq \tau_1, \\ \dot{\varepsilon}(\tau) &= \dot{\varepsilon}(0)e^{-\tau} + (A_1 - \sigma_{01}\xi_0)(1 - e^{-\tau}); \quad 0 < \tau \leq \tau_1 \end{aligned} \right\} \quad (126)$$

oraz

$$\left. \begin{aligned} \varepsilon(\tau) &= \varepsilon(0) + \dot{\varepsilon}(0)(1 - e^{-\tau}) + (A_1 - \sigma_{01}\xi_0)(\tau + e^{-\tau} - 1) + \\ &+ 2\sigma_{01}\xi_0[(\tau - \tau_1) + e^{-(\tau - \tau_1)} - 1], \\ \dot{\varepsilon}(\tau) &= \dot{\varepsilon}(0)e^{-\tau} + (A_1 - \sigma_{01}\xi_0)(1 - e^{-\tau}) + 2\sigma_{01}\xi_0[1 - e^{-(\tau - \tau_1)}] \end{aligned} \right\} \quad \tau_1 \leq \tau \leq \tau_2. \quad (127)$$

Równania wyjściowe mają wobec tego postać następującą:

$$\left. \begin{aligned} \varepsilon(\tau_1) &= \varepsilon(0) + \dot{\varepsilon}(0)(1 - e^{-\tau_1}) + (A_1 - \sigma_{01}\xi_0)(\tau_1 + e^{-\tau_1} - 1), \\ \dot{\varepsilon}(\tau_1) &= \dot{\varepsilon}(0)e^{-\tau_1} + (A_1 - \sigma_{01}\xi_0)(1 - e^{-\tau_1}) \end{aligned} \right\} \quad (128)$$

oraz

$$\begin{aligned} \varepsilon(\tau_2) &= \varepsilon(0) + \dot{\varepsilon}(0)(1 - e^{-\tau_2}) + (A_1 - \sigma_{01}\xi_0)(\tau_2 + e^{-\tau_2} - 1) + \\ &+ 2\sigma_{01}\xi_0[(\tau_2 - \tau_1) + e^{-(\tau_2 - \tau_1)} - 1] = 0, \\ \dot{\varepsilon}(\tau_2) &= \dot{\varepsilon}(0)e^{-\tau_2} + (A_1 - \sigma_{01}\xi_0)(1 - e^{-\tau_2}) + 2\sigma_{01}\xi_0[1 - e^{-(\tau_2 - \tau_1)}] = 0. \end{aligned} \quad (129)$$

Wyznaczając z równań (129) warunki początkowe $\varepsilon(0)$ i $\dot{\varepsilon}(0)$ jako funkcje czasów τ_1 i τ_2 i podstawiając je do równań (128) otrzymamy

$$\left. \begin{aligned} \varepsilon(\tau_1) &= -(A_1 + \sigma_{01}\xi_0)[(\tau_2 - \tau_1) + e^{\tau_2 - \tau_1} - 1], \\ \dot{\varepsilon}(\tau_1) &= -(A_1 + \sigma_{01}\xi_0)(e^{\tau_2 - \tau_1} - 1). \end{aligned} \right\} \quad (130)$$

Rugując z równań (130) różnicę $\tau_2 - \tau_1$, traktowaną jako parametr, otrzymujemy równanie krzywej przełączania:

$$\varepsilon + \dot{\varepsilon} + (A_1 + \sigma_{01}\xi_0) \ln \left(1 - \frac{\dot{\varepsilon}}{A_1 + \sigma_{01}\xi_0} \right) = 0. \quad (131)$$

Odejmując drugie z równań (130) od pierwszego otrzymamy

$$\tau_2 - \tau_1 = - \frac{\varepsilon(\tau_1) + \dot{\varepsilon}(\tau_1)}{A_1 + \sigma_{01}\xi_0}. \quad (132)$$

Z uwagi na to, że $\tau_2 - \tau_1 > 0$ i $|A_1| < \xi_0$, ze (132) wynika, że

$$\sigma_{01} = -\operatorname{sgn} [\varepsilon(\tau_1) + \dot{\varepsilon}(\tau_1)] \quad (133)$$

i równanie (131) można napisać w postaci

$$F(\varepsilon, \dot{\varepsilon}, \dot{x}_0) = \varepsilon + \dot{\varepsilon} + [A_1 - \xi_0 \operatorname{sgn} (\varepsilon + \dot{\varepsilon})] \ln \left[1 - \frac{\dot{\varepsilon}}{A_1 - \xi_0 \operatorname{sgn} (\varepsilon + \dot{\varepsilon})} \right] = 0. \quad (134)$$

Podobnie jak w rozpatrywanym poprzednio serwomechanizmie z ograniczonym momentem obrotowym, krzywa przełączania (134) jest w przypadku ogólnym kwazistacjonarna i zależy od parametru sygnału wejściowego — prędkości $\dot{x}_0 = A_1$.

lub przy $\sigma_{01} = -1$ i $\dot{\varepsilon}(0) > A_1 + \xi_0$,

ponieważ [zob. (126)]

$$\operatorname{sgn} [\dot{\varepsilon} - (A_1 - \sigma_{01} \xi_0)] = \operatorname{sgn} [\dot{\varepsilon}(0) - (A_1 - \sigma_{01} \xi_0)]. \quad (136)$$

Tory fazowe (135) dążą w każdym przypadku do asymptot

$$\dot{\varepsilon} = A_1 - \sigma_{01} \xi_0 = \text{const}, \quad (137)$$

które reprezentują ustaloną prędkość silnika wykonawczego.

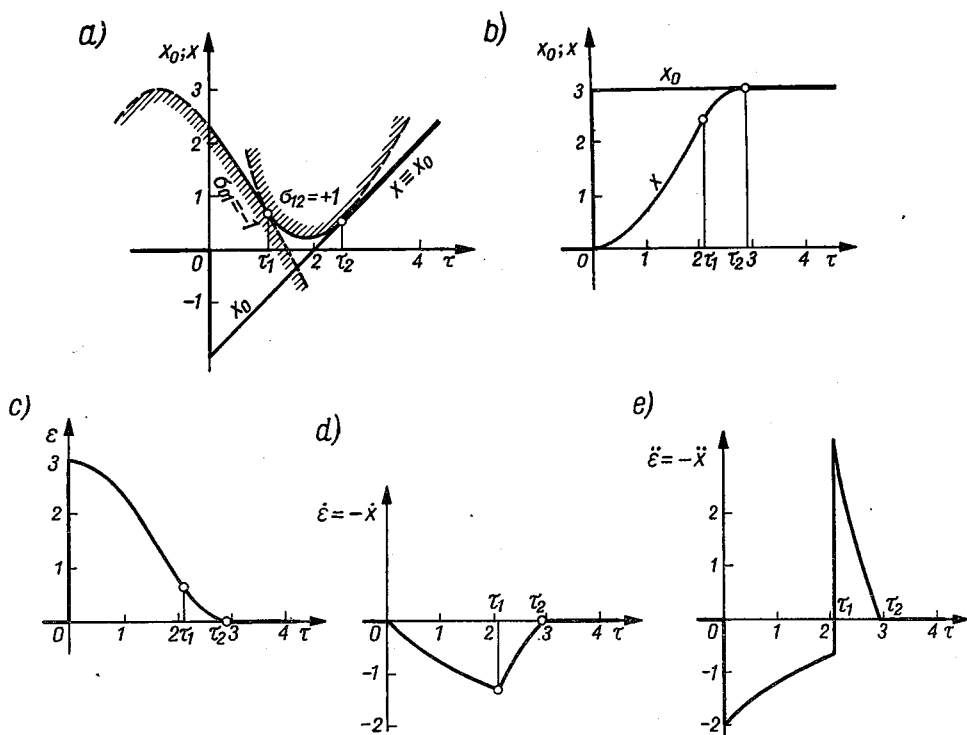
Odcinek toru fazowego w przedziale $\tau_1 \leq \tau \leq \tau_2$, ($\sigma_{12} = -\sigma_{01}$) jest oczywiście identyczny z odcinkiem odpowiedniej gałęzi krzywej przełączania (136), której równanie można napisać w postaci następującej:

$$\varepsilon + \dot{\varepsilon} + (A_1 + \xi_0) \ln \left(1 - \frac{\dot{\varepsilon}}{A_1 + \xi_0} \right) = 0 \quad (138)$$

przy $\operatorname{sgn} (\varepsilon + \dot{\varepsilon}) = -1$; ($\sigma_{01} = +1$),

$$\varepsilon + \dot{\varepsilon} + (A_1 - \xi_0) \ln \left(1 - \frac{\dot{\varepsilon}}{A_1 - \xi_0} \right) = 0 \quad (139)$$

przy $\operatorname{sgn} (\varepsilon + \dot{\varepsilon}) = +1$; ($\sigma_{01} = -1$).



Rys. 22. Przebiegi optymalne w serwomechanizmie z ograniczeniem

$$|\ddot{x} + \dot{x}| \leq \xi_0$$

a) $x(0) = 2,3$; $\dot{x}(0) = -1,2$; $A_0 = -2$; $A_1 = 1$; $\xi_0 = 2$;

b) $\div$ e) $x(0) = \dot{x}(0) = 0$; $A_0 = \dot{\varepsilon}(0) = 3$; $A_1 = \dot{\varepsilon}(0) = 0$; $\xi_0 = 2$

tor M_2N_2 , gdy $\dot{\varepsilon}(0) = \dot{\varepsilon}_g(0) = -\xi_0$,

tor M_3N_3 , gdy $\dot{\varepsilon}(0) < -\xi_0$.

Analogiczne trzy przypadki można wyróżnić w odniesieniu do gałęzi L'' .

Drugi i trzeci przypadek praktycznie nie występują w rzeczywistych układach technicznych, bowiem trudno się spodziewać, aby prędkość początkowa silnika napędowego była równa lub większa od jego prędkości ustalonej. Jeżeli $A_1 \neq 0$ i krzywa przełączania jest kwazistacjonarna, to prędkością początkową graniczną $\dot{\varepsilon}_g(0)$ jest asymptota $\dot{\varepsilon} = A_1 - \xi_0$ (lub $\dot{\varepsilon} = A_1 + \xi_0$).

Struktura układu opiera się na funkcji F , która określa sygnał operacyjny z uwzględnieniem związku (140). Zatem

$$x_F = F(\varepsilon, \dot{\varepsilon}, \dot{x}_0) = \varepsilon + \dot{\varepsilon} + (A_1 - \xi_0 \operatorname{sgn} \dot{\varepsilon}) \ln \left(1 - \frac{\dot{\varepsilon}}{A_1 - \xi_0 \operatorname{sgn} \dot{\varepsilon}} \right). \quad (143)$$

Łatwo zauważyć, że

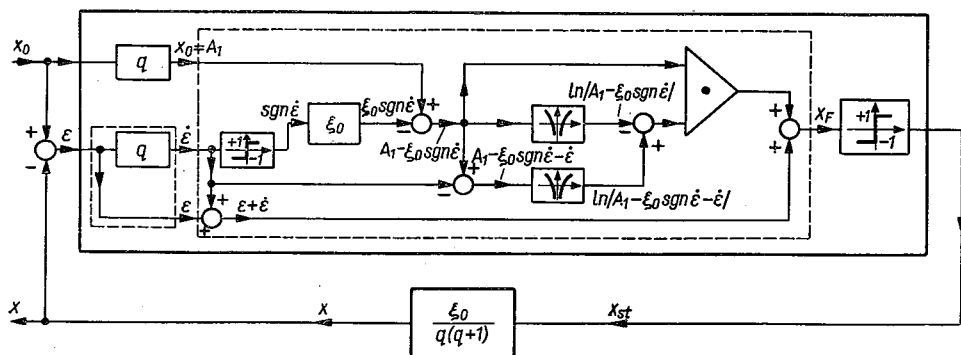
$$\operatorname{sgn} F = \sigma. \quad (144)$$

Funkcja F traci ciągłość w punktach $\dot{\varepsilon} = 0$ i dzięki temu argument logarytmu jest zawsze dodatni. Dlatego też funkcja F jest określona i analityczna w całej płaszczyźnie $(\varepsilon, \dot{\varepsilon})$ z wyjątkiem punktów leżących na osi ε . Z przyczyn wymienionych poprzednio, korzystnie jest przedstawić sygnał operacyjny w postaci

$$x_F = \varepsilon + \dot{\varepsilon} + (A_1 - \xi_0 \operatorname{sgn} \dot{\varepsilon}) [\ln |A_1 - \xi_0 \operatorname{sgn} \dot{\varepsilon} - \dot{\varepsilon}| - \ln |A_1 - \xi_0 \operatorname{sgn} \dot{\varepsilon}|], \quad (145)$$

w którym oczywiście argumenty logarytmów zastąpiono ich wartościami bezwzględny.

Schemat blokowy układu jest przedstawiony na rysunku 24. Do organu operacyjnego zaznaczonego na rysunku linią przerywaną, zostają wprowadzone sygnały wyjściowe z organów prognozy i rejestrującego: $\dot{x}_0 = A_1$, $\dot{\varepsilon}$, ε . Organ operacyjny formujący sygnał x_F nie wymaga spe-

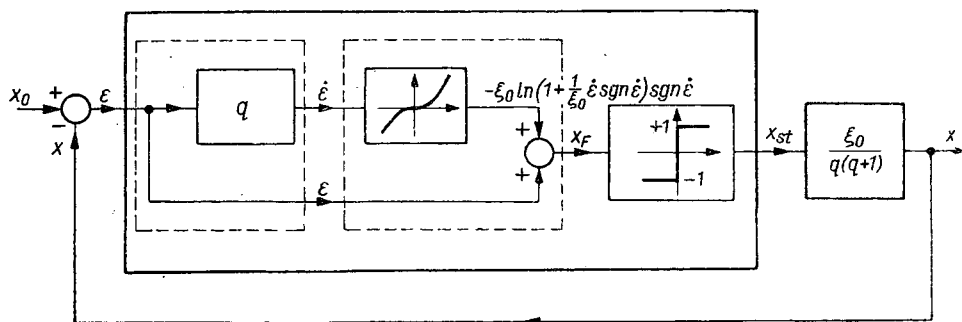


Rys. 24. Schemat blokowy optymalnego serwo mechanizmu przekąźnikowego z ograniczeniem typu $|\ddot{x} + \dot{x}| \leq \xi_0$; $x_0(\tau) = (A_0 + A_1 \tau) \cdot 1(\tau)$

cjalnego omówienia. Jest on znacznie bardziej złożony od analogicznego organu w serwomechanizmie z ograniczonym momentem obrotowym. Zawiera on aż pięć sumatorów i dwa nieliniowe przetworniki logarytmiczne¹⁾, których realizacja techniczna napotyka na poważne trudności.

Jeżeli serwomechanizm ma odtwarzać sygnał skokowy (141), to sygnał operacyjny, a więc i struktura organu operacyjnego znacznie upraszczają się. Wówczas, zgodnie z (143), otrzymujemy

$$x_F = F(\varepsilon, \dot{\varepsilon}) = \varepsilon + \dot{\varepsilon} - \xi_0 \ln \left(1 + \frac{1}{\xi_0} \dot{\varepsilon} \operatorname{sgn} \dot{\varepsilon} \right) \operatorname{sgn} \dot{\varepsilon}. \quad (146)$$



Rys. 25. Struktura serwomechanizmu optymalnego z ograniczeniem

$$|\ddot{x} + \dot{x}| \leq \xi_0; \quad x_0(\tau) = A_0 \cdot 1(\tau)$$

Odpowiedni schemat blokowy jest przedstawiony na rys. 25. Przetwornik nieliniowy w torze sygnału pochodnej uchybu jest określony charakterystyką zgodną z zależnością

$$u = v - \xi_0 \ln \left(1 + \frac{1}{\xi_0} v \operatorname{sgn} v \right) \operatorname{sgn} v. \quad (147)$$

Charakterystyka ta jest symetryczna i dla małych wartości v jest zbliżona do charakterystyki z rys. 17.

Struktura serwomechanizmu z rys. 24 jest praktycznie identyczna ze strukturą serwomechanizmu z rys. 19. Oczywiście i w tym przypadku jest możliwe zastąpienie toru sygnału pochodnej uchybu odpowiednim wewnętrznym sprzężeniem zwrotnym, analogicznie do schematu przedstawionego na rys. 20. Należy jednak zwrócić uwagę na to, że pomimo identycznych struktur obu serwomechanizmów realizacja przetwornika nieliniowego o charakterystyce (147) jest znacznie trudniejsza od realizacji przetwornika o charakterystyce (115), oraz na to, że serwomechanizm z silnikiem sterowanym napięciowo o strukturze przedstawionej na rys. 25 jest w stanie odtwarzać optymalnie jedynie sygnały skokowe.

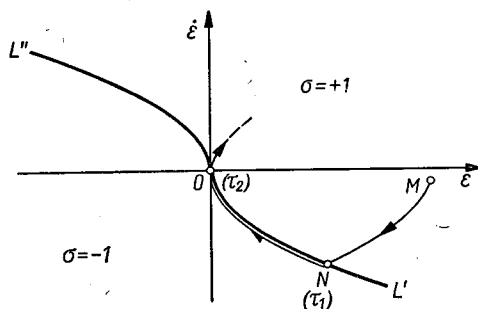
¹⁾ W sensie $u = \ln |v|$.

4. PRZEBIEGI USTALONE W OPTYMALNYCH SERWOMECHANIZMACH PRZĘKAŹNIKOWYCH

W publikacjach zagranicznych dotyczących optymalnych układów regulacji automatycznej, a w szczególności optymalnych serwomechanizmów przekąźnikowych, poświęcono stosunkowo niewiele uwagi analizie przebiegów ustalonych. Zagadnienie to wymaga jednak ścisłego rozpatrzenia nawet w przypadku układów idealnych, tj. układów z przekąźnikiem bez strefy martwej i histerezy i z czasem zadziałania i zwalniania równym zeru. Rozpatrzenie „mechanizmu” przebiegu ustalonego w układzie idealnym stanowić będzie przyczynek do analizy przebiegów w układach realizowanych technicznie.

Rozpatrzmy dla przykładu zagadnienie stanu ustalonego w optymalnym serwomechanizmie przekąźnikowym z ograniczonym momentem obrotowym silnika wykonawczego.

W przypadku odtwarzania sygnału typu (106) lub (107) zagadnienie jest proste. Mechanizm zjawiska jest przedstawiony poglądowo na rys. 26.



Rys. 26. Stan ustalony w optymalnym
serwomechanizmie przekąźnikowym
 $A_2 = 0$

Punkt opisujący biegnąc po torze parabolicznym, leżącym na przykład w obszarze $\sigma = -1$ nieskończenie blisko gałęzi L' osiąga w chwili τ_2 nieskończenie małe otoczenie początku układu współrzędnych — $\epsilon(\tau_2) = \dot{\epsilon}(\tau_2) = 0$ — i przechodzi (tor zaznaczony linią przerywaną) przez krzywą przełączania. Jednakże przejście do obszaru $\sigma = +1$ nie następuje. Gdy punkt opisujący znajdzie się na samej krzywej przełączania sygnał operacyjny $x_{F_1}(\tau_2) = 0$, przekąźnik idealny zwalnia, a ponieważ $\ddot{x}_0 = 0$ i w układzie nie występuje tłumienie $\left[K(q) = \frac{\xi_0}{q^2} \right]$, więc wyłączony w chwili τ_2 silnik bądź zatrzymuje się (jeżeli $A_1 = 0$), bądź też ustala swoją prędkość (jeżeli $A_1 \neq 0$). Zatem również dla $\tau > \tau_2$ $\epsilon(\tau) = \dot{\epsilon}(\tau_2) = 0$ i punkt opisujący zatrzymuje się w początku układu współrzędnych na krzywej przełączania. Wniosek ten wynika również bezpośrednio z równania dynamiki układu

[por. (50)]. Warunek równowagi układu drugiego rzędu ma postać

$$\left. \begin{aligned} \frac{d\varepsilon}{d\tau} &= 0, \\ \frac{d\dot{\varepsilon}}{d\tau} &= 0, \end{aligned} \right\} \quad (148)$$

Podstawiając do równania (50) wyrażenie (112) oraz $x = x_0 - \varepsilon$ i zakładając $A_2 = 0$ otrzymamy

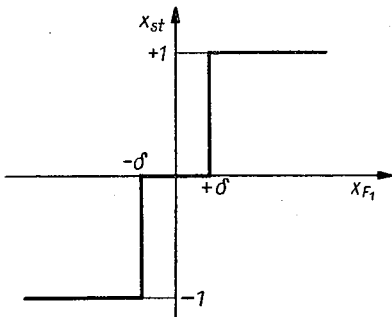
$$\frac{d^2\varepsilon}{d\tau^2} = -\xi_0 \operatorname{sgn} F_1(\varepsilon, \dot{\varepsilon}) = -\xi_0 \operatorname{sgn}(2\xi_0\varepsilon + \dot{\varepsilon}^2 \operatorname{sgn} \dot{\varepsilon}). \quad (149)$$

Zatem zgodnie ze (148) punktem równowagi układu jest początek układu współrzędnych $\varepsilon = \dot{\varepsilon} = 0$.

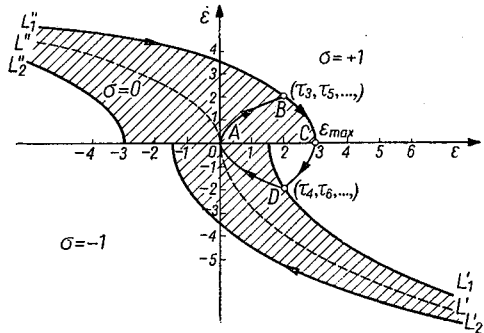
Zagadnienie komplikuje się, jeżeli $A_2 \neq 0$, tzn. jeżeli prędkość sygnału wejściowego nie jest stała, lecz wzrasta liniowo. Wówczas równanie (149) ma postać:

$$\frac{d^2\varepsilon}{d\tau^2} = -\xi_0 \operatorname{sgn}[2\varepsilon(\xi_0 - 2A_2 \operatorname{sgn} \dot{\varepsilon}) + \dot{\varepsilon}^2 \operatorname{sgn} \dot{\varepsilon}] + 2A_2. \quad (150)$$

W tym przypadku warunek równowagi (150) nie jest spełniony w żadnym punkcie płaszczyzny fazowej $(\varepsilon, \dot{\varepsilon})$, bowiem prawa strona wyraże-



Rys. 27. Charakterystyka przekąznika ze strefą martwą o szerokości 2δ



Rys. 28. Cykl graniczny w optymalnym serwomechanizmie przekąznikowym
 $2A_2 = 1$; $\xi_0 = 3$; $\delta = 12$

nia (150) jest zawsze różna od zera. Załóżmy, podobnie jak poprzednio, że punkt opisujący osiągnąwszy nieskończenie małe otoczenie początku układu współrzędnych przeszedł przez krzywą przełączania do obszaru o przeciwnym znaku σ , na przykład $\sigma = +1$, gdyż na samej krzywej przełączania nie mógł się w tym przypadku zatrzymać. W celu ujednoznacznienia rozważań załóżmy również, że od chwili $\tau = \tau_2$ zastąpiliśmy przekąznik idealny przekąznikiem ze strefą martwą (zob. tablica 1) o charakterystyce przedstawionej na rys. 27.

Przy takim założeniu na płaszczyźnie fazowej pojawiają się dwie pary gałęzi krzywej przełączania, odpowiadające wartościom progowym przekąźnika $\pm \delta$ (rys. 28). Zgodnie ze (114) otrzymamy:

$$x_{F_1} = 2\varepsilon(\xi_0 - 2A_2 \operatorname{sgn} \dot{\varepsilon}) + \dot{\varepsilon}^2 \operatorname{sgn} \dot{\varepsilon} = \pm \delta. \quad (151)$$

Zatem para gałęzi L'_1 i L'_2 przy $\operatorname{sgn} \dot{\varepsilon} = -1$ jest opisana równaniem

$$\varepsilon = \frac{\dot{\varepsilon}^2}{2(2A_2 + \xi_0)} \pm \frac{\delta}{2(2A_2 + \xi_0)}, \quad (152)$$

zaś para L''_1 i L''_2 przy $\operatorname{sgn} \dot{\varepsilon} = +1$ —

$$\varepsilon = \frac{\dot{\varepsilon}^2}{2(2A_2 - \xi_0)} \mp \frac{\delta}{2(2A_2 - \xi_0)}, \quad (153)$$

a związek (111) przybiera postać zgodną z charakterystyką przekąźnika:

$$\sigma = \frac{1}{2} [\operatorname{sgn}(x_{F_1} - \delta) + \operatorname{sgn}(x_{F_1} + \delta)]. \quad (154)$$

Gałęzie L'_1 , L'_2 i L''_1 , L''_2 wyodrębniają zatem na płaszczyźnie fazowej obszar $\sigma = 0$ odpowiadający strefie martwej przekąźnika i oddzielającej od siebie obszary $\sigma = +1$ i $\sigma = -1$. „Szerokości” obszaru martwego $\sigma = 0$ przy $\dot{\varepsilon} = 0$ są różne w górnej $\operatorname{sgn} \dot{\varepsilon} = +1$ i dolnej $\operatorname{sgn} \dot{\varepsilon} = -1$ półpłaszczyznach, co wynika z zastąpienia funkcji F funkcją F_1 . Warto dodać, że przy zastosowaniu funkcji F szerokości te są równe, a dla ujednoznacznienia rozważań nie wystarczy zastąpić przekąźnika idealnego przekąźnikiem ze strefą martwą i konieczne staje się wprowadzenie przekąźnika ze strefą martwą i histerezą. Zagadnienie to zostanie omówione później.

Ogólna postać równania toru fazowego w przedziale $\tau_k \leq \tau \leq \tau_{k+1}$, gdzie τ_k — chwila przejścia przez jedną z gałęzi krzywej przełączania, jest następująca [por. (105)]:

$$\varepsilon = \frac{\dot{\varepsilon}^2}{2(2A_2 - \sigma_{k,k+1}\xi_0)} + \left[\varepsilon(\tau_k) - \frac{\dot{\varepsilon}^2(\tau_k)}{2(2A_2 - \sigma_{k,k+1}\xi_0)} \right]. \quad (155)$$

Tak więc, w przedziale $\tau_2 \leq \tau \leq \tau_3$ $\varepsilon(\tau_2) = \dot{\varepsilon}(\tau_2) = 0$, $\sigma_{23} = 0$, silnik jest wyłączony i

$$\varepsilon = \frac{\dot{\varepsilon}^2}{2 \cdot 2A_2}. \quad (156)$$

Jeżeli $A_2 > 0$, to w chwili τ_3 tor (156) przetnie gałąź L''_1 i przekąźnik zadziała przy wartości $x_{F_1} = +\delta$. Ze związków (153) i (156) otrzymamy:

$$\left. \begin{aligned} \varepsilon(\tau_3) &= \frac{\delta}{2\xi_0}, \\ \dot{\varepsilon}(\tau_3) &= +\sqrt{\frac{2A_2\delta}{\xi_0}} \end{aligned} \right\} \quad (157)$$

Zatem w przedziale $\tau_3 \leq \tau \leq \tau_4$ $\sigma_{34} = +1$ i równanie toru fazowego przybiera postać:

$$\varepsilon = \frac{\dot{\varepsilon}^2}{2(2A_2 - \xi_0)} - \frac{\delta}{2(2A_2 - \xi_0)}. \quad (158)$$

Równanie (158) w obszarze $\dot{\varepsilon} > 0$ jest identyczne z równaniem gałęzi L_1'' (153) i w tym obszarze przekąznik pracuje na progu wzbudzenia.

W chwili τ_4 tor (158) przetnie gałąź L_1' i przekąznik zwolni również przy wartości $x_{F_1} = +\delta$. Ze związków (152) i (158) otrzymamy:

$$\left. \begin{aligned} \varepsilon(\tau_4) &= \frac{\delta}{2\xi_0}, \\ \dot{\varepsilon}(\tau_4) &= -\sqrt{\frac{2A_2\delta}{\xi_0}} \end{aligned} \right\} \quad (159)$$

Wobec tego w przedziale $\tau_4 \leq \tau \leq \tau_5$ $\sigma_{45} = 0$ silnik jest ponownie wyłączony i równanie (155) przybiera postać:

$$\varepsilon = \frac{\dot{\varepsilon}^2}{2 \cdot 2A_2}, \quad (160)$$

a więc identyczną jak w przedziale $\tau_2 \leq \tau \leq \tau_3$. Jest to oczywiste z racji symetrii warunków granicznych (157) i (159). Zatem tor fazowy w przedziale $\tau_4 \leq \tau \leq \tau_5$ przechodzi przez punkt $\varepsilon(0) = \dot{\varepsilon}(0) = 0$, współrzędne kolejnego włączenia $[\varepsilon(\tau_5), \dot{\varepsilon}(\tau_5)]$ są identyczne ze (157) i stan ustalony w układzie jest reprezentowany przez stabilny cykl graniczny ABCD. Stabilność cyklu jest oczywista i nie wymaga uzasadnienia.

Amplituda oscylacji wynosi [zob. (158)]

$$\varepsilon_{max} = \frac{\delta}{2(\xi_0 - 2A_2)}, \quad (161)$$

a okres

$$T = \int_{ABCD} \frac{d\varepsilon}{\dot{\varepsilon}} = \int_{\varepsilon(0)}^{\varepsilon(\tau_3)} \frac{d\varepsilon}{\dot{\varepsilon}} + \int_{\varepsilon(\tau_3)}^{\varepsilon_{max}} \frac{d\varepsilon}{\dot{\varepsilon}} + \int_{\varepsilon_{max}}^{\varepsilon(\tau_4)} \frac{d\varepsilon}{\dot{\varepsilon}} + \int_{\varepsilon(\tau_4)}^{\varepsilon(0)} \frac{d\varepsilon}{\dot{\varepsilon}}. \quad (162)$$

Podstawiając do wyrażenia (162) związki (156)÷(161) otrzymamy

$$T = 2 \left(\sqrt{\frac{\delta}{2A_2\xi_0}} + \frac{1}{\xi_0 - 2A_2} \sqrt{\frac{2A_2\delta}{\xi_0}} \right), \quad (163)$$

Jeżeli $A_2 < 0$, to rozumowanie jest analogiczne, a wartości (161) i (163) nie ulegają zmianie.

Cechą charakterystyczną rozpatrywanego stanu ustalonego jest periodyczna styczność przebiegów x_0 i x , odpowiadająca każdorazowemu przejściu punktu opisującego przez początek układu współrzędnych. Jednocześnie warto zwrócić uwagę na poślizgowy charakter pracy przekąznika. W stanie ustalonym, jak łatwo zauważyć, sygnał operacyjny jest zawsze

jednakowego znaku (w naszym przypadku przy $A_2 > 0$ również $x_{F_1} > 0$) i przekładnik pracuje w otoczeniu tylko jednej wartości progowej (w rozpatrywanym przypadku $x_{F_1} = +\delta$). W konwencjonalnych serwomechanizmach przekładnikowych taki charakter pracy może mieć miejsce tylko przy spełnieniu pewnych warunków, jak mały moment bezwładności mas wirujących, odpowiednio dobrany luz mechaniczny itp. [2].

Jeżeli w przekładniku nieidealnym szerokość strefy martwej maleje nieograniczenie, tj. przekładnik zbliża się do idealnego, to gałęzie L'_1 i L'_2 oraz L''_1 i L''_2 dążą do odpowiednich gałęzi L' i L'' , a

$$\left. \begin{aligned} \lim_{\delta \rightarrow 0} \varepsilon_{max} &= 0, \\ \lim_{\delta \rightarrow 0} T &= 0. \end{aligned} \right\} \quad (164)$$

Jeżeli zatem $A_2 \neq 0$, to — w przypadku granicznym — punkt opisujący po osiągnięciu w chwili τ_2 nieskończenie małego otoczenia początku układu współrzędnych będzie oscylował w otoczeniu tego punktu z częstotliwością nieskończenie wielką i nieskończenie małą amplitudą. Dlatego też przebiegi przedstawione na rysunkach 10 i 11 przy $\tau > \tau_2$ należy traktować w sensie wartości średnich.

Podobny charakter stanu ustalonego ma oczywiście optymalny serwomechanizm przekładnikowy z silnikiem sterowanym napięciowo, z tym, że role przyspieszenia i prędkości sygnału wejściowego zostaną odpowiednio zastąpione przez prędkość i skok położenia.

5. ZAKOŃCZENIE

Praca niniejsza nie wyczerpuje oczywiście zagadnień związanych z optymalnymi serwomechanizmami przekładnikowymi drugiego rzędu. Wydaje się konieczne przeprowadzenie szczegółowej analizy tych układów bez założeń upraszczających. Najbardziej istotne, nie opracowane dotychczas zagadnienia, wydają się być następujące: analiza wpływu parametrów, technicznie realizowalnych elementów różniczkujących, przetworników nieliniowych i przekładników na dynamikę rozpatrzonych układów optymalnych. Również celowe będzie opracowanie metod syntezy i analizy optymalnych serwomechanizmów przekładnikowych uwzględniających obciążenie niedynamiczne części wykonawczej i występujące w niej ograniczenia prędkości. Interesujące byłoby przy tym ustalenie możliwości częściowej kompensacji wpływu nieidealnych elementów układu.

Niezależnie od powyższych uwag wydaje się niezbędne przeprowadzenie gruntownej analizy porównawczej serwomechanizmów konwencjonalnych z optymalnymi i wyciągnięcie wniosków odnośnie do stosowności tych ostatnich, zarówno z punktu widzenia dynamiki układów, jak i ich ekonomiki.

WYKAZ LITERATURY

1. Coales J. F., Noton A. R.: An On-Off Servo Mechanisms with Predieted Change-Over. — Proceedings of IEE, t. 103, cz. B, Nr 10, 1956.
2. Cypkin J. Z.: Teorja rieliejnych sistiem awtomatyczskowo riegulirowanija. — Gosudarstwennoje izdatielstwo tiechniko-tieorieticzeskoj litieratury, Moskwa 1955.
3. Feldbaum A. A.: Układy elektryczne regulacji automatycznej. — PWT, Warszawa 1958.
4. Feldbaum A. A.: K woprosu o sintezie optimalnych sistiem awtomatyczskowo riegulirowanija. — Trudy wtorowo wsiesojuznowo sowieszczanija po teorii awtomatyczskowo riegulirowanija, t. II, Izdatielstwo AN SSSR, Moskwa 1955.
5. Feldbaum A. A.: Sintez optimalnych sistiem s pomoszczu fazowowo prostranstwa. — Awtomatika i tielemiechanika, t. XVI, Nr 2, 1955.
6. Findeisen W.: Dobór elementu wykonawczego w serwomechanizmie. Referat na Krajowej Konferencji Automatyki, Warszawa. — Archiwum Automatyki i Telemechaniki, t. 4, zeszyt 1, Warszawa 1959.
7. Gosiewski A.: Niektóre zagadnienia optymalnych serwomechanizmów przekąźnikowych. Referat na Krajowej Konferencji Automatyki, Warszawa. — Archiwum Automatyki i Telemechaniki, t. III, zeszyt 4, Warszawa 1958.
8. Hopkin A. M.: A Phase-Plane Approach to the Compensation of Saturating Servomechanisms. — Transactions of AIEE, t. 70, cz. I, 1951.
9. Kalman R. E.: Analysis and Design Principles of Second and Higher Order Saturating Servomechanisms. — Transactions of AIEE, t. 74, cz. II, 1955.
10. Kazda L. F.: Errors in Relay Servo Systems. — Transactions of AIEE, t. 72, cz. II, 1953.
11. Kazda L. F., Bogner I.: An Investigation of the Switching Criteria for Higher Order Contactor Servomechanisms. — Transactions of AIEE, t. 73, cz. II, 1954.
12. Krasowski A. A.: Integralnyje ocenki. — Osnovy awtomatyczskowo riegulirowanija, pod red. W. W. Sołodownikowa, Maszgiz, Moskwa 1954.
13. Krassowski N. N.: K teorii optimalnowo riegulirowanija. — Awtomatika i Tieliemiechanika, Nr 11, 1957.
14. Kulikowski R.: Optymalne przebiegi w układach automatycznej regulacji. — Archiwum Automatyki i Telemechaniki, t. I, zeszyt 1-2, Warszawa 1956.
15. Lierner A. J.: Postrojenije bystrodejstwujuuszczich sistiem awtomatyczskowo riegulirowanija pri ograniczenii znaczenij koordinat nieregulirujemowo obiekta. — Trudy wtorowo wsiesojuznowo sowieszczanija po teorii awtomatyczskowo riegulirowanija, t. II, Izdatielstwo AN SSSR, Moskwa 1955.
16. Neiswander R. S., Mac Neal R. H.: Optimization of Nonlinear Control Systems by Means of Nonlinear Feedbacks. — Transactions of AIEE, t. 72, cz. II, 1953.
17. Niejmark J. I.: O pieriodiczeskich dwizenijach rieliejnych sistiem. Pamiati Aleksandra Aleksandrowicza Andronowa. — Izdatielstwo AN SSSR, Moskwa 1955.
18. Silva L. M.: Predictor Servomechanism. — Transactions of IRE, t. CT-1, 1954.
19. Silva L. M.: Predictor Control Optimizes Control-System Performance. — Transactions of ASME, t. 77, Nr 8, 1955.
20. Smith O. J.: Feedback Control Systems. — Mc Graw-Hill Book Co., Inc., New York 1958.

21. Sołodownik W. W.: Sintez korriektirujuszczich ustrojstw sistiem awtomaticheskowo riegulirovaniya. — Osnovy awtomaticheskowo riegulirovaniya, pod red. W. W. Sołodownikowa, Maszgiz, Moskwa 1954.
22. Stout T. M.: Switching Errors in an Optimum Relay Servomechanisms. — Proceedings of National Electronic Conference, t. 9, 1953.
23. Stout T. M.: Effects of Friction in an Optimum Relay Servomechanism. — Transactions of AIEE, t. 72, cz. II, 1953.
24. Stout T. M.: On the Comparison of Linear and Nonlinear Servomechanisms Response. Transactions of IRE, t. CT-1, 1954.
25. Truxal J. G.: Automatic Feedback Control Systems Synthesis. — Mc Graw-Hill Book Co., Inc., 1955.
26. Tsien H. S.: Engineering Cybernetics. — Mc Graw-Hill Book Co., Inc., New York 1955.
27. Wiener N.: Extrapolation, Interpolation and Smoothing of Stationary Time Series. — J. Wiley, New York 1954.

A. ГОСЕВСКИ

ОПТИМАЛЬНЫЕ РЕЛЕЙНЫЕ СЛЕДЯЩИЕ СИСТЕМЫ ВТОРОГО ПОРЯДКА

Резюме

В работе представлены и применены элементы усистематизированной теории оптимальных систем автоматического регулирования — систем характеризующихся самым кратким временем переходного процесса — с целью синтеза и анализа основных типов релейных следящих систем второго порядка: с исполнительным двигателем управляемым напряжением и током.

Теоретические рассуждения обоснованы на нескольких упрощающих положениях, идеализирующих рассматриваемые следящие системы, не искажающих однако качественного характера рассматриваемых явлений.

Анализ систем произведён на основании математического аппарата фазовой плоскости, обращая одновременно особое внимание на вопросы до сих пор недостаточно разработанные в литературе (напр. вопрос установившегося режима), а также на физический смысл математического изложения.

A. GOSIEWSKI

OPTIMUM SECOND ORDER RELAY SERVOMECHANISMS

Summary

The elements of systemized theory of the time-optimized automatic control systems i.e. systems distinctive by the shortest response time, and their applications to synthesis and analysis of the basic types of the second order relay servomechanisms with current or voltage controlled servomotor are considered in the paper.

Actually the theoretical considerations are based on a few simplifying assumptions which, though, idealize the examined servomechanisms don't distort however investigated phenomena qualitatively.

The analysis of systems has been carried out with the help of the mathematical instrument of the phase plane. Simultaneously a special attention has been paid to the problems poorly enlightened in the literature (such as the problem of steady state) and to the physical interpretation of the mathematical description.

JANUSZ KACPROWSKI, TADEUSZ UTHKE

Akustyczna metoda rezonansowa analizy mieszanin gazowych

Rękopis dostarczono 8.7.1960

Przedmiotem pracy jest rozwinięcie akustycznej metody rezonansowej pomiaru składu objętościowego mieszanin gazowych. Kryterium pomiarowym metody jest prędkość dźwięku w mieszaninie gazowej, jako parametr, który przy danych prędkościach dźwięku w gazach składowych jest funkcją składu objętościowego mieszaniny.

W teoretycznej części pracy omówiono fizyczne podstawy metody i przeanalizowano jej dokładność oraz zakresy stosowalności z punktu widzenia fizycznej struktury gazów składowych oraz wpływu czynników ubocznych na wynik pomiaru. Szczegółowo przedyskutowano warunki pracy rezonansowego falowodu akustycznego, zastosowanego jako układ do wyznaczania prędkości dźwięku na podstawie pomiaru częstotliwości wzbudzonych w układzie drgań elektrycznych.

W części eksperymentalnej opisano przebieg i podano wyniki pomiarów, których celem było doświadczalne sprawdzenie słuszności założeń i rozważań teoretycznych, a następnie omówiono zaprojektowany i wykonany w Zakładzie Badania Drgań IPPT model urządzenia pomiarowego przeznaczonego do wyznaczania składu objętościowego mieszanin gazowych stosowanych w anestezji. Opracowany przyrząd ma jednak charakter uniwersalny i po wprowadzeniu odpowiednich zmian może być z powodzeniem stosowany przy podobnych pomiarach w technice, przemyśle, klimatyzacji itp.

1. WSTĘP

Zagadnienie analizy ilościowej mieszanin gazowych występuje w wielu procesach technologicznych i produkcyjnych. Stosowane w tym celu klasyczne metody pomiarowe są na ogół dość uciążliwe i pracochłonne i wymagają stosowania kosztownych aparatur. Brak jest metod prostych, pozwalających na przeprowadzanie szybkich, lecz jednocześnie dostatecznie dokładnych pomiarów technicznych.

W praktyce często zachodzi konieczność przeprowadzania analizy ilościowej mieszaniny dwóch znanych gazów o różnych parametrach fizycznych. Analizę taką można wykonać mierząc parametry wypadkowe mieszaniny. Pomiary tego typu są uzasadnione przede wszystkim ogólnymi tendencjami do instrumentalizacji i automatyzacji technicznych metod pomiarowych.

Bardzo wygodną, zwłaszcza w przypadku konieczności ciągłej kontroli mieszaniny gazowej o szybko zmieniającym się składzie, przepływającej ze zmienną prędkością, okazała się metoda oparta na pomiarze parametrów akustycznych mieszaniny. Parametrem takim, dogodnym z punktu widzenia techniki pomiarowej, jest prędkość dźwięku, jako wielkość fizyczna, którą można mierzyć dokładnie za pomocą stosunkowo prostych metod pomiarowych. Analiza składu objętościowego mieszaniny gazowej sprowadza się zatem pośrednio do wyznaczania prędkości dźwięku w ośrodku gazowym, oczywiście tylko w tym przypadku, gdy prędkości dźwięku w gazach wchodzących w skład mieszaniny różnią się między sobą odpowiednio; przypadki takie często zachodzą w praktyce (tablica 1).

Tablica 1

Prędkość fali dźwiękowej w gazach i parach według Beranka [1] i Jeżewskiego [2]. Temperatura 0°, ciśnienie 760 Tr (z wyjątkiem par, dla których prędkość dźwięku podana jest przy ich prężności w stanie nasycenia). Uszeregowanie wg prędkości dźwięku

Lp.	Ośrodek	Symbol chem.	Prędkość dźwięku c [m/sec]	Gęstość ρ [g/l]	Współczynnik $n = \frac{c_p}{c_v}$	Wartości obliczone	
						Względna prędkość dźwięku $\frac{c_{p, w}}{c} = \frac{c_1}{c_2}$	Ilość wykrywana metodą rezonansową α_s [%]
1	Para eteru	$(C_2H_5)_2O$	179,2	—	—	1,850	0,17
2	Dwusiarczek węgla	CS_2	189	—	—	1,752	0,19
3	Chlor	Cl	205,3	3,214	1,355 (15°C)	1,615	0,25
4	Para alkoholu etylowego	C_2H_6O	230,6	—	1,13 (90°C)	1,437	0,39
5	Dwutlenek węgla	CO_2	258,0	1,9769	1,304 (15°C)	1,283	0,6
6	Podtlenek azotu	N_2O	261,8	1,9778	1,32 (15°C)	1,265	0,7
7	Etylen	C_2H_4	317	1,2604	1,255 (15°C)	1,045	4,5
8	Tlen	O_2	317,2	1,42904	1,401 (15°C)	1,045	4,5
9	Argon	A	319	1,7837	1,668 (15°C)	1,038	5
10	Tlenek azotu	NO	325	—	1,400 (15°C)	1,020	10
11	Powietrze suche (bez CO_2)	—	331,45	1,29283	1,40	1,000	—
12	Azot	N_2	337	1,25055	1,404 (15°C)	0,984	11
13	Tlenek węgla	CO	337,1	1,2504	1,404 (15°C)	0,984	11
14	Para wodna	H_2O	401 (OPC) 404,8 (100°C)	—	— 1,324 (100°C)	0,826 0,819	1,3 1,2
15	Amoniak	NH_3	415	0,7710	1,310 (15°C)	0,798	1,1
16	Metan	CH_4	432	0,7168	1,3	0,767	1,0
17	Neon	Ne	435	0,90035	1,64	0,762	1,0
18	Gaz świetlny	—	490,4	—	—	0,677	0,7
19	Hel	He	970	0,17847	1,17	0,342	0,5
20	Wodór	H	1269,5	0,08988	1,420 (15°C)	0,261	0,43

Ze wszystkich znanych metod wyznaczania względnych zmian prędkości dźwięku, wywołanych w omawianym przypadku zmianami składu objętościowego mieszaniny znanych gazów, najbardziej dokładne są metody rezonansowe, oparte na zależności

$$c = \lambda \cdot f, \quad (1)$$

w której:

c — jest prędkością dźwięku w ośrodku,

λ — jest długością fali dźwiękowej,

f — jest częstotliwością drgań akustycznych.

Zakładając stałość jednego z parametrów ($\lambda = \text{const}$, lub $f = \text{const}$) można traktować względne zmiany drugiego parametru (odpowiednio $\frac{\Delta f}{f}$ lub $\frac{\Delta \lambda}{\lambda}$) jako pośrednią miarę zmiany prędkości dźwięku $\frac{\Delta c}{c}$, zależnej od chwilowego składu objętościowego badanej mieszaniny, oczywiście przy zachowaniu stałych wartości pozostałych parametrów fizycznych pomiaru, przede wszystkim temperatury.

Najbardziej dogodnym rezonansowym układem akustycznym służącym do wyznaczania zależności (1) okazał się falowód akustyczny w stanie obustronnego zamknięcia [3], [4], [5], [6].

2. OGÓLNY OPIS REZONANSOWEJ METODY POMIAROWEJ

2.1. Prędkość dźwięku w mieszaninie dwóch gazów

Kryterium pomiarowym metody jest wypadkowa prędkość dźwięku c , będąca funkcją prędkości dźwięku w gazach składowych oraz funkcją proporcji objętościowych mieszaniny.

Dla fali płaskiej w gazach obowiązuje zależność:

$$c = \sqrt{\frac{P \cdot \kappa}{\rho}}, \quad (2)$$

gdzie:

P — ciśnienie statyczne,

$\kappa = \frac{c_p}{c_v}$ — stosunek ciepła właściwego gazu przy stałym ciśnieniu do jego ciepła właściwego przy stałej objętości,

ρ — gęstość gazu.

Odpowiednie prędkości dźwięku w gazach składowych wynoszą:

$$c_1 = \sqrt{\frac{P_1 \cdot \kappa_1}{\rho_1}} \quad \text{oraz} \quad c_2 = \sqrt{\frac{P_2 \cdot \kappa_2}{\rho_2}}. \quad (2a, 2b)$$

Pomiar przeprowadza się przy stałym ciśnieniu $P_1 = P_2 = P$. Niewielkie granice zmienności współczynnika κ dla szeregu gazów ważnych z punktu

widzenia metody pomiarowej (zob. tablica 1) pozwalają również założyć jego stałość: $\kappa_1 = \kappa_2 = \kappa$. Przy tych założeniach prędkość dźwięku jest wyłącznie funkcją wypadkowej gęstości mieszaniny.

Oznaczamy objętość i masę całkowitą mieszaniny przez V i m , oraz objętości i masy gazów składowych odpowiednio przez V_1 , m_1 i V_2 , m_2 . Wówczas

$$V = V_1 + V_2, \quad (3)$$

oraz

$$m = m_1 + m_2. \quad (4)$$

Wprowadzając współczynniki zawartości objętościowej $\alpha_1 = \frac{V_1}{V}$, oraz $\alpha_2 = \frac{V_2}{V}$ można określić gęstość wypadkową jako

$$\varrho = \frac{m}{V} = \frac{m_1 + m_2}{V} = \frac{\varrho_1 \cdot V_1 + \varrho_2 \cdot V_2}{V} = \alpha_1 \cdot \varrho_1 + \alpha_2 \cdot \varrho_2. \quad (5)$$

Ponieważ

$$\varrho_1 = \frac{P \cdot \kappa}{c_1^2} \quad \text{oraz} \quad \varrho_2 = \frac{P \cdot \kappa}{c_2^2}, \quad (6a, 6b)$$

przeto

$$\varrho = P \cdot \kappa \left(\frac{\alpha_1}{c_1^2} + \frac{\alpha_2}{c_2^2} \right). \quad (7)$$

Po wstawieniu otrzymanej wartości do równania (2) otrzymuje się wyrażenie:

$$c = \frac{1}{\sqrt{\frac{\alpha_1}{c_1^2} + \frac{\alpha_2}{c_2^2}}}. \quad (8)$$

2.2. Możliwości pomiarowe metody

Czułość metody pomiarowej wyznaczają dwa niezależne warunki:

$$\frac{c_1}{c_2} \neq 1 \quad \text{oraz} \quad 0 < \alpha_2 \leq 1.$$

Wychodząc z ogólnej postaci wzoru (1), określającego zasadę pomiaru prędkości dźwięku c metodą rezonansową, oraz zakładając kolejno stałość jednego z parametrów $f = \text{const}$ lub $\lambda = \text{const}$, otrzymuje się związki

$$\frac{dc}{c} = \frac{d\lambda}{\lambda} \bigg|_{f = \text{const}}, \quad (9a)$$

oraz

$$\frac{dc}{c} = \frac{df}{f} \bigg|_{\lambda = \text{const.}} \quad (9b)$$

które przedstawiają dwa możliwe warianty rozwiązań technicznych metody pomiaru względnych zmian prędkości dźwięku $\frac{dc}{c}$. Ze względów natury metrologicznej wygodniej jest oprzeć pomiar $\frac{dc}{c}$ na pomiarze względnych zmian częstotliwości $\frac{df}{f}$. Przyjmując, że bezpośredni pomiar techniczny względnych zmian częstotliwości rzędu 1% nie następuje trudności, można określić czułość metody pomiarowej jako równą $\frac{\Delta c}{c_1} = 1\%$. W oparciu o tę wartość, przyjętą chwilowo jako orientacyjną, można wyznaczyć w sposób analityczny możliwości pomiarowe metody przyjmując jako wartości graniczne:

a) minimalną zawartość objętościową jednego z gazów składowych przy założonych wartościach c_1 i c_2 ;

b) minimalną różnicę prędkości dźwięku w gazach składowych $|c_1 - c_2|$ przy założonej proporcji objętościowej a_1 lub a_2 mieszaniny.

Przedstawiając wzór (8) w postaci przyrostowej

$$\Delta c = c - c_1 = \frac{1}{\sqrt{\frac{a_1}{c_1^2} + \frac{a_2}{c_2^2}}} - c_1, \quad (10)$$

a następnie przechodząc do postaci zredukowanej

$$\frac{\Delta c}{c_1} = \frac{c - c_1}{c_1} = \frac{1}{\sqrt{a_1 + a_2 \left(\frac{c_1}{c_2}\right)^2}} - 1 = \frac{1}{\sqrt{1 + \left[\left(\frac{c_1}{c_2}\right)^2 - 1\right] a_2}} - 1, \quad (11)$$

można wyznaczyć kolejno minimalną zawartość a_2 gazu dodatkowego o prędkości dźwięku c_2 w mieszaninie, której podstawowym składnikiem jest gaz o prędkości dźwięku c_1 , przy założonym stosunku $\frac{c_1}{c_2} = \text{const.}$,

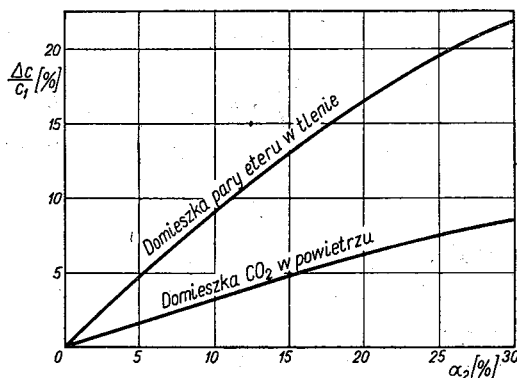
a następnie niezbędny stosunek prędkości gazów składowych $\frac{c_1}{c_2}$, przy założonej proporcji objętościowej mieszaniny, wyrażonej zawartością gazu domieszanego (a_2).

Na rys. 1 przedstawiono wykres zależności (11) dla mieszanin: tlen — eter oraz powietrze — dwutlenek węgla. Przyjętej czułości pomiaru $\frac{\Delta c}{c_1} = 1\%$ odpowiadają minimalne wykrywalne zawartości objętościowe domieszki:

$$\alpha_2 = 2,9\% \text{ CO}_2 \text{ w powietrzu} \quad \left(\frac{c_1}{c_2} = 1,28 \right),$$

$$\alpha_2 = 1,1\% \text{ pary eteru w tlenie} \quad \left(\frac{c_1}{c_2} = 1,77 \right).$$

Czułość metody wzrasta w miarę wzrostu wartości $|c_1 - c_2|$.



Rys. 1. Zmiany prędkości dźwięku w funkcji składu objętościowego mieszaniny gazowej

Własne prace badawcze wykazały, że w warunkach technicznego analizatora gazu można przy użyciu metody rezonansowej wykrywać taką ilość domieszki gazowej, która wywołuje zmianę prędkości dźwięku o wartość $\frac{\Delta c}{c_1} = 0,2\%$. Odpowiadającymi tej czułości najmniejszymi wykrywalnymi w powietrzu ilościami różnych gazów uzupełniono tablicę 1.

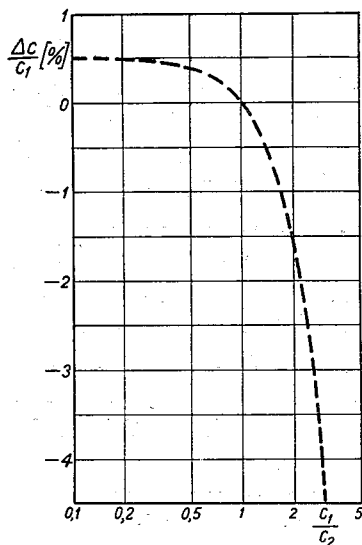
Na rys. 2 przedstawiono wykres ilustrujący czułość metody w funkcji stosunku $\frac{c_1}{c_2}$. Przeliczenie tego wykresu oraz przeliczenie umieszczone w tablicy 1 wykonano przy zastosowaniu wzoru uproszczonego, który wyprowadzono ze wzoru (11) przy założeniu $\alpha_2 \ll 1$:

$$\frac{\Delta c}{c_1} \approx \frac{1}{2} \left[1 - \left(\frac{c_1}{c_2} \right)^2 \right] \alpha_2. \quad (12)$$

Wykreślony na rys. 2 przebieg obrazuje w sposób jaskrawy dodatkowe ograniczenie możliwości pomiarowych metody wskutek asymptotycznego zakrzywienia charakterystyki w kierunku malejącego stosunku $\frac{c_1}{c_2}$:

$$\lim_{\frac{c_1}{c_2} \rightarrow 0} \frac{\Delta c}{c_1} \bigg|_{\alpha_2 = 1\%} = 0,5\%.$$

Wskutek tego dla $\frac{c_1}{c_2} \ll 1$ pomiar jest praktycznie niemożliwy. Przyczyna wynika z zależności $c = \sqrt{\frac{P \cdot \kappa}{\rho}}$. Gęstość gazu o dużej prędkości dźwięku jest niewielka, niewielki więc dodatek gazu o bardzo małej gę-



Rys. 2. Względna zmiana prędkości dźwięku w mieszaninie dla różnych stosunków prędkości c_1/c_2 w gazach składowych. Zawartość dodatku $\alpha_2 = 1\%$ const

stości nie zmienia w sposób zasadniczy gęstości wypadkowej. W rezultacie nawet dla bardzo małych wartości $\frac{c_1}{c_2}$ uzyskana zmiana prędkości dźwięku $\frac{\Delta c}{c_1}$ przy założonej czułości pomiaru $\frac{\Delta c}{c_1} = 1\%$ jest niemierzalna.

W zakresie, w którym $\frac{c_1}{c_2} > 1$, sytuacja jest daleko korzystniejsza, gdyż czułość metody wzrasta ze wzrostem stosunku $\frac{c_1}{c_2}$. Przyjętej czułości pomiaru $\frac{\Delta c}{c_1} = 1\%$ odpowiada tu dopuszczalna wartość stosunku $\frac{c_1}{c_2} \geq 1,73$. Realizacja pomiaru mieszanin o $0 < \frac{c_1}{c_2} < 1,73$ wymaga zwiększenia czułości pomiaru $\left(\frac{\Delta c}{c_1} < 1\%\right)$.

Reasumując można stwierdzić, że w warunkach technicznych stosowność metody jest ograniczona głównie do pomiaru składu objętościowego mieszanin, w których gaz dodatkowy charakteryzuje się mniejszą prędkością dźwięku niż gaz podstawowy ($c_1 > c_2$). Realizacja pomiaru składu obję-

tościowego mieszanin w przypadku gdy $c_1 < c_2$ również wymaga zwiększenia czułości pomiaru zmian prędkości dźwięku.

2.3. Analiza wpływów ubocznych na prędkość dźwięku

Należy przede wszystkim sprawdzić celowość stosowania urządzeń stabilizujących warunki pomiarowe, analizując wpływ wszystkich możliwych parametrów ubocznych na dokładność pomiaru. Dotychczasowe zastosowania metod rezonansowych analizy objętościowej mieszanin gazowych obejmują pomiary zawartości domieszek gazowych w powietrzu (lub w tlenie) pod ciśnieniem atmosferycznym, trzeba więc przeprowadzić rozważania nad możliwymi w tym przypadku wpływami temperatury, ciśnienia statycznego i domieszki przypadkowej w postaci pary wodnej.

Wpływ temperatury na bezwzględną wartość prędkości dźwięku jest dość duży, jednak przy pomiarach względnej zmiany prędkości dźwięku istotne są jedynie wahania temperatury w czasie trwania pomiaru. Odpowiedni uchyb przy założeniu niewielkich zmian temperatury wynosi:

$$\frac{\Delta c}{c_1} \approx \frac{\Delta t}{2(273 + t_p)}, \quad (13)$$

gdzie:

Δt — przyrost temperatury [$^{\circ}\text{C}$],

t_p — temperatura początkowa [$^{\circ}\text{C}$].

Jeżeli przyjmiemy dopuszczalną wartość uchybu $\frac{\Delta c}{c_1} \leq 0,1\%$ oraz $t_p = 20^{\circ}\text{C}$, to dopuszczalne wahanie temperatury wyniesie:

$$\Delta t \leq \frac{\Delta c}{c_1} \cdot 2(273 + t_p) = 10^{-3} \cdot 2 \cdot 293^{\circ}\text{C} \approx 0,6^{\circ}\text{C}.$$

Zmiany ciśnienia statycznego, nawet w dość dużych granicach, nie mają wpływu na prędkość dźwięku. Ponieważ gęstość gazu ρ jest proporcjonalna do ciśnienia, więc prędkość dźwięku $c = \sqrt{\frac{P \cdot \kappa}{\rho}}$ jest niezależna od ciśnienia. Odstępstwa od tej zależności występują dopiero przy ciśnieniach bardzo dużych (ponad 100 at).

Wpływ obecności pary wodnej na prędkość dźwięku w przyjętych warunkach jest niewielki. Zawartość objętościową a_2 pary wodnej w powietrzu określa wzór:

$$a_2 = \frac{P_{pn}}{P} W, \quad (14)$$

gdzie:

P_{pn} — ciśnienie pary nasyconej w danej temperaturze,

P — sumaryczne ciśnienie mieszaniny,

W — wilgotność względna.

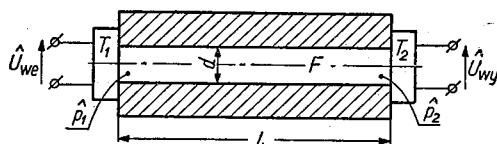
Dla $t_p = 20^\circ\text{C}$, $P_{pn} = 17,5$ Tr. Zakładając $P = 760$ Tr oraz $\frac{\Delta c}{c_1} \leq 0,1\%$ można obliczyć dopuszczalne wahanie wilgotności korzystając ze wzorów (5) i (7):

$$\frac{P_{pn}}{P} W = \frac{2 \Delta c}{c_1} \cdot \frac{1}{1 - \left(\frac{c_1}{c_2}\right)^2}, \quad (15)$$

$$\Delta W \leq \frac{P}{P_{pn}} \cdot \frac{2 \Delta c}{c_1} \cdot \frac{1}{1 - \left(\frac{c_1}{c_2}\right)^2} \frac{760}{17,5} \cdot 2 \cdot 10^{-3} \cdot \frac{1}{1 - (0,826)^2} = 27 \%$$

2.4. Metoda rezonansowa pomiaru prędkości dźwięku

Pomiar przeprowadza się z wykorzystaniem jako rezonatora falowodu akustycznego zamkniętego obustronnie dużymi opornościami akustycznymi czołowych powierzchni dwóch przetworników odwracalnych (rys. 3).



Rys. 3. Rezonator falowodowy

F — wnęka falowodu, T_1 — przetwornik nadawczy, T_2 — przetwornik odbiorczy

Z teorii falowodu bezstratnego o długości l i oporności akustycznej charakterystycznej $Z_c = \frac{\rho_0 c}{S}$ (gdzie S — przekrój rury), zamkniętego opornością akustyczną $\hat{Z}_k$, wynika wyrażenie na oporność wejściową falowodu [7]:

$$\hat{Z}_p = Z_c \frac{\hat{Z}_k + j Z_c \operatorname{tg} \beta l}{Z_c + j \hat{Z}_k \operatorname{tg} \beta l}. \quad (16)$$

Dla falowodu o długości $l = \frac{\lambda}{2} n$, gdzie $n = 0, 1, 2, 3 \dots$,

$$\operatorname{tg} \beta l = \operatorname{tg} \frac{2\pi}{\lambda} \cdot \frac{\lambda}{2} n = \operatorname{tg} (\pi \cdot n) = 0$$

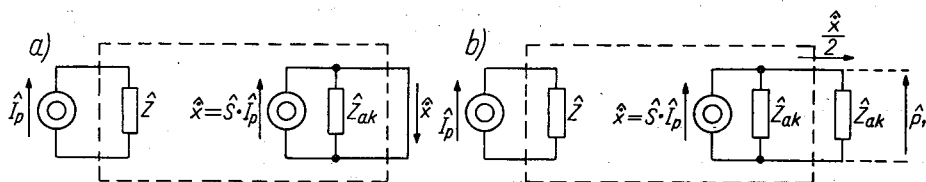
i wówczas

$$\hat{Z}_p = \hat{Z}_k. \quad (17)$$

Oporność wejściowa falowodu równa się jego oporności obciążenia. Falowód bezstratny o długości $n \cdot \frac{\lambda}{2}$ jest więc transformatorem aktycznym o przekładni 1 : 1. Tłumienie wprowadzone do obwodu przez falowód będzie więc dla tego przypadku głównie funkcją skuteczności przetworników.

Aby obliczyć tłumienie wprowadzane przez cały omawiany czwórnik elektryczno-akustyczno-elektryczny, należy wprowadzić następujące założenia:

1) przetworniki są identyczne, odwracalne, posiadające charakter czwornika nieprzezroczystego¹⁾ o oporności strony akustycznej $\hat{Z}_{ak}$ i strony elektrycznej $\hat{Z}$, o danej skuteczności odbiorczej w stanie otwarcia zacisków elektrycznych $\hat{S}$;



Rys. 4. Schemat zastępczy przetwornika nadawczego, nieprzezroczystego, zasilanego ze źródła o stałej wydajności prądowej

a) Stan zwarcia akustycznego, b) Stan zamknięcia impedancją równą wewnętrznej impedancji akustycznej przetwornika.

2) obowiązuje teoria wzajemności bazująca na założeniu równości skuteczności nadawczej i odbiorczej [1] przetworników odwracalnych:

$$\frac{\hat{x}}{\hat{I}_p} = \frac{\hat{E}}{\hat{p}}, \quad (18)$$

gdzie:

- $\hat{x}$ — akustyczna prędkość objętościowa wytwarzana przez przetwornik pobudzany od strony elektrycznej prądem $\hat{I}_p$,
- $\hat{I}_p$ — prąd elektryczny pobudzający przetwornik,
- $\hat{E}$ — siła elektromotoryczna na zaciskach przetwornika pobudzanego od strony akustycznej ciśnieniem $\hat{p}$,
- $\hat{p}$ — ciśnienie akustyczne działające na przetwornik,
- $\hat{S}$ — skuteczność odbiorcza przetwornika w stanie otwarcia zacisków elektrycznych.

Lewa strona równości (18) $\left(\frac{\hat{x}}{\hat{I}_p} \right)$ jest skutecznością nadawczą prze-

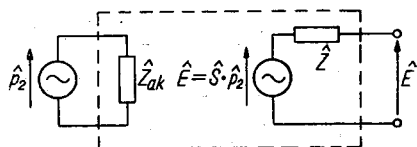
¹⁾ Warunek ten oznacza nieuwzględnianie oporności ruchowej przetwornika.

twornika w stanie zwarcia strony akustycznej, zaś prawa $\left(\frac{\hat{E}}{\hat{p}}\right)$ jest skutecznością odbiorczą w stanie otwarcia zacisków elektrycznych.

Odpowiednie układy zastępcze przetwornika przedstawiono na rys. 4 i 5.

Ze wzoru (17) wynika, że ciśnienie akustyczne $\hat{p}_2$ działające na membranę przetwornika odbiorczego we wnętrzu falowodu jest równe ciśnieniu $\hat{p}_1$ wytwarzanemu przez membranę przetwornika nadawczego (rys. 3).

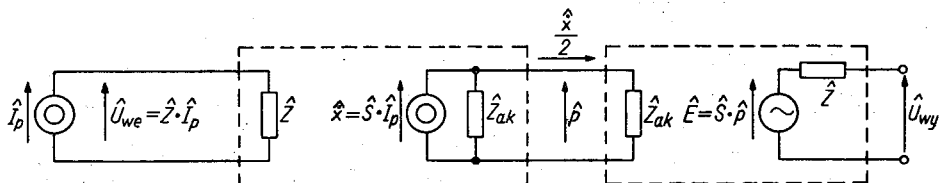
Wniosek ten pozwala na utworzenie pełnego schematu zastępczego falowodu z przetwornikami (rys. 6). Wynikają z niego następujące zależności:



Rys. 5. Schemat zastępczy przetwornika odbiorczego, nieprzezroczystego, zasilanego ze źródła o stałym ciśnieniu akustycznym

ponieważ

$$\hat{p} = \frac{\hat{x}}{2} \hat{Z}_{ak} \quad (19)$$



Rys. 6. Schemat zastępczy układu dwóch jednakowych przetworników odwracalnych, nieprzezroczystych, sprzężonych bezpośrednio ze sobą po stronie akustycznej

oraz

$$\hat{x} = \hat{S} \cdot \hat{I}_p, \quad (20)$$

przeto

$$\hat{p} = \frac{1}{2} \hat{S} \cdot \hat{I}_p \cdot \hat{Z}_{ak} \quad (21)$$

i

$$\hat{U}_{wy} = \hat{E} = \hat{S} \cdot \hat{p} = \frac{1}{2} \hat{S}^2 \cdot \hat{I}_p \cdot \hat{Z}_{ak}. \quad (22)$$

Przechodząc na zasilanie ze źródła o stałej wydajności napięciowej $\hat{U}_{we} = \hat{Z} \cdot \hat{I}_p$ otrzymamy ostatecznie:

$$\hat{U}_{wy} = \frac{1}{2} \hat{S}^2 \frac{\hat{Z}_{ak}}{\hat{Z}} \hat{U}_{we} \quad (23)$$

i przenoszenie napięciowe układu dla stanu otwarcia wyniesie:

$$\frac{\hat{U}_{wy}}{\hat{U}_{we}} = \frac{1}{2} \hat{S}^2 \frac{\hat{Z}_{ak}}{\hat{Z}}. \quad (24)$$

2.5. Warunki pobudzania falowodu do drgań

Ze względów metrologicznych i konstrukcyjnych najkorzystniejszy jest przypadek drgań półfalowych. W celu wyeliminowania możliwości wzbudzania się drgań innego rodzaju należy wprowadzić następujące warunki dodatkowe.

1. Umieszczenie falowodu w pętli sprzężenia zwrotnego ograniczy możliwość generacji drgań do przypadku $l = (2n + 1) \frac{\lambda}{2}$, gdzie $n = 0, 1, 2, 3$ itd., ponieważ drgania mogą się wzbudzać tylko przy warunkach fazowych spełniających warunki generacji.

2. Wybranie odpowiedniego stosunku długości do szerokości falowodu. Kompromisem pomiędzy spełnieniem warunków właściwej pracy falowodu przy częstotliwości rezonansu półfalowego, a niespełnieniem ich dla nieparzystych wielokrotności tej częstotliwości jest stosunek $\frac{l}{d} = 5$.

Wówczas, dla $n = 0$, $\frac{\lambda}{d} = 10$.

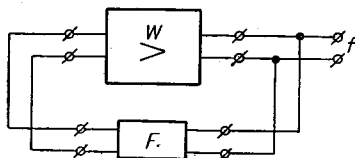
3. Zastosowanie przetworników i wzmacniacza wzbudzającego o ograniczonym pasmie przenoszenia w zakresie dużych częstotliwości.

Stosując w sposób racjonalny wymienione warunki można, jak się okazuje w praktyce, bardzo łatwo spełnić założenia pomiarowe. W dalszych rozważaniach zostaną pominięte możliwości wzbudzenia się drgań innych poza drganiem określonym warunkiem: $n = 0$ i $l = \frac{\lambda}{2}$.

Stąd warunek na częstotliwość rezonansową:

$$f_r = \frac{c}{2l}. \quad (25)$$

Obwód generacji falowodu można przedstawić jako obwód złożony z czwórnika o charakterystyce rezonansowej i wzmacniacza (rys. 7).



Rys. 7. Obwód generacyjny falowodu
W — wzmacniacz, F — czwórnik falowodowy

Z warunku liniowej teorii generacji wynikają zależności:

$$\hat{m} \cdot \hat{k} = 1, \quad (26)$$

a więc

$$m \cdot k = 1, \quad (26a)$$

oraz

$$\varphi_m + \varphi_k = 0, \quad (26b)$$

gdzie:

$\hat{m} = m \cdot e^{j\varphi_m}$ — wartość zespolona współczynnika sprzężenia zwrotnego,

$\hat{k} = k \cdot e^{j\varphi_k}$ — wartość zespolona wzmocnienia.

Przesunięcie fazowe w czwórniku sprzężenia zwrotnego składa się w tym przypadku z przesunięcia φ_f we wnęce falowodu oraz z przesunięcia φ_p wniesionego przez przetworniki:

$$\varphi_m = \varphi_f + \varphi_p. \quad (27)$$

Można założyć, że wzmacniacz nie przesuwają fazy, co praktycznie jest łatwo zrealizować w zakresie średnich częstotliwości akustycznych, a więc

$$f_k = 0, \quad (28)$$

i warunki generacji przybiorą postać:

$$m \cdot k = 1, \quad (29a)$$

oraz

$$\varphi_m = \varphi_k + \varphi_p = 0. \quad (29b)$$

Każde zniekształcenie fazy φ_p wprowadzone przez przetworniki musi być skompensowane w falowodowym czwórniku rezonansowym przez jego oddziaływanie fazowe φ_f :

$$\varphi_f = -\varphi_p. \quad (30)$$

Przesunięcie fazowe wniesione przez obwód rezonansowy można wyrazić wzorem:

$$\varphi_f = \arctg(Q \cdot \nu), \quad (31)$$

gdzie:

Q — dobroć obwodu,

$\nu \approx \frac{2(f - f_r)}{f_r}$ — odstrojenie względne (wartość przybliżona dla małych odstrojeń).

W związku z tym

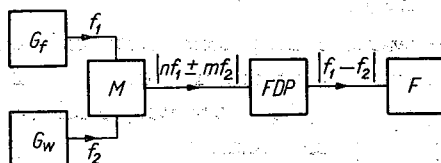
$$\nu = \frac{-\operatorname{tg} \varphi_p}{Q}. \quad (32)$$

Wtrącenie składowej urojonej do obwodu generacyjnego wywiera wpływ na warunek fazy, czyli na częstotliwość [8].

W konkretnym przypadku

$$\hat{m} = \frac{\hat{U}_{wy}}{\hat{U}_{we}} = \frac{1}{2} \hat{S}^2 \cdot \frac{\hat{Z}_{ak}}{\hat{Z}} \quad (33)$$

Przyjmując, że $\hat{S}$ jest zawsze rzeczywiste, oraz eliminując wpływ $\hat{Z}$ na przebiegi (co zawsze można zrealizować, gdy przetworniki będą pracować na bardzo duże oporności elektryczne), stwierdzamy istotny wpływ składowej urojonej oporności akustycznej $\hat{Z}_{ak}$ przetworników na częstotliwość generacji. Chcąc więc, aby częstotliwość generacji była bliska częstotli-



Rys. 8. Schemat blokowy układu pomiaru częstotliwości

G_f — generator faliwodowy, G_w — generator wzorcowy, M — modulator, FDP — filtr dolnoprzepustowy, F — częstotliwościomierz wskazówkowy

wości rezonansowej faliwodu, należy stosować przetworniki o czysto stratnościowym charakterze oporności akustycznej oraz faliwody o jak największej dobroci Q , co stwierdzić można na podstawie wzoru (32).

Ze względu na swój charakter techniczny pomiar częstotliwości wzbudzonych drgań musi się odbywać mało dokładną metodą wychyłową. Dokładność tę, wynoszącą $2 \div 5\%$, można zwiększyć o jeden rząd wielkości przez obniżenie częstotliwości metodą heterodynowania. W tym celu układ pomiarowy należy zaopatrzyć w generator pomocniczy o dużej stałości częstotliwości drgań oraz w modulator i filtr dolnoprzepustowy (rys. 8).

3. BADANIA EKSPERYMENTALNE

3.1. Opis aparatury eksperymentalnej do wyznaczania zawartości pary eteru w powietrzu

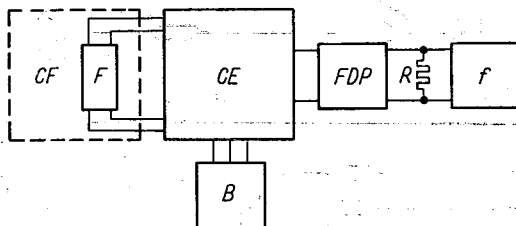
Eksperymentalne sprawdzenie zależności wynikających z założeń teoretycznych opisywanej metody pomiarowej wykonano w Zakładzie Badań Drgań I.P.P.T. na przykładzie mieszaniny pary eteru i powietrza, przy ciśnieniu atmosferycznym, w temperaturze otoczenia. O wyborze zastosowanego rodzaju mieszaniny zdecydowały następujące przesłanki:

- 1) duża czułość pomiaru wynikająca z korzystnego stosunku prędkości c_1 dźwięku w powietrzu do prędkości c_2 dźwięku w parze eteru $\left(\frac{c_1}{c_2} = 1,85\right)$;
- 2) łatwość określenia zawartości pary eteru w naczyniu o znanej objętości przez pomiar masy odparowanego eteru;

3) korzystne właściwości fizyczne eteru: duża prężność pary nasyconej (432 Tr w temperaturze 20°C) i jej znaczna dyfuzyjność;

4) prostota i niewielki koszt aparatury, a w szczególności jej części dozującej mieszaninę gazową;

5) możliwość wykorzystania doświadczeń pomiarowych do prac nad zastosowaniem metody w anestezji przy określaniu zawartości pary eteru w tlenie, w którym prędkość dźwięku jest zaledwie o 4,3% mniejsza od prędkości dźwięku w powietrzu.



Rys. 9. Schemat blokowy aparatury doświadczalnej

CF — człon fizyczny, F — falowód, CE — człon elektryczny, B — bateria zasilająca, FDP — filtr dolnoprzepustowy, R — opór dopasowujący, f — miernik częstotliwości

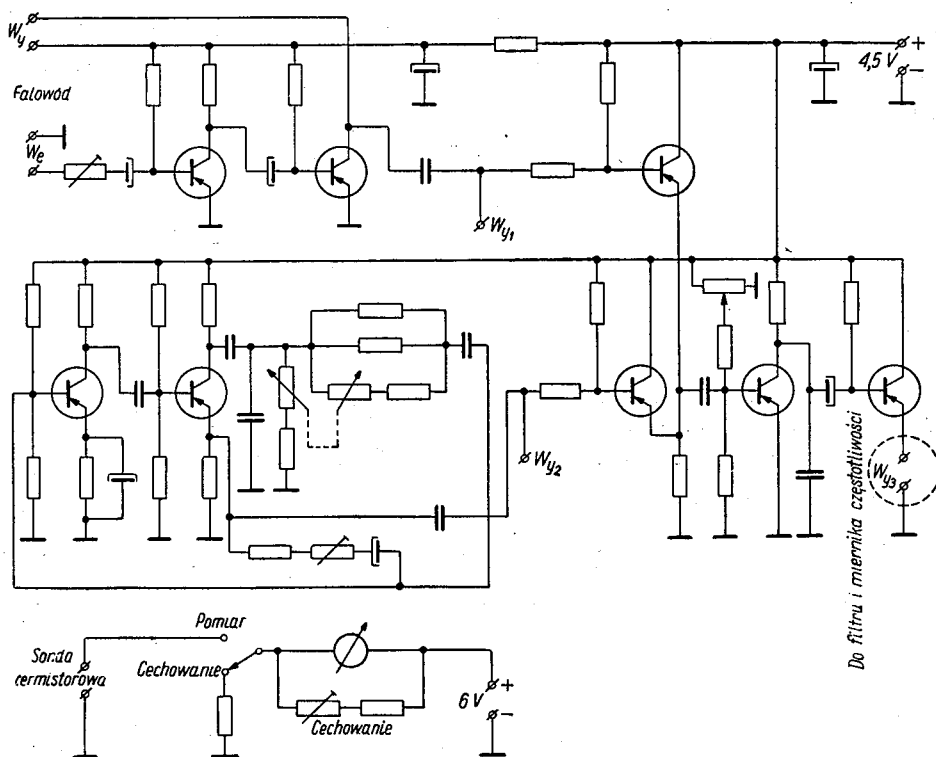
Aparatura eksperymentalna służąca do opisywania celów musi zawierać oprócz falowodu akustycznego dwa dodatkowe człony. Człon pierwszy, określany dalej mianem członu fizycznego, służy do sporządzania mieszaniny gazowej o znanej proporcji objętościowej gazów składowych i przekazywania jej do wnęki falowodu, zaś człon drugi — elektryczny, wykorzystany jest do pomiaru częstotliwości rezonansowej falowodu (rys. 9).

Zastosowany w doświadczeniach falowód został wykonany w formie grubościennej rury mosiężnej o wewnętrznej średnicy $d = 12$ mm i długości $l = 82$ mm (rys. 3), zamkniętej obustronnie miniaturowymi przetwornikami magnetycznymi typu PH firmy Philips. Temperaturę wewnątrz falowodu kontrolowano termistorem miniaturowym typu ZE 3 produkcji Zakładu Elektroniki I.P.P.T. Doprowadzenie mieszaniny gazowej realizowano poprzez dwa otwory wykonane w połowie długości falowodu.

W przypadku badania mieszaniny tlen-eter zadowalające wyniki uzyskano przy użyciu członu fizycznego aparatury w postaci hermetycznego klosza szklanego, w którego wnętrzu umieszczono falowód, eliminując w ten sposób konieczność wykonania zewnętrznych połączeń rurowych obiegu zamkniętego mieszaniny gazowej. Oprzyrządowanie członu fizycznego stanowił barometr rtęciowy o zakresie -250 Tr $\div$ $+250$ Tr, laboratoryjny termometr rtęciowy o zakresie $0 \div 50^\circ\text{C}$, wentyl przyciskowy do wyrównywania ciśnienia oraz pompka tłokowa ssąca, służąca do obniżenia ciśnienia wewnątrz klosza w stosunku do ciśnienia atmosferycznego; czynność ta jest konieczna do odparowania eteru. Kilka dodatkowych otworów w pokrywie klosza, zamykanych korkami gumowymi, przezna-

czono do spełniania czynności dodatkowych, jak wkraplanie eteru i wentylacja.

Człon elektryczny obejmował następujące podzespoły: układ zasadniczy, filtr dolnoprzepustowy oraz miernik częstotliwości. Schemat układu



Rys. 10. Schemat zasadniczego układu członu elektrycznego aparatury eksperymentalnej

zasadniczego przedstawia rys. 10. Układ ten zawiera dwustopniowy wzmacniacz wzbudzający drgania w obwodzie falowodu, generator RC, stopnie separujące, modulator i wzmacniacz końcowy, pracujący w układzie z uziemionym kolektorem na oporność wejściową użytego filtra dolnoprzepustowego. Wszystkie stopnie wykonano jako tranzystorowe z wykorzystaniem tranzystorów typu TC-11 i T-16 produkcji krajowej.

Generator RC zrealizowano jako układ przestrajany w niewielkim zakresie częstotliwości (1800÷2200 Hz). Jego stabilność została sprawdzona i okazała się dobra.

W płytę układu zasadniczego wmontowano termometr termistorowy, przeskalowany według laboratoryjnych termometrów rtęciowych w zakresie 14÷26°C.

Jako miernik częstotliwości zastosowano miernik wskazówkowy typu LMC-2 firmy „Elektra” o uchybie mniejszym od $\pm 3\%$.

3.2. Badanie wpływów parametrów ubocznych na odczyt zmian częstotliwości

Badanie wpływów parametrów ubocznych wykonano w celu umożliwienia późniejszego oszacowania błędów pomiaru.

Pomiar wpływu temperatury przeprowadzono pod ciśnieniem $P_0 = 757$ Tr, przy wilgotności $W = 40\%$, w zakresie temperatur $18 \div 24^\circ\text{C}$. Wpływ miał w tym zakresie przebieg liniowy i wynosił:

$$\gamma_t = \frac{\Delta f}{f_1} = 0,20 \cdot \Delta t, \quad (34)$$

gdzie:

γ_t — uchyb wywołany zmianą temperatury [%],

Δt — przyrost temperatury $^\circ\text{C}$.

Przyrost $\frac{\Delta f}{f_1} = 0,1\%$ został więc uzyskany przy $\Delta t = 0,5^\circ\text{C}$, co jest w przybliżeniu zgodne z wynikiem teoretycznym $\Delta t = 0,6^\circ\text{C}$ (rozdz. 2.3).

Pomiar wpływu ciśnienia statycznego przeprowadzono w warunkach: $W = 42\%$, $t = 22,1^\circ\text{C}$, zmieniając ciśnienie w zakresie $760 \div 560$ Tr. Wpływ miał również przebieg liniowy w tym zakresie ciśnień:

$$\gamma_p = \frac{\Delta f}{f_1} = -0,0048 \cdot \Delta P, \quad (35)$$

gdzie:

γ_p — uchyb wywołany zmianą ciśnienia [%],

ΔP — przyrost ciśnienia [Tr].

Wpływ ten pozwala określić wielkości błędu popełnianego przez założenie, że $\frac{\Delta c}{c_1} = \frac{\Delta f}{f_1}$, gdyż zmiana ciśnienia nie może mieć wpływu na prędkość dźwięku.

Pomiar wpływu wilgotności względnej przeprowadzono w warunkach: $P_0 = 757$ Tr, $t = 19,4^\circ\text{C}$, w zakresie wilgotności $0 \div 40\%$. Osuszenie powietrza przeprowadzono za pomocą pięciotlenku fosforu P_2O_5 , kontrolę zaś wilgotności początkowej za pomocą higrometru włosowego. Przyjęto wpływ jako liniowy:

$$\gamma_w = \frac{\Delta f}{f_1} = 0,05 \cdot \Delta W, \quad (36)$$

gdzie:

γ_w — uchyb wywołany zmianą wilgotnością [%],

ΔW — przyrost wilgotności względnej [%].

Przyrost $\frac{\Delta f}{f_1} = 0,1\%$ zachodził dla $\Delta W = 20\%$, a więc wpływ okazał się większy od wyniku teoretycznego $\left(\frac{\Delta f}{f_1} = 0,1\% \text{ dla } \Delta W = 27\% \right)$, co mogło być wywołane małą dokładnością wskazań higrometru.

Dokładność pomiaru częstotliwości obliczono. Najmniej dokładnym użytym zakresem miernika częstotliwości był zakres 0–250 Hz.

$$\gamma_{f_1} = \frac{\Delta f}{f_1} = \frac{250 \cdot 2,5}{2000} \approx 0,31\%, \quad (37)$$

gdzie γ_{f_1} — uchyb pomiaru częstotliwości.

Stabilność generatora RC została sprawdzona przez długotrwałą obserwację wskazań częstotliwości zdudnionej w ustalonych warunkach fizycznych. Maksymalny uchyb wynikający z niestabilności generatora RC, generatora „falowodowego” i miernika częstotliwości wynosił:

$$\gamma_{f_2} = \frac{\Delta f}{f_1} = \pm 0,05\%. \quad (38)$$

3.3. Przebieg pomiaru

Dozowanie przeprowadzono za pomocą strzykawki lekarskiej, ważonej przed napełnieniem i po napełnieniu dawką eteru etylowego, przy użyciu wagi laboratoryjnej. Po wytworzeniu w kloszu członu fizycznego podciśnienia rzędu — 200 Tr wstrzykiwano odważoną dawkę eteru przez ściankę rurki gumowej połączonej z kloszem. Dzięki istnieniu podciśnienia uzyskiwano całkowite wyssanie eteru ze strzykawki po kilku ruchach jej tłoczkiem. Po odparowaniu eteru wewnątrz klosza wyrównywano ciśnienie do poziomu ciśnienia zewnętrznego za pomocą wentyla przyciskowego.

Pomiary przeprowadzono w kilku seriach w ustalonych warunkach ciśnienia, temperatury i wilgotności dla jednej serii, ustawiając przed każdorazowym dawkowaniem eteru częstotliwość f_2 generatora RC na taką wartość, wyższą od częstotliwości f_1 drgań w pętli falowodu, aby uzyskać dla czystego powietrza pewną stałą, początkową wartość częstotliwości różnicowej $\Delta f = |f_1 - f_2|$. Manipulacja ta okazała się konieczna ze względu na skłonność generatorów do wzajemnej synchronizacji oraz ze względu na pewne, niewielkie zmiany temperatury między poszczególnymi dawkowaniami.

Po wstrzyknięciu i odparowaniu dawki eteru obserwowano gwałtowny wzrost częstotliwości różnicowej i jej asymptotyczne przejście do stanu ustalonego, któremu odpowiadał równomierny rozkład przestrzenny pary eteru w kloszu i we wnęce falowodu. Czas pełnej dyfuzji pary wynosił w opisanym aparaturze około 10 min dla mniejszych dawek, a do 20 min dla dawek większych. Zawartość a_2 pary eteru przeliczono ze wzoru wynikającego z równań gazowych [9]:

$$a_2 = \frac{P_{pn}}{e_{pn} P_0 V} m_2. \quad (39)$$

w którym:

P_{pn} — ciśnienie pary nasyconej eteru etylowego w danej temperaturze [Tr],

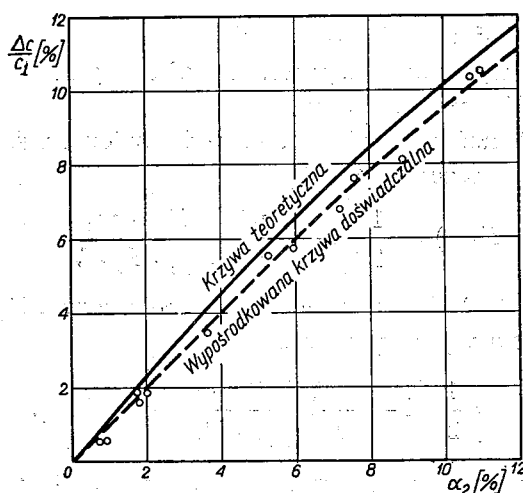
ρ_{pn} — gęstość pary nasyconej eteru etylowego w danej temperaturze $\left[\frac{\text{g}}{\text{cm}^3} \right]$,

P_0 — bezwzględna wartość ciśnienia statycznego w kloszu [Tr],

V — efektywna objętość całego członu fizycznego [cm^3],

m_2 — masa eteru [g].

Poszczególne serie pomiarowe wykazały dość dobrą zgodność wyników w postaci wspólnej krzywej cechowania, przedstawionej na rys. 11, obra-



Rys. 11. Krzywa cechowania aparatury eksperymentalnej

zującej względny przyrost częstotliwości drgań w pętli falowodu (któremu odpowiada przyrost $\frac{\Delta c}{c_1}$ prędkości dźwięku) w funkcji zawartości objętościowej α_2 pary eteru w powietrzu. Krzywa doświadczalna przebiega nieco poniżej krzywej wyliczonej ze wzoru teoretycznego (11), co wskazuje na istnienie bliżej nieznanego błędu systematycznego. Źródło tego błędu może kryć się w pominięciu w rozważaniach w rozdz. 2.2 wpływu wielkości α na wypadkową prędkości dźwięku lub może wynikać z niedoskonałości aparatury pomiarowej. Jego wielkość nie ma istotnego wpływu na dokładność pomiaru, wskazuje jednak na konieczność doświadczalnego cechowania aparatów tego typu.

Wielkość rozrzutów punktów pomiarowych wokół wyposrodkowanej krzywej cechowania pozwala oszacować uzyskiwalną przy istniejącym wyposażeniu dokładność pomiaru na $\Delta \alpha_2 \approx 0,2\%$ zawartości pary eteru w powietrzu dla $0,5\% < \alpha_2 < 11\%$.

3.4. Analiza błędów pomiarowych

W celu uproszczenia rozważań oznaczamy wynik konkretnego pomiaru jako

$$\frac{\Delta f}{f_1} = \xi. \quad (40)$$

Wynik pomiaru jest funkcją zawartości domieszki, wpływów ubocznych i wskazań przyrządów pomiarowych:

$$\xi = F(a_2, \gamma_i) = F(a_1) + \sum_{i=1}^n \gamma_i, \quad (41)$$

gdzie:

a_2 — zawartość domieszki,

γ_i — uchyby wywołane czynnikami ubocznymi i niedokładnością przyrządów pomiarowych.

Wielkość $F(a_2)$ można określić przez podstawienie wartości a_2 ze wzoru (39) do wzoru uproszczonego (12):

$$F(a_2) = \frac{1}{2} \left[1 - \left(\frac{c_1}{c_2} \right)^2 \right] \frac{P_{pn} \cdot m_2}{\varrho_{pn} \cdot P_o \cdot V}. \quad (42)$$

Po wstawieniu (42) do wyrażenia na ξ ze wzoru (41) otrzymuje się ogólne wyrażenie na wynik pomiaru

$$\xi = \frac{1}{2} \left[1 - \left(\frac{c_1}{c_2} \right)^2 \right] \frac{P_{pn} \cdot m_2}{\varrho_{pn} \cdot P_o \cdot V} + \sum_{i=1}^n \gamma_i. \quad (43)$$

Przeprowadzając analizę otrzymanego wzoru według zasad rachunku błędów dochodzi się do wyrażenia na całkowity uchyb pomiaru:

$$\frac{\Delta \xi}{\xi} = \frac{\Delta m_2}{m_2} + \frac{\Delta V}{V} + \frac{\Sigma \gamma_i}{\xi}. \quad (44)$$

Wzór ten nie uwzględnia uchybów pochodzących od danych liczbowych stałych fizycznych $c_1, c_2, P_{pn}, \varrho_{pn}$ oraz uchybu wskazania barometru — P_o .

Przyjmując następujące wahania warunków pomiarowych: $t = 0,2^\circ\text{C}$, $\Delta P = 2 \text{ Tr}$, $\Delta W = 2\%$ i wykonując przeliczenia według wzoru (34), (35) i (36), otrzymano następujące wartości uchybów od wpływów ubocznych:

$$\gamma_t = 0,04\%, \quad \gamma_p = 0,01\%, \quad \gamma_w = 0,01\%.$$

Stąd:

$$\sum \gamma_i = \gamma_t + \gamma_p + \gamma_w + \gamma_{f1} + \gamma_{f2} = 0,04 + 0,01 + 0,01 + 0,31 + 0,05 = 0,42\%$$

niezależnie od punktu pomiarowego (należy pamiętać, że jest to wartość bezwzględna błędu).

W tabelicy 2 podano całkowite przeliczenie błędu pomiarowego dla kilku punktów pomiarowych.

Tablica 2

Przykładowe przeliczenie błędu pomiarowego

m_2	g	$0,100 \pm 0,010$	$0,200 \pm 0,010$	$0,400 \pm 0,010$	$0,600 \pm 0,010$
α_2	%	1,97	3,58	7,2	10,7
ξ	%	15,6	3,41	6,7	10,3
$\sum \gamma_i$	%	04,2	0,42	0,42	0,42
$\frac{\Delta m_2}{m_2}$	%	10,0	5,0	2,5	1,7
$\frac{\Delta v}{v}$	%	0,85	0,85	0,85	0,85
$\frac{\sum \gamma_i}{\xi}$	%	27	12	6,5	4,1
$\frac{\Delta \xi}{\xi}$	%	37,9	17,9	9,9	6,7
$\Delta \xi$	%	0,6	0,6	0,7	0,7
ξ	%	$1,6 \pm 0,6$	$3,4 \pm 0,6$	$6,7 \pm 0,7$	$10,3 \pm 0,7$

Bezwzględny uchyb pomiaru okazał się w przybliżeniu stały, jego wielkość zawiera się w granicach $0,6 \div 0,7\%$ w zakresie do $\alpha_2 = 10,7\%$ zawartości pary eteru w powietrzu. Względny uchyb pomiaru $\frac{\Delta \xi}{\xi} \leq 10\%$ został uzyskany dopiero dla $\alpha_2 > 7\%$. Najmniejszą ilość wykrywalną przy pomocy opisanego urządzenia można oszacować na $\alpha_2 \approx 1\%$ pary eteru w powietrzu.

Największe błędy zostały wprowadzone w pomiarze częstotliwości ($\gamma_{f1} = 0,31\%$) oraz w pomiarze masy eteru ($\Delta m = 0,010$ g). Odkładne uchyby są o rząd wielkości większe od pozostałych.

Wskazuje to na konieczność położenia nacisku w dalszych badaniach na pomiar zawartości par eteru w powietrzu, użycie dokładniejszej wagi laboratoryjnej ($\Delta m_2 \leq 0,0010$ g) oraz zwiększenie dokładności pomiaru częstotliwości.

Zwiększenie dokładności pomiaru częstotliwości można uzyskać dwiema drogami:

- 1) przez zastosowanie miernika o większej dokładności,
- 2) przez przejście na mniejszy zakres pomiarowy.

Sposób pierwszy jest trudny do zrealizowania miernikiem wychyłowym, którego dokładność wynosi w najlepszym przypadku $1 \div 2\%$. Możliwe byłoby zastosowanie licznika elektronowego, co jednak zwiększyłoby wielokrotnie koszt urządzenia.

Na uwagę zasługuje sposób drugi. Użyty w powyższych pomiarach miernik „Elektra” posiada niższy zakres pomiarowy, $0 \div 50$ Hz, przy użyciu którego uchyb γ_{f1} pomiaru częstotliwości zmniejszyłby się do wartości

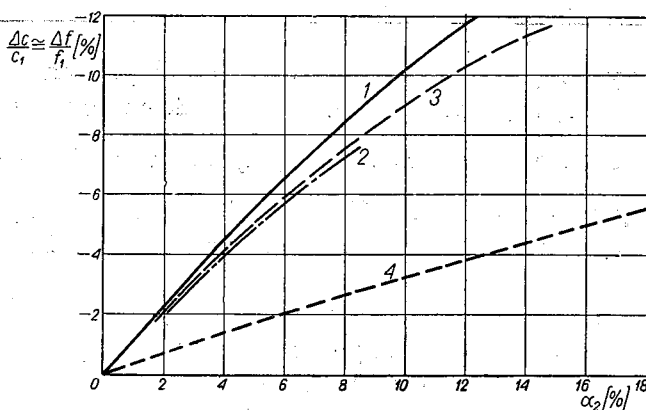
$$\gamma_{f1} = \frac{\Delta f}{f_2} = \frac{50 \cdot 25}{2000} = 0,0625\%,$$

pięciokrotnie mniejszej od poprzednich, przy założonej dokładności wskazań $2,5\%$. Zakres pomiaru zmniejszy się odpowiednio do wartości $\frac{\Delta f_{max}}{f_1} = \frac{50}{2000} = 2,5\%$, co odpowiada maksymalnej zawartości przy eteru w powietrzu $\alpha_2 max = 2,5\%$.

Reasumując należy stwierdzić, że zakres pomiarowy powinien być jak najmniejszy; w przypadku przeprowadzenia pomiarów w szerokich granicach należy zastosować przełączanie zakresów.

3.5. Podsumowanie badań eksperymentalnych

W pracowni Elektroakustyki Zakładu Badania Drgań I.P.P.T. przeprowadzono poza opisanymi pomiarami pomiary dodatkowe z trzema wariantami aparatury pomiarowej, w celu ustalenia wymagań dotyczą-



Rys. 12. Krzywe cechowania trzech wariantów aparatury eksperymentalnej (oznaczenia krzywych podane są w tekście)

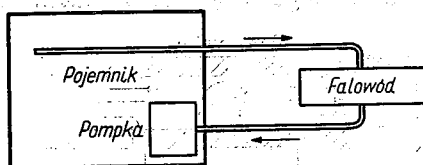
cych konstrukcji i sposobu cechowania aparatury użytkowej do pomiaru zawartości pary eteru w powietrzu lub w tlenie. Wyniki poszczególnych serii pomiarowych obrazuje rodzina krzywych na wykresie, rys. 12.

Krzywa teoretyczna została tu oznaczona cyfrą 1. Ilustruje ona granicę najwyższej osiągalnej czułości pomiaru. Krzywa 2 przedstawia wyniki cechowania aparatury opisanej w rozdziale 3 niniejszej pracy, zaś krzywa 3 i 4 aparatury zmodyfikowanej.

W nowej wersji aparatury dozowanie eteru zrealizowano poprzez umieszczenie dawek eteru w małych probówkach o pojemności ok. 3 ml,

szczelnie zakorkowanych, ważonych, a następnie tłuczonych wewnątrz klosza, co poważnie zwiększyła dokładność cechowania aparatury. Udoskonalono też sposób odparowania eteru i podawania mieszaniny do wnętrza falowodu. Przekonstruowany i uszczelniony falowód umieszczono poza kloszem, realizując obieg zamknięty badanej mieszaniny poprzez podwójne połączenie rurowe i miniaturową pompką odśrodkową umieszczoną wewnątrz klosza (rys. 13). Dzięki wprowadzeniu obiegu ciągłego zbliżono się do warunków eksploatacyjnych, w których opisana metoda pomiarowa może mieć największe zastosowanie.

Najbardziej zbliżoną do krzywej teoretycznej jest krzywa 3. Uzyskano ją przez zastosowanie w zmodernizowanej aparaturze przetworników typu PH firmy Philips. Dzięki zastosowanym udoskonaleniom rozrzuty punk-



Rys. 13. Obieg mieszaniny gazowej w zmodernizowanej aparaturze do cechowania

tów cechowania zmalały znacznie, co sugeruje możliwość zmniejszenia uchybu bezwzględnego pomiaru do wartości $\Delta\xi \leq 0,2\%$ zawartości pary eteru w powietrzu.

Dla zobrazowania konieczności właściwego doboru przetworników elektroakustycznych przeprowadzono na zmodernizowanej aparaturze pomiary, których wynikiem jest krzywa 4. Zastosowanie jako przetworników słuchawek telefonicznych typu CB-48 dających duże zniekształcenia fazowe wpłynęło na trzykrotne zmniejszenie czułości i znaczne zwiększenie rozrzutów.

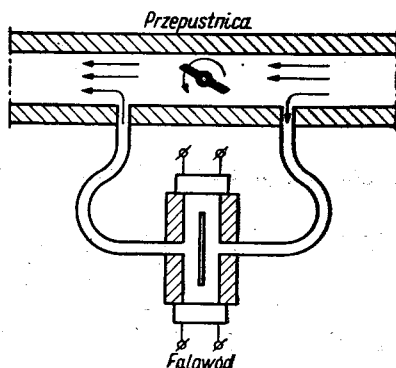
3.6. Wytyczne do budowy miernika składu objętościowego anestetycznych mieszanin gazowych

Doświadczenia uzyskane w Zakładzie Badania Drgań przy pracach nad rezonansową metodą analizy mieszanin gazowych pozwalają na stwierdzenie możliwości skonstruowania użytkowego miernika zawartości objętościowej pary eteru w tlenie, podtlenku azotu w tlenie oraz dwutlenku węgla w tlenie. Spotykane w kraju aparatury do znieczulania ogólnego (anestezji) nie posiadają miernika dającego bezpośredni obraz składu objętościowego stosowanych mieszanin gazowych. Najczęściej stosowanymi miernikami przez takie firmy, jak niemiecka „Dräger” lub czeska „Chirana” są przepływomierze aerodynamiczne pozwalające z grubsza oszacować ilości podtlenku azotu lub dwutlenku węgla w tlenie. Zawartość pary eteru nie jest w ogóle mierzona, lecz tylko ustawiana doświadczalnie.

Omawiany miernik może mieć przede wszystkim duże zastosowanie jako uzupełnienie istniejących aparatów do znieczulania ogólnego z obiegiem określonym gazów. Zasada działania miernika wymaga separacji domieszek zniekształcających wynik pomiaru, co jest spełnione w tego typu aparaturach w pochłaniaczach pary wodnej i dwutlenku węgla.

Miernik w postaci eksploatacyjnej powinien stanowić zwartą całość w niewielkiej obudowie stawianej na stole, z wlotami rurowymi do włączania w obieg gazowy. Duże opory aerodynamiczne stawiane przez obwód gazowy miernika nie umniejszają jego stosowności; obwód ten może być bocznikowany przez przelot o regulowanej przepustowości (rys. 14). Bardzo ważne jest zapewnienie łatwości dezynfekcji obwodu gazowego.

Wymagania dotyczące miernika powinny iść głównie w kierunku jego pewności działania. Z tego względu zasilacz członu elektrycznego



Rys. 14. Przepustnica regulowana, umożliwiająca zaopatrywanie falowodu w próbki mieszaniny gazowej w warunkach pomiaru technicznego

aparatury powinien być stabilizowany. Ze względu na duże temperatury sal operacyjnych (ok. 37°C) wydaje się konieczne zrezygnowanie ze stosowania tranzystorów w członie elektrycznym i zastąpienie ich lampami elektronowymi, ogólny schemat blokowy członu elektrycznego pozostanie jednak bez zmian, taki jak na rys. 8.

Najbardziej odpowiedzialną, wymagającą dobrego zaprojektowania częścią aparatury jest częstotliwościomierz wskazówkowy. Ze względu na duże zakresy pomiaru należy zrealizować go z całkowaniem częstotliwości w układzie lampowym, bez przekaźnika.

4. OPIS WYKONANEGO PROTOTYPU EKSPLOATACYJNEGO ANALIZATORA ETTEROWEGO

4.1. Założenia wstępne projektu prototypu

W celu zorientowania się w wymaganym zakresie pomiaru wykonano przybliżone przeliczenie średniej zawartości pary eteru w mieszaninie anestetycznej tlen-eter. Przyjmując, że pacjent wykonuje 16 oddechów

na minutę o przeciętnej objętości 400 cm^3 i zakładając, że w ciągu minuty zostaje odparowany w obiegu okrężnym i całkowicie wchłonięty 1 g eteru na minutę otrzymuje się:

$$V = 16 \cdot 400 = 6,4 \cdot 10^3 \text{ cm}^3,$$

$$a_2 = \frac{P_{pn}}{Q_{pn}} \cdot \frac{m_2}{p \cdot V} = 2,4 \cdot 10^5 \frac{1,00}{0,755 \cdot 10^3 \cdot 6,4 \cdot 10^3} \approx 4\%.$$

Zakładając, że wartość a_2 może być dwukrotnie większa, przyjęto jako największy zakres pomiarowy $a_2 = 0 \div 15\%$. Do dokładnego pomiaru mniejszych zawartości wykonano drugi, czulszy zakres pomiaru $0 \div 3\%$. Oprócz wymienionych zrealizowano trzeci zakres pomiarowy jako rezerwowy.

Projekt wykonano w oparciu o przebadaną zasadę heterodynową pomiaru częstotliwości, według schematu blokowego z rys. 8. W trosce o miniaturyzację urządzenia zastosowano w nim łatwo dostępne próżniowe lampy miniaturowe, ograniczając ilość typów do czterech: EF85, ECC81, EAA91 i AZ21. Stabilizację napięć anodowych zrealizowano z użyciem lamp 85A1 oraz CF-4C.

4.2. Opis układu analizatora

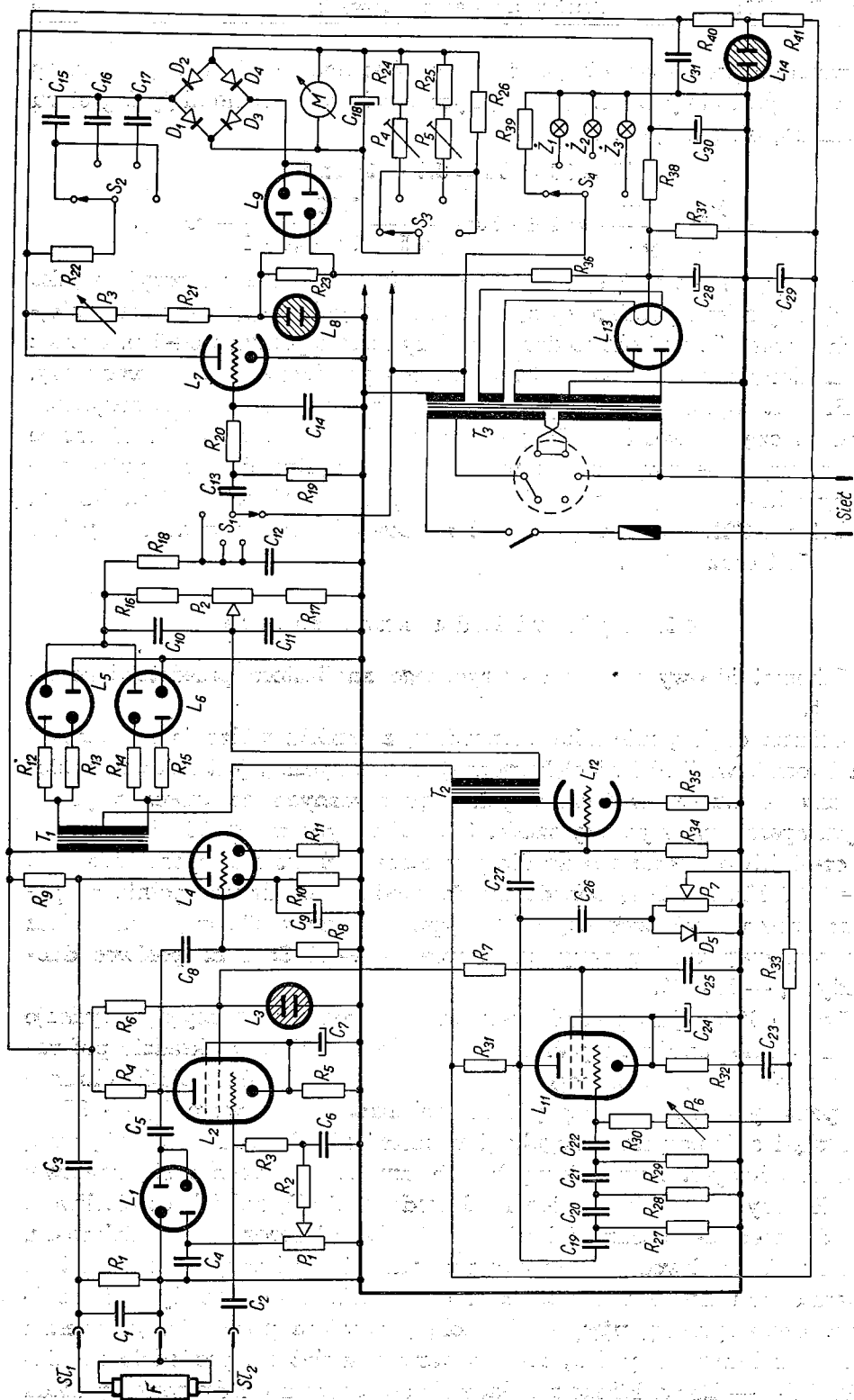
Schemat ideowy układu elektrycznego analizatora przedstawiono na rys. 15.

Główną częścią miernika, stanowiącą z punktu widzenia automatyki przetwornik wielkości nieelektrycznej w elektryczną, jest falowód umieszczony ze względów termicznych i praktycznych na trzech wtykach z tyłu aparatury. Wymiary wnętrza falowodu wynoszą: $\phi 16 \times 82 \text{ mm}$, zaś jej częstotliwość rezonansowa dla powietrza suchego o temperaturze 20°C $f_1 = 2000 \text{ Hz}$. Funkcję pobudzania falowodu F do drgań spełnia za pośrednictwem słuchawek Sl_1 i Sl_2 (typu PH firmy Philips) wzmacniacz dwustopniowy zrealizowany na lampie L_2 — EF85 i na połówce duotriody L_4 — ECC81.

Lampa EF85 jest pentodą regulowaną; minus regulacyjny uzyskuje ona z własnej anody za pośrednictwem prostownika w układzie podwajacza napięcia zrealizowanego na duodiodzie L_1 — EAA91. Siatka ekranowa pentody otrzymuje napięcie ze stabilizatora jonowego L_3 — 85A1.

Dzięki automatycznej regulacji wzmocnienia drgania zawierają małą ilość składowych harmoniczných, ich amplituda zmienia się zaledwie o 3 dB przy zmianach tłumienia falowodu o 12 dB, zaś ich częstotliwość jest dostatecznie stabilna i bliska częstotliwości rezonansowej falowodu.

Drgania w obwodzie falowodu wzmacnia i przekazuje do modulatora separator w układzie wzmacniacza z drugą połówką duotriody L_4 i transformatorem symetryzującym T_1 . Rolę generatora pomocniczego spełnia generator łańcuchowy RC, zrealizowany również na pentodzie regulacyjnej L_{11} — EF85. Minus regulacyjny dostarcza z anody pentody dioda



Rys. 15. Schemat ideowy prototypu eksploatacyjnego analizatora eterowego

germanowa D_5 typu DOG21. Ekran lampy jest zasilany poprzez układ odsprzęgający R_{33} i C_{25} ze stabilizatora L_3 .

Drgania generatora utrzymywane są przez układ automatycznej regulacji na granicy ich powstawania, co stabilizuje ich amplitudę i częstotliwość.

Częstotliwość generacji jest regulowana potencjometrem P_6 (pokrętko „Zdudnianie” na ścianie czołowej przyrządu) w granicach $f_2 = 1900 \div 2100$ Hz. Drgania przekazywane są do obwodu modulatora poprzez separator z triodą L_{12} — 1/2 ECC81 i z transformatorem T_2 .

Zdudnianie drgań o częstotliwościach f_1 i f_2 przeprowadzono w układzie modulatora kołowego zrównoważonego, zrealizowanego na duodiodach L_5 i L_6 typu EAA91. Wybór modulatora miał na celu otrzymanie na jego wyjściu tylko częstotliwości $\Delta f = (f_1 \pm f_2)$, z wyeliminowaniem częstotliwości f_1 i f_2 . W wyniku otrzymano 40 dB odstępu sygnału $\Delta f = (f_1 - f_2)$ od szumów, po zastosowaniu trójczłonowego filtra dolnoprzepustowego RC.

W rozwiązaniu modulatora istotny jest jego układ wyjściowy.

Aby modulator jako układ produkujący powolne przebiegi w zakresie $f = 0 \div 600$ Hz mógł spełnić swoje zadanie, było konieczne zastąpienie symetrycznego transformatora wyjściowego dzielnikiem oporowym R_{16} i R_{17} symetryzowanego potencjometru P_2 , co wywołuje duże straty energetyczne dla źródła pomocniczego (generatora RC). Straty te udało się wybitnie zmniejszyć i usprawnić działanie modulatora poprzez umieszczenie w dzielniku oporowym R_{16} i R_{17} pierwszego elementu filtra dolnoprzepustowego w formie pojemności składającej się z szeregowo połączonych kondensatorów C_{10} i C_{11} . Układ ten powoduje znaczne zwiększenie prądu polaryzującego diody z transformatora T_2 oraz — jako zjawisko uboczne — pewne zmniejszenie oporności również dla transformatora T_1 .

Sygnał różnicowy przyłożony jest do siatki ogranicznika amplitudowego anodowo-siatkowego z triodą L_7 — 1/2 ECC81. Amplituda wyjściowa jest regulowana w niewielkich granicach za pomocą potencjometru P_3 na płycie czołowej przyrządu (pokrętko „Cechowanie”). Zasilanie odbywa się stabilizatora jonowego L_8 typu CT-4C.

Częstotliwość Δf sygnału różnicowego jest mierzona jako efekt całkowania przez mikroamperomierz M typu MEL 100P firmy A-3 impulsów różniczkowych powstających w wyniku ładowania kondensatora C_{15} , C_{16} lub C_{17} , w zależności od zakresu pomiarowego. Zastosowano układ prostowniczy Graetza z diodami germanowymi D_1 , D_2 , D_3 i D_4 , typu DOG21. Aby uniezależnić się od oddziaływania szumów, wśród których przeważają sygnały o częstotliwości stosunkowo dużej $f_1 + f_2 = 4000$ Hz, znacznie zmniejszających czułość analizatora, powstała konieczność zastosowania w obwodzie różniczkującym miernika częstotliwości szeregowego ogranicznika amplitudowego „górnoprzepustowego” zrealizowanego z użyciem

duodiody L_8 —EAA91. Trzecim organem manipulacyjnym na czołowej płycie analizatora jest przełącznik zakresów. Posiada on cztery styki ruchome oznaczone na schemacie literami S_1 , S_2 , S_3 i S_4 .

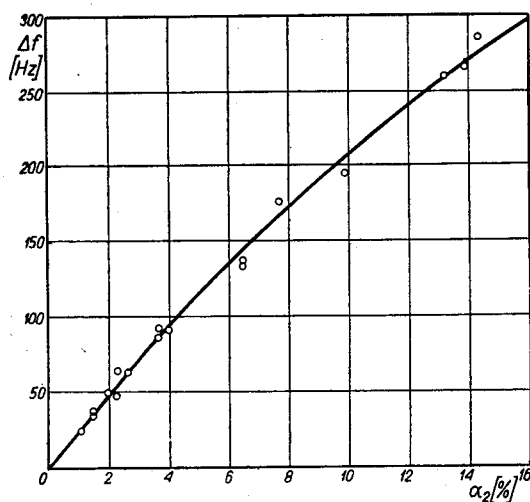
Styk S_1 jest umieszczony w obwodzie siatki ogranicznika i przełącza ją ze źródła o wzorcowej częstotliwości (pozycja „Cechowanie”) na modulator (pozycje „Pomiar”). Styki S_2 i S_3 przełączają elementy miernika częstotliwości. Ostatni ze styków — S_4 spełnia rolę pomocniczą. Przełącza on żarówki sygnalizacyjne Z_1 , Z_2 i Z_3 , których zastosowanie uniemożliwia popełnienie pomyłki przy zdalnej obserwacji miernika.

Analizator posiada 3 zakresy pomiarowe, z których jeden jest zakresem rezerwowym, nie jest uwidoczniiony na skali przyrządu i może być wykorzystany do pomiaru innego typu gazu. Pozostałe zakresy to zakres $0 \div 15\%$ zawartości pary eteru, nieco zagęszczony przy końcu skali, oraz zakres $0 \div 3\%$, całkowicie liniowy i pokrywający się z pierwotną skalą mikroamperomierza $0 \div 150 \mu A$. Zakresowi temu odpowiadają częstotliwości różnicowe $\Delta f = 0 \div 72 \text{ Hz}$. Na tym zakresie umieszczono na skali przyrządu kreskę odpowiadającą częstotliwości 50 Hz, przykładanej do miernika częstotliwości w pozycji „Cechowanie”.

4.3. Charakterystyka własności użytkowych analizatora

Największa czułość miernika odpowiada zawartości $0,2\%$ pary eteru. Dokładność odczytu wynosi $0,05\%$ (0,5 działki) na zakresie $0 \div 3\%$ oraz $0,5\%$ (0,5 działki) na zakresie $0 \div 15\%$.

Analizator charakteryzuje się bardzo szybkim pomiarem dzięki zastosowaniu małych stałych czasu w obwodzie automatycznej regulacji wzmac-



Rys. 16. Krzywa cechowania analizatora eterowego

niacza w pętli falowodu oraz w obwodzie całkującym. Wskazanie mikroamperomierza ustala się w czasie ok. 1 sec., co w praktyce pozwala na ciągły pomiar strugi gazu w falowodzie, a jest rewelacją w porównaniu z metodami katalitycznymi analizy gazowej. Cechowanie analizatora wykonano wg krzywej teoretycznej; przebieg tej krzywej z naniesionymi punktami pomiarów kontrolnych obrazuje wykres na rys. 16.

Jako wadę urządzenia należy wymienić konieczność stabilizacji termicznej badanej mieszaniny oraz konieczność zastosowania regulatora prędkości przepływu przy zastosowaniu analizatora do znieczulenia ogólnego w obiegu okrężnym gazów.

5. ZAKOŃCZENIE

Opisane przykłady medycznych zastosowań rezonansowej metody analizy gazowej nie wyczerpują jej perspektyw użytkowych. Możliwości innych zastosowań są bardzo szerokie. Na pierwszy plan wysuwają się potrzeby automatyzacji przemysłu chemicznego. Sanitarna kontrola jakości powietrza służącego do oddychania jest również problemem ważnym z punktu widzenia bezpieczeństwa i higieny pracy; problem ten rozwiązać można przez automatyzację urządzeń klimatyzacyjnych. Całkowicie realne wydają się także metody kontroli zawartości gazów wybuchowych w kopalniach oparte na opisanej zasadzie.

W tablicy 3 wyszczególniono najważniejsze możliwości zastosowań metody rezonansowej.

Tablica 3
Perspektywy zastosowań rezonansowej metody
analizy gazowej

Dziedzina zastosowania	Gaz wykrywany i mierzony
Medycyna	eter, dwutlenek węgla, podtlenek azotu
Kontrola procesów przemysłowych	dwusiarczek węgla, chlor, amoniak, wodór,
Bezpieczeństwo i higiena pracy	dwutlenek węgla, metan, gaz świetlny, wodór
Meteorologia	para wodna, wodór

Wykonany miernik eksploatacyjny może być użyty do wszystkich zastosowań medycznych. Adaptacja do pomiaru podtlenku azotu czy CO₂ polegałaby na przecechowaniu analizatora. W niektórych przypadkach, jak np. przy pomiarze zawartości CO₂ w powietrzu wydychanym, anali-

zator musiałby być zaopatrzony w urządzenia stabilizujące warunki pomiarowe takie, jak wilgotność i temperatura.

Pozostałe zastosowania rezonansowej metody analizy gazowej wymagają nowych rozwiązań układowych, różniących się między sobą precyzją działania i stopniem automatyzacji. Przewiduje się następujące kierunki rozbudowy analizatora.

1. Zastosowanie dwóch identycznych, samozdudniających się falowdowych obwodów drgań.

2. Zastosowanie układu przekąźnikowego reagującego na zmiany częstotliwości lub amplitudy sygnału wyjściowego, wywołane zmianą składu objętościowego badanej mieszaniny.

3. Zaopatrzenie układu w urządzenia alarmowe.

4. Przystosowanie układu do sterowania urządzenia samopiszącego.

5. Przystosowanie urządzenia do współpracy z analizatorami innych typów w celu rozszerzenia możliwości pomiarowych.

6. Zwiększenie czułości metody przez daleko idącą stabilizację warunków pomiarowych.

7. Budowę analizatorów przenośnych w oparciu o układy tranzystorowe.

Pomiary przeprowadzone w ramach badań eksperymentalnych wykazały, że para eteru wnosi bardzo duże tłumienie fali dźwiękowej do wnętrza falowodu, co sugeruje możliwość realizacji pomiaru zawartości par w gazach o właściwościach zbliżonych do własności gazów szlachetnych przy pomocy pomiaru parametru tłumienia fali dźwiękowej lub ultradźwiękowej. Badania w tym kierunku musiałyby obejmować poza badaniem wielkości tłumienia w funkcji zawartości domieszki pary w gazie przy ustalonej częstotliwości drgań również badania wielkości tłumienia w funkcji częstotliwości drgań przy ustalonej zawartości domieszki.

Należy się spodziewać, że otrzymamy w przeprowadzonych badaniach przyrost tłumienia po domieszananiu eteru do powietrza zwiększy się wraz z częstotliwością drgań. Proponowany kierunek badań wymagałby poparcia wyników doświadczalnych głęboką analizą teoretyczną zjawiska propagacji fali dźwiękowej w parach.

*Instytut Podstawowych Problemów Techniki PAN
Zakład Badania Drgań — Pracownia Elektroakustyki*

WYKAZ LITERATURY

1. Beranek L. L.: Acoustic measurements. — J. Wiley, New York—London 1949.
2. Jeżewski M., Kalisz K.: Tablice wielkości fizycznych. — P.W.N., Warszawa 1957.
3. Faulconer, Clarke, Osterberg: Proc. Staff. Meet., Mayo Clin., t. 18, 89, 1943.

4. Faulconer: Anesthesiology, t. 10, 1, 1949.
5. Wilson J. E.: The Review of Scientific Instruments, t. 25, nr 9, 1954, s. 927—928.
6. Stott F. D.: The Review of Scientific Instruments, t. 28, nr 11, 1957, s. 914—915.
7. Żyszkowski Z.: Podstawy elektroakustyki. — P.W.T., Warszawa 1953.
8. Groszkowski J.: Wytwarzanie drgań elektrycznych. P.W.T., Warszawa 1958.
9. Jeżewski M.: Fizyka. P.W.N., Warszawa 1953.

Я. КАЦПРОВСКИ, Т. УТГЭ

АКУСТИЧЕСКИЙ РЕЗОНАНСНЫЙ МЕТОД АНАЛИЗА ГАЗОВЫХ СМЕСЕЙ

Резюме

Предметом работы является развитие акустического метода измерения объёмного состава газовых смесей. Проблема количественного анализа газовых смесей часто встречается в практике во многих технологических и производственных процессах. Применяемые для этих целей измерительные методы являются в общем затруднительными и трудоёмкими и они требуют применения весьма дорогих комплектов аппаратуры. Отсутствуют методы простые, позволяющие проводить быстро, но одновременно в достаточной степени точно технические измерения, особенно в случае газовой смеси с быстро изменяющимся объёмным составом протекающей с переменной скоростью.

В резюмированной работе, в качестве измерительного параметра принята скорость звука в газовой смеси, объёмный состав которой должен быть определён. Обозначая коэффициенты объёмного содержания составляющих газов со скоростями звука c_1 и c_2 в виде $\alpha_1 = \frac{v_1}{v_2}$ и $\alpha_2 = \frac{v_2}{v}$, получается зависимость:

$$\frac{\Delta c}{c_1} = \frac{1}{\sqrt{1 + \left[\left(\frac{c_1}{c_2} \right)^2 - 1 \right] \alpha_2}} - 1, \quad (1)$$

которая в случаях, когда $\alpha_2 \ll 1$ приводится к упрощённой формуле:

$$\frac{\Delta c}{c_1} \approx \frac{1}{2} \left[1 - \left(\frac{c_1}{c_2} \right)^2 \right] \alpha_2. \quad (2)$$

На основании зависимости (1) или (2) можно определить процентное объёмное содержание α_2 добавочного газа со скоростью звука c_2 в смеси, в которой главной составляющей является газ со скоростью звука c_1 .

Из анализа формулы (2) следует, что в случаях когда $\frac{c_1}{c_2} \leq 0,5$ чувствительность метода не зависит от значения отношения $\frac{c_1}{c_2}$ и определена выражением:

$$\alpha_2 \approx 2 \frac{\Delta c}{c_1}. \quad (3)$$

Но в тех случаях, когда $\frac{c_1}{c_2} > 1$, чувствительность метода растёт вместе с ростом отношения $\frac{c_1}{c_2}$. Исходя из вышеуказанного, в условиях технических

измерений, соответственный диапазон применения метода определяется зависимостью $c_1 > c_2$. Применение метода в случаях обнаруживания примесей газа со скоростями звука $c_2 \geq 2c_1$ требует соответственного увеличения чувствительности измерения $\frac{\Delta c}{c_1}$.

Анализируя влияние посторонних факторов на точность измерения получено, что при предположенной значении погрешности измерения $\frac{\Delta c}{c_1} = 0.1\%$ допускаемые колебания температуры составляют $\Delta t \approx \pm 0,6^\circ\text{C}$, допускаемые изменения содержимого водяного пара $\Delta W \approx \pm 27\%$, а допускаемое значение изменения статического давления $\Delta P_0 = \pm 20 \text{ Tr}$.

В качестве схемы для измерения относительных изменений скорости звука $\frac{\Delta c}{c_1}$ применён акустический волновод длиной $l = \text{const}$, закрытый с обоих сторон обратимыми электроакустическими преобразователями магнитного типа. Волновод, включённый в петлю обратной связи усилителя низкой частоты, вызывает возникание в системе электрических колебаний с частотой f , которая при данной скорости звука c_1 — соответствует длине волны $\lambda = \frac{c_1}{f}$. Условие возбуждения колебаний соблюдается только лишь тогда, когда $l = (2n + 1) \frac{\lambda}{2}$: таким образом по мере относительных изменений частоты $\frac{\Delta c}{c_1} = f(\alpha_2)$, вызванных изменением объёмного состава газовой смеси, появляются относительные изменения частоты генерированных колебаний: $\frac{\Delta c}{c_1} = \frac{\Delta f}{f_1} \Big|_{\lambda = \text{const}}$, измеряемые с помощью электронного частотомера. Предварительные экспериментальные исследования были проведены в лабораторной схеме, принимая в качестве газовой смеси пар эфира в воздухе. Разработанный метод измерения даёт возможность достижения точности $\pm 0,2\%$ при определении содержимого пара эфира в воздухе для $0,5\% < \alpha_2 < 11\%$.

На основании результатов экспериментальных исследований, которые подтвердили правильность теоретических рассуждений и доказали полную возможность технического осуществления разработанного метода, была запроектирована и изготовлена модель измерительного устройства для определения содержания эфира в анестетической смеси кислород — эфир. Измерение частоты производится по методу смешивания колебаний частотой в f_1 , генерированных в волноводе, с колебаниями частотой в f_2 производимыми вспомогательным генератором R-C, обладающим значительной стабильностью частоты. Этот метод значительно повышает точность измерений.

Изготовленный анализатор предназначен в основном для целей объёмного анализа анестетических смесей (кислород — эфир). Анализатор имеет два измерительных диапазона, соответствующих содержимому пара эфира $\alpha_2 = 0 \div 3\%$ и $\alpha_2 = 0 \div 15\%$. Чувствительность измерения соответствует содержимому пара эфира $\alpha_2 = 0,2\%$. Точность составляет $0,05\%$ в диапазоне $\alpha_2 = 0 \div 3\%$, и $0,5\%$ в диапазоне $\alpha_2 = 0 \div 15\%$.

Изготовленный анализатор имеет универсальный характер и, после введения незначительных изменений, он может быть успешно использован в медицине [$\text{C}_2\text{H}_5\text{O}$; CO_2 ; N_2O], при контроле процессов в промышленности [H_2S ; Cl ; NH_3 ; H_2], в метеорологии [H_2O ; H_2] и при контроле безопасности труда, например в шахтах [CO_2 ; CH_4 ; H_2 ; светильный газ и др.].

J. KACPROWSKI, T. UTHKE

ACOUSTIC RESONANCE METHOD FOR ANALYSIS OF GASEOUS MIXTURES

Summary

The subject of this paper is the development of the acoustic method for measuring the volume composition of different gaseous mixtures. This problem is a very important one in many technological as well as industrial processes. The methods used for the quantitative analysis of gaseous mixtures are in general difficult and require rather expensive and complicated instrumentation. There is a lack of simple but accurate methods, which could be used in technical measurements, especially of those gaseous mixtures, which quickly change their volume composition and/or their flow-velocity.

The method described is based on the measurement of the velocity of sound in the mixture under investigation. Designing the velocities of sound in two different gases c_1 and c_2 , as well as their volume coefficients $\alpha_1 = \frac{v_1}{v}$ and $\alpha_2 = \frac{v_2}{v}$, respectively, the following equation can be formulated:

$$\frac{\Delta c}{c_1} = \frac{1}{\sqrt{1 + \left[\left(\frac{c_1}{c_2} \right)^2 - 1 \right] \alpha_2}} - 1. \quad (1)$$

If $\alpha_2 \ll 1$, equation (1) can be approximated to the form:

$$\frac{\Delta c}{c_1} \approx \frac{1}{2} \left[1 - \left(\frac{c_1}{c_2} \right)^2 \right] \alpha_2. \quad (2)$$

Equations (1) or (2) can be used for the determination of the volume coefficient α_2 of the additional gas (velocity of sound c_1) in the mixture, the main component of which is the gas of sound-velocity c_2 .

The discussion of equation (2) shows, that as long as $\frac{c_1}{c_2} \ll 0,5$, the sensitivity of the method is independent of the ratio $\frac{c_1}{c_2}$ as in this case

$$\alpha_2 \approx 2 \frac{\Delta c}{c_1}. \quad (3)$$

If, however, $\frac{c_1}{c_2} > 1$, the sensitivity increases as the ratio $\frac{c_1}{c_2}$ increases. It follows, that the useful range of application of this method is limited to the special case, when $c_1 > c_2$: the velocity of sound in the additional gas should be higher than the sound velocity in the gas which is the main component of the mixture, as the sensitivity in measuring the incremental ratio $\frac{\Delta c}{c_1}$ must be increased in another case.

It has been proved that designing the accuracy of the measurement $\frac{\Delta c}{c_1} = 0,1\%$, the maximum admissible variations of environment condition are the following: for temperature — $\Delta t \approx \pm 0,6^\circ\text{C}$, for static pressure — $\Delta P_0 \approx \pm 20 \text{ Tr}$, and for humidity water vapor — $\Delta W \approx \pm 27\%$.

The measurement of the incremental velocity of sound $\frac{\Delta c}{c_1}$ is being achieved by means of an acoustic tube (wave guide) of the length $l = \text{const}$, closed at both ends

with reversible electroacoustic transducers (magnetic type used in commercial hearing-aids). This arrangement forms the feed-back loop of the low-frequency amplifier. If $l = (2n + 1) \frac{\lambda}{2}$ the feed-back gets positive and electrical oscillations of frequency $f_1 = \frac{c_1}{\lambda}$ arise. In this connection the relative variation of the sound-velocity $\frac{\Delta c}{c_1}$ in the mixture under investigation, caused by the variation of its volume composition $\alpha_2 \left[\frac{\Delta c}{c_1} = f(\alpha_2) \right]$ can be simply detected by the direct measurement of the relative variation of frequency: $\frac{\Delta c}{c_1} = \frac{\Delta f}{f_1} \Big|_{\lambda = \text{const}}$. For the measurement of frequency variation an electronic frequency meter can be used.

The preliminary experimental tests have been carried on in laboratory equipment using the gaseous mixture of ether-vapor and air. For the values of volume coefficient α_2 of ether between 0,5% and 11% the accuracy was better than $\pm 0,2\%$.

An experimental analyser has been built for the analysis of anesthetic gaseous mixtures (e.g. oxygen-ether). The measurement of frequency variations $\frac{\Delta f}{f_1}$ is carried on by means of an auxiliary stable R—C oscillator. The frequency f_2 of its oscillations is being mixed with the frequency f_1 generated in the feed-back amplifier and the difference $(f_1 - f_2)$ is measured with an electronic meter; in this way the accuracy can be highly increased.

The analyser has two measuring ranges: a narrow one for α_2 from 0 to 3%, and a broad one, for α_2 from 0 to 15%. The sensitivity is about 0,2% of ether in oxygen volume. The accuracy is better than $\pm 0,05\%$ in the range $\alpha_2 = 0 \div 3\%$, and better than $\pm 0,5\%$ in the range $\alpha_2 = 0 \div 15\%$.

The analyser, although primarily designed for oxygen-ether mixture, can be used after re-calibration for measuring of other gaseous mixtures in medicine [e.g. $(C_2H_5)_2O$; CO_2 ; N_2O], in meteorology [H_2O ; H_2], in controlling industrial processes [e.g. H_2S ; Cl ; NH_3 ; H_2] or in controlling safety conditions in mines, factories etc. [CO_2 ; CH_4 ; lighting gas].

STANISŁAW BELLERT, HENRYK DZIATLIK, MIECZYŚLAW KOWALSKI

Model elektryczny kanału wodnego

Rękopis dostarczono 3.6.1960

W artykule rozpatrzono analogie między uproszczonymi równaniami Saint-Venanta obrazującymi dynamikę ruchu wody w kanale wodnym i równaniami elektrycznymi linii długiej. Na podstawie zaobserwowanych zależności rozpatrzono techniczną możliwość realizacji modelu elektrycznego kanału wodnego o przekroju trapezowym i o małym nachyleniu dna, modelującego uproszczone równania Saint-Venanta o z góry zadanych współczynnikach. Za pomocą takiego modelu można by badać przebieg fali powodziowej w rzece lub w kanale w przypadku małych prędkości i małych przyspieszeń ruchu wody, co jest zwykle spełnione w rzekach o małym nachyleniu dna (dla rzek nizinnych).

1. WSTĘP

Dynamikę ruchu wody w rzece lub w kanale wodnym można opisać za pomocą równań różniczkowych Saint Venanta. Ze względu na istnienie pewnych analogii tych równań z równaniami telegrafistów, które opisują przebiegi prądów i napięć w torze telekomunikacyjnym, nasuwa się możliwość zbudowania elektrycznego urządzenia analogowego symulującego dynamikę ruchu wody w rzece. Zagadnieniu opracowania odpowiedniego programu dla takiego urządzenia poświęcona jest właśnie niniejsza praca.

2. ANALOGIE MIĘDZY RÓWNANIAMİ SAINT-VENANTA A RÓWNANIAMİ UKŁADÓW ELEKTRYCZNYCH

2.1. Ustalenie ogólnych zależności

Dynamikę ruchu wody w rzece lub w kanale wodnym określają następujące równania różniczkowe

$$i_0 - \frac{\partial h}{\partial s} = \frac{u^2}{k^2 a} + \frac{1}{g} \left(\frac{\partial u}{\partial t} + u \frac{\partial u}{\partial s} \right), \quad (1)$$

$$\frac{\partial q}{\partial s} + \frac{\partial \omega}{\partial t} = 0, \quad (2)$$

gdzie:

- i_0 — spadek dna koryta rzeki,
- h — wysokość (mierzona od dna) zwierciadła wody danego przekroju poprzecznego rzeki,
- u — prędkość ruchu wody,
- k, a — współczynnik Chezy i promień hydrauliczny,
- g — przyspieszenie ziemskie,
- q — przepływ (ilość wody przepływająca w ciągu 1 sek przez przekrój poprzeczny rzeki),
- ω — powierzchnia przekroju poprzecznego rzeki,
- s — droga,
- t — czas.

Aby ustalić równoważność powyższych równań z równaniami układu elektrycznego, powyższe równania przekształcamy w ten sposób, by występowały w nich jako zmienne niewiadome wielkości h i q .

Biorąc mianowicie pod uwagę, że prędkość u jest związana z wielkościami q i ω następującym oczywistym związkiem

$$u = \frac{q}{\omega}$$

oraz uwzględniając, że q i ω są następującymi funkcjami czasu:

$$q = f_1[s(t), t], \quad \omega = f_2[s(t), t],$$

otrzymamy

$$\frac{du}{dt} = \frac{\dot{q}\omega - q\dot{\omega}}{\omega^2} = \frac{1}{\omega} \dot{q} - \frac{q}{\omega^2} \dot{\omega},$$

oraz

$$\begin{aligned} \dot{q} &= \frac{\partial q}{\partial s} \frac{ds}{dt} + \frac{\partial q}{\partial t} = \frac{q}{\omega} \cdot \frac{\partial q}{\partial s} + \frac{\partial \omega}{\partial t}, \\ \dot{\omega} &= \frac{\partial \omega}{\partial s} \cdot \frac{q}{\omega} + \frac{\partial \omega}{\partial t}, \end{aligned}$$

a zatem

$$\frac{du}{dt} = \frac{q}{\omega^2} \cdot \frac{\partial q}{\partial s} + \frac{1}{\omega} \cdot \frac{\partial q}{\partial t} - \frac{q^2}{\omega^3} \cdot \frac{\partial \omega}{\partial s} - \frac{q}{\omega^2} \cdot \frac{\partial \omega}{\partial t}.$$

Uwzględniając powyższe wyrażenie w pierwszym równaniu Saint-Venanta otrzymamy

$$i_0 - \frac{\partial h}{\partial s} = \frac{q^2}{k^2 \omega^2 a} + \frac{1}{\omega g} \cdot \frac{\partial q}{\partial t} + \frac{q}{g \omega^2} \cdot \frac{\partial q}{\partial s} - \frac{q^2}{g \omega^3} \cdot \frac{\partial \omega}{\partial s} - \frac{q}{g \omega^2} \cdot \frac{\partial \omega}{\partial t}.$$

Możemy następnie z równania powyższego wyrugować z wyrażeń $\frac{\partial \omega}{\partial t}$ i $\frac{\partial \omega}{\partial s}$ zmienną ω . Biorąc mianowicie pod uwagę następujące zależności

$$\frac{\partial \omega}{\partial t} = \frac{d\omega}{dh} \cdot \frac{\partial h}{\partial t} = b(h) \frac{\partial h}{\partial t},$$

$$\frac{\partial \omega}{\partial s} = \frac{d\omega}{dh} \cdot \frac{\partial h}{\partial s} = b(h) \frac{\partial h}{\partial s},$$

otrzymamy

$$i_0 - \frac{\partial h}{\partial s} = \frac{q^2}{k^2 \omega^2 a} + \frac{1}{\omega g} \cdot \frac{\partial q}{\partial t} + \frac{q}{g \omega^2} \cdot \frac{\partial q}{\partial s} - \frac{q^2}{g \omega^3} \cdot b(h) \cdot \frac{\partial h}{\partial s} - \frac{q}{g \omega^2} b(h) \frac{\partial h}{\partial t}.$$

Korzystając następnie z równania (2) zastępujemy wyrażenie $\frac{\partial q}{\partial s}$ przez $\frac{\partial h}{\partial t}$ według wzoru

$$\frac{\partial q}{\partial s} = - \frac{d\omega}{dh} \cdot \frac{\partial h}{\partial t} = -b(h) \frac{\partial h}{\partial t}.$$

Po wykonaniu elementarnych przekształceń otrzymujemy ostatecznie następujące równania różniczkowe:

$$i_{01} - \frac{\partial h}{\partial s} = A q + B \frac{\partial q}{\partial t} + C \frac{\partial h}{\partial t}, \quad (3)$$

$$- \frac{\partial q}{\partial s} = D \frac{\partial h}{\partial t}, \quad (4)$$

gdzie:

$$A = \frac{g \cdot \omega}{k^2 a (g \omega^3 - q^2 b)} q,$$

$$B = \frac{\omega^2}{g \omega^3 - q^2 b},$$

$$C = \frac{-2 b \omega}{g \omega^3 - q^2 b} q,$$

$$D = b,$$

$$i_{01} = \frac{g \omega^3}{g \omega^3 - q^2 b} i_0.$$

Równania powyższe stanowią układ przekształconych równań Saint-Venanta, w których rolę zmiennych niewiadomych grają funkcje q i h . Należy zwrócić uwagę na fakt, że współczynniki powyższych równań nie są stałymi liczbowymi, lecz skomplikowanymi funkcjami zmiennych q i h .

Można wykazać, że punkty nieciągłości funkcji A , B , C i D wyznaczają tzw. krytyczne wartości głębokości h lub przepływu q wody. Dla wartości tych energia właściwa rozważanego „zwilżonego” przekroju koryta

osiąga wartość minimalną. Można również wykazać, że dla wartości h_{kr} i q_{kr} po ustaleniu się przebiegów nachylenie zwierciadła rzeki w danym punkcie wynosi: $\frac{dh}{ds} = \infty$.

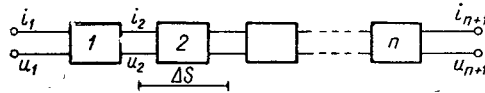
Dla przypadku $h > h_{kr}$ i $q < q_{kr}$ przepływ wody jest spokojny, natomiast dla $h < h_{kr}$ i $q > q_{kr}$ mamy do czynienia z przepływem burzliwym. O tym, czy w danym przypadku obowiązują pierwsze czy też drugie nierówności, decyduje przede wszystkim nachylenie dna koryta rzeki.

Przyjmując jako założenie, że wysokość h poziomu wody w rzece będzie rozpatrywana w modelu elektrycznym przez napięcie u , a przepływ q — przez prąd i , otrzymujemy następujące równanie różniczkowe określające elektryczny układ modelujący:

$$E_{01} - \frac{\partial u}{\partial s} = A(i, u)i + B(i, u)\frac{\partial i}{\partial t} + C(i, u)\frac{\partial u}{\partial t}, \quad (5)$$

$$-\frac{\partial i}{\partial s} = D(u)\frac{\partial u}{\partial t}. \quad (6)$$

Ponieważ zakładamy, że model elektryczny będzie zbudowany ze skończonej ilości elementów skupionych, więc równania powyższe mogą



Rys. 1. Linia o postaci łańcucha czwórników

być zrealizowane jedynie w sposób przybliżony. Zakładamy mianowicie, że układem modelującym jest linia o postaci łańcucha czwórników wskazanych na rys. 1.

Ponumerowanie powyższych czwórników jest równoznaczne z założeniem, że element drogi Δs odwzorowanej przez układ jest równy jedności.

Wobec powyższego przyrost napięcia $-\frac{\partial u}{\partial s}$ i przyrost prądu $-\frac{\partial i}{\partial s}$ na „ n -tym” czwórniku zastąpimy różnicami $i_n - i_{n+1}$, $u_n - u_{n+1}$

$$-\frac{\partial i}{\partial s} \sim i_n - i_{n+1},$$

$$-\frac{\partial u}{\partial s} \sim u_n - u_{n+1}.$$

W ten sposób otrzymujemy następujące równania różnicowo-różniczkowe:

$$u_n - u_{n+1} = -E_{01} + A i_n + B \frac{di_n}{dt} + C \frac{du_{n+1}}{dt}, \quad (7)$$

$$i_n - i_{n+1} = D \frac{du_{n+1}}{dt}. \quad (8)$$

2.2. Badanie szczególnych własności równań Saint-Venanta

Wykażemy, że równania powyższe nie mogą być modelowane za pomocą układu biernego, energetycznie symetrycznego. Jeżeli mianowicie dla uproszczenia rozumowania założymy, że współczynniki powyższych równań są liczbami, wówczas pochodne

$$\frac{di_n}{dt} \text{ i } \frac{du_{n+1}}{dt}$$

można będzie zastąpić wyrażeniami

$$sI_n \text{ i } sU_{n+1},$$

gdzie s jest liczbą zespoloną, a I_n , U_{n+1} są obrazami Laplace'owskimi prądu i_n i napięcia u_{n+1} . Po wykonaniu elementarnych przekształceń otrzymamy więc następujące równania liniowe:

$$U_n + E_{01} = I_n Z_{11} - I_{n+1} Z_{12}, \quad (9)$$

$$U_{n+1} = I_n Z_{21} - I_{n+1} Z_{22}, \quad (10)$$

gdzie:

$$Z_{11} = A \frac{C}{D} + Bs + \frac{1}{Ds},$$

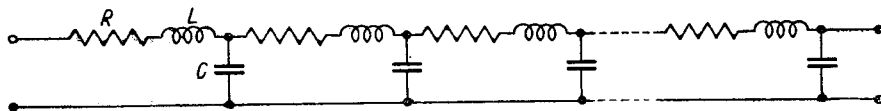
$$Z_{12} = \frac{C}{D} + \frac{1}{Ds},$$

$$Z_{21} = Z_{22} = \frac{1}{Ds}.$$

Weźmy pod uwagę macierz Z układu równań (9) i (10):

$$Z = \begin{bmatrix} Z_{11} & Z_{12} \\ Z_{21} & Z_{22} \end{bmatrix}$$

Ponieważ macierz Z jest macierzą niesymetryczną ($Z_{12} \neq Z_{21}$), więc na mocy znanego twierdzenia teorii liniowych układów elektrycznych



Rys. 2. Układ elektryczny modelujący równania (11)

otrzymujemy wniosek, że równania (9), (10) są nierealizowalne za pomocą energetycznie symetrycznego układu biernego.

Można zauważyć, że warunek symetrii macierzy Z byłby spełniony w przypadku, gdy współczynnik C równania (7) jest równy zeru

$$C = 0.$$

Wówczas otrzymalibyśmy

$$Z_{12} = Z_{21} = Z_{22} \frac{1}{Ds}, \quad (11)$$

$$Z_{11} = A + Bs + \frac{1}{Ds}.$$

W tym przypadku układem modelującym powyższe równania byłby układ zbudowany z czwórników wskazanych na rys. 2.

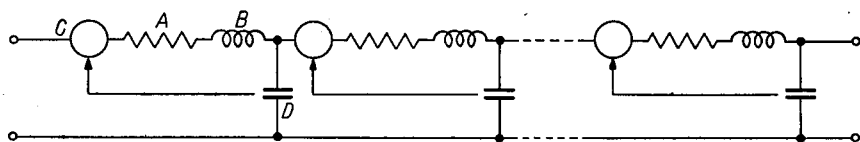
Należy jednak podkreślić, że przyjęte powyżej założenie

$$C = 0$$

jest w ogólnym przypadku niedopuszczalne.

2.3. Ustalenie struktury układu modelującego pełne równanie Saint-Venanta

Badając możliwości realizacji układu elektrycznego modelującego „pełne” równania różniczkowe (7) i (8) dochodzimy do wniosku, że układem takim jest układ łańcuchowego połączenia czwórników o strukturze



Rys. 3. Struktura układu modelującego pełne równanie Saint-Venanta

wskazanej na rys. 3. Czwórniki te zawierają oprócz elementów biernych element sterowany, którego zadaniem jest odtworzenie wyrazu

$$C \frac{d u_{n+1}}{d t}$$

w równaniu różniczkowym (7).

W układzie wskazanym na rys. 3 założono, że oporność czynna r jest znacznie mniejsza od oporności pojemnościowej $\left| \frac{1}{Ds} \right|$, dzięki czemu napięcie sterujące element C jest w przybliżeniu proporcjonalne do pochodnej napięcia u_{n+1} .

2.4. Ustalenie struktury układu elektrycznego modelującego uproszczone równania Saint-Venanta

Weźmy pod uwagę równanie Saint-Venanta (1) i oznaczmy w tym równaniu poszczególne wyrazy jako I_1, I_2, I_3 :

$$I_1 = \frac{u^2}{k^2 a},$$

$$I_2 = \frac{1}{g} \frac{\partial u}{\partial t},$$

$$I_3 = \frac{1}{g} u \frac{\partial u}{\partial t}.$$

Wyrazom powyższym możemy nadać pewną interpretację fizyczną, wyraz I_1 mianowicie można traktować jako składnik nachylenia fali wynikający ze straty energii na tarcie. Wyraz I_2 można traktować jako składnik nachylenia fali wynikający ze zmiany (w czasie) prędkości w danym przekroju rzeki. Wyraz I_3 można na koniec traktować jako składnik nachylenia fali wynikający z przyrostu prędkości przy przejściu fali z jednego przekroju rzeki do drugiego. Składniki I_1 , I_2 , I_3 są oczywiście proporcjonalne do sił, jakie są zużyte odpowiednio na tarcie, na zmianę prędkości w czasie i na zmianę prędkości w funkcji drogi.

Zastanówmy się obecnie, w jakim stosunku względem siebie pozostają wyrazy I_1 , I_2 i I_3 .

Trzeba podkreślić, że wzajemny stosunek powyższych wyrazów może być bardzo różny w zależności od warunków, w jakich odbywa się ruch fali. Na przykład w przypadku gdy mamy do czynienia z małą prędkością początkową i równocześnie z nagłą zmianą przepływu lub z nagłą zmianą poziomu wody, wyrazy inercyjne I_2 i I_3 wskutek powstawania dużych sił inercji będą znacznie przewyższały wyraz I_1 , który związany jest ze stratą energii na tarcie. W tym przypadku możemy, nie wprowadzając praktycznie niemal żadnego błędu, pominąć w równaniu (1) Saint-Venanta wyraz I_1 . Otrzymalibyśmy wówczas równania cieczy idealnej, pozbawionej tarcia, a postać tych równań byłaby następująca:

$$i_0 - \frac{\partial h}{\partial s} = \frac{1}{g} \left(\frac{\partial h}{\partial t} + u \frac{\partial h}{\partial s} \right),$$

$$\frac{\partial q}{\partial s} + \frac{\partial \omega}{\partial t} = 0,$$
(12)

W przypadku fali powodziowej sytuacja jest wręcz odwrotna. Fala powodziowa ma mianowicie bardzo spłaszczony zarys.. Wystarczy — na potwierdzenie powyższej tezy — przytoczyć fakt, że nawet w przypadku bardzo gwałtownych powodzi stosunek wysokości fali do jej długości nigdy nie przekracza dziesięciotysięcznych części. Z tego względu nie może być w tym przypadku mowy o dużych zmianach prędkości i o dużych wartościach wyrazów inercyjnych I_2 i I_3 . Według Delemé dla stanu powodziowego jednej z rzek pomierzone wartości wyrazów I_2 i I_3 wynoszą:

$$I_2 = 0,0000285,$$

$$I_3 = 0,0000176,$$

oraz

$$I = I_1 + I_2 + I_3 = 0,0041.$$

Wobec powyższego nachylenie fali spowodowane stratą energii na tarcie w stosunku do całkowitego nachylenia fali wynosi

$$\frac{I_1}{I} 100\% \approx 99\%. \quad (13)$$

Nie ulega wątpliwości, że w tym przypadku można bez szkody dla dokładności wyników praktycznych pominąć wyrazy inercyjne I_2 i I_3 . Otrzymujemy więc dla tego przypadku następujący układ równań uproszczonych:

$$\left. \begin{aligned} i_0 - \frac{\partial h}{\partial s} &= \frac{u^2}{k^2 a}, \\ \frac{\partial q}{\partial s} + \frac{\partial \omega}{\partial t} &= 0. \end{aligned} \right\} \quad (14)$$

Występujące w pierwszym z powyższych równań współczynniki k oraz a są funkcjami niewiadomej h , a mianowicie:

$$k \approx \frac{h^{1/6}}{n}, \quad a \approx \beta h, \quad (15)$$

gdzie n jest współczynnikiem chropowatości koryta oraz β jest współczynnikiem zależnym od kształtu koryta. W przybliżeniu można przyjąć, że β jest wielkością stałą, niezależną od h .

Uwzględniając w równaniach (14) oczywiste związki

$$\left. \begin{aligned} q &= u \omega, \\ h &= f(\omega) \end{aligned} \right\} \quad (16)$$

oraz uwzględniając zależności (15), otrzymamy ostatecznie następujące równania:

$$\left. \begin{aligned} i_0 - \frac{\partial h}{\partial s} &= \frac{n^2}{\beta \omega^2 h^{4/3}} q^2, \\ -\frac{\partial q}{\partial s} &= b(h) \frac{\partial h}{\partial t}, \end{aligned} \right\} \quad (17)$$

gdzie

$$b(h) = \frac{d\omega}{dh}. \quad (18)$$

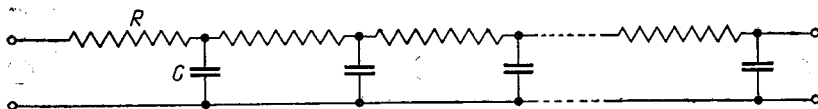
Przyjmując jako założenie, że wysokość poziomu wody w rzece będzie reprezentowana na modelu elektrycznym przez napięcie u , a przepływ

wody q — przez prąd i , otrzymamy następujące równania różniczkowe określające elektryczny układ modelujący:

$$\left. \begin{aligned} e_0 - \frac{\partial u}{\partial s} &= \frac{n^2 i}{\beta \omega^2 u^{4/3}} i, \\ - \frac{\partial i}{\partial s} &= D(u) \frac{\partial u}{\partial t}. \end{aligned} \right\} \quad (19)$$

Można zauważyć, że równania (19) są równoważne z równaniami toru elektrycznego o elementach R , C , a mianowicie z następującymi równaniami:

$$\begin{aligned} e_0 - \frac{\partial u}{\partial s} &= R i, \\ - \frac{\partial i}{\partial s} &= C \frac{\partial u}{\partial t}. \end{aligned}$$



Rys. 4. Struktura układu modelującego uproszczone równania Saint-Venanta

Płyńie stąd wniosek, że równania (19) można modelować za pomocą układu elektrycznego zbudowanego z łańcuchowego połączenia czwórników o strukturze wskazanej na rys. 4.

3. UKŁADY REALIZUJĄCE WSPÓŁCZYNNIKI RÓWNAŃ SAINT-VENANTA

3.1. Wstęp

W rozdziale tym rozszerzymy techniczne możliwości realizacji elementów elektrycznych odtwarzających współczynniki równań Saint-Venanta. Rozważania ograniczymy jedynie do przypadku uproszczonych równań Saint-Venanta oraz do przypadku równania opisującego krzywą spiętrzenia.

Celem rozważań przedstawionych w niniejszym rozdziale jest wskazanie niektórych możliwości technicznej realizacji elementów wchodzących w skład modelujących układów elektrycznych. Trzeba podkreślić, że oczywiście można wskazać bardzo wiele sposobów realizacji powyższych elementów: rozwiązanie podane w referacie prawdopodobnie jest jednak najbardziej racjonalne.

3.2. Ustalenie postaci funkcji określających elementy układu elektrycznego modelującego uproszczone równania Saint-Venanta

Według schematu przedstawionego na rys. 4, w układzie elektrycznym modelującym uproszczone równania Saint-Venanta występują w zasadzie

dwa elementy, a mianowicie element oporowy $R(i, u)$ oraz element pojemnościowy $D(u)$. Występujące również w tym układzie źródło siły elektromotorycznej e_0 nie będzie przedmiotem naszych rozważań.

W celu zorientowania się w charakterze zależności $R(i, u)$ i $D(u)$ od prądu i napięcia przeprowadzimy następujące rozumowanie.

Rozważmy dwa następujące przypadki szczególne, a mianowicie przypadek kanału o korycie prostokątnym i przypadek kanału o korycie trapezowym.

a) Przypadek kanału o korycie prostokątnym

W tym przypadku przekrój czynny ω kanału wyraża się wzorem

$$\omega = \begin{cases} b_0 \cdot h, & \text{gdy } h \leq H, \\ \infty, & \text{gdy } h > H, \end{cases}$$

gdzie b_0 jest szerokością koryta kanału i H jest głębokością kanału. Wówczas

$$D(u) \sim \frac{d}{dh} = \begin{cases} b_0 = \text{const}, & \text{gdy } h \leq H, \\ \infty, & \text{gdy } h > H, \end{cases}$$

oraz

$$\omega^2 \sim \begin{cases} b_0^2 U^2, & \text{gdy } U \leq U_0, \\ \infty, & \text{gdy } U > U_0, \end{cases}$$

i wobec powyższego otrzymujemy następujące równania (20), słuszne dla przypadku napięć U nie przekraczających wartości granicznej U_0 , odpowiadającej głębokości kanału.

$$\left. \begin{aligned} e_0 - \frac{\partial U}{\partial s} &= \frac{n^2}{\beta \cdot b^2 U^{10/3}} i^2, \\ - \frac{\partial i}{\partial s} &= b_0 \frac{\partial U}{\partial t}. \end{aligned} \right\} \quad (20)$$

b) Przypadek kanału o korycie trapezowym

W tym przypadku przekrój czynny ω kanału wyraża się wzorem

$$\omega = \begin{cases} b_0 \cdot h + m h^2, & \text{gdy } h \leq H, \\ \infty, & \text{gdy } h > H, \end{cases}$$

gdzie:

b — szerokość dna kanału,

m — nachylenie skarpy,

H — głębokość kanału.

Wówczas

$$D(u) \sim \frac{d\omega}{dh} \sim \begin{cases} b_0 + 2mU, & \text{gdy } U \leq U_0, \\ \infty, & \text{gdy } U > U_0, \end{cases}$$

oraz

$$\omega^2 \sim \begin{cases} b_0^2 U^2 + 2 b_0 m U^3 + m^2 U^4, & \text{gdy } h \leq H, \\ \infty, & \text{gdy } h > H \end{cases}$$

i wobec powyższego otrzymujemy następujące równanie (21), słuszne dla napięć U nie przekraczających wartości granicznej U_0 , odpowiadającej głębokości kanału

$$\begin{aligned} e_0 - \frac{\partial U}{\partial s} &= \frac{n^2}{\beta (b_0^2 + 2 b_0 m U + m^2 U^2) U^{10/3}} i^2, \\ - \frac{\partial i}{\partial s} &= (b_0 + 2 m U) \frac{\partial U}{\partial t}. \end{aligned} \quad (21)$$

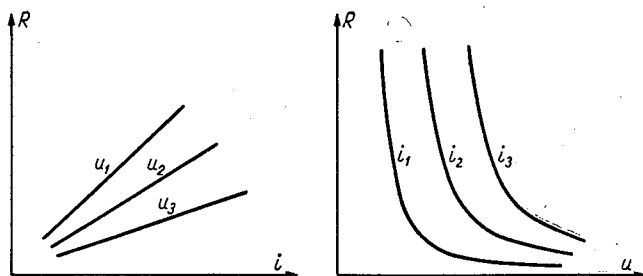
Według wzorów (20) dla przypadku kanału o korycie prostokątnym otrzymamy następujące funkcje określające elementy $R(i, u)$ i $D(u)$ modelu

$$\left. \begin{aligned} R(i, u) &= \frac{n^2}{\beta b_0^2 U^{10/3}} i = \lambda \frac{i}{U^{10/3}}, \\ D(u) &\neq b_0 = \text{const}, \end{aligned} \right\} \quad (22)$$

gdzie:

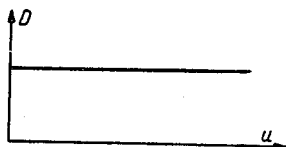
$$\lambda = \frac{n^2}{\beta b_0^2}.$$

Wykresy funkcji $R(i, u)$, $D(u)$ podane są na rysunkach 5 i 6.



Rys. 5. Charakter przebiegu funkcji $R(i, u)$ określającej element oporowy w modelu kanału wodnego o przekroju prostokątnym

Na podstawie równań (21) możemy w sposób analogiczny wyznaczyć funkcje określające elementy modelu kanału o przekroju trapezowym.



Rys. 6. Charakter przebiegu funkcji $D(u)$ określającej element pojemnościowy w modelu kanału wodnego o przekroju prostokątnym

Otrzymamy mianowicie następujące wyniki

$$R(i, u) = \frac{n^2 i}{\beta (b_0^2 + 2 b_0 m U + m^2 U^2) U^{10/3}} = \lambda \frac{i}{W(u) \cdot U^{10/3}},$$

gdzie:

$$\lambda = \frac{n^2}{\beta},$$

$$W(u) = b_0^2 + 2 b_0 m U + m^2 U^2 = a_0 + a_1 U + a_2 U^2;$$

$$D(u) = b_0 + 2 m U = b_0 + b_1 U, \quad \text{gdzie} \quad b_1 = 2 m.$$

Założmy, że współczynniki n , b , m posiadają następujące wartości liczbowe:

$$n = 0,013; \quad b = 5 m; \quad m = 3.$$

Funkcje $R(i, u)$, $D(u)$ dla rozważanego przypadku są określone przez wzory

$$R(i, u) = \frac{169 \cdot 10^{-6}}{25 + 30 U + 9 U^2} \cdot \frac{i}{U^{10/3}} [k_\Omega \Omega], \quad (23a)$$

$$D(u) = (5 + 6 U) [k_F F], \quad (23b)$$

gdzie k_Ω i k_F są umownie przyjętymi współczynnikami odwzorowania.

Przebieg funkcji (23a) i (23b) przedstawiony jest na rys. 5 i 6.

3.3. Sposób realizacji elementów układu modelującego uproszczone równanie Saint-Venanta

W przypadku kanału prostokątnego pojemność $D(u)$ czwórników w układzie modelującym jest wielkością stałą, niezależną od napięcia. Realizacja tej pojemności nie przedstawia więc żadnego problemu.

W przypadku kanału o przekroju nieprostokątnym, a w szczególności w przypadku kanału o przekroju trapezowym pojemność ta jest wielkością nieliniową. Pojemności nieliniowe można realizować za pomocą specjalnych materiałów dielektrycznych, które nazywamy segnetodielektrykami. W szczególności do segnetodielektryków zaliczamy tytanian baru BaTiO_3 . Przebieg charakterystyki stałej dielektrycznej tytanianu baru w funkcji natężenia pola elektrycznego posiada charakter podobny do przebiegu charakterystyki przenikalności magnetycznej materiałów ferromagnetycznych w funkcji natężenia pola magnetycznego. Przez łączenie równoległe lub szeregowo elementów pojemnościowych o powyższych dielektrykach oraz liniowych elementów pojemnościowych można by realizować z góry zadane charakterystyki nieliniowych pojemności w czwórnikach modelujących ruch wody w rzece. Należy zauważyć, że charakterystyka stałej dielektrycznej w funkcji temperatury dla omawianych dielektryków jest niemal płaska w zakresie temperatur od 0 do 80 °C. Płyńie stąd wniosek, że wpływ temperatury na pracę tego ro-

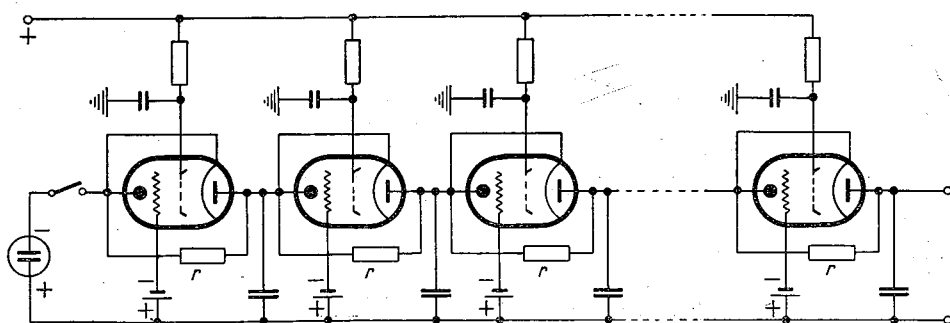
dzaju elementów jest nieznaczny. Zagadnienie realizacji odpowiedniej charakterystyki pojemności nieliniowej z kondensatorów o dielektrykach z tytanianu baru stanowi samo w sobie trudny problem techniczny.

Oprócz wskazanej powyżej możliwości realizacji pojemności nieliniowych za pomocą segnetodielektryków istnieją oczywiście również inne możliwości. Między innymi pojemności nieliniowe można realizować za pomocą układów zawierających lampy reaktancyjne. Trzeba jednak przy tym wziąć pod uwagę fakt, że w ten sposób realizowane pojemności spełniałyby stawiane im warunki tylko w przypadku przebiegów zmiennych zawierających się w określonym pasmie częstotliwości.

Dla realizacji pojemności nieliniowej można by też korzystać ze specjalnych układów tryggerowych, automatycznie przełączających pojemności odpowiadające kilku „napięciom progowym”. W ten sposób można by otrzymać pojemność o charakterystyce schodkowej.

Spśród wymienionych powyżej metod realizacji pojemności nieliniowej najbardziej prostym rozwiązaniem, a zatem najbardziej realnym, jest rozwiązanie oparte na wykorzystaniu elementów pojemnościowych wykonanych z segnetodielektryków, np. z tytanianu baru.

Drugim elementem wchodzącym w skład czwórników modelujących jest element oporowy $R(i, u)$. Trudność realizacji tego elementu polega na tym, że zależy on nie tylko od przepływającego przez niego prądu, lecz również od napięcia panującego na kondensatorze $D(u)$. Charakterystyki prądowo-napięciowe $R(i, u)$ są wskazane na rys. 5. W najprostszy sposób element posiadający ten rodzaj charakterystyki można zrealizować przy wykorzystaniu odpowiedniej lampy elektronowej, a mianowicie pentody napięciowej.



Rys. 7. Pełny układ modelujący uproszczone równania Saint-Venanta

Na przebieg charakterystyk pentody można by w pewnym stopniu wpływać przez bocznikowanie jej odpowiednim oporem liniowym lub ewentualnie nieliniowym. Poważną wadą w powyższy sposób wykonanych elementów $R(i, u)$ jest trudność uzyskania zadanych charakterystyk. Poza tym elementy w ten sposób wykonane w zasadzie nie umożliwiają zmiany

swoich charakterystyk, a zatem mogą być stosowane tylko i wyłącznie w modelu przeznaczonym do z góry zadanego kanału. Elementy te nie mogłyby więc być stosowane w modelu uniwersalnym, przeznaczonym do badań dowolnych rzek. Model elektryczny zrealizowany w układzie zawierającym wyżej omawiane elementy oporowe $R(i, u)$ jest wskazany na rys. 7.

4. UKŁAD MODELUJĄCY KRZYWĄ SPIĘTRZENIA

Ruch nieustalony wody w rzece w ogólnym przypadku opisują następujące równania różniczkowe:

$$\begin{aligned} i_0 - \frac{\partial h}{\partial s} &= \frac{v^2}{k^2 a} + \frac{1}{g} \left(\frac{\partial v}{\partial t} + v \frac{\partial v}{\partial s} \right), \\ \frac{\partial q}{\partial s} + \frac{\partial \omega}{\partial t} &= 0, \end{aligned} \quad (24)$$

gdzie:

v — prędkość ruchu ośrodka,

q — przepływ,

h — głębokość wody,

ω — przekrój czynny koryta.

Założmy, że istnieją następujące granice:

$$\lim_{t \rightarrow \infty} q(s, t) = q_\infty(s), \quad (a)$$

$$\lim_{t \rightarrow \infty} h(s, t) = h_\infty(s). \quad (b) \quad (25)$$

Funkcję $h_\infty = f(q_\infty)$, określoną dla ustalonych wartości zmiennej s ($s = s_0$), nazywamy krzywą konsumpcyjną w punkcie s_0 . Funkcję $h_\infty(s) = f(s)$ nazywamy zaś krzywą spiętrzenia.

Ze względu na założone istnienie granicy (a), dla każdej wartości zmiennej s zachodzi związek

$$\lim_{t \rightarrow \infty} |q(s, t + \Delta t) - q(s, t)| = 0,$$

zatem

$$\lim_{t \rightarrow \infty} \frac{\partial}{\partial t} q(s, t) = \lim_{\substack{t \rightarrow \infty \\ \Delta t \rightarrow 0}} \frac{q(s, t + \Delta t) - q(s, t)}{\Delta t} = 0$$

oraz z uwagi na związek

$$v = \frac{q}{\omega},$$

$$\lim_{t \rightarrow \infty} \frac{\partial}{\partial t} v(s, t) = 0.$$

Ze względu na założenie istnienia granicy (b) dla każdej wartości zmiennej s , zachodzi równość

$$\lim_{t \rightarrow \infty} \frac{\partial}{\partial t} h(s, t) = 0.$$

Zatem z powodu jednoznaczności funkcji $\omega = f(h)$ mamy

$$\frac{\partial \omega}{\partial t} = 0.$$

Z równania ciągłości otrzymujemy więc natychmiast

$$\frac{\partial q_{\infty}}{\partial s} = 0,$$

czyli

$$q_{\infty} = \text{const.}$$

Według powyższych rozważań dla badanego przypadku granicznego otrzymujemy następujące równanie różniczkowe zwyczajne określające przy założonych warunkach brzegowych krzywą spiętrzenia $h_{\infty}(s)$:

$$i_0 - \frac{d h_{\infty}}{d s} = \frac{v_{\infty}^2}{k^2 a} + \frac{v_{\infty}}{g} \cdot \frac{d v_{\infty}}{d s}; \quad (26a)$$

(dla uproszczenia zapisu, w dalszych rozważaniach będziemy pomijali indeksy ∞ pisząc po prostu h, v, ω, q).

Równanie (26a) można również zapisać w następującej postaci równoważnej:

$$i_0 = \frac{v^2}{k^2 a} + \frac{d}{d s} \left(h + \frac{v^2}{2g} \right). \quad (26b)$$

Wyrażenie pod znakiem różniczki, a więc

$$h + \frac{v^2}{2g}$$

jest sumą energii potencjalnej i kinetycznej w badanym przekroju rzeki w odniesieniu do jednostki ciężaru wody. Wyrażenie

$$\frac{v^2}{k^2 a}$$

jest energią strat na tarcie przypadającą na jednostkę długości drogi i jednostkę ciężaru wody. W hydrotechnice wyrażenie

$$E = h + \frac{v^2}{2g}$$

jest znane pod nazwą energii właściwej przekroju (zob. [1]).

Warto zauważyć, że istnieje tu prosta analogia z obwodem elektrycznym zawierającym opornik, cewkę indukcyjną i kondensator, a więc zawierającym elementy R , L , C . Moc zużyta w powyższym obwodzie wyraża się mianowicie następującym równaniem różniczkowym

$$P = RI^2 + \frac{d}{dt} \left(\frac{LI^2}{2} + \frac{CU^2}{2} \right) \quad (27)$$

analogicznym do równania (26a). Odpowiednikiem energii zużytej na tarcie przy przepływie wody jest więc w powyższym obwodzie elektrycznym strata energii na ciepło wyrażająca się wielkością RI^2 , odpowiednikiem zaś sumy energii potencjalnej i kinetycznej wody jest energia elektryczna nagromadzona w polu elektrycznym kondensatora i energia magnetyczna nagromadzona w polu magnetycznym cewki indukcyjnej, wyrażające się wielkością

$$\frac{CU^2}{2} + \frac{LI^2}{2}.$$

Dzięki istnieniu powyższych analogii między równaniami (26a) i (27) istnieje teoretyczna możliwość modelowania krzywej spiętrzenia za pomocą obwodu elektrycznego zawierającego elementy R , L , C . Oczywiście w tym przypadku zmienną drogi s w układzie hydrotechnicznym reprezentowałaby zmienna czasu t w układzie elektrycznym. Krzywa spiętrzenia $h = f(s)$ byłaby w powyższym układzie elektrycznym reprezentowana przez funkcję napięcia na kondensatorze $U = f(t)$.

Modelowanie krzywej spiętrzenia za pomocą powyższego obwodu elektrycznego kryje w sobie jednak duże trudności techniczne ze względu na konieczność realizacji nieliniowych elementów biernych, np. nieliniowej indukcyjności.

Z powyższych względów krzywą spiętrzenia dogodniej prawdopodobnie będzie modelować za pomocą układu blokowego realizującego nieliniowe równanie różniczkowe. W tym celu przeprowadzimy następujące rozumowanie:

Weźmy pod uwagę równanie różniczkowe

$$i_0 - \frac{dh}{ds} = \frac{v^2}{k^2 a} + \frac{v}{g} \cdot \frac{dv}{ds},$$

określające krzywą spiętrzenia $h(s)$.

Dla ujawnienia w powyższym równaniu zmiennej h uwzględnimy związek:

$$v = \frac{q}{\omega},$$

ponieważ ω jest jednoznaczna funkcją zmiennej h , a ponadto w ogólnym przypadku zależy od zmiennej s ; otrzymamy więc:

$$i_0 - \frac{dh}{ds} = \frac{q^2}{k^2 a \omega^2} + \frac{q^2}{g \omega} \cdot \frac{d}{ds} \cdot \left(\frac{1}{\omega} \right).$$

Następnie, ponieważ

$$\left(\frac{1}{\omega} \right) = - \frac{1}{\omega^2} \cdot \frac{d\omega}{ds},$$

więc

$$i_0 - \frac{dh}{ds} = q^2 \left[\frac{1}{k^2 a \omega^2} - \frac{1}{g \omega^3} \cdot \frac{d\omega}{ds} \right],$$

uwzględniając zaś następujący oczywisty związek:

$$\frac{d\omega}{ds} = \frac{d\omega}{dh} \cdot \frac{dh}{ds} = b \frac{dh}{ds},$$

gdzie zgodnie z poprzednio przyjmowanym oznaczeniem

$$b = \frac{d\omega}{dh},$$

będzie

$$i_0 - \frac{dh}{ds} = q^2 \left(\frac{1}{k^2 a \omega^2} - \frac{b}{g \omega^2} \cdot \frac{dh}{ds} \right).$$

Równanie krzywej spiętrzenia można zatem zapisać w następującej postaci

$$A(h) \frac{dh}{ds} + B(h) = i_0, \quad (28)$$

gdzie:

$$A(h) = \frac{1 - b(h)}{g \omega^3(h)} q^2,$$

$$B(h) = \frac{q^2}{k^2 a \omega^2(h) \cdot h}.$$

Równanie (28) określa krzywą spiętrzenia przy założonych warunkach brzegowych, jest ono nieliniowym równaniem różniczkowym pierwszego rzędu.

W przypadku kanału o korycie prostokątnym mamy:

$$\omega = b_0 \cdot h,$$

gdzie $b_0 = b$ jest szerokością kanału. Zatem w tym przypadku

$$A(h) = 1 - \frac{q^2}{g b_0^2} \cdot \frac{1}{h^3},$$

$$B(h) = \frac{q^2}{k^2 a b_0^2} \cdot \frac{1}{h^3},$$

uwzględniając zaś wyprowadzone poprzednio zależności (15)

$$k \approx \frac{h^{1/6}}{n}, \quad a \approx \beta h \approx h$$

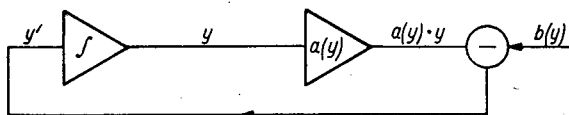
otrzymamy ostatecznie dla przypadku koryta o przekroju prostokątnym

$$A(h) = 1 - \frac{q^2}{g b_0^2} \cdot \frac{1}{h^3},$$

$$B(h) = \frac{n^2 q^2}{b_0^2} \cdot \frac{1}{h^{4/3}},$$

gdzie n jest współczynnikiem chropowatości koryta.

Równanie (28) jako równanie różniczkowe zwyczajne pierwszego rzędu można zrealizować za pomocą następującego układu zawierającego człon całkujący i dwa generatory funkcyjne (rys. 8).



Rys. 8. Układ modelujący krzywą spiętrzenia

W układzie wskazanym na rysunku 8 mamy trzy niezależne elementy, a mianowicie człon całkujący, człon przekazujący o nieliniowym współczynniku przenoszenia $a(y)$, węzeł sumujący.

Ze względu na sposób połączenia powyższych elementów otrzymamy równanie

$$y' = b(y) - a(y) \cdot y,$$

czyli otrzymamy następujące równanie różniczkowe:

$$y' + a(y) \cdot y = b(y), \quad (29)$$

identyczne z równaniem (28) określającym krzywą spiętrzenia; o ile założymy

$$y \sim h, \quad s \sim t,$$

$$a(y) = \frac{B}{A} = \frac{n^2 g \cdot q}{g b_0^2 y^3 - q^2} \cdot \frac{1}{y^{4/3}}, \quad (30a)$$

$$b(y) = \frac{i_0}{A} = \frac{g \cdot b_0^2 y^3}{g b_0^2 y^3 - q^2} i_0. \quad (30b)$$

5. WNIOSKI KOŃCOWE

Streszczając rozważania podane w poprzednich rozdziałach należy stwierdzić, co następuje:

1. Istnieje teoretyczna i techniczna możliwość realizacji układu elektrycznego modelującego uproszczone równanie Saint-Venanta o z góry

zadanych współczynnikach. Za pomocą takiego układu można by badać przebieg fali powodziowej w rzece lub w kanale w przypadku małych prędkości i małych przyspieszeń ruchu wody. Dla spełnienia powyższych warunków jest między innymi rzeczą konieczną, aby spadek dna koryta rzeki lub kanału był niewielki i znacznie mniejszy od spadku krytycznego. Według przeprowadzonych wstępnych obliczeń warunek ten jest spełniony w przypadku rozważanego przez nas kanału.

2. Dla określenia krzywej spiętrzenia, a więc krzywej swobodnej powierzchni wody w kanale przy ruchu ustalonym, konieczne jest opracowanie specjalnego układu modelującego, realizującego równanie krzywej spiętrzenia. Układ ten byłby układem uniwersalnym i służyłby do badania krzywych spiętrzenia przy dowolnych założeniach.

3. Istnieje również możliwość realizacji układu modelującego pełne równanie Saint-Venanta o z góry zadanych współczynnikach. Realizacja takiego układu kryje w sobie jednak znacznie większe trudności techniczne niż realizacja układu modelującego równanie Saint-Venanta uproszczone.

4. Dokładność uzyskiwanych na modelu elektrycznym wyników zależy od całego szeregu przyczyn wprowadzających nieuniknione błędy. Błędy powyższe można podzielić na następujące kategorie.

4.1. Błąd wprowadzony przez równanie różniczkowe opisujące dynamikę ruchu wody. Błąd ten wynika z idealizacji zjawiska, którą w sposób nieunikniony wprowadza się przy ustalaniu odpowiedniego równania różniczkowego.

4.2. Błąd wprowadzony przez niedokładne określenie współczynników równań różniczkowych. Błąd ten jest uwarunkowany dokładnością pomiarów parametrów charakteryzujących kanał lub rzekę (współczynnika chropowatości, nachylenia dna, współczynnika skarp itp.).

4.3. Błędy wprowadzone przez układ modelujący. Błędy te można z kolei podzielić na następujące grupy.

a) błąd wynikający z założenia, że układ modelujący równanie różniczkowe cząstkowe jest zbudowany z elementów skupionych i posiada skończoną liczbę tych elementów (w przypadku układów modelujących nieustalony ruch wody),

b) błędy wynikające z niedokładności realizacji charakterystyk poszczególnych elementów układu modelującego,

c) błędy wynikające z niedokładności pomiaru przebiegów elektrycznych odwzorowujących na modelu wielkości hydrotechniczne,

d) błędy wynikające z idealizacji warunków początkowych i warunków granicznych,

e) błędy przypadkowe mające swe źródło w różnorodnych przyczynach i wynikające np. z niestałości w czasie elementów układu modelującego i z wpływu czynników zakłócających (np. szumy).

Należy zauważyć, że spośród wszystkich wymienionych wyżej grup błędów decydujące znaczenie będą niewątpliwie miały błędy wynikające z niedokładnego określenia współczynników odwzorowywanych równań.

Błędy te są uwarunkowane dokładnością wykonanych w terenie pomiarów wielkości hydrotechnicznych.

Grupy błędów ujętych w punkcie 4.1 i 4.2 nie zależą oczywiście od układu modelującego, natomiast niemal wyłącznie one decydują o dokładności wyników uzyskiwanych na układzie modelującym. Dla zilustrowania powyższego wystarczy chociażby zauważyć, że dokładność określonych na podstawie pomiarów parametrów hydrotechnicznych w najlepszym razie jest rzędu kilku procent, a przeciętnie rzędu kilkunastu lub nawet kilkudziesięciu procent. Dokładność natomiast błędów ujętych w grupie 3, a więc błędów wprowadzanych przez układ modelujący — przy obecnym stanie techniki tego rodzaju układów — może być zawarta poniżej jednego procenta.

Z powyższych rozważań dotyczących charakteru i źródeł błędów wynika ważny wniosek, że nie byłoby rzeczą celową dążenie do opracowania bardzo dokładnego układu modelującego.

Dokładność samego układu modelującego miałyby bowiem jedynie minimalny wpływ na dokładność uzyskiwanych wyników. Natomiast koszt realizacji modelu elektrycznego w sposób niewspółmierny rośnie przy zwiększaniu wymagań co do jego dokładności.

Należy przy tym zaznaczyć, że wyniki uzyskane na elektrycznym układzie modelującym co najmniej nie będą gorsze od wyników uzyskiwanych przy analitycznych obliczeniach.

5. Największą zaletą modelu elektrycznego byłaby możliwość szybkiego badania całego szeregu wariantów postawionego zagadnienia, bez konieczności przeprowadzenia żmudnych obliczeń; badania przeprowadzone na modelu umożliwiłyby więc dokonanie wyboru optymalnych (pod względem technicznym i ekonomicznym) rozwiązań.

6. Istnieje możliwość realizacji „uniwersalnego” układu modelującego ruch wody, a więc układu, który można by w sposób stosunkowo prosty i szybki przystosować do badania dowolnej istniejącej rzeki lub kanału.

Dodatek

Obliczenie krytycznej wartości nachylenia dna kanału
(według podręcznika: J. J. Agroskin „Gidrawlika”)

1) Nachylenie krytyczne dna możemy obliczyć ze wzoru

$$I_{kr} = \frac{g X_{kr}}{\alpha C_{kr} B_{kr}}.$$

Wielkości X_{kr} , C_{kr} i B_{kr} są funkcjami głębokości krytycznej h_{kr} , którą

obliczamy z wartości krytycznej $h_{kr\ pr}$ dla kanału o przekroju prostokątnym wg wzoru

$$h_{kr} = \left(1 - \frac{\sigma}{3} + 0,1056 \sigma^2\right) h_{kr\ pr}.$$

We wzorze tym $h_{kr\ pr} = 0,482 \sqrt[3]{q^2}$ [m].

Przyjmując $q = 30$ m³/sek otrzymamy

$$h_{kr\ pr} = 0,482 \sqrt[3]{30^2} \text{ m} = 4,7 \text{ m}.$$

Wielkość σ obliczamy ze wzoru

$$\sigma = \frac{m h_{kr\ pr}}{b_0}.$$

Ponieważ $m = 3$, $b = 5$ m,
więc

$$\sigma = \frac{3 \cdot 4,7}{5} = 2,8.$$

Podstawiając znalezione wartości σ i $h_{kr\ pr}$ do wzoru znajdujemy h_{kr} dla przekroju trapezowego

$$h_{kr} = \left(1 - \frac{2,8}{3} + 0,105 \cdot 2,8^2\right) \cdot 4,7 = 4,2 \text{ m}.$$

2) Obwód zwilżony krytyczny X_{kr} obliczamy ze wzoru

$$x_{kr} = b_0 + 2 h_{kr} \sqrt{1 + m^2} = 5 + 2 \cdot 4,2 \sqrt{1 + 3^2} = 31,5 \text{ m}.$$

Szerokość krytyczną zwierciadła otrzymamy ze wzoru

$$B_{kr} = b_0 + 2 m h_{kr} = 5 + 2 \cdot 3 \cdot 4,2 = 30,2 \text{ m}.$$

Współczynnik krytyczny C_{kr} obliczamy ze wzoru Manninga dla danego $n = 0,013$

$$C_{kr} = \frac{h_{kr}^{1/6}}{n} = \frac{4,2^{1/6}}{0,013} = 98,1 \text{ m}^{1/2} \text{ sec}^{-1}.$$

3) Podstawiając obliczone wartości do wzoru przytoczonego na wstępie znajdujemy ostatecznie krytyczne nachylenie dna ($g = 9,81$ m/sec², $\alpha = 1,1$)

$$I_{kr} = \frac{9,81 \cdot 31,5}{1,1 \cdot 98,1^2 \cdot 30,2} \approx 1 \cdot 10^{-3}.$$

4) Ponieważ dla badanego kanału mamy dane nachylenie $J = 0,2 \cdot 10^{-3}$, więc jest ono pięciokrotnie mniejsze od nachylenia krytycznego

$$J \ll J_{kr}.$$

С. БЕЛЛЕРТ, Х. ДЗЯТЛИК, М. КОВАЛЬСКИ

ЭЛЕКТРИЧЕСКАЯ МОДЕЛЬ ВОДЯНОГО КАНАЛА

Резюме

В статье рассмотрено подобие упрощенных уравнений Сент-Венанта определяющих динамику движения воды в канале и уравнений электрических однородных линий.

На основании выступающих аналогий соотношений проанализирована техническая возможность реализации электрической модели водяного канала с поперечным сечением в виде трапеции и с малым падением дна, соответствующей упрощенным уравнениям Сент-Венанта с заранее заданными коэффициентами. При помощи такой модели можно было бы исследовать движение фронтовой волны наводнения в реке или канале в случае выступления малых скоростей и ускорений, что обычно имеет место в реках с малым падением дна (на низменностях).

S. BELLERT, H. DZIATLIK, M. KOWALSKI

ELECTRIC MODEL OF WATER CHANNEL

Summary

The paper is concerned with the examination of analogy between Saint-Venant's simplified equations portraying dynamics of water motion in a channel and those of the long electric lines.

Departing from the observed dependencies the technical feasibility of providing the electric model of water channel of the trapeziform cross-section and the slight bottom slop which modulates Saint-Venant's simplified equations with the coefficients given a priori, is being discussed.

Such a model enables to examine the flow of the flood wave either in a river or channel, especially if the small rate of flow and small accelerations of water in motion are involved what is the case in rivers with the slight bottom slope (shallow dale rivers).

ALOJZY SPICHAŁSKI

Wpływ skaz na rozkład pola w szynie kolejowej i wytyczne dla defektoskopii magnetycznej¹⁾

Rękopis dostarczono 8.7.1959

4. WŁAŚCIWOŚCI DEFEKTOSKOPII MAGNETYCZNEJ W PORÓWNANIU Z PRĄDOWĄ

Defektoskopia magnetyczna o polu wzdłużnym wymaga znacznie lżejszego i tańszego urządzenia pomiarowego niż defektoskopia prądowa. Osobliwościami metody magnetycznej w odróżnieniu od prądowej są: zmienność przenikalności magnetycznej, przenikalność cienkiej skazy w postaci szczeliny powietrznej oraz przenikalność środowiska otaczającego szynę ze skazą. Ze względu na trudność doświadczalnego zbadania oddzielnie poszczególnych wpływów zastosowano tu zmodyfikowaną metodę numeryczną.

Wpływ zmiennej przenikalności magnetycznej okazał się wysoce korzystny w porównaniu z polem prądowym przy założeniu, że indukcja w szynie przekracza nieco wartość odpowiadającą maksymalnej przenikalności magnetycznej stali szynowej.

Przewodność magnetyczna cienkiej szczeliny zmniejsza wprawdzie stopień niejednostajności pola. Zmniejszenie to nie przekracza jednak 50% przy skazach małych, poniżej 15% przekroju sztaby o grubości większej niż 0,1 mm.

Przenikalność magnetyczna otaczającego środowiska powoduje wypychanie linii indukcji na zewnątrz szyny.

Powstaje dopełniające zewnętrzne pole magnetyczne, uzależnione od rozkładu stopnia niejednostajności pola przy powierzchni. Strukturę takiego pola wyznaczono dla skaz wzorcowych w główce oraz w stopce szyny. Składową poziomą i pionową zewnętrznego pola dopełniającego przedstawiono wykreślnie w zależności od odległości.

4.1. Wstęp

W celu wykrywania skaz w szynach kolejowych wykorzystuje się zjawisko deformacji pola prądowego lub magnetycznego w sąsiedztwie skazy. Jeśli przy powierzchni szyny z dala od skazy gęstość prądu wy-

1) Całość pracy składa się z pięciu części. W zeszycie poprzednim zostały opublikowane pierwsze trzy, mianowicie: 1. Analityczne wyznaczanie pola w sztabie prostokątnej z cienką szczeliną poprzeczną, 2. Numeryczne wyznaczanie struktury pola w sztabie prostokątnej przy skazie poprzecznej w położeniu dowolnym oraz 3. Doświadczone odwzorowanie pola w szynie kolejowej S49 w zależności od wielkości i położenia skazy.

nosi σ_{∞} , a nad skazą σ , to stosunek $\varepsilon_{\sigma} = \frac{\sigma - \sigma_{\infty}}{\sigma_{\infty}}$ stanowi stopień niejednostajności pola prądowego i od niego głównie zależy wielkość impulsu czujnika wykrywającego. Uzyskanie jednak dostatecznej czułości wykrywania w metodzie prądowej wymaga dużego prądu w szynie, powyżej 2 kA, co jest zasadniczą wadą defektoskopii prądowej.

W defektoskopii magnetycznej natomiast stosunkowo łatwo uzyskać można wystarczającą czułość wykrywania przy indukcji w szynie rzędu 1 Wb/m². Defektoskopia magnetyczna w odróżnieniu od prądowej musi uwzględnić dodatkowo następujące właściwości pola magnetycznego.

a) Przenikalność stali szynowej w przeciwieństwie do jej przewodności elektrycznej zmienia się silnie w zależności od natężenia pola magnetycznego.

b) Nie można pomijać w rozważaniach przenikalności przestrzeni objętej skazą ani otoczenia szyny.

c) Maksymalna wartość impulsu czujnika głównie cewkowego zależy od jego posuwu, a zatem od prędkości jazdy defektoskopu szynowego.

4.2. Wpływ zmiennej przenikalności magnetycznej na rozkład pola

Zwykła stal szynowa martenowska zawiera ok. 0,6% węgla. Cechuje ją stosunkowo niska przenikalność magnetyczna zarówno początkowa, jak i maksymalna. Dla szyny typu S49 zmierzono podstawową krzywą magnesowania $B = f(H)$ i obliczono względną przenikalność magnetyczną ze wzoru:

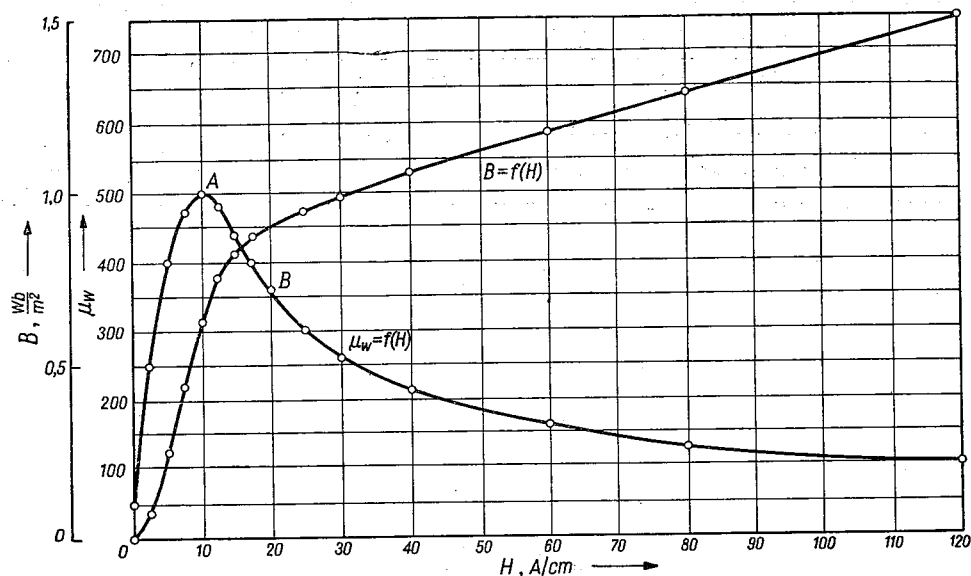
$$\mu_w = \frac{\mu}{\mu_0} = \frac{B \cdot 10^5}{H \cdot 4 \pi}$$

dla B w Wb/m² i H w A/cm.

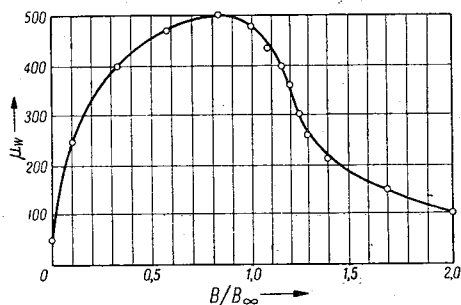
Tablica 1

H A/cm	B Wb/m ²	μ_w	B/B_{∞}
0	0	50	0
2,5	0,078	250	0,104
5,0	0,25	400	0,334
7,5	0,44	470	0,586
10,0	0,625	500	0,835
12,5	0,75	480	1,00
15,0	0,82	435	1,09
17,5	0,87	400	1,16
20,0	0,90	360	1,20
25,0	0,94	300	1,25
30,0	0,98	260	1,30
40,0	1,05	210	1,40
120,0	1,50	100	2,0
187,5	1,875	80	2,5

Z dala od skazy przyjęto uzasadnioną dalej wartość indukcji $B_{\infty} = 0,75$ Wb/m² i podano stosunek B/B_{∞} . Wyniki zestawiono w tablicy 1 i na rysunkach 1 i 2.



Rys. 1. Krzywa magnesowania i przenikalność względna szyny kolejowej



Rys. 2. Przenikalność magnetyczna stali szynowej w zależności od stosunku B/B_{∞} dla $B_{\infty} = 0,75$ Wb/m²

Zmiana przenikalności magnetycznej μ ma identyczny wpływ na przebieg pola jak zmiana przewodności elektrycznej w polu prądowym. Istnieje bowiem całkowita analogia obu tych pól. Wyrażając natężenie pola elektrycznego przez $\bar{K}$, a magnetycznego przez $\bar{H}$ możemy wypisać zasadnicze równania wektorowe dla obu pól:

$$\text{rot } \bar{K} = 0, \quad \text{div } \bar{\sigma} = 0 \quad \bar{\sigma} = \gamma \cdot \bar{K}$$

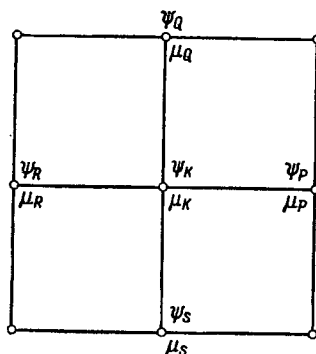
$$\text{rot } \bar{H} = 0, \quad \text{div } \bar{B} = 0 \quad \bar{B} = \mu \cdot \bar{H}.$$

Pole magnetyczne określone funkcją strumienia ψ jest bezźródłowe, ale o zmiennej przenikalności magnetycznej. Zastępując rzeczywisty przekrój szyny zastępczym przekrojem prostokątnym, możemy podać podstawowe równanie różniczkowe cząstkowe dla takiego pola płaskiego:

$$\frac{\partial^2 \psi}{\partial x^2} + \frac{\partial^2 \psi}{\partial y^2} - \frac{1}{\mu} \left[\frac{\partial \mu}{\partial x} \cdot \frac{\partial \psi}{\partial x} + \frac{\partial \mu}{\partial y} \cdot \frac{\partial \psi}{\partial y} \right] = 0,$$

gdzie ψ oznacza funkcję strumienia magnetycznego.

Wzór ten i jego dalsze przekształcenia opierają się na analogii do podobnego wzoru dla pola prądowego o zmiennej przewodności elek-



Rys. 3. Fragment siatki prostokątnej do wyznaczania pola przy uwzględnieniu zmiennej przenikalności magnetycznej

trycznej i stanowią treść poprzedniego artykułu¹⁾. Po wprowadzeniu siatki prostokątnej powyższe równanie różniczkowe daje się przekształcić na równanie różnicowe.

$$p = \frac{\Sigma \psi - 4 \psi_k}{4} - \frac{1}{4} \left[\Delta_x \psi \frac{\Delta_x \mu}{\mu} + \Delta_y \psi \frac{\Delta_y \mu}{\mu} \right] = 0$$

lub w skrócie $p = p_\psi + p_\mu = 0$.

Oznaczenia. Stosownie do rys. 3 przedstawiającego fragment siatki prostokątnej mamy następujące oznaczenia:

$$p_\psi = \frac{\Sigma \psi - 4 \psi_k}{4} = \frac{\psi_P + \psi_Q + \psi_R + \psi_S - 4 \psi_k}{4},$$

$$\Delta_x \psi = \frac{\psi_P - \psi_R}{2}, \quad \Delta_y \psi = \frac{\psi_Q - \psi_S}{2},$$

$$\Delta_x \mu = \frac{\mu_P - \mu_R}{2}, \quad \Delta_y \mu = \frac{\mu_Q - \mu_S}{2}.$$

1) Zob. „Rozprawy Elektrotechniczne” zeszyt 4/1960.

Wartość funkcji strumienia ψ dla poszczególnych węzłów siatki obliczać należy ze wzoru (1). Nasuwa się tu następujący tok obliczenia.

a) Najprzód oblicza się równanie uproszczone $p = p_\psi = 0$. Jest to zasadnicze pole regularne dla określonej wielkości i położenia skazy w szynie o stałej przenikalności. Przyjęto skazę względnej wielkości $h_s/h = 1/8$ umieszczoną w dolnej części szyny o przekroju zastępczym. Dokładny rozptyw pola dla tego wypadku obliczono analitycznie, numerycznie i sprawdzono doświadczalnie w wannie elektrolitycznej. Stanowi to treść rozważań autora nad polem prądowym.

b) Pole zasadnicze przyjmuje się jako pierwsze przybliżenie szukanego pola magnetycznego. We wszystkich węzłach obliczamy teraz przyrosty strumienia w kierunku osi x i y , mianowicie $\Delta_x \psi$ i $\Delta_y \psi$, oraz składowe indukcji magnetycznej $B_x = \frac{\Delta_y \psi}{a}$, $B_y = -\frac{\Delta_x \psi}{a}$. Całkowita indukcja wyniesie: $B = \sqrt{B_x^2 + B_y^2}$. W powyższych obliczeniach przyjęto dowolnie za jednostkę $a = 1$ bok siatki o długości ósmej części wysokości szyny, za całkowity strumień w szynie $\Phi_c = 10\,000$ jednostek, co odpowiada indukcji w szynie z dala od skazy $B_\infty = 1250$ jednostek. Ze stosunku B/B_∞ możemy teraz na podstawie rysunku 2 poszukać odpowiadającej mu wartości przenikalności magnetycznej $\mu = f\left(\frac{B}{B_\infty}\right)$ i wyznaczyć przyrosty tej przenikalności $\Delta_x \mu$ i $\Delta_y \mu$. Zgodnie z wzorem (1) poprawka od zmiany przenikalności magnetycznej wyniesie:

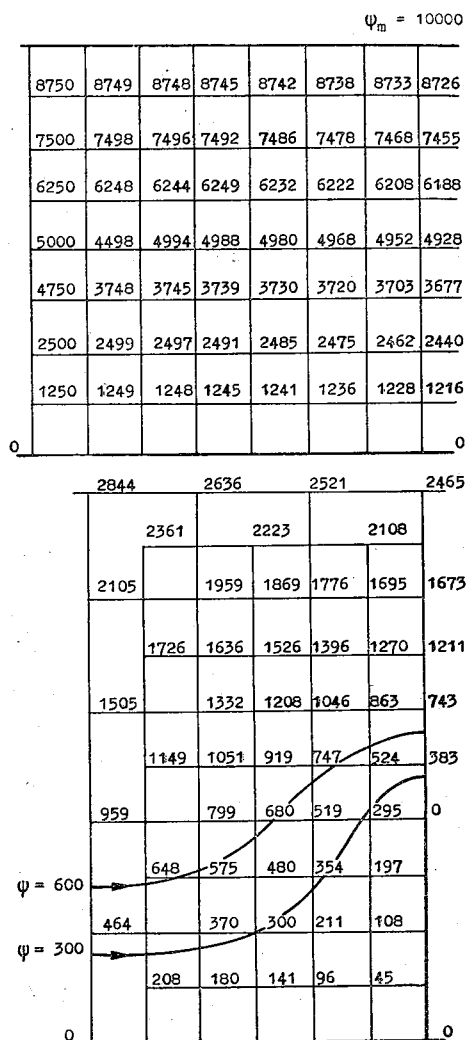
$$p_\mu = -\frac{1}{4} \left[\Delta_x \psi \frac{\Delta_x \mu}{\mu} + \Delta_y \psi \frac{\Delta_y \mu}{\mu} \right].$$

c) Zauważamy w trakcie obliczeń, że zmiana funkcji strumienia ψ_k o jednostkę $+1$ powoduje zmianę poprawki strumienia p_ψ o -1 , gdy tymczasem poprawka p_μ ulega znacznie mniejszym zmianom nawet w pobliżu skazy, gdzie zachodzi zwykle szybka zmiana indukcji i przenikalności magnetycznej. W takim przypadku po obliczeniu poprawek przenikalności najdogodniej wyznaczyć pole dopełniające od tych poprawek i dodać je do pierwotnego pola zasadniczego. Otrzymane w ten sposób pole będzie drugim dokładniejszym przybliżeniem szukanego pola w szynie. Przyjmując do obliczeń przyrostów strumienia i przenikalności nowe wartości strumienia otrzymalibyśmy także nowe wartości poprawek p'_μ . W naszym jednak przypadku drugie przybliżenie stanowi dostatecznie dokładny obraz pola magnetycznego w szynie.

Rysunki 4 i 5 przedstawiają obliczony w podany sposób obraz pola magnetycznego w szynie w bliższym i dalszym sąsiedztwie skazy. Jako

pierwsze przybliżenie przyjęto tu pole o stałej przenikalności z uwzględnieniem poprawek od niedostatecznego zagęszczenia siatki, a następnie dodano pole dopełniające od poprawek p_μ . Poprawka wypadkowa tylko w kilku punktach przekracza nieco wartość 1, a suma poprawek wypadkowych w całym obszarze jest praktycznie równa zeru.

Interesuje nas głównie wartość strumienia ψ_K w bezpośrednim sąsiedztwie górnej krawędzi szyny, gdyż na tej podstawie będziemy mogli wyznaczyć składową poziomą indukcji B_x przy powierzchni oraz stopień niejednostajności pola dla danej skazy. Zakładając na powierzchni



Rys. 4. Fragment pola w bezpośrednim otoczeniu skazy oraz w pobliżu umownego brzegu przy uwzględnieniu zmiennej przenikalności magnetycznej

$\psi_m = 10000$ otrzymamy:

$$B_x = \frac{\psi_m - \psi_K}{a}, \quad B_{\infty} = \frac{\psi_m - \psi_{K\infty}}{a}.$$

Stąd stopień niejednostajności pola dla górnej powierzchni:

$$\varepsilon_B = \frac{B_x - B_{\infty}}{B_{\infty}} = \frac{\psi_{K\infty} - \psi_K}{\psi_m - \psi_{K\infty}}. \quad (2)$$

Dla dolnej powierzchni przy założeniu $\psi_S = 0$ otrzymamy:

$$\varepsilon_B = \frac{\psi_K - \psi_{K\infty}}{\psi_{K\infty}}. \quad (3)$$

W obu przypadkach zmniejszenie funkcji strumienia w sąsiedztwie powierzchni wywołuje zwiększenie bezwzględnej wartości stopnia niejedno-

$\psi_m = 10000$

			9352	9349	9346	9345	9349	9343
8726	8718	8708		8697	8693	8689	8687	8685
			8054	8046	8038	8032	8028	8027
7455	7438	7417		7394	7383	7374	7368	7366
			6755	6740	6725	6712	6704	6701
6188	6162	6126		6085	6065	6047	6034	6030
			5447	5431	5401	5376	5358	5351
4928	4892	4842		4772	4735	4699	4671	4661
			4164	4118	4068	4016	3972	3956
3677	3637	3572		3468	3402	3327	3258	3228
			2898	2831	2844	2636	2521	2465
2440	2303	2337		2210	2105	1959	1776	1673
			1689	1615	1505			
1216	1194	1150		1053	959			
			551	517	464			

$\psi = 0$

Rys. 5. Fragment pola w dalszym sąsiedztwie skazy z uwzględnieniem zmiennej przenikalności magnetycznej

stajności pola, który dla skazy w położeniu dolnym jest dodatni na górnej powierzchni, a ujemny na dolnej.

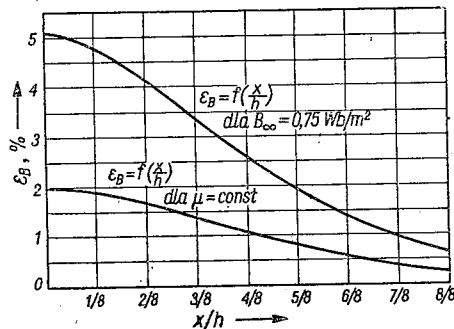
Takie zmniejszenie funkcji ψ_K może nastąpić wtedy, gdy pojawi się przewaga ujemnych poprawek od zmiany przenikalności magnetycznej. Znak i wartość tych poprawek zależy od położenia i wielkości skazy, przede wszystkim jednak od doboru indukcji B_∞ w szynie. We wzorze (1) dla skazy jak na rys. 4 przyrosty strumienia $\Delta_x\psi$, $\Delta_y\psi$ są dodatnie. Z wyjątkiem najbliższego otoczenia skazy idąc w kierunku osi x i y napotykamy rozrzedzenie linii indukcji, jak to widać na rysunku 4. Takiemu rozrzedzeniu powinny towarzyszyć dodatnie przyrosty przenikalności magnetycznej $\Delta_x\mu$, $\Delta_y\mu$. Wówczas bowiem poprawka p_μ będzie ujemna. Przypadek taki zajdzie wtedy, jeśli indukcja B_∞ będzie nieco większa od tej, której odpowiada największa wartość przenikalności magnetycznej, czyli na podstawie rysunku 1 powinno być B_∞ większe od $0,625 \text{ Wb/m}^2$. Należy przypuszczać, że najkorzystniejsze warunki osiągniemy przyj-

Tablica 2

Stopień niejednostajności pola w %

x/h	0	1/8	2/8	3/8	4/8	5/8	6/8	7/8	8/8
Powierzchnia górna									
$\mu = \text{const}$	1,96	1,88	1,68	1,39	1,10	0,83	0,61	0,43	0,30
$\mu = f(B)$	5,1	4,8	4,2	3,4	2,6	1,9	1,36	0,96	0,64
Powierzchnia dolna									
$\mu = \text{const}$	-100	-28,5	-10,0	-4,5	-2,5	-1,5	-1,0	-0,6	-0,4
$\mu = f(B)$	-100	-40,8	-17,3	-8,0	-4,5	-2,7	-1,8	-1,1	-0,7

mując punkt pracy pomiędzy ekstremalnym punktem A i punktem przecięcia B. Na tej podstawie obraliśmy punkt pracy dla $B_\infty = 0,75 \text{ Wb/m}^2$ i $\mu_w = 480$. Ostateczny wybór punktu pracy można ustalić jedynie na drodze doświadczalnej.

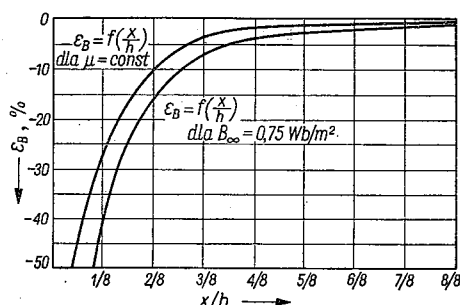


Rys. 6. Stopień niejednostajności pola na górnej powierzchni sztaby dla zmiennej i stałej przenikalności magnetycznej

Wzory (2) i (3) pozwolą nam wyznaczyć stopień niejednostajności pola przy założeniu stałej i zmiennej przenikalności magnetycznej, a następnie porównać otrzymane wyniki.

Stopień niejednostajności pola na górnej i dolnej powierzchni szyny w zależności od oddalenia w kierunku poziomym $\varepsilon_B = f\left(\frac{x}{h}\right)$ dla stałej i zmiennej przenikalności magnetycznej przedstawiono w tablicy 2 oraz na rysunkach 6 i 7.

Okazuje się, że na górnej powierzchni stopień niejednostajności pola wzrósł około 2,5 razy, a na dolnej około 1,5 razy w porównaniu do przypadku dla $\mu = \text{const}$. Prócz tego nastąpiło pewne nieznaczne zmniejszenie



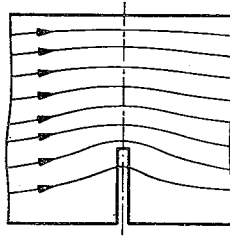
Rys. 7. Stopień niejednostajności pola na dolnej powierzchni sztaby dla zmiennej i stałej przenikalności magnetycznej

średniej rozpiętości skazy. Należy podkreślić, że dobór niewłaściwego punktu pracy, np. na wznoszącej się części charakterystyki $\mu = f(B)$, spowodowałby przewagę poprawek dodatnich i w rezultacie zmniejszenie się stopnia niejednostajności pola w porównaniu z przypadkiem dla $\mu = \text{const}$. Nie można też znaleźć sposobu, który by doświadczalnie, np. drogą odwzorowania jak dla pola prądowego, pozwolił sprawdzić dodatni efekt zmiennej przenikalności magnetycznej otrzymany drogą bardzo uciążliwych obliczeń. W praktyce ten dodatni efekt może być przytłumiony przy uwzględnianiu przenikalności cienkiej szczeliny powietrznej stanowiącej skazę.

4.3. Wpływ przenikalności magnetycznej przestrzeni objętej skazą

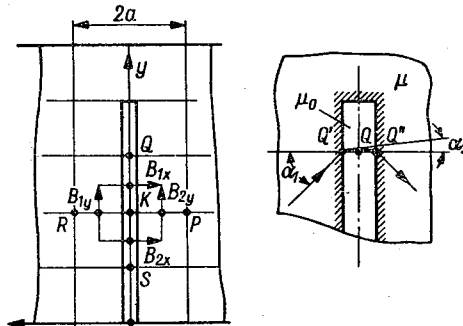
Pęknięcia poprzeczne jako skazy zmęczeniowe posiadają grubość zaledwie kilku mikronów. Grupują się one jednak zwykle w pobliżu drobnych nieciągłości, np. wypełnień zużłowych, i dlatego jako skazę zastępczą przyjmujemy początkowo skazę grubości około 0,2 mm o przenikalności względnej $\mu_w = 1$.

Linie indukcji, jak pokazano na rysunku 8, przenikają częściowo skażę, skutkiem czego pole nad nią bardziej zbliża się do równomiernego, co pociąga niekorzystne zmniejszenie stopnia niejednostajności pola przy



Rys. 8. Przebieg linii strumienia z uwzględnieniem przenikalności magnetycznej szczeliny

powierzchni szyny. Zachodzi obawa, że przenikalność stosunkowo cienkiej skaży w poważnym stopniu obniża wartość metody magnetycznej w porównaniu do prądowej, tym bardziej, że doświadczalne sprawdzenie tego zjawiska utrudnia możliwość wykonania tak cienkiej sztucznej szczeliny.



Rys. 9. Fragment siatki w otoczeniu szczeliny powietrznej

Numeryczna metoda wyznaczania pola pozwoli nam rozwiązać powyższe zagadnienie. Na granicy dwóch ośrodków o różnych przenikalnościach magnetycznych następuje załamanie linii indukcji, zachowują one jednak ciągłość. Dla dużej przenikalności względnej linie indukcji przebiegają w szczelinie powietrznej prawie poziomo stosownie do znanego wzoru:

$\mu_w = \frac{\operatorname{tg} \alpha_1}{\operatorname{tg} \alpha_2}$. Wówczas, jak to pokazano na rysunku 9, zwłaszcza dla cienkiej szczeliny powietrznej, wartość funkcji strumienia będzie jednakowa w poziomym przekroju tej szczeliny, czyli $\psi_Q = \psi_{Q'} = \psi_{Q''}$. Będzie więc:

$$B_{1x} = \frac{\psi_Q - \psi_K}{a}, \quad B_{1y} = \frac{\psi_R - \psi_K}{a},$$

$$B_{2x} = \frac{\psi_K - \psi_S}{a}, \quad B_{2y} = \frac{\psi_K - \psi_P}{a},$$

gdzie a oznacza bok siatki.

Z warunku bezwirowości: $\oint \vec{H} \cdot d\vec{l} = 0$ otrzymujemy na obwodzie kwadratu o boku a następujące równanie:

$$\frac{(B_{1x} - B_{2x})(a - \delta)}{\mu} + \frac{(B_{1y} - B_{2y})a}{\mu} + \frac{(B_{1x} - B_{2x})\delta}{\mu_0} = 0.$$

Po podstawieniach i pominięciu $\frac{\delta}{a}$ wobec 1 dla cienkiej szczeliny równanie powyższe da się sprowadzić do postaci:

$$\frac{\Sigma \psi - 4 \psi_K}{4} + \frac{\psi_Q + \psi_S - 2 \psi_K}{4} \mu_w \frac{\delta}{a} = 0 \quad (4)$$

lub w skrócie:

$$p = p_\psi + p_\delta = 0. \quad (4a)$$

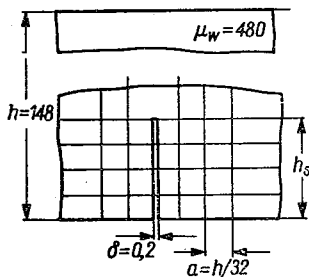
Nieco inaczej wypadnie dla węzła na górnym skraju szczeliny, mianowicie:

$$\frac{\Sigma \psi - 4 \psi_K}{4} + \frac{\psi_S - \psi_K}{4} \mu_w \frac{\delta}{a} = 0 \quad (5)$$

lub w skrócie:

$$p = p_\psi + p'_\delta = 0. \quad (5a)$$

Stosownie do rysunku 10 przyjmujemy teraz wartości: względną przenikalność magnetyczną stali szynowej $\mu_w = 480$, grubość szczeliny po-



Rys. 10. Położenie szczeliny powietrznej w sztabie ze stali szynowej

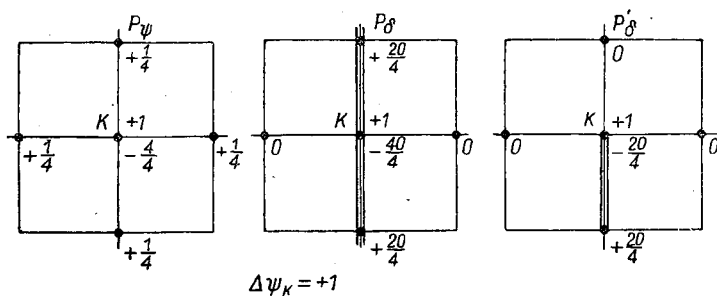
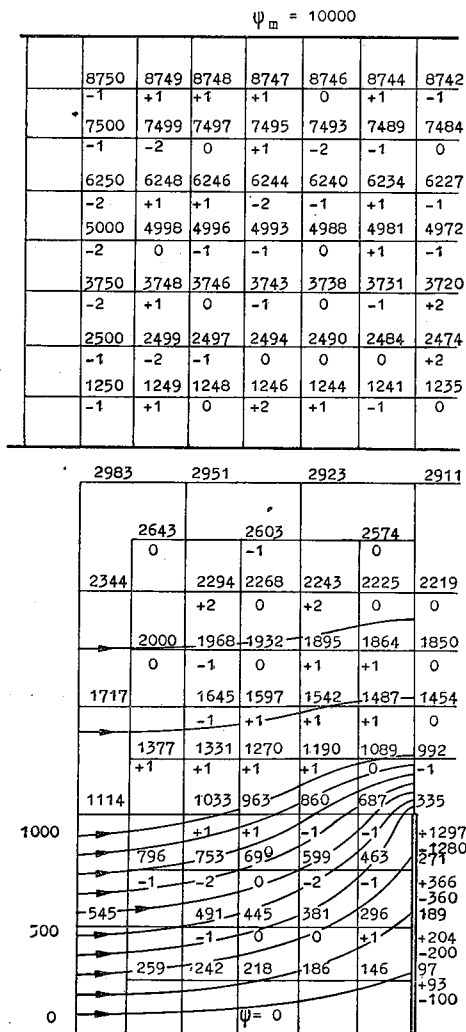
wietrznej — $\delta = 0,2$ mm, bok siatki $a = h/32 = 148/32 = 6,64$ mm.

Dla powyższych danych wyrażenie: $\mu_w \frac{\delta}{a} = 480 \frac{0,2}{6,64} \approx 20$.

Poprawki szczelinowe wyniosą odpowiednio:

$$p_\delta = \frac{\psi_Q + \psi_S - 2 \psi_K}{4} 20, \quad p'_\delta = \frac{\psi_S - \psi_K}{4} 20.$$

Na rysunku 11 pokazano, w jakim stopniu wzrośnie poprawka p_ψ i poprawki szczelinowe w otaczających węzłach dla przyrostu strumienia $\Delta \psi_K = +1$. Poprawki szczelinowe w znacznie większym stopniu zmieniają się tu przy wzroście strumienia niż poprawka p_ψ . One to głównie decydują o rozkładzie pola w pobliżu skazy.

Rys. 11. Wpływ przyrostu strumienia $\Delta\psi_K = +1$ na przyrost poprawek

Rys. 12. Fragment pola w bezpośrednim otoczeniu skazy oraz w pobliżu umownego brzegu przy uwzględnieniu przenikalności szczeliny powietrznej. Przy poprawkach opuszczono mianownik równy 4

W celu wyznaczenia tego rozkładu najkorzystniej jest przyjmować pewne pole wyjściowe, np. jednostajne, i obliczać oraz rozdzielać poprawki według schematu na rysunku 11 tak, aby spełniały się dla szcze-

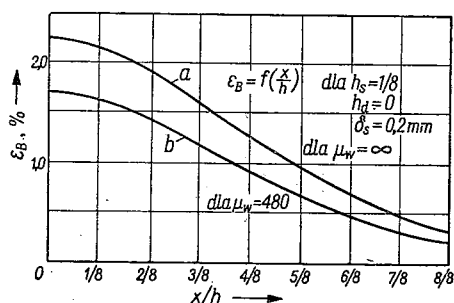
$\Psi_{II} = 10000$

		9367 -1	9366 0	9365 +2	9365 0	9365 -1	9365 -1
8742	8739	8735	8732	8731	8730	8729	8729
-1	-1	+2	0	-1	-1	+1	0
		8100 +2	8098 0	8096 0	8094 +1	8093 0	8093 -2
7484	7478	7471	7464	7461	7458	7456	7455
0	-1	-1	-1	-1	0	+1	+1
		6834 -1	6829 +1	6825 -1	6821 0	6818 +1	6817 +1
6227	6217	6206	6194	6188	6183	6179	6178
-1	+2	0	-1	0	-1	+1	-1
		5566 -1	5558 -1	5550 -1	5543 -1	5538 -1	5536 0
4972	4959	4942	4921	4910	4900	4893	4890
-1	0	0	-1	0	+1	0	+2
		4298 0	4284 -1	4269 -1	4255 -2	4244 +1	4240 +1
3720	3705	3682	3647	3626	3605	3589	3583
+2	0	-2	-2	0	+1	0	-2
		3036 0	3012 +1	2983 +2	2951 +2	2924 +2	2912 +2
2474	2459	2432	2383	2344	2294	2243	2219
+2	0	+1	-1	+1			
		1795 +2	1764 +1	1717 0	1645		
1235	1225	1205	1162	1114	1033		
0	-1	-1	0	+1			
		591 +2	574 +2	545 -1	491		

δ

$\Psi = 0$

Rys. 13. Fragment pola w dalszym sąsiedztwie skazy z uwzględnieniem przenikalności szczeliny powietrznej. Przy poprawkach opuszczono mianownik równy 4



Rys. 14. Stopień niejednostajności pola na górnej powierzchni sztaby bez uwzględnienia (a) i z uwzględnieniem (b) przenikalności szczeliny powietrznej

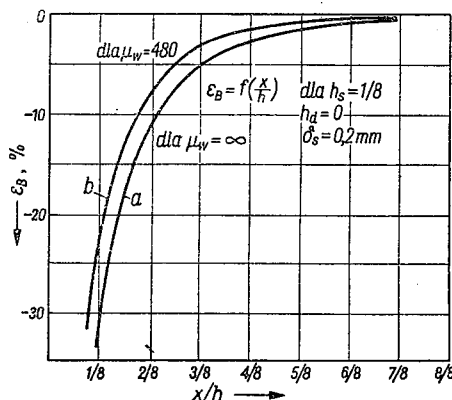
liny równania (4) i (5). Poza szczeliną przyjmujemy pole regularne o stałej przenikalności:

$$p = p_{\psi} = 0.$$

Obliczony w ten sposób przebieg funkcji strumienia w sąsiedztwie skazy podaje rysunek 12, a dla dalszego otoczenia skazy rysunek 13. Obok węzłów w szczelinie wpisano wartości poprawek szczelinowych. Poprawki wypadkowe, z wyjątkiem szczeliny, nie przekraczają nigdzie wartości $4/4$.

Stopień niejednostajności pola dla górnej i dolnej powierzchni sztaby obliczono ze wzorów (2) i (3) oraz pokazano na rysunkach 14 i 15.

Na skutek przenikalności magnetycznej szczeliny powietrznej stopień niejednostajności pola ε_{δ} zmalał w naszym przypadku do wartości $0,7 \varepsilon$ na górnej oraz do $0,75 \varepsilon$ na dolnej powierzchni sztaby, gdzie ε wyraża



Rys. 15. Stopień niejednostajności pola na dolnej powierzchni sztaby bez uwzględnienia (a) i z uwzględnieniem (b) przenikalności szczeliny powietrznej

stopień niejednostajności w polu prądowym względnie magnetycznym, ale bez uwzględnienia przenikalności szczeliny powietrznej. Ogólnie możemy napisać: $\varepsilon_{\delta} = k \cdot \varepsilon$.

Okazuje się, że dla stałej wielkości stosunkowo cienkiej szczeliny powietrznej mnożnik szczelinowy k rośnie praktycznie w jednakowym stopniu zarówno przy wzroście względnej przenikalności stali μ_w , jak i względnej szerokości szczeliny $\frac{\delta}{h}$, czyli:

$$k = f\left(\mu_w \frac{\delta}{h}\right).$$

W naszym przypadku mieliśmy: $\mu_w \frac{\delta}{h} = 480 \frac{0,2}{148} = 0,65$ oraz $k = 0,75$.

Po obliczeniu dwóch dalszych punktów oraz wartości granicznych dla $k = 0$ i $k = \infty$ otrzymujemy wyniki ujęte w tablicy 3.

Tablica 3

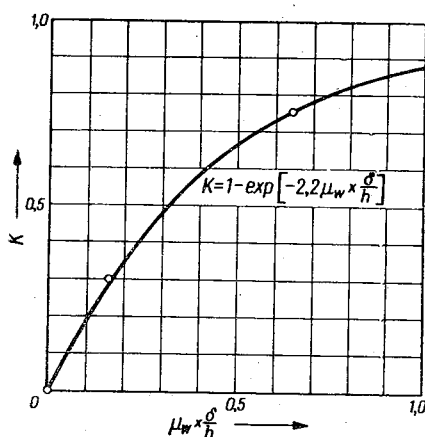
$\mu_w \frac{\delta}{h}$	0	0,16	0,32	0,65	∞
Mnożnik szczelinowy k	0	0,30	0,51	0,75	1,0

Wartości k otrzymane w tablicy 3 dają się dosyć dobrze aproksymować wzorem empirycznym:

$$k = 1 - \exp \left[-2,22 \mu_w \frac{\delta}{h} \right]. \quad (6)$$

Widać to na rysunku 16, gdzie obok punktów z tablicy 3 narysowano przebieg funkcji (6). Zmniejszenie stopnia niejednostajności pola do połowy wg wzoru (6) otrzymuje się dla $\mu_w \frac{\delta}{h} = 0,31$. Odpowiada temu przy wysokości sztaby $h = 148$ mm $\mu_w = 480$, $\delta = 0,1$ mm lub $\mu_w = 240$, $\delta = 0,2$ mm.

We wzorze $\varepsilon_\delta = k \cdot \varepsilon$ podstawmy zależność (6) oraz jak w polu prądowym $\varepsilon = C_s \left(\frac{h_s}{h} \right)^2$, gdzie współczynnik oddziaływania skazy C_s zależy od jej położenia i daje się wyznaczyć pomiarowo w wannie elektrolitycz-



Rys. 16. Mnożnik szczelinowy w zależności od względnej przenikalności i grubości szczeliny. $k = f \left(\mu_w \frac{\delta}{h} \right)$

nej. Stopień niejednostajności pola z uwzględnieniem przenikalności szczeliny powietrznej wyrazi się wówczas jako:

$$\varepsilon_\delta = C_s \left(\frac{h_s}{h} \right)^2 \left[1 - \exp \left(-2,2 \mu_w \frac{\delta}{h} \right) \right]. \quad (6a)$$

Współczynnik 2,2 ważny jest w powyższym wzorze empirycznym dla skazy wielkości względnej $h_s/h = 1/8$, położonej u dołu sztaby. Okazuje się, że współczynnik ten w małym jedynie stopniu zależy od położenia skazy, a w dosyć znacznym stopniu rośnie przy jej zmniejszeniu. W tych samych przeto warunkach stopień niejednostajności pola dla skazy mniejszej zmaleje w mniejszym stopniu niż dla skazy dużej.

Z dotychczasowych rozważań wynika następujący wniosek. W polu magnetycznym z jednej strony zmienność przenikalności magnetycznej szyny μ_w przy właściwie dobranej indukcji B_∞ zwiększa stopień niejednostajności pola przy powierzchni, z drugiej jednak strony stopień ten maleje wskutek przenikalności magnetycznej skazy w postaci cienkiej szczeliny powietrznej. Wpływy te częściowo się znoszą. W rezultacie dla skazy o grubości powyżej 0,1 mm zmniejszenie stopnia niejednostajności pola przy właściwie dobranej indukcji jest nieznaczne, a wyniki otrzymane dla pola prądowego w wannie elektrolitycznej są w przybliżeniu ważne także i dla pola magnetycznego. Dopiero przy większych szczelinach, dla h_s/h powyżej $1/8$, a także dla cieńszych od 0,1 mm, pole magnetyczne daje wyraźnie gorsze wyniki od pola prądowego. Przypadek taki nie ma jednak praktycznego znaczenia wobec łatwości wykrywania skaz dużych i rzadkości skaz cieńszych od 0,1 mm bez dodatkowych jam usadowych czy też wtrąceń żuźlowych.

4.4. Wpływ przenikalności magnetycznej przestrzeni otaczającej szynę

Dotychczas we wszystkich rozważaniach zakładaliśmy, że strumień magnetyczny nie wychodzi poza powierzchnię szyny, czyli wartość funkcji strumienia dla powierzchni szyny jest stała, np. $\psi_m = 10\,000$ jednostek. W rzeczywistości strumień wychodzi z szyny na powierzchnię i stwarza pole zewnętrzne, które oddziałuje na czujnik do wykrywania skaz.

Założenie $\psi_m = 10\,000 = \text{const}$ spełnia się wyłącznie w miejscu z dala od skazy. Natomiast w sąsiedztwie skazy stosownie do rysunku 17 funkcja strumienia spełniać musi pewne równanie różnicowe wynikające z warunku bezwirowości pola, w szczególności:

$$\frac{1}{\mu_0} \left[B_1 \frac{a}{2} + B_2 a + B_3 \frac{a}{2} + \frac{B_3}{\mu_w} \cdot \frac{a}{2} + \frac{B_4}{\mu_w} a + \frac{B_1}{\mu_w} \cdot \frac{a}{2} \right] = 0.$$

Podstawiając:

$$B_1 = \frac{\psi_K - \psi_P}{a}, \quad B_2 = \frac{\psi_K - \psi_Q}{a},$$

$$B_3 = \frac{\psi_K - \psi_R}{a}, \quad B_4 = \frac{\psi_K - \psi_S}{a}$$

otrzymamy po drobnych przekształceniach:

$$\frac{\psi_P + \psi_R + 2\psi_Q - 4\psi_K}{4} + \frac{1}{\mu_w} \frac{\psi_P + \psi_R + 2\psi_S - 4\psi_K}{4} = 0 \quad (7)$$

lub w skrócie

$$p = p_p + p_{Fe} = 0. \quad (7a)$$

Z równania (7) możemy obliczyć wartość funkcji strumienia na powierzchni szyny ψ_K przyjmując początkowy rozkład pola tak w szynie, jak i na zewnątrz, a następnie pola te skorygować. Jako pole początkowe w szynie najpraktyczniej przyjąć pole zniekształcone obecnością skazy, ale dla założenia, że strumień nie wychodzi na zewnątrz. Za pole początkowe zewnętrzne najdogodniej uważać pole jednostajne o takim samym natężeniu jak w szynie, ale w miejscu dalekim od skazy. Dla początkowego rozkładu pola przyjęto więc następujące założenia:

$$\psi_P = \psi_R = \psi_K, \quad (8)$$

$$\frac{\psi_Q - \psi_K}{a} = \frac{B_{x\infty}}{\mu_w}, \quad (9)$$

$$\psi_K - \psi_S = a B_{\infty} [1 + \varepsilon_B]. \quad (10)$$

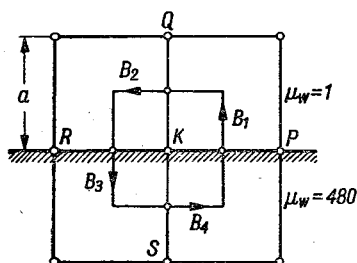
Po podstawieniu zależności (8) do wzoru (7) będzie:

$$p = \frac{2\psi_Q - 2\psi_K}{4} + \frac{1}{\mu_w} \cdot \frac{2\psi_S - 2\psi_K}{4}.$$

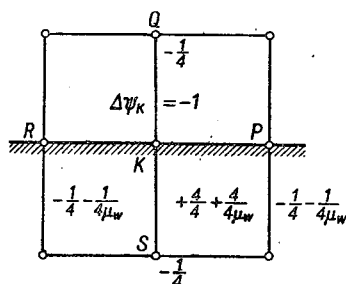
Podstawiając zaś dalej zależność (9) i (10) otrzymamy ostatecznie:

$$p = -\frac{B_{\infty}}{2\mu_w} a \cdot \varepsilon_B. \quad (11)$$

Na powierzchni szyny wystąpią zatem poprawki, które należy odpowiednio rozdzielić, stosownie do rysunku 18. Dla $\Delta\psi_K = -1$ przyro-



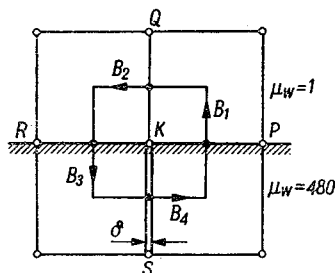
Rys. 17. Fragment siatki pola magnetycznego przy powierzchni szyny w zasięgu skazy



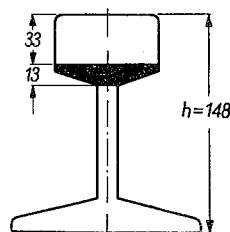
Rys. 18. Wpływ zmiany strumienia $\Delta\psi_K = -1$ na przyrost poprawek przy powierzchni szyny

sty poprawek na powierzchni szyny określa wzór (7), a przyrosty w punktach Q i S są takie jak w polu regularnym. Widać od razu, że dla dużej przenikalności względnej, jak np. dla $\mu_w = 480$, rozdział poprawek nie różni się prawie od rozdziału w polu regularnym. Małe bowiem poprawki

dotatkowe na powierzchni szyny nie wpływają praktycznie na pole wewnątrz szyny. Pole dodatkowe od poprawek powierzchniowych stanowi więc pole dopełniające dla pola stałego, które przyjęliśmy za pole początkowe. Pole zewnętrzne jest więc superpozycją dwóch pól stałego i dopełniającego od poprawek powierzchniowych. Jeśli pole stałe pochodzące od ustroju magnesującego nie będzie się zmieniać w czasie wykrywania skaz, to efekt wykrywania zależeć będzie jedynie od przebiegu zewnętrznego pola dopełniającego, które jest funkcją poprawek węzłowych związanych z rozkładem stopnia niejednostajności pola na powierzchni szyny.



Rys. 19. Fragment siatki dla skazy przy górnej powierzchni szyny



Rys. 20. Szyna typu S49. Skaza wzorcowa w główce

W szczególnym przypadku skaza może graniczyć z powierzchnią szyny, jak to przedstawia rysunek 19. Całka okrężna $\oint \bar{H} d\bar{l}$ powiększy się wówczas o człon $\frac{B_4 \cdot \delta}{\mu_o}$, a równanie różnicowe dla węzła K wyniesie:

$$\frac{\psi_P + \psi_R + 2\psi_Q - 4\psi_K}{4} + \frac{1}{\mu_w} \frac{\psi_P + \psi_R + 2\psi_S - 4\psi_K}{4} + \frac{\delta}{a} \cdot \frac{\psi_S - \psi_K}{4} = 0 \quad (12)$$

lub w skrócie:

$$p_p + p_{Fe} + p_\delta'' = 0. \quad (12a)$$

Możemy teraz zająć się wyznaczeniem zewnętrznego pola dopełniającego od skazy wzorcowej położonej w dolnej części główki szyny. Wymiary i rozmieszczenie skazy wzorcowej podaje rysunek 20. Skaza ta o grubości powyżej 0,1 mm, zajmująca około 10% poprzecznego przekroju szyny oznaczana jest w skrócie symbolem „A”. Ponieważ dla prawidłowo dobranej indukcji magnetycznej w szynie dodatni wpływ zmiennej przenikalności magnetycznej kompensuje ujemny wpływ przenikalności magnetycznej szczeliny powietrznej o grubości ponad 0,1 mm, więc możemy przyjąć, że stopień niejednostajności pola jest dla skazy „A” prawie taki sam jak w polu prądowym.

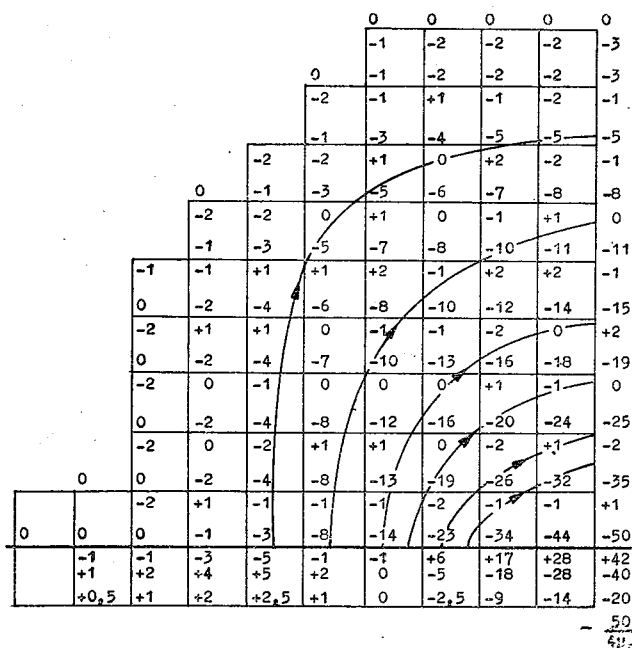
Odwzorowanie skazy wzorcowej „A” w wannie elektrolitycznej pozwala wyznaczyć stopień niejednostajności pola na górnej powierzchni

szyny dla różnych odległości od płaszczyzny skazy. Zależność tę w postaci $\varepsilon_B = f\left(\frac{x}{h}\right)$ wraz z zaokrąglonymi wartościami poprawek powierzchniowych zawiera tabela 4.

Tabela 4

x/h	0	1/16	2/16	3/16	4/16	5/16	6/16	7/16	8/16
$\varepsilon_B = 10^{-2} \times$	3,24	2,22	1,42	0,40	0	-0,2	-0,4	-0,3	-0,2
$p = \frac{1}{480} \times$	$\frac{-40}{4}$	$\frac{-28}{4}$	$\frac{-18}{4}$	$\frac{-5}{4}$	0	$\frac{+2}{4}$	$\frac{+5}{4}$	$\frac{+4}{4}$	$\frac{+2}{4}$

W obliczeniach dogodnie jest rozważać jedynie pole dopełniające od poprawek powierzchniowych. Pole to nakłada się na pole stałe od ustroju magnesującego na zewnątrz szyny w siatce podwójnie zagęszczonej dla



Rys. 21. Przebieg pola dopełniającego na zewnątrz szyny dla skazy wzorcowej w dolnej części główki (skaza „A”). Opuszczono mianowniki nad szyną dla wartości strumienia: 480, dla poprawek: $4 \cdot 480$, w szynie dla poprawek 4

$a = 0,5$, $\mu_w = 480$, $B_\infty = 1250$. Poprawka powierzchniowa stosownie do wzoru (11) wyniesie:

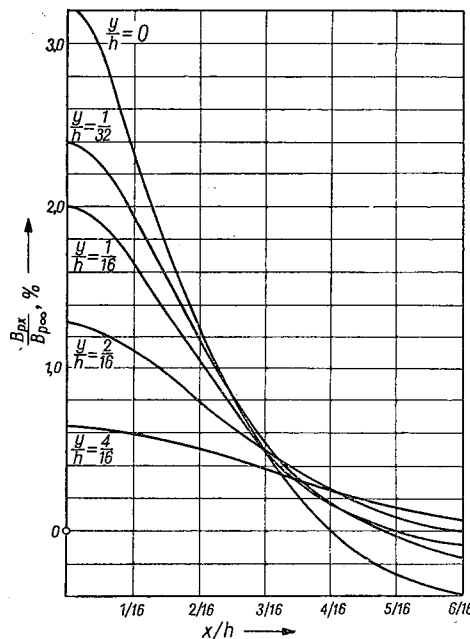
$$p = -\frac{1250}{2 \cdot 480} 0,5 \cdot \varepsilon_B = \frac{312,5}{480} \varepsilon_B.$$

Przebieg pola dopełniającego wewnątrz i na zewnątrz szyny pokazany jest na rysunku 21. Pole dopełniające wewnątrz szyny zostało określone

na podstawie odwzorowania elektrolitycznego. Pole regularne, dopełniające na zewnątrz szyny wynika z rozdziału poprawek powierzchniowych. Z powodu niewygody pisania zbyt małych, ułamkowych liczb opuszczono mianownik $\mu_w = 480$ dla wartości strumienia oraz $4 \cdot \mu_w$ dla poprawek na zewnątrz szyny.

Funkcja strumienia na powierzchni spełnia równanie różnicowe (7). Wartości strumienia w szynie wymagają wprowadzenia poprawek, które jednak w naszym przypadku są pomijalnie małe. Największa poprawka w szynie dla węzła mającego $\psi = -20$ wyniesie bowiem $-\frac{50}{4 \cdot 480}$. Pole zewnętrzne nie wpływa przeto praktycznie na pole w szynie, a linie indukcji są prostopadłe do powierzchni.

Funkcja strumienia dopełniającego na zewnątrz szyny pozwala obliczyć składową poziomą i pionową zewnętrznego pola dopełniającego dla skazy wzorcowej „A”. Wprowadzając oznaczenia $\Delta_y \psi = \psi_Q - \psi_K$ i $\Delta_x \psi = \psi_p - \psi_K$ otrzymamy:



Rys. 22. Stosunkowa wartość składowej poziomej dopełniającego pola zewnętrznego

$$\frac{B_{px}}{B_{p\infty}} = f\left(\frac{x}{h}\right)$$

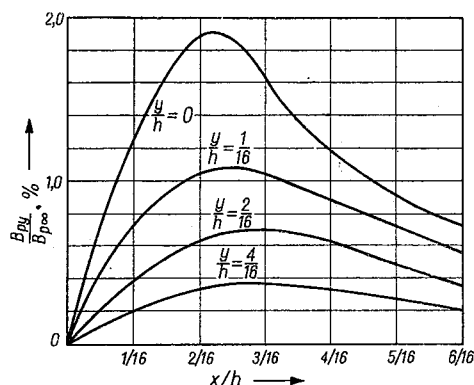
$$B_{px} = \frac{\Delta_y \psi}{a} \quad \text{oraz} \quad B_{py} = \frac{\Delta_x \psi}{a}.$$

W odniesieniu do indukcji przy powierzchni szyny poza zasięgiem

skazy równej $B_{p\infty} = \frac{B_\infty}{\mu_w} = \frac{1250}{480}$ mamy dla siatki o boku $a = 0,5$:

$$B_{px}/B_{p\infty} = \frac{\Delta_y \psi \cdot 480}{625} \quad \text{oraz} \quad B_{py}/B_{p\infty} = \frac{\Delta_x \psi \cdot 480}{625}. \quad (13)$$

Obliczone z wzoru (13) składowe pola dopełniającego przedstawiają rysunki 22 i 23. Stosunkowa wartość składowej poziomej tuż przy po-



Rys. 23. Stosunkowa wartość składowej pionowej dopełniającego pola zewnętrznego

$$\frac{B_{py}}{B_{p\infty}} = f\left(\frac{x}{h}\right)$$

wierzchni szyny wynosi tyle co stopień zniekształcenia pola. Na wysokości od powierzchni równej odległości od górnej krawędzi skazy, to jest zgodnie z rys. 20 dla $y/h = 33/148 = 0,22$, obie składowe maleją w przybliżeniu 4-krotnie w stosunku do swych wartości przy powierzchni szyny.

W odległości poziomej od skazy składowa pozioma osiąga zero tym bliżej, im bliższa jest ona powierzchni szyny, a następnie przybiera wartości ujemne. W miejscu, gdzie składowa pozioma maleje do połowy, to jest dla średniej rozpiętości skazy, mamy w przybliżeniu punkt przegięcia tej krzywej, a składowa pionowa pola dopełniającego posiada tam maksymalną wartość.

Powyższy rozkład pola zewnętrznego, dopełniającego otrzymaliśmy zakładając, że jest ono polem płaskim. W rzeczywistości pole to poza najbliższym otoczeniem powierzchni szyny zatracą stopniowo charakter płaski przyjmując symetrię osiową szyny. W związku z tym należy przypuszczać, że składowe zewnętrznego pola dopełniającego będą nieco silniej malały przy oddaleniu od szyny, niż to wynika z rysunków 22 i 23.

4.5. Uwagi końcowe

Rozpatrywane osobliwości pola magnetycznego dotyczące cienkich skaz poprzecznych występują także przy skazach objętościowych. W szczególności ważny jest tu dobór właściwej indukcji w szynie.

Elementarna rurka o przekroju s i długości l znajdująca się w szynie doznaje korzystnego dla defektoskopii przyrostu oporu magnetycznego przy pojawieniu się skazy, jeśli będzie spełniony warunek:

$$\frac{l}{s} \left(\frac{1}{\mu} - \frac{1}{\mu_{\infty}} \right) > 0. \quad (14)$$

Nastąpi to wtedy, jeśli indukcja B_{∞} w zdrowej szynie będzie równa indukcji B_A , której odpowiada maksymalna wartość przenikalności magnetycznej. Przy pojawieniu się skazy zarówno w obszarze zagęszczenia, jak i rozrzedzenia pola będzie wtedy zawsze $\mu < \mu_A$ i spełni się warunek (14).

Okazuje się jednak, że wobec przewagi obszaru zagęszczenia największy stopień niejednostajności pola osiąga się przy $B_{\infty} > B_A$ i to w większym jeszcze stopniu dla skaz owalnych niż cienkich poprzecznych.

Niekorzystny wpływ przenikalności przestrzeni objętej skazą jest dla skaz objętościowych bez znaczenia. W porównaniu z cienką skazą poprzeczną o tej samej powierzchni rzutu na płaszczyznę pionową skaza objętościowa posiada zwykle nieco większy stopień niejednostajności pola przy powierzchni, większy zasięg, ale mniejszą stromość zewnętrznych składowych pola dopełniającego. Wobec tego czujnik cewkowy jest na ogół czulszy dla cienkich skaz poprzecznych niż dla skaz objętościowych, gdyż reaguje na stromość pola.

WYKAZ LITERATURY

1. Bozort R. M.: Ferromagnetism, New York 1951.
2. Markuszewicz M., Mierzyjewski A.: Materiały magnetyczne. W.G.H., Katowice 1954.
3. Hugh H. Skilling: Fale elektromagnetyczne. P.W.N., Warszawa 1954.
4. Kifer I. I., Pantiuszyn W. S.: Ispitanja fierromagnitnykh matieriałow. G.E.I., Leningrad 1955.
5. Tamm I. E.: Osnovy teorii elektricestwa. Moskwa 1949.
6. Deutch M.: La voiture de contrôle des rails de la SNCF. Révue générale des chemins de fer. Maj 1954.

5. PRACA CZUJNIKA CEWKOWEGO W POLU MAGNETYCZNYM PRZY DEFEKTOSKOPII SZYNOWEJ

W metodzie magnetycznej o strumieniu wzdłużnym czujnik wykrywający znajduje się przy powierzchni szyny, w środku pod elektromagnesem. Podlega on różnym zakłóceniom zewnętrznym. Najmniej ulega zakłóceniom ekranowany elektrycznie i magnetycznie cewkowy czujnik różnicowy bez rdzenia o cewkach poziomych. Reaguje on na stromość składowej poziomej zewnętrznego pola dopełniającego, jakie występuje w sąsiedztwie skazy.

Na skazach sztucznych sprawdzono czułość wykrywania, dobór optymalnej indukcji w szynie, wpływ względnej wielkości skazy i jej położenia w głowce

szyny. Podano wytyczne dla uzyskania równomiernej czułości wykrywania skaz w główce.

Doświadczenia na torowisku pozwoliły wprowadzić podział skaz na przypadkowe, powtarzalne i niebezpieczne. Określono wskaźniki niejednorodności główki i gęstości rozmieszczenia skaz w główce. Impuls większy od impulsu skazy wzorcowej w główce trafiał się średnio 4 razy na kilometr. Skaz przekraczających 10% przekroju szyny było jednak znacznie mniej. Przyczyną tego niepożądanego zjawiska była dosyć znaczna nierównomierność wykrywania skaz czujnika cewkowego.

Przeprowadzono krytykę czujnika cewkowego i podano wytyczne dla budowy ustroju magnesującego, sygnalizacji i kontroli wykrywania skaz.

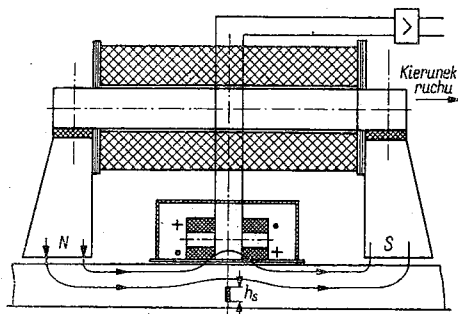
5.1. Wstęp

Najważniejszym elementem defektoskopu szynowego jest czujnik wykrywający. W celu poznania warunków jego prawidłowej pracy należy wybrać jeden z najprostszych typów czujnika, np. cewkowy, i rozpatrzyć jego działanie dla określonych skaz w szynie, zarówno sztucznych, jak i naturalnych. Następnie dopiero będzie można podać wytyczne dla czujnika ulepszanego.

W polu magnetycznym zarówno zasilanie szyny strumieniem magnetycznym, jak i sposób wykrywania skaz upraszczają się wielokrotnie w stosunku do pola prądowego, a skazy poprzeczne o grubości powyżej 0,1 mm przy prawidłowo dobranej indukcji w szynie wywołują nie mniejszy stopień niejednorodności pola niż w polu prądowym.

5.2. Zakłócenie w pracy czujnika

W czynnym polu między biegunami elektromagnesu jak na rys. 1 istnieje przy powierzchni szyny zewnętrzne pole główne $B_{p\infty} = B_{\infty}/\mu_w$.



Rys. 1. Umieszczenie i osłona czujnika cewkowego

W zasięgu skazy nakłada się na nie zewnętrzne pole dopełniające o składowej poziomej:

$$B_{px} = B_{p\infty} \cdot \varepsilon_B \quad (1)$$

gdzie:

B_{∞} — indukcja w szynie poza zasięgiem skazy,

ε_B — stopień niejednostajności pola na skutek obecności skazy,

$$\varepsilon_B = \frac{B - B_{\infty}}{B_{\infty}},$$

μ_w — względna przenikalność stali szynowej, średnio około 480.

Rozkład stopnia niejednostajności pola wyznacza się przez odwzorowanie skazy w wannie elektrolitycznej. Można także wykonać sztuczną skazę w szynie i określić rozkład niejednostajności pola doświadczalnie, za pomocą próbnej zwojniczki połączonej z galwanometrem. Przesunięcie zwojniczki poza zasięg skazy wywołuje wychylenie galwanometru proporcjonalne do składowej poziomej pola dopełniającego przy powierzchni szyny: $\alpha = c \cdot B_{px}$. Półobrót zaś zwojniczki powoduje wychylenie proporcjonalne do podwójnej wartości indukcji całkowitego pola zewnętrznego:

$$\alpha_{\infty} = 2 \cdot c [B_{py} + B_{px}].$$

Po uwzględnieniu wzoru (1) otrzymujemy:

$$\varepsilon_B = \frac{2 \cdot \alpha}{\alpha_{\infty} - 2\alpha}. \quad (2)$$

Czujnik powinien reagować jedynie na zewnętrzne pole dopełniające wywołane skazą. Wszelkie zmiany pola głównego pochodzące od wstrząsów, zmian strumienia rozproszenia ustroju magnesującego itp. powodują niekorzystne zakłócenia. Wady tej nie posiada wprawdzie praca na remanencie magnetycznym, ale trudno wytworzyć w czasie kontroli bieżącej wzdlużny remanent w szynie. Po przesunięciu elektromagnesu powstaje jedynie remanent pionowy, zamykający się przez powietrze od główki do stopki szyny. Taki remanent nadaje się do wykrywania pęknięć w płaszczyźnie poziomej, ale zawodzi przy wykrywaniu wąskich, a niebezpiecznych pęknięć poprzecznych, istotnych dla defektoskopii szynowej.

Szkodliwy wpływ zakłóceń od zmiany strumienia głównego oraz strumienia rozproszenia daje się zmniejszyć w następujący sposób:

a) Należy stosować możliwie duży opór magnetyczny na drodze strumienia głównego przez wprowadzenie szczeliny powietrznej, ale nie między nabiegunkami i szyną, tylko między pieńkami i jarzmem. Zmniejsza się wówczas zarówno zmiany pola głównego przy wstrząsach, jak i zmniejszy się wpływ strumienia rozproszenia wskutek oddalenia poza sąsiedztwo czujnika.

b) Zamiast jednej zwojniczki należy dać dwie i połączyć je w przeciwnym kierunku, czyli różnicowo. Zmiana pola głównego indukuje w nich dwie przeciwnie skierowane, znoszące się wzajemnie SEM. Jedynie zmiana pola dopełniającego w zasięgu skazy wywołuje w nich impuls napięciowy. Zwojniczki te ułożone poziomo reagują na składową poziomą

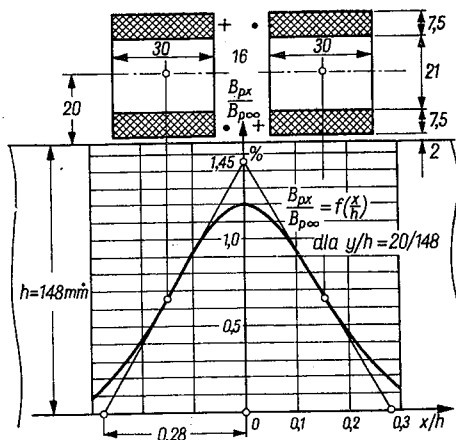
dopełniającego pola zewnętrznego B_{px} jak na rysunku 1, ułożone zaś pionowo reagują na składową pionową B_{py} .

Dalsze źródła zakłóceń mogą pochodzić od wstrząsów czujnika w czasie ruchu, jego oddziaływania na pole w szynie oraz wpływów zewnętrznych, jak obce pola elektryczne bądź magnetyczne. Zwykły czujnik cewkowy bez rdzenia żelaznego nie wpływa na rozkład pola przy powierzchni szyny. Jest on również mało wrażliwy na wstrząsy. Wpływy zewnętrzne zaś dają się odgrodzić przez ekranowanie czujnika od spodu blachą z metalu niemagnetycznego. Osłonę boczną i górną stanowi blacha o dużej przenikalności magnetycznej, ekranująca zarówno elektrycznie, jak i magnetycznie.

Zastosowanie powyższych środków zaradczych pokazanych na rysunku 1 uwalnia praktycznie czujnik od niebezpiecznego wpływu zakłóceń.

5.3. Zasada czujników cewkowych

Rozpatrzmy najpierw prosty czujnik cewkowy, różnicowy, reagujący na składową poziomą zewnętrznego pola dopełniającego wywołanego ska-



Rys. 2. Poziome umieszczenie cewek w polu dopełniającym nad szyną dla skazy u dołu główki (skaza „A”)

zą w główce szyny. Czujnik składa się z dwóch cewek od licznika indukcyjnego dwusystemowego na napięcie 380 V. Wymiary cewki widoczne są na rys. 2.

Dane cewki: ilość zwojów $z = 17\,000$, średnica drutu $d = 0,09$ mm, oporność czynna $R = 4,3$ k Ω , średni przekrój cewki $q = 6,15$ cm².

Obie cewki ułożone poziomo posuwają się nad szyną w odległości ok. 2 mm. Odstęp osi cewek nad szyną wynosi 20 mm, co w stosunku do wysokości szyny stanowi $20/148 = 0,135$.

Największą wartość *SEM* indukowanej w cewkach osiągamy w chwili, gdy środki ich znajdują się równocześnie w miejscach największej stromości składowej poziomej B_{px} . Jeśli znany jest rozkład stopnia niejednostajności pola na powierzchni szyny $\varepsilon_B = f\left(\frac{x}{h}\right)$ oraz indukcja w szynie B_∞ , to metoda numeryczna pozwala wyznaczyć funkcję strumienia i indukcję B_{px} w dowolnym punkcie dopełniającego pola zewnętrznego.

Na rysunku 2 przedstawiono stosunek $B_{px}/B_{p\infty} = f\left(\frac{x}{h}\right)$ w płaszczyźnie poziomej, przechodzącej przez oś cewek czujnika i to dla skazy wzorcowej „A”, stanowiącej 10% przekroju szyny, i położonej w dolnej części główki. Odległość punktów przegięcia składowej $B_{px}/B_{p\infty}$, równa odległości środków cewek, wynosi tu 46 mm.

W czasie ruchu indukuje się w cewce *I SEM*: $e_I = \frac{d}{dt} [\Phi_Z]$. Przyjmując, że średnia indukcja w cewce równa się indukcji w jej środku, a czujnik porusza się z prędkością $v = \frac{dx}{dt}$, otrzymamy $e_I = z \cdot q \cdot v \left[\frac{dB_{px}}{dx} \right]_s$. Pochodną $\frac{dB_{px}}{dx}$ w środku cewek oblicza się z wykresu $\frac{B_{px}}{B_{p\infty}} = f\left(\frac{x}{h}\right)$ na rysunku 2 jako $\left[\frac{dB_{px}}{dx} \right]_s = \frac{1,45 \cdot 10^{-2} \cdot B_{p\infty}}{0,28 \cdot h}$. Przy wzroście skazy, zachowującej swoje położenie, pochodna ta rośnie — jak wynika z doświadczenia — proporcjonalnie do maksymalnej wartości stopnia niejednostajności pola ε_{B0} , który dla skazy „A” wynosi $3,24 \cdot 10^{-2}$. Będzie więc po podstawieniu $B_{p\infty} = \frac{B_\infty}{\mu_0}$

$$e_I = z \cdot q \cdot v \cdot \frac{1,45 \cdot 10^{-2} \cdot B_{p\infty}}{0,28 \cdot h} \cdot \frac{\varepsilon_{B0}}{3,24 \cdot 10^{-2}} = 1,6 \frac{z \cdot q \cdot v \cdot B_\infty \cdot \varepsilon_{B0}}{h \cdot \mu_0}.$$

W drugiej cewce indukuje się *SEM* e_{II} , równa co do wartości poprzedniej, ale odwrotnego znaku. Na wejściu do wzmacniacza występuje różnica obu *SEM*, a więc suma ich bezwzględnych wartości:

$$e = k_x \frac{z \cdot q \cdot v \cdot B_\infty}{h \cdot \mu_0} \varepsilon_{B0}. \quad (3)$$

Współczynnik k_x charakteryzuje zasięg składowej poziomej pola zewnętrznego w środku obu cewek. W naszym przypadku, gdy środki cewek znajdują się w punktach przegięcia składowej B_{px} , otrzymujemy $k_x = 3,2$.

Jeśli grubość skazy przekracza 0,1 mm, to przy prawidłowym doborze indukcji w szynie B_∞ skaza wywołuje w polu magnetycznym taki sam

lub większy stopień niejednostajności niż w polu prądowym. Stopień niejednostajności pola ε_{B0} będzie jak w polu prądowym proporcjonalny do kwadratu względnej wielkości skazy:

$$\varepsilon_B = C_s \cdot \varepsilon_s^2, \quad (4)$$

gdzie:

C — współczynnik oddziaływania skazy, zależny od jej położenia. Dla skazy wzorcowej w dolnej części główki $C = 3,24$,

ε_s — względna wielkość skazy w stosunku do wysokości szyny,

$$\varepsilon_s = \frac{s}{h}.$$

Po podstawieniu wyrażenia (4) do wzoru (3) otrzymamy:

$$e = k_x \cdot C_s \frac{z \cdot q \cdot v \cdot B_{\infty}}{h \cdot \mu_w} \varepsilon_s^2. \quad (5)$$

Przy dalszych doświadczeniach zakładamy, że prędkość posuwu wózka z czujnikiem równa się $v = 1 \text{ m/sec}$. Przyjmując poza tym z i q jak w rozdziale 5.3 oraz $B_{\infty} = 0,75 \text{ Wb/m}^2$, $\mu_w = 480$, $h = 0,148 \text{ m}$ otrzymamy:

$$e = k_x \cdot C_s \frac{17000 \cdot 6,15 \cdot 10^{-4} \cdot 1 \cdot 0,75}{0,148 \cdot 480} \varepsilon_s^2,$$

$$e = k_x \cdot C_s \cdot 0,11 \cdot \varepsilon_s^2. \quad (6)$$

Dla skazy wzorcowej „A” $k_x = 3,2$, $C_s = 3,24$, zatem SEM indukowana w czujniku o cewkach poziomych wyniesie

$$e = 1,14 \varepsilon_s^2 \text{ woltów}. \quad (7)$$

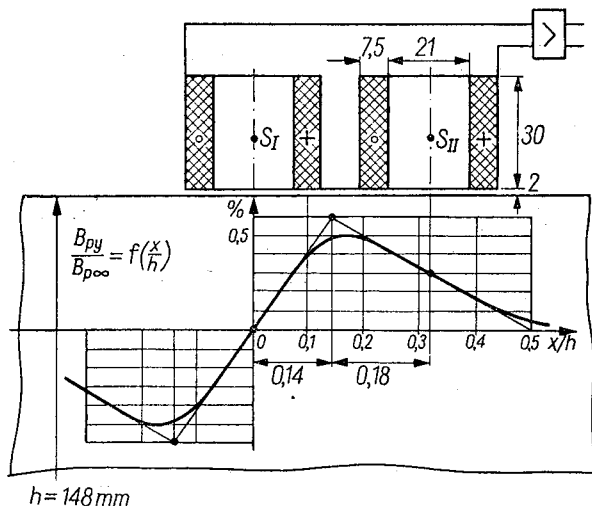
Te same co poprzednio cewki ustawione pionowo nad szyną stanowią prosty czujnik indukcyjny dla składowej pionowej pola zewnętrznego B_{py} . Przebieg tej składowej wyznacza się na podstawie siatki pola analogicznie jak poprzednio dla B_{px} . Na rysunku 3 przedstawiono umieszczenie czujnika i stosunek: $\frac{B_{py}}{B_{p\infty}} = f\left(\frac{x}{h}\right)$ w odległości środka cewek od powierzchni szyny, czyli dla $y/h = 17/148 = 0,115$. Największą wartość SEM czujnika osiągamy w chwili, gdy pierwsza cewka znajdzie się w środku nad skazą, a druga w punkcie przegięcia na opadającej gałęzi składowej pionowej pola zewnętrznego B_{py} , to jest w odstępnie 47 mm od pierwszej. SEM indukowane w cewkach wynoszą:

$$e_I = z \cdot q \cdot v \left[\frac{dB_{py}}{dx} \right]_{s_I}, \quad e_{II} = z \cdot q \cdot v \left[\frac{dB_{py}}{dx} \right]_{s_{II}}.$$

Na wejściu wzmacniacza mamy różnicę obu tych SEM, więc:

$$e = z \cdot q \cdot v \left[\left(\frac{dB_{py}}{dx} \right)_{s_I} - \left(\frac{dB_{py}}{dx} \right)_{s_{II}} \right].$$

Pochodne indukcji w środku cewek odczytuje się z rysunku 3 dla stopnia niejednostajności pola $3,24 \cdot 10^{-2}$. Przy wzroście skazy zachowującej swoje położenie pochodne indukcji wzrosną jak poprzednio w stosunku $\varepsilon_B/3,24 \cdot 10^{-2}$.



Rys. 3. Pionowe umieszczenie cewek w polu dopełniającym nad szyną dla skazy u dołu główki (skaza „A”)

Będzie zatem:

$$\left(\frac{dB_{py}}{dx} \right)_{S_I} = + \frac{0,6 \cdot 10^{-2}}{0,14 \cdot 3,24 \cdot 10^{-2}} \cdot \frac{B_{\infty} \cdot \varepsilon_{B0}}{h \cdot \mu_w} = 1,32 \frac{B_{\infty} \cdot \varepsilon_{B0}}{h \cdot \mu_w},$$

$$\left(\frac{dB_{py}}{dx} \right)_{S_{II}} = - \frac{0,6 \cdot 10^{-2}}{0,36 \cdot 3,24 \cdot 10^{-2}} \cdot \frac{B_{\infty} \cdot \varepsilon_{B0}}{h \cdot \mu_w} = -0,52 \frac{B_{\infty} \cdot \varepsilon_{B0}}{h \cdot \mu_w}.$$

Na wejściu wzmacniacza różnica obu SEM wyniesie:

$$e = k_y \frac{z \cdot q \cdot v \cdot B}{h \cdot \mu_w} \varepsilon_{B0}. \quad (8)$$

Współczynnik k_y charakteryzuje zasięg składowych pionowych zewnętrznego pola dopełniającego w środku cewek. Jest on nieco mniejszy niż analogiczny współczynnik dla składowej poziomej i wynosi $k_y = 1,84$.

5.4. Próby wykrywania skaz sztucznych

Próby laboratoryjne przeprowadzono czujnikiem o cewkach poziomych z prędkością posuwu ok. 1 m/sec, a wartość indukowanej SEM mierzono oscylografem katodowym. Najprzód sprawdzono słuszność wzoru (7), wyrażającego wielkość SEM indukowanej w czujniku. W tym celu

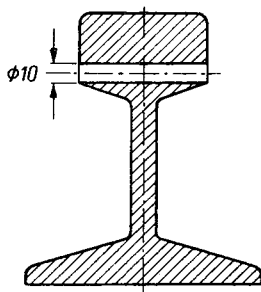
w dolnej części główki wywiercono otwór $d = 10$ mm jak na rys. 4. Względna wielkość skazy wyniesie wówczas:

$$\varepsilon_s = \frac{1 \cdot 6,8}{62,5} = 10,9 \cdot 10^{-2}.$$

Ze wzoru (7) mamy $e = 1,14 \cdot 10,9^2 \cdot 10^{-4} = 13,6 \cdot 10^{-3}$ V, a z pomiarów:

$$e_{pom} = 10,8 \text{ mV}.$$

Obliczona wg wzoru (7) wartość impulsu ważna jest dla cienkiej skazy poprzecznej. Z powodu trudności wykonania takiej skazy zastąpiono ją



Rys. 4. Skaza w postaci otworu $d = 10$ mm w dolnej części główki szyny S49

skązą okrągłą. Z jednej strony należało się spodziewać dla skazy okrągłej większej nieco wartości impulsu przy prawidłowym doborze indukcji B_∞ , gdyż strumień prawie całkowicie omija tak grubą skazę. Z drugiej jednak strony kołowy kształt skazy łagodzi stromość składowej poziomej pola, a rozszerza jej średni zasięg. Tymczasem impuls czujnika cewkowego zależy jedynie od stromości pola, a nie od zasięgu skazy.

W celu sprawdzenia, jaki wpływ wywiera kształt szczeliny, wywiercono jeden nad drugim dwa otwory o średnicy $d = 5$ mm w tym samym co poprzednio miejscu. Otrzymano nieco większą niż poprzednio wartość impulsu $e_{pom} = 11,6$ mV. Różnicę w stosunku do wzoru (7) należy tłumaczyć niewypełnieniem założeń koniecznych dla tego wzoru.

Ważnym z kolei zagadnieniem jest dobór najkorzystniejszej indukcji w szynie B_∞ . Stosownie do wzoru (3) pożądane jest przekroczenie indukcji odpowiadającej największej przenikalności magnetycznej, tak aby przy dalszym wzroście indukcji przenikalność magnetyczna malała. Wartość indukowanego impulsu czujnika nie może być jednak jedynym wskaźnikiem dobroci wykrywania. Wraz ze wzrostem tego impulsu wzrastają zwykle także i wszelkie zakłócenia. Należy raczej dla określonego typu skazy obrać taką wartość indukcji, aby uzyskać największą wartość stopnia niejednostajności pola.

Indukcja we wzorze (3) nie wpływa praktycznie na współczynnik stromości pola k_x , a stopień niejednostajności wyniesie:

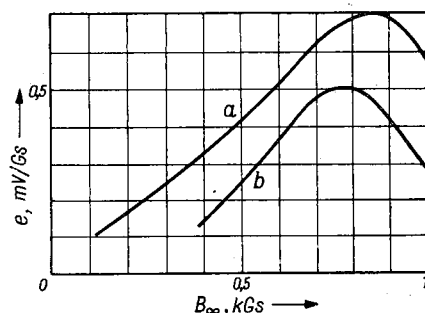
$$\varepsilon_{B0} = \frac{e \cdot \mu_w}{B_\infty} \cdot \frac{h}{z \cdot q \cdot v \cdot k_x} = \frac{e}{B_{p\infty}} \cdot \frac{h}{z \cdot q \cdot v \cdot k_x}.$$

Najkorzystniejsze warunki wykrywania wystąpią przy indukcji B_∞ , dla której ułamek $\frac{e \cdot \mu_w}{B_\infty} = \frac{e}{B_{p\infty}}$ osiągnie największą wartość. Doświadczenia przeprowadzono dla skazy okrągłej jak na rysunku 4. Dla różnych prądów zasilania elektromagnesu mierzono zwojniczką próbną składową poziomą pola zewnętrznego $B_{p\infty} = \frac{B_\infty}{\mu_w}$ w miejscu oddalonym od skazy, a wartość impulsu e oscylografem katodowym. Natężenie pola przy po-

Tablica 1

$B_{p\infty}$ Gs	H A/cm	B_∞ Gs	e mV	$e/B_{p\infty}$ mV/Gs	e mV	$\frac{e}{B_{p\infty}}$
			Otwór $d=10\text{mm}$		Otwór zatkany	
6,2	5,0	2500	1,22	0,2	—	—
9,4	7,5	4400	3,2	0,34	1,7	0,18
12,5	10,0	6250	6,9	0,55	5,0	0,4
15,6	12,5	7500	10,5	0,67	7,8	0,5
21,9	17,5	8700	15,3	0,70	9,8	0,45
50,0	40,0	10500	22,5	0,45	13,0	0,26

wierzchni, a więc także i w szynie, oblicza się jako: $H = 0,8 \cdot B_{p\infty}$ dla H w A/cm i $B_{p\infty}$ w Gs. Indukcję w szynie określamy na podstawie krzywej magnesowania dla stali szynowej. Doświadczenie powtórzono dla tej samej skazy, ale zatkaney kołkiem ze stali szynowej tak, by luz wynosił ok. 0,1 mm. Otrzymane wartości zestawiono w tablicy 1. Pomiary dla zatkanego otworu zostały wykonane po rozmagnesowaniu szyny. Zależ-



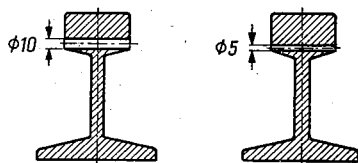
Rys. 5. Zależność: $\frac{e}{B_{p\infty}} = f(B_\infty)$ dla otworu zwykłego $d = 10$ mm — krzywa a, i zatkanego kołkiem stalowym — krzywa b

ność: $\frac{e}{B_{p\infty}} = f(B_{\infty})$ przedstawiono wykreślnie dla obu wypadków na rysunku 5.

Z pomiarów wynika, że najkorzystniejsza wartość indukcji wynosi 8500 Gs dla otworu $d = 10$ mm oraz 8000 Gs przy zatkaniu tego otworu kołkiem. Przez włożenie kołka impuls czujnika dla $B_{\infty} = 7500$ Gs zmniejszył się do ok. 0,74 pierwotnej wartości.

Pomiary w polu prądowym wykazały, że dla stałego położenia skazy stopień niejednostajności pola, a zatem i wartość impulsu czujnika rośnie z kwadratem względnej wielkości skazy stosownie do wzoru (4), o ile ta względna wartość nie przekroczy 15%. Dla skaz owalnych za wielkość skazy przyjmuje się jej rzut na przekrój poprzeczny szyny. Analogiczną zależność przyjmowaliśmy również dla pola magnetycznego, którą teraz sprawdzimy doświadczalnie.

W celu kontroli tej zależności wywiercono w główce szyny dwa otwory jak na rys. 6 o średnicach $d_1 = 10$ mm i $d_2 = 5$ mm. Otrzymano wartości

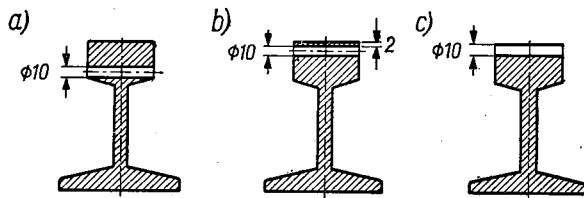


Rys. 6. Otwory różnej średnicy w dolnej części główki

impulsów $e_1 = 10,7$ mV oraz $e_2 = 2,7$ mV. Stosunek $e_1/e_2 = 10,7/2,7 = 3,97$.

Zarówno ten, jak i inne pomiary potwierdzają, że zależność kwadratowa jak we wzorze (4) utrzymuje się także i w polu magnetycznym.

W miarę przesuwania skazy z dolnej części główki ku górze rośnie silnie współczynnik oddziaływania skazy C_s , natomiast maleje współczynnik k_x , ponieważ odstęp punktów przecięcia składowej poziomej pola



Rys. 7. Skazy sztuczne w różnych częściach główki:

a) Otwór w dolnej części główki, b) Otwór przy powierzchni szyny, c) Nacięcie główki

zewnętrznej B_{px} zmniejsza się poniżej średniej odległości cewek. W celu porównania czułości wykrywania w górnej i dolnej części główki wykonano następujące doświadczenie. Wywiercono otwór o średnicy $d = 10$ mm w dolnej i górnej części główki. Brzeg górnego otworu znalazł się w odległości 2 mm od powierzchni szyny. Prócz tego nacięto szynę od główki

na głębokość 10 mm. Pomiary dla tych trzech przypadków (jak na rysunku 7) dały następujące wartości impulsów czujnika:

$$e_1 = 10,4 \text{ mV}; \quad e_2 = 17,1 \text{ mV}; \quad e_3 = 23,0 \text{ mV}.$$

Stosunkowo duża różnica wartości na rysunku 7b i 7c tłumaczy się zarówno zmianą kształtu szczeliny, jak i bocznikującym działaniem stali pomiędzy otworem a powierzchnią szyny. Okazuje się, że dla czujnika indukcyjnego o cewkach poziomych czułość wykrywania skaz powierzchniowych jest ok. dwukrotnie większa niż czułość wykrywania skaz położonych w dolnej części główki. W innych doświadczeniach dla skaz mniejszych wzrost czułości wykrywania skaz powierzchniowych był nawet nieco większy. Zjawisko powyższe jest — jak się to dalej okaże — wybitnie niekorzystne dla selektywności wykrywania.

5.5. Granica i selektywność wykrywania skaz

Dla czujnika indukcyjnego o cewkach poziomych przy wykrywaniu skaz w dolnej części główki ważny jest wzór (7). Napięciu $e = 1 \text{ mV}$, które daje się bez trudu wzmocnić, odpowiada tu skaza 3-procentowa ($\epsilon_s = 2,95 \cdot 10^{-2}$). Niestety w czasie ruchu defektoskopu powstają najrozmaitsze zakłócenia, które mogą zatrzeć impuls wywołany badaną skazą.

W zależności od wielkości impulsu zakłócenia dają się podzielić następująco:

a) Zakłócenia przypadkowe. Zakłócenia te o impulsie poniżej $\pm 0,5 \text{ mV}$ są bardzo częste i pochodzą głównie od wstrząsów i drgań mechanicznych wywołujących zmiany pola elektromagnetycznego. Dzięki starannemu ekranowaniu elektrycznemu i magnetycznemu poziom napięciowy tych zakłóceń nie przekraczał $0,5 \text{ mV}$, co odpowiada skazie $1,5\%$ przekroju szyny.

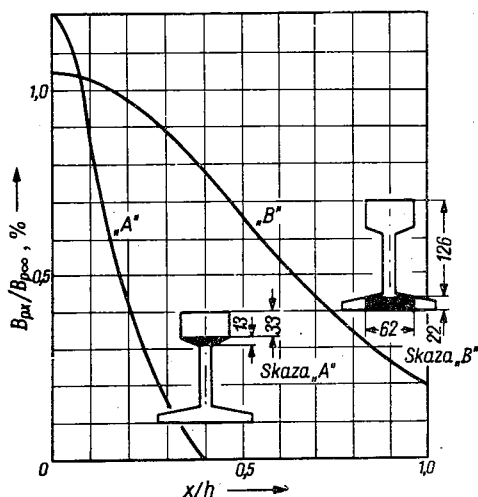
b) Zakłócenia powtarzalne. Zakłócenia te powodują szybki skok plamki oscylografu, co wskazuje na ich powierzchniowe pochodzenie. Wywołują je drobne defekty powierzchniowe, zmiany struktury szyny, wtrącenia żuźlowe itp. Powtarzają się one nieregularnie, ale zawsze w tych samych miejscach. Na ogół trudno określić ich przyczynę, jeśli napięcie czujnika nie przekracza 3 mV , co odpowiada skazie $5,1\%$ przekroju szyny.

c) Skazy zwykłe i niebezpieczne. Po przekroczeniu wartości impulsu 3 mV już częściej, ale nie zawsze, da się określić przyczynę impulsu. W przypadku cienkich pęknięć poprzecznych przesunięcie cięcia szyny o niespełna 1 mm nie wykryje już pęknięcia, które może być nawet bardzo niebezpieczne. Normalnie przyczyną impulsu bywają tu mniejsze jamy usadowe lub pęknięcia, a także większe niejednorodności struktury, zwłaszcza przy powierzchni starych, zużytych szyn.

Granica niebezpieczeństwa skazy daje się określić dopiero po długich próbach eksploatacyjnych. Na podstawie dotychczasowych prób i doświadczeń zagranicznych, zwłaszcza francuskich, za niebezpieczną uznaje

się skazę stanowiącą 10‰ całkowitego przekroju szyny. Skaza taka umieszczona w dolnej części główki daje dla naszego czujnika impuls napięciowy ok. 10,5 mV.

Przeprowadzając próby na skazach sztucznych przekonaliśmy się, że czujnik cewkowy jest około dwukrotnie czulszy dla skaz położonych w górnej części główki niż w dolnej. Przy ocenie skaz na podstawie indukowanej SEM w czujniku cewkowym skazy w górnej części będą przejawiskrawione i niesłusznie mogą być zaliczone do skaz niebezpiecznych. W naszym czujniku odległość środków cewek równa się odległości punktów przegięcia składowej pola zewnętrznego B_{px} , pochodzącej od 10-procentowej skazy wzorcowej „A” znajdującej się w dolnej części główki. Ze wzrostem odstepu cewek, zwłaszcza przy równoczesnym zwiększaniu



Rys. 8. Składowa pozioma B_{px} pola dopełniającego w odległości 20 mm od powierzchni szyny dla skaz wzorcowych „A” i „B”. $\frac{B_{px}}{B_p} = f\left(\frac{x}{h}\right)$

ich długości i oddalaniu od szyny, czułość wykrywania maleje silniej dla skaz w górnej części główki niż w części dolnej. Dla określonego typu skazy, np. dla dwóch otworów o średnicy po 5 mm umieszczonych nad sobą raz w dolnej, drugi raz w górnej części główki, udawało się uzyskiwać silnie obniżoną, ale prawie jednakową czułość wykrywania w obu przypadkach przez odpowiednie rozmieszczenie cewek.

W miarę posuwania się skazy od główki poprzez szyjkę ku stopce rośnie odległość punktów przegięcia składowej poziomej dopełniającego pola zewnętrznego i maleje jej stromość. Rysunek 8 przedstawia obok skazy wzorcowej „A” skazę wzorcową „B” umieszczoną w stopce i stanowiącą 20‰ przekroju szyny. Stopień niejednostajności pola przy powierzchni szyny dla powyższych skaz znany jest z odwzorowania w wannie elektrolitycznej.

Rozkład pola w określonej odległości od szyny, np. 20 mm, daje się wyznaczyć rachunkowo metodą numeryczną. Okazuje się, że pole od skazy w stopce maleje znacznie wolniej z odległością od powierzchni szyny niż dla skazy w główce. W odległości bliskiej wysokości szyny ($h = 148$ mm) składowa pozioma indukcji B_{px} jest około czterech razy mniejsza niż przy powierzchni. W miarę oddalania się od powierzchni indukcja maleje w przybliżeniu odwrotnie proporcjonalnie do kwadratu odległości górnej krawędzi skazy od powierzchni szyny. Indukcja zewnętrznego pola dopełniającego w odległości y od szyny wynosi więc:

$$\text{dla skazy „A” } B_1 = \frac{\varepsilon_1 \cdot B_\infty}{\mu_w} \left(\frac{33}{33 + y} \right)^2,$$

$$\text{dla skazy „B” } B_2 = \frac{\varepsilon_2 \cdot B_\infty}{\mu_w} \left(\frac{126}{126 + y} \right)^2.$$

Liczby 33 i 126 wyrażają tu odległości w mm górnej krawędzi skazy od powierzchni szyny, stopnie niejednostajności wynoszą $\varepsilon_1 = 3,24 \cdot 10^{-2}$ oraz $\varepsilon_2 = 1,4 \cdot 10^{-2}$. Po podstawieniu i podzieleniu otrzymamy:

$$B_2/B_1 = 6,3 \left(\frac{33 + y}{126 + y} \right)^2. \quad (9)$$

Ze wzoru (9) widać, że w odległości $y = 28$ mm od powierzchni szyny obie składowe są praktycznie równe. Skaza „B”, dwukrotnie większa, posiada jednak prawie 4-krotnie większą rozpiętość, a jej składowa pola prawie 4-krotnie mniejszą stromość. Zwykły czujnik cewkowy reagujący jedynie na stromość pola o odstępie cewek 16 mm, jak na rysunku 2, wykazuje impuls ok. 10-krotnie mniejszy, a nawet przy najkorzystniejszym odstępie cewek ok. 60 mm impuls jego jest jeszcze prawie 4-krotnie mniejszy, niż to miało miejsce dla skazy „A”. Dla skaz w stopce o długim, ale płaskim zasięgu działania czujnik cewkowy nie daje zadowalających rezultatów. Dopiero dodanie integratora poprawia nieco efekt wykrywania skaz w stopce.

5.6. Próby eksploatacyjne z dotychczasowym czujnikiem cewkowym

Próby eksploatacyjne przeprowadzono za pomocą prototypu defektoskopu laboratoryjnego z czujnikiem cewkowym jak na rysunku 2, przy szybkości jazdy 2,5 m/sec. Otrzymane wartości impulsów jako proporcjonalne do prędkości przeliczono na prędkość laboratoryjną 1 m/sec.

Na każdym metrze bieżącym szyny czujnik daje pewien powtarzalny, maksymalny impuls. Niech średnia wartość tych maksymalnych impulsów na określonym odcinku szyny wynosi e_{sr} , a impuls od skazy wzorcowej „A” w dolnej części główki niech wyniesie e_A . Wtedy stosunek $\zeta = e_{sr}/e_A$ jako średni impuls względny jest wskaźnikiem niejednorodności główki szyny.

Jeśli bowiem czujnik wykazuje jednakową czułość dla górnej i dolnej części główki, to ze wzoru (7) otrzymamy:

$$\zeta_{sr} = \frac{\sum \varepsilon_s^2}{n \cdot \varepsilon_{sA}^2}, \quad (10)$$

gdzie:

- ε_s — maksymalna względna wartość skazy na długości jednego metra,
- ε_{sA} — względna wielkość skazy wzorcowej „A”,
- n — ilość metrów bieżących rozpatrywanego odcinka szyn.

Średni impuls względny nie zależy tu od budowy czujnika, a jedynie od niejednorodności struktury oraz wielkości i gęstości rozłożenia skaz w główce. W rzeczywistości przedstawiając dla czujnika cewkowego wzór (7) w postaci: $e = k \cdot \varepsilon_s^2$, mamy dla skazy w dolnej części główki $k_2 = 1,14$, a dla górnej części główki $k_1 \approx 2,0$. Przy równomiernym rozkładzie skaz na górną i dolną część główki otrzymamy wówczas jako pozorny względny impuls średni:

$$\zeta'_{sr} = \frac{k_1 + k_2}{2 \cdot k_2} \zeta_{sr} = 1,38 \zeta_{sr}.$$

Na skutek nierównomiernej czułości czujnika cewkowego nastąpiło pozorne zwiększenie wskaźnika niejednorodności główki. Zależy on więc w tym wypadku także od budowy czujnika. Badany przez nas odcinek szyn, wykazujący stosunkowo dużo braków, dawał średni impuls $e_{sr} = 2,5$ mV, $e_A = \text{ok. } 11,4$ mV, więc $\zeta'_{sr} = 0,219$. Rzeczywisty zaś wskaźnik niejednorodności główki można oszacować jako $\zeta_{sr} = 0,219/1,38 = 0,158$.

Drugim wskaźnikiem przy defektoskopii szynowej może być wskaźnik gęstości rozmieszczenia skaz w główce. Niech $\zeta = e/e_A$ wyraża pewną względną wartość impulsu i niech na odcinku n metrów występuje N impulsów większych od ζ , przy czym wykluczamy impulsy na stykach. Gęstością rozmieszczenia skaz nazywamy stosunek $\sigma = N/n$. Małe skazy występują gęsto, duże zaś rzadko. Zależność $\sigma = f(\zeta)$ jest funkcją monotonicznie malejącą, odbiegającą od normalnego rozkładu Gaussa. Na sprawdzanym odcinku szyn zanotowano gęstość rozmieszczenia skaz jak w tablicy 2.

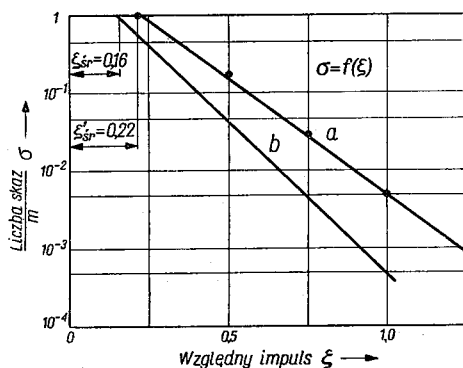
Tablica 2

Względny impuls ζ	0,219	0,5	0,75	1,0	1,25
Gęstość rozmieszczenia skaz $\sigma \left(\frac{\text{ilość skaz}}{\text{m}} \right)$	1	170	30	5	0,9
				$\times 10^{-3}$	

Na rysunku 9 widać, że punkty pomiarowe w skali logarytmicznej dla σ leżą w przybliżeniu na linii prostej i dają się wyrazić równaniem:

$$\sigma = \sigma_0 \exp[-C'(\xi - \xi_{sr})], \quad (11)$$

gdzie C' jest pozornym wskaźnikiem gęstości rozmieszczenia skaz w główce. W naszym przypadku było $C' = 6,8$, $\sigma_0 = 1$. Przekroczenie impulsu od skazy wzorcowej „A” trafiało się na badanym odcinku około 5 razy na kilometr ($\sigma = 5 \cdot 10^{-3}$). Niektóre z tych skaz jako powierzchniowe i wi-



Rys. 9. Gęstość rozmieszczenia skaz na określonym odcinku toru w zależności od przekraczanego impulsu czujnika: a — dla czujnika zwykłego, b — dla czujnika o jednakowej czułości w główce

doczne nie zostały uznane za niebezpieczne. Inne niewidoczne okazały się po przecięciu szyny większymi jamami usadowymi i pęknięciami. Czujnik o równomiernej czułości wykrywania w główce wykazywałby taką samą gęstość rozmieszczenia skaz dla impulsów zmniejszonych w stosunku $\xi_{sr}/\xi_{sr}' = 0,158/0,219$. Rzeczywisty wskaźnik gęstości rozmieszczenia skaz w główce wyniesie więc $C = 9,35$. Skaza rzeczywiście równorzędna pod względem wielkości skazie wzorcowej „A” trafia się na badanym odcinku tylko około czterech razy na 10 kilometrów, czyli prawie 10 razy rzadziej niż wykazuje to czujnik o czułości nierównomiernej. Na rysunku 8 linia b) przedstawia zależność $\sigma = f(\xi)$ dla $C = 9,35$. Otrzymany rezultat uzasadnia dążność do budowy czujników o równomiernej czułości w główce szyny.

Na podstawie dotychczasowych prób eksploatacyjnych można już wyciągnąć pewne wnioski dotyczące czujnika cewkowego. Czujnik ten w stosunkowo prosty i łatwy sposób wykrywa skazy w główce szyny przekraczające 10% przekroju poprzecznego. Pozwala on także dzięki zwiększeniu długości cewek i odpowiedniemu rozmieszczeniu ich uzyskać prawie równomierną czułość wykrywania dla całej główki. Jest on poza tym mało wrażliwy na wstrząsy mechaniczne, gdyż nie wpływa na przebieg pola dopełniającego od skazy.

Obok prostej budowy i działania czujnik indukcyjny cewkowy posiada jednak następujące wady:

a) Stosownie do wzoru (6) wartość impulsu napięciowego czujnika zależy nie tylko od wielkości skazy, ale i od prędkości wykrywania. Prędkość ta powinna być stała, albo też trzeba regulować wzmocnienie impulsu w zależności od prędkości jazdy.

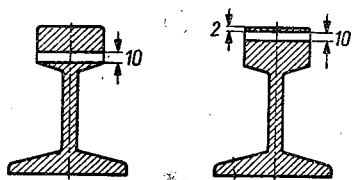
b) Czujnik reaguje jedynie na stromość składowej poziomej zewnętrznego pola dopełniającego pochodzącego od skazy, jest więc z natury mało wrażliwy na skazy w stopce o płaskiej składowej pola, które mają znaczną rozpiętość, a stanowią poważne niebezpieczeństwo dla ruchu.

c) Czujnik cewkowy, zwłaszcza gdy znajduje się w pewnym oddaleniu od powierzchni, w celu uzyskania równomiernej czułości wymaga dużej liczby stosunkowo cienkich zwojów. Posiada on wskutek tego dużą oporność czynną i podlega łatwo zakłóceniom pola elektrycznego.

Próby wykrywania skaz w stopce szyny czujnikiem cewkowym nie dały zadowalającego rezultatu. Do wykrywania skaz w stopce konieczny jest specjalny czujnik reagujący na maksymalną wartość pola B_{pz} oraz jego rozpiętość. Przewiduje się więc dwa czujniki: pierwszy dla skaz w główce i drugi, mało wrażliwy na skazy w główce, ale za to uczulony na skazy w stopce. Ostateczny wybór i konstrukcja czujnika zależy od zbadania innych typów czujnika, jak np. transformatorowego o pełnym lub niepełnym rdzeniu żelaznym. Przede wszystkim zbadać należy czujnik polaryzowany polem skazy w układzie wzmacniacza magnetycznego zasilanego prądem zmiennym o częstotliwości akustycznej. Działanie takiego czujnika nie zależy od szybkości wykrywania, tylko od indukcji i zasięgu pola dopełniającego od skazy. Dobór odpowiednich czujników jest zagadnieniem oddzielnym, podobnie jak budowa ustroju magnesującego i sygnalizacja skaz.

5.7. Wytyczne dla sygnalizacji i ustroju magnesującego

Sygnalizacja przewiduje zróżnicowanie skaz pod względem położenia i wielkości. Zróżnicowania skaz na skazy w główce i stopce dokonuje



Rys. 10. Sztuczna skaza wzorcowa w postaci dwóch otworów $\phi = 5$ mm umieszczonych nad sobą

się poprzez dwa oddzielne czujniki selektywne z przynależnymi wzmacniaczami. Granicą podziału dla skaz dużych i małych jest skaza wzorcowa stanowiąca 10% przekroju szyny, położona w dolnej części główki, możli-

wie cienka. Ze względu na trudność wykonania skazy sztucznej grubości 0,1 mm przyjmuje się umownie za skazę graniczną skazę złożoną z dwóch otworów $d = 5$ mm położonych jeden nad drugim i to raz w dolnej, a drugi raz w górnej części główki, jak na rysunku 10. Takie dwie niedaleko siebie położone skazy pozwalają nie tylko określać stosunkową wartość impulsu czujnika, lecz również sprawdzić równomierność czułości wykrywania. Skazy o względnym impulsie w przedziale $\zeta = 0,75$ do 1,0 uważać możemy za małe, a powyżej tego przedziału za duże. Z rysunku 9 widać, że dla naszego czujnika skazy małe trafiają się około 25 razy na kilometr, a skazy duże około 5 razy. Czujnik o równomiernej czułości daje znacznie mniejszą gęstość rozmieszczenia. Przy takiej sygnalizacji dwustopniowej mała skaza wyzwać może jedynie sygnał lampkowy, a duża oprócz tego sygnał lampkowy i akustyczny, połączony z natryskiem farby na szynę. Natrysk odmiennego koloru farby dla skaz małych wydaje się niecelowy. Natomiast obok rejestracji skaz na taśmie z zapisem ciągłym pożądanym jest oscylograf do pomiaru impulsów, przełączalny na oba toki. Oscylograf taki jest najlepszą kontrolą działania aparatury wykrywającej, a poza tym pozwala określić na pewnym odcinku wskaźnik jednolitości szyn i gęstość rozmieszczenia skaz jako wielkości charakteryzujących stan danego odcinka szyn.

Zaznaczone farbą skazy duże powinny być powtórnie zbadane i to możliwie inną metodą. Jeszcze w roku 1954 dodatkowa kontrola zarejestrowanych skaz odbywała się kosztowną metodą pola prądowego. Obecnie wprowadza się do tego celu defektoskop ultradźwiękowy. Załóżmy, że prawdopodobieństwo nieszkodliwości wśród zarejestrowanych skaz wynosi przy defektoskopii magnetycznej $p_m = 0,3$, a w ultradźwiękowej $p_u = 0,2$. Wówczas wypadkowa prawdopodobieństwo nieszkodliwości skaz po kontroli ultradźwiękowej magnetycznie zarejestrowanych skaz wynosi: $p_w = p_m \cdot p_u = 0,06$, a zatem prawdopodobieństwo niebezpieczeństwa 0,94, co kwalifikuje szyny do wymiany. Wadą defektoskopii ultradźwiękowej jest trudność wykrywania skaz w ruchu, np. przy równoczesnym połączeniu obu defektoskopów na wózku pomiarowym.

5.8. Uwagi końcowe

Dotychczasowe rozważania i próby wykazały, że defektoskopia magnetyczna jest prostsza, tańsza i pewniejsza od prądowej. Wśród różnych odmian tej metody najlepsze wyniki daje magnesowanie wzdlużne z czujnikiem umieszczonym w środku biegunów. Stosunkowo łatwo wykrywa się tą metodą jamy i pęknięcia poprzeczne o grubości zastępczej powyżej 0,1 mm występujące w główce. Obok niewrażliwości na wstrząsy i zakłócenia główną zaletą czujnika jest możliwie równomierna czułość wykrywania skaz w całej główce. Czujnik cewkowy w zasadzie nie nadaje

się do wykrywania skaz w stopce. Czujnik na tego rodzaju skazy powinien reagować nie na stromość, lecz na wielkość i zasięg składowej poziomej zewnętrznego pola dopełniającego. W zasadzie jednak wykrywanie skaz w stopce jest możliwe. W celu zwiększenia pewności wykrywania skaz jest rzeczą celową sprawdzać wykryte skazy dodatkowo defektoskopem ultradźwiękowym.

*Politechnika Gdańska
Katedra Miernictwa Elektrycznego*

WYKAZ LITERATURY

1. Dayton I. E., Schoemaker F. C., Morzley R. F.: The Measurement of Two-Dimensional Fields. Part II. — Rev. scien. Instrument, t. 25, nr 5, maj 1954, s. 484.
2. Deutsch M.: La voiture de contrôle des rails de la SNCF. — Revue générale des chemins de fer, maj 1954.
3. Elmore W. C., Garrett M. W.: Measurement of Two-Dimensional Fields. Part I. — Rev. scien. Instrument, t. 25, nr 5, maj 1954, s. 480.
4. Jeżewski, Szklarski L., Kawecki: Badania magnetycznych pól rozproszenia pochodzących od uszkodzeń w linach drucianych w związku z elektromagnetyczną metodą badania lin. — Arch. Elektrot., II, nr 2, 1954, s. 127.
5. Karpow F. M.: Wremiennaja instrukcja po uchodu i obskuzыwaniu welo-diefektoskopow sistemu izoprietatiela F. M. Karpowa. — G.T.Ž.I., Moskwa 1941.
6. Kuryłowicz J.: Sonda permalojowa w zastosowaniu do badań materiałów magnetycznych. — Pomiarы, Automatyka, Kontrola, nr 1, 1957, s. 2.
7. Kuhne R.: Magnetfeldmessung mit Eisenkernmagnetometer nach dem Oberwellenverfahren. — A.T.M., 1952, V, 392-1.
8. Samow A.: Wagon defektoskop. — Żel. Transport, sierpień 1956.

А. СПИХАЛЬСКИ

ВЛИЯНИЕ ИЗЪЯНОВ НА РАСПРЕДЕЛЕНИЕ ПОЛЯ В ЖЕЛЕЗНОДОРОЖНОЙ РЕЛЬСЕ И УКАЗАНИЯ ДЛЯ МАГНИТНОЙ ДЕФЕКТОСКОПИИ (части 4 и 5)

4. СВОЙСТВА МАГНИТНОЙ ДЕФЕКТОСКОПИИ ПО СРАВНЕНИЮ С ТОКОВОЙ

Резюме

Магнитная дефектоскопия с продольным полем требует значительно более лёгкую и дешёвую измерительную аппаратуру, чем токовая дефектоскопия. Особенности магнитного метода отличающимися от токового являются: изменяемость магнитной проницаемости, проницаемость узкого дефекта в виде воздушной щели, а также проницаемость среды окружающей рельсы с дефектом. Из-за трудности экспериментального исследования каждого влияния в отдельности применен здесь, соответственно приспособленный, нумерический метод.

Влияние переменной магнитной проницаемости оказалось весьма полезным по сравнению с токовым полем, принимая положение, что индукция в рельсе немного превышает значение, соответствующее максимальной проницаемости магнитной рельсовой стали.

Магнитная проводимость узкой щели сокращает степень нестационарности поля. Сокращение это не превышает однако 50% при незначительных дефектах, ниже 15% сечения стержня толщиной более чем в 0,1 мм.

Магнитная проницаемость окружающей среды приводит к вытеснению линий индукции вне рельсы.

Возникает дополнительное внешнее магнитное поле, зависимое от распределения степени неоднородности поля при поверхности. Структура такого поля определена при наличии образцовых дефектов в головке, а также в основе. Вертикальная и горизонтальная составляющие дополнительного поля представлены графически в зависимости от расстояния.

5. РАБОТА ИНДУКТИВНОГО ДАТЧИКА В МАГНИТНОМ ПОЛЕ ПРИ РЕЛЬСОВОЙ ДЕФЕКТОСКОПИИ

В магнитном методе с продольным потоком обнаруживающий датчик находится при поверхности рельсы в середине под электромагнитом. Он подвергается действию различных внешних помех. Меньше всего чувствителен по отношению к помехам индуктивный разностный датчик с электрическим и магнитным экранами без сердечника с горизонтальными катушками. Он реагирует на наклон горизонтальной составляющей внешнего дополнительного поля, какое выступает в близости дефекта.

На искусственных дефектах была проведена проверка чувствительности обнаружения, выбор оптимальной индукции в шине, влияние относительной величины дефекта и его положение в головке рельсы. Приведены указания для получения равномерной чувствительности обнаружения дефектов в головке.

Эксперименты произведенные на пути дали возможность введения классификации дефектов на случайные, повторяемые и опасные. Определены показатели неоднородности головки и плотности распределения дефектов в головке. Импульс больше импульса образцового дефекта в головке случался в среднем 4 раза на километр. Дефекты превышающие 10% сечения шины были обнаружены в значительно меньшем количестве. Причиной этого нежелательного явления была довольно значительная неравномерность обнаружения дефектов индуктивного датчика.

Произведена критика индуктивного датчика и приведены указания для построения намагничивающей системы сигнализации и контроля обнаружения дефектов.

A. SPICHALSKI

INFLUENCE OF FLAWS ON THE FIELD DISTRIBUTION IN THE RAILWAY BAR AND DIRECTIONS FOR MAGNETICAL FLAW DETECTION (parts 4 and 5)

Summary

4. PROPERTIES OF MAGNETIC DEFECTOSCOPY COMPARED WITH DETECTION BY CURRENT METHOD

The magnetic defectoscopy implementing lengthwise field requires a much lighter and cheaper arrangement if compared with the current defectoscopy. Such peculiarities as varying magnetic permeability, permeability through a thin flaw of a kind of air-crack and the permeability through medium surrounding the rail

with a flaw, differ the magnetic method from that of current. Due to the difficulty associated with the experimental investigation of the influences exercised by the particular factors separately a modified numeric method was used.

Under assumption that the induction in rail slightly exceeds the value corresponding to the maximum magnetic permeability in the steel rail, the effect of varying magnetic permeability turn out to be more valuable than the current field.

Though the magnetic permeability of the thin crack reduces the degree of the field non-uniformity this however does not exceed 50 percent for the small cracks below 15 percent of the cross-section when the crack is thicker than 0,1 mm.

The magnetic permeability of the surrounding medium forces the induction lines out of the rail.

The additional magnetic outer field which is then emerging depends on the distribution of degree of the field non-uniformity close to the surface. For the standard flaws detected in the rail head and base the structure of such a field has been determined. The horizontal and vertical components of the additional outer field are presented in function of the distance graphically.

5. MAGNETIC COIL DETECTION PRACTICE IN DEFECTOSCOPY OF RAILS

In magnetic flaw detection method with the longitudinal flux the coil detector is placed close to the rail surface right in the middle under the electro-magnet. It should be pointed out that the coil detector is liable to the diverse outer disturbances though the differential coil detector being shielded electrically and magnetically proved to be the least effected. Its response to the slope of the horizontal component of the additional external field emerging in the vicinity of flaw.

The sensitivity of detection, the selection of the optimum induction in the rail, the effect of the relative magnitude of flaw and the flaw position in the rail head have been checked making use of artificially made flaws. The indications as to the achieving of the uniform sensitivity in detecting the flaws in the rail head are given.

The experiments carried out for the railway track showed that the detected flaws may be subdivided into the casual, repetitive and dangerous type. The non-uniformity coefficient of the rail head and the density of the distribution of flaws within the head have been determined. Four times per kilometer as an average occurred the impulse stronger than that in the standard flaw of the rail head.

It would be worthwhile to notice that the flaws which exceeded 10 percent of the rail cross-section have occurred in much lesser amount.

This undesirable phenomenon is caused by considerable irregularity of the magnetic flaw detection with the coil detector.

The critical observations about the magnetic coil detector are made and the indications showing the way to the construction of the magnetizing, monitoring and controlling device are given.

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 VII · ZESZYT 2

K W A R T A L N I K

PAŃSTWOWE WYDAWNICTWO NAUKOWE
WARSZAWA 1961

RADA REDAKCYJNA

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

KOMITET REDAKCYJNY

Redaktor Naczelny
PROF. WITOLD NOWICKI

Sekretarz
JULIUSZ MIERZEJEWSKI

ADRES REDAKCJI

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

PAŃSTWOWE WYDAWNICTWO NAUKOWE
Warszawa, Miodowa 10

Nakład 640 (502 + 138)	Oddano do składania 15. III. 1961
Ark. wyd. 11, ark. druk. 9 + 3 wkl.	Podpisano do druku we wrześniu 1961
Papier druk. sat. kl. III, 80 g, 70×100	Druk ukończono we wrześniu 1961
Zamówienie nr 496/61	Cena zł 25.— S-34

Drukarnia im. Rewolucji Październikowej, Warszawa

TADEUSZ SCHWARTZ

Wyznaczanie wielkości charakteryzujących stalownicze urządzenia łukowe

Rękopis dostarczono 8.7.1960

Wychodząc z założenia, że prawidłowe stany pracy elektrycznego pieca łukowego są możliwe do uzyskania przy znajomości poszczególnych, charakteryzujących piec wielkości, podano sposoby ich wyznaczania. Omówiono wyznaczanie czasu roztopiania złomu, zużycia właściwego energii elektrycznej, przelotności pieca, oporności R i X obwodu zasilającego i wyznaczanie roboczego zakresu prądowego pieca. Do wyznaczania wymienionych wielkości wybrano metody możliwie proste, oparte na nieskomplikowanych pomiarach, możliwych do wykonania w elektrostalowniach przez załogi hutnicze.

1. WSTĘP

Operując w dalszym ciągu terminem „stalownicze urządzenie łukowe” będziemy mieli na myśli zestaw trójfazowy złożony z transformatora i pieca (typu Héroulta) o wyprawie zasadowej, przeznaczony do pełnego¹⁾ cyklu produkcji stali szlachetnych.

Urządzenie takie mimo swej prostoty konstrukcyjnej podlega złożonym prawom elektrotermicznym i jedynie w rękach obsługi obznajmionej z tymi prawami może służyć za precyzyjne narzędzie produkcyjne.

W cyklu pracy takiego urządzenia jako najważniejszą fazę roboczą wyróżnia się okres roztopiania złomu stalowego, ponieważ w okresie tym jest pochłaniana maksymalna część energii elektrycznej pobieranej przez urządzenie w trakcie całego cyklu pracy. Tak na przykład przy wytopie stali w czterotonowym piecu roztopianie wsadu może trwać jedną godzinę i czterdzieści minut przy średnim poborze mocy wynoszącym 1200 kW, utlenianie — 45 minut przy poborze mocy 680 kW, a pozostałe fazy przez 30 minut przy poborze mocy 600 kW. Z tego wynika, że przeszło 70% pobranej w ciągu całego cyklu energii zużywa się w trakcie roztopiania. Dlatego też znajomość stanów pracy urządzenia możliwych do utrzymania w trakcie roztopiania złomu jest podkreślana w wielu pracach naukowych dotyczących pieców łukowych jako szczególnie ważna. Ograniczając podane dalej uwagi do tej właśnie fazy roboczej trzeba podkreślić, że wskazania odnoszące się do niej dają się rzutować na następne fazy, objęte

1) W odróżnieniu od cyklu niepełnego, stosowanego w procesach duplex.

ogólną nazwą rafinacji stali, to znaczy, że dają one orientację w zakresie całego cyklu produkcyjnego.

Okres roztopiania wsadu w kadzi piecowej pochłania energię, której 85 ... 90% doprowadza się do urządzenia w postaci elektrycznej, szczególnie interesująca jest zatem umiejętność ustalania prawidłowych parametrów elektrycznych tego okresu, wpływających na wartości wskaźników, które decydują o sposobie prowadzenia procesu technologicznego. Na wskaźniki te zresztą wpływają również, jak wykazemy, parametry cieplne.

Ponieważ okres roztopiania ma zasadniczo na celu rozgrzanie znajdującego się w kadzi metalu do temperatury wyższej od punktu topliwości, przeto mogłoby się pozornie wydawać, że aby zredukować do minimum straty energii odprowadzanej w postaci ciepła do otoczenia z gorących ścian pieca, powinno się w tym okresie dążyć zawsze do osiągnięcia najkrótszego czasu trwania roztopiania wsadu drogą doprowadzania do pieca możliwie dużej mocy. W rzeczywistości jednak sprawa nie przedstawia się tak prosto. Powodem jej skomplikowania jest zjawisko rozbieżności dwu podstawowych wskaźników pracy pieca, którymi jak wiadomo [4] są:

zużycie właściwe energii elektrycznej, e , kWh/t,
przelotność urządzenia, g , t/h.

Rozbieżność tych wskaźników polega na tym, że prowadząc roztopianie z możliwie największą prędkością uzyskiwaną w rezultacie pobierania maksymalnej mocy uzyskuje się wprawdzie największą przelotność, otrzymując tym samym największą produkcję, bynajmniej jednak nie osiąga się wówczas najmniejszego zużycia energii na jednostkę ciężarową produktu wyrażoną w kilowatogodzinach na tonę.

Najkorzystniejsze warunki energetyczne i najkorzystniejsze warunki urobkowe osiąga się w stalowniczym urządzeniu łukowym przy dwu zupełnie różnych elektrycznych stanach pracy urządzenia. Z tego punktu widzenia operowanie terminem „optymalny stan pracy pieca łukowego” nie jest godne zalecenia, może bowiem wywołać nieporozumienie w przypadku braku definicji tego terminu. Za optymalny stan pracy pieca łukowego może być uznany zarówno stan elektryczny, w którym osiąga się minimalne zużycie właściwe, e_{min} , jak również stan gwarantujący maksymalną przelotność, g_{max} , ponieważ decydujący jest cel, a cel ten wytyczają względy ekonomiczne i koniunkturalne. W niektórych okresach gospodarczych mogą przeważać dążenia do intensyfikacji produkcji z aprobatą wzrostu kosztów energii spowodowanych większym zużyciem właściwym, w innych, na przykład przy wyraźnych niedoborach energii elektrycznej, mogą się okazać decydujące tendencje oszczędzania energii elektrycznej, prowadzące do konieczności zachowania minimalnego zużycia jednostkowego energii elektrycznej, dające się zrealizować tylko przy przelotności mniejszej od największej możliwej do uzyskania na danym urządzeniu.

Na tle wskazanej rozbieżności dwu najważniejszych wskaźników eksploatacyjnych charakteryzujących produkcję elektrostali staje się zrozumiała potrzeba opanowania umiejętności wyznaczania stanów pracy pieca łukowego, związana ściśle z koniecznością wyznaczania wielkości charakteryzujących stalownicze urządzenie łukowe.

2. OGÓLNE WYTYCZNE ROZTAPIANIA ZŁOMU

Przed opisem sposobów badania stanów pracy pieca łukowego i sposobów wyznaczania jego parametrów wydaje się celowe zwrócenie uwagi na pewne ogólne zalecenia mające na celu zachowanie prawidłowości procesu roztapiania wsadu w stalowniczym piecu łukowym.

Pomijając wskazówki dotyczące złomu, metod jego ładowania i sposobów wypełniania nim kadzi piecowej, co — wydaje się — jest dość powszechnie znane, do wskazań ogólnych zaliczają się następujące wskazania.

W opisach technologii roztapiania złomu w stalowniczym piecu łukowym można się spotkać z uwagami wskazującymi na to, że roztapianie rozpoczyna się przy niższym wtórnym napięciu transformatora piecowego od napięcia najwyższego możliwego do zastosowania. Praca przy obniżonym napięciu, odpowiadającym 60...75% mocy znamionowej, przewidywana jest na czas około 10...15 minut, po czym wymaga się podniesienia elektrod i zmiany stopnia napięciowego na najwyższy.

Rozpoczynanie roztapiania przy obniżonym napięciu piecowym jest motywowane zarówno warunkami termicznymi, stwarzanymi przez początkowo nie osłonięty łuk elektryczny, jak też i warunkami elektrycznymi, wywołanymi małym stopniem zjonizowania przestrzeni łuku, co — jak wiadomo — wpływa ujemnie na jego stabilność. Są jednak i argumenty przeciwne. Mianowicie, w praktyce, w piecach o pojemności wsadowej od 10 do 40 ton w okresie roztapiania złomu elektrody pograżają się we wsadzie, zależnie od stopnia jego zwartości, a więc zależnie od jego ciężaru objętościowego, z prędkością wynoszącą 2,5...6,0 m/h. W takich warunkach już po upływie pięciu minut od początku roztapiania elektroda jest zanurzona w wytworzonym przez nią kraterze na głębokość 200...500 mm. Gdy zagłębienie elektrod we wsadzie wyniesie 100...200 mm, promieniowanie łuku nie jest już niebezpieczne dla sklepienia i ścian kadzi piecowej. Początkowo łuk jest bowiem krótki, krater zaś ciasny. Stwarza to dobre warunki ekranowania promieniowania. Te przesłanki wskazują na niecelowość obniżania napięcia na początku okresu roztapiania wsadu. Są także dalsze, przemawiające za tym samym poglądem. Początkową niestabilność łuku elektrycznego wywołuje wysoki potencjał jonizacji atmosfery łuku. Wprowadzenie w obszar łuku związków o niskim potencjale jonizacji, na przykład wapna lub innych związków wapnia, wpływa dodatnio na stabilność łuku i pozwala go wydłużyć, co

wpływa korzystnie na proces roztopiania metalu, likwiduje dymienie pieca i ułatwia pracę samoczynnych regulatorów elektrod.

Omówione względy powodują, że w nowych poglądach przeważa tendencja unikania obniżonego napięcia na początku okresu roztopiania wsadu, co jest uzasadnione zwłaszcza w odniesieniu do urządzeń stalowniczych wyposażonych w nowoczesne, szybko działające regulatory elektrod.

Postępowe załogi elektrostalowni wyposażonych w piece z nowoczesnymi regulatorami elektrod prowadzą roztopianie od początku przy najwyższym wtórnym napięciu transformatora piecowego. W rezultacie stosowania takiej metody można uzyskać skrócenie czasu roztopiania wsadu o 3...5%, ponieważ rozporządza się wówczas większą ilością energii i ponadto likwiduje się stratę czasu około dwu minut wywoływaną przełączaniem napięcia.

Przy dążeniu do skracania okresu roztopiania wsadu jest celowe utrzymywanie jak najdłużej pracy pieca przy najwyższym stopniu napięciowym. Zbyt niemu przedłużaniu tego stanu w okresie roztopiania stoją jednak na przeszkodzie względy termiczne.

Promieniowanie łuku na sklepienie i na ściany kadzi staje się bowiem niebezpieczne wówczas, gdy skutkiem postępującego procesu roztopiania złomu odkrywają się nad kąpielą łuki elektryczne. Temperatura powierzchni wewnętrznych kadzi nad zwierciadłem wsadu zaczyna wtedy szybko wzrastać, a to pociąga za sobą konieczność obniżenia mocy łuku za pośrednictwem zmiany stopnia napięcia. Przy dążeniu do skracania okresu roztopiania obniżenie napięcia powinno nastąpić jak najpóźniej, jednak z zachowaniem bezpieczeństwa termicznego. Decydująca jest w tym przypadku temperatura topliwości materiału ogniotrwałego i z tego względu materiały ogniotrwałe powinny być jak najbardziej wytrzymałe termicznie. Na wybór prawidłowego momentu obniżenia napięcia ma wpływ temperatura panująca w najgorętszym miejscu wyprawy pieca. Operowanie średniówkami nie zawsze jest dopuszczalne, ponieważ w przypadku asymetrii elektrycznej pieca łuki pieców średniej i dużej mocy mogą wydzielać ilości ciepła różniące się o 15...30%. Z tych względów racjonalny wybór momentu obniżenia napięcia w końcowym okresie roztopiania wsadu wymaga najczęściej oparcia się na wynikach pomiarów temperatury.

Im jest wyższa dopuszczalna temperatura ścian i sklepienia pieca, tym dłużej można korzystać z najwyższego stopnia napięciowego, wprowadzając większe ilości energii do wsadu i skracając tym samym czas roztopiania wsadu.

Gdy temperatura w najbardziej nagrzanej części wyprawy piecowej dochodzi do dopuszczalnej granicznej wartości, która na przykład dla dynasu wynosi 1700°C, wówczas skokowe obniżenie napięcia piecowego zmienia kształt krzywej wzrostu temperatury w ten sposób, że znów po upływie pewnego interwału czasowego występuje potrzeba dalszego obni-

zenia napięcia. Nie jest przy tym obojętne stopniowanie napięcia. Warunki pracy są tym bardziej korzystne, im większa ilość stopni napięcia jest do dyspozycji. Posługiwanie się tylko dwoma stopniami napięcia wydłuża czas roztopienia. Operowanie 4...5 stopniami napięcia, ze względu na większe zbliżenie takiej zmienności do ciągłej regulacji mocy, pozwala skrócić czas roztopienia o około 3...4 minut.

Po omówieniu podanych wyżej wskazań ogólnych przejdziemy z kolei do zobrazowania przebiegów zmienności dwu wymienionych na wstępie wskaźników eksploatacyjnych rozpoczynając od przelotności.

3. CHARAKTERYSTYKA CZASU ROZTAPIANIA WSADU

Przelotność w okresie roztopienia wsadu jest ilością roztopionego metalu przypadającą na jednostkę czasu. Z definicji tej wielkości wynika, że przelotność jest funkcją czasu, w ciągu którego odbywa się przebieg roztopienia wsadu załadowanego do kadzi piecowej. Przy tej samej pojemności wsadowej pieca przelotność jest tym większa, im krótszy jest czas roztopienia — i odwrotnie. Z tego względu wyznaczanie czasu roztopienia wsadu w postaci charakterystyki jego zmienności ma istotne znaczenie dla oceny przelotności, która interesuje użytkownika pieca łukowego.

Czas trwania roztopienia wsadu można wyznaczyć dzieląc wartość energii elektrycznej zużytej w tym czasie przez wartość mocy pobieranej w tym czasie. Jeżeli moc pobierana w tym czasie ma kilka różnych wartości, działania należy wykonać dla każdej z nich, rozbijając jednocześnie całkowity okres czasu roztopienia na interwały, w których przypadają poszczególne wartości mocy, albo też można operować średnią wartością mocy dla całego okresu roztopienia.

Zasada wyznaczania czasu roztopienia nie ulegnie zmianie, jeżeli w dalszym ciągu przyjmiemy dla całego okresu roztopienia T jednakową wartość mocy pobieranej przez piec.

Przy wyznaczaniu czasu T najwygodniej jest operować energią zużytą wyłącznie na nagrzewanie wsadu, czyli tak zwaną energią *użyteczną* (inaczej: energią teoretyczną), wyrażającą ilość kilowatogodzin potrzebną na roztopienie jednej tony złomu określonego gatunku w warunkach idealnych, czyli bez strat ciepła. Nazywając tę wielkość *teoretycznym zużyciem właściwym*, oznaczmy ją przez e_t .

Jeżeli oznaczmy przez Δt całkowity przyrost temperatury wsadu w trakcie okresu jego roztopienia w $^{\circ}\text{C}$, wówczas odejmując od końcowej temperatury t_k początkową temperaturę t_p wsadu otrzymujemy $\Delta t = t_k - t_p$. Jeżeli ponadto oznaczmy średnie ciepło właściwe wsadu przez c_s , w $\text{kcal/kg}^{\circ}\text{C}$, wówczas teoretyczne zużycie właściwe możemy obliczyć na podstawie prostej zależności:

$$e_t = 1,16 \cdot c_s \cdot \Delta t \quad [\text{kWh/t}]. \quad (1)$$

Najczęściej w praktyce elektrostalowniczej wartość e_t jest przyjmowana jako równa 340 ... 360 kWh/t, z tym że wartość niższa jest częściej używana.

Mnożąc zużycie e_t przez ciężar G ton wsadu piecowego otrzymujemy teoretyczną ilość energii potrzebną na roztopienie tej ilości:

$$E_t = e_t \cdot G \quad [\text{kWh}], \quad (2)$$

a dzieląc z kolei E_t przez fazową moc użyteczną P_u [kW] otrzymujemy czas roztopiania wsadu

$$T = \frac{E_t}{3 P_u}. \quad (3)$$

Jednak moc użyteczna P_u jest mocą, jaka pozostaje do dyspozycji po odjęciu od mocy P_1 łuku elektrycznego mocy P_c kW straconej w postaci ciepła odpływającego od pieca do otoczenia

$$P_u = P_1 - P_c, \quad (4)$$

wobec czego poszukiwany czas roztopiania wsadu może być wyrażony jako

$$T = \frac{E_t}{3(P_1 - P_c)} [\text{h}]. \quad (5)$$

Z podanego wyrażenia (5) widać, że charakterystykę czasu roztopiania można łatwo uzyskać na podstawie charakterystyk mocy łuku i mocy strat cieplnych pieca, mianowicie dzieląc stałą wartość E_t przez różnicę rzędnych krzywych P_1 i P_c . Podzielenie przez 3 wynika stąd, że wartości P_u , P_1 , P_c traktuje się jako odnoszące się do jednej fazy pieca. W praktyce za P_c przyjmuje się średnią wartość strat cieplnych wyznaczonych drogą pomiarów, do czego jeszcze wrócimy w dalszym ciągu, i zakłada się, że ta średnia nie ulega zmianie w trakcie okresu roztopiania wsadu.

Krzywą P_1 w zależności od prądu I łuku elektrycznego można wykreślić na podstawie zależności

$$P_1 = I \sqrt{U^2 - I^2 X^2} - I^2 R, \quad (6)$$

w której:

U — wartość skuteczna napięcia piecowego, w V, stała dla wybranego zacze pu napięciowego,

R — stała oporność czynna, w omach, obwodu zasilającego piec,

X — stała oporność indukcyjna, w omach, obwodu zasilającego piec.

Wszystkie wielkości traktowane są jako fazowe.

Do wykreślenia krzywej $P_1 = f(I)$, jak z tego widać, konieczne jest zmierzenie wielkości R i X , do czego także jeszcze wrócimy w dalszym ciągu.

Z zależności (5) wynika bezpośrednio, że najkrótszy czas roztopiania uzyskuje się przy takim prądzie łuku, przy którym różnica $P_1 - P_c$ staje się największa, a więc wówczas, gdy moc pieca (łuku) osiąga wartość największą. Wartość tego prądu można wyznaczyć wobec tego z warunku:

$$\frac{dP_1}{dI} = 0, \quad (7)$$

z którego, po wykonaniu działań, otrzymuje się wyrażenie na prąd I_{max} gwarantujący najkrótszy czas roztopiania w postaci

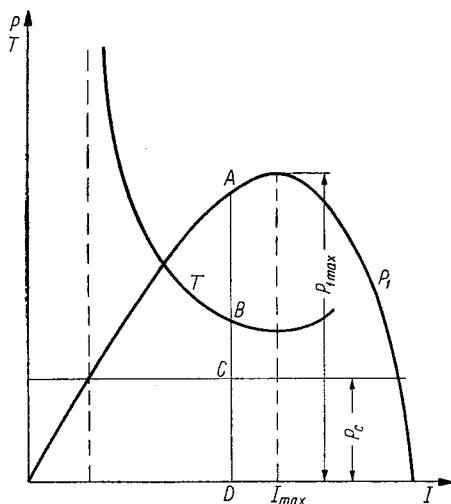
$$I_{max} = \frac{1}{\sqrt{2}} \cdot \frac{U}{X} \sqrt{1 - \frac{R}{\sqrt{R^2 + X^2}}}. \quad (8)$$

Ze wzoru (8) wynika, że prąd I_{max} może być obliczony na podstawie tego wzoru jedynie wówczas, gdy znane są:

U — napięcie fazowe zasilające piec (wtórne napięcie fazowe transformatora, w V),

R i X — fazowe oporności, odpowiednio czynna i bierna, w omach, obwodu zasilającego łuk elektryczny.

Na rys. 1 jest pokazany charakter przebiegu zależności $T = f(I)$. Wyznaczenie rzędnej BD krzywej obrazującej zmiany czasu T wymaga od-



Rys. 1. Charakterystyka czasu roztopiania wsadu, $T = f(I)$

jęcia od rzędnej AD , wyrażającej moc łuku, rzędnej CD , wyrażającej straty ciepłne. Tak otrzymany dzielnik i dzielna E_t dają iloraz będący miarą czasu T wyrażonego przez rzędną BD .

Wartość liczbowa najkrótszego czasu T_{min} , osiąganego zawsze przy prądzie I_{max} , zależy przy określonej maksymalnej mocy łuku P_{1max} , jak wynika ze wzoru (5), od wielkości mocy strat ciepłnych P_c . Im mniejsze

są te straty P_c , tym jest mniejsza wartość minimalnego czasu T_{min} . Widać stąd, jak ważną w praktyce jest dbałość o zachowanie podczas pracy pieca jak najmniejszych strat cieplnych.

Oprócz metody analitycznego wyznaczania prądu I_{max} , opartej na zależności (8), istnieją także jeszcze inne. W następnym rozdziale przedstawimy jedną z nich, szczególnie prostą.

4. WYZNACZANIE PRĄDU I_{max} ZA POMOCĄ ZWARCIA

Wyznaczanie prądu I_{max} , przy którym osiąga się najkrótszy czas roztopiania wsadu, na podstawie pomiaru dokonanego w czasie zwarcia, jest szczególnie interesujące z tego względu, że metoda ta nie posilkuje się wielkościami R i X oporności fazowych, które, jak już wiadomo, są niezbędne przy obliczaniu I_{max} za pomocą wzoru (8).

Postępowanie przy wyznaczaniu tą metodą [3] prądu I_{max} jest dość proste i może być wobec tego wykonane w każdej elektrostalowni bez uciekania się do pomocy z zewnątrz. Należy mianowicie wyznaczyć w stanie zwarcia moc, prąd i napięcie. Próba może być łatwo wykonana po roztopieniu metalu, a więc w stanie, który pozwala zanurzyć końce elektrod w kąpeli metalowej na czas kilku sekund, wystarczający do dokonania odczytu na przyrządach pomiarowych. Bliższe szczegóły dotyczące warunków zwarcia będą podane w jednym z następnych punktów, poniższe uwagi ograniczymy wobec tego tylko do zależności prowadzących do określenia wartości prądu I_{max} .

Zmierzone wielkości fazowe związane są ze sobą, jak wiadomo, zależnością

$$P_z = U \cdot I_z \cdot \cos \varphi_z, \quad (9)$$

z której przy znanych wielkościach mocy zwarcia P_z , prądu zwarcia I_z oraz napięcia U , oblicza się współczynnik mocy zwarcia

$$\cos \varphi_z = \frac{P_z}{U \cdot I_z}$$

i kąt przesunięcia fazowego przy zwarcu

$$\varphi_z = \arccos \frac{P_z}{U \cdot I_z}; \quad (10)$$

w dalszym ciągu należy oprzeć się na zależnościach odnoszących się do maksymalnej mocy łuku. Moc łuku o oporności r omów występującego w obwodzie o stałych opornościach: czynnej R i biernej X , można wyrazić w postaci wzoru

$$P_1 = I^2 \cdot r = \frac{U^2 \cdot r}{r^2 + 2R \cdot r + Z_z^2}, \quad (11)$$

w którym

$$Z_z = \sqrt{R^2 + X^2} \quad (12)$$

oznacza oporność pozorną obwodu zasilającego łuk.

Narzucając funkcji (11) warunek

$$\frac{dP_1}{dr} = 0,$$

przy którym występuje maksymalna moc łuku $P_{1\max}$, otrzymuje się po wykonaniu działania ten sam warunek w postaci zależności

$$r = Z_z \quad (13)$$

oznaczający, że maksymalna moc łuku występuje wówczas, gdy oporność r łuku elektrycznego staje się równa oporności obwodu zasilającego ten łuk. Warunek (13) można zobrazować za pomocą wykresu wektorowego pokazanego na rys. 2.

Na wykresie tym wektor U oznacza napięcie zasilające, które podczas pracy urządzenia łukowego rozkłada się, jak wiadomo, na spadek napięcia występujący na oporności r łuku, wynoszący $I \cdot r$ i będący w fazie z prądem I łuku elektrycznego oraz na spadek napięcia występujący na oporności Z_z i przesunięty względem prądu o kąt φ_z . Te dwa spadki napięcia są zobrazowane na wykresie wektorami OA i OB . Przy spełnionym warunku $r = Z_z$ mamy także spełniony warunek $I \cdot r = I \cdot Z_z$ i wtedy kąty $\sphericalangle COA$ i $\sphericalangle COB$ są sobie równe, a ponieważ

$$\sphericalangle BOA = \varphi_z$$

wobec tego

$$\sphericalangle COA = \varphi_0 = \frac{1}{2} \sphericalangle BOC = \frac{1}{2} \varphi_z,$$

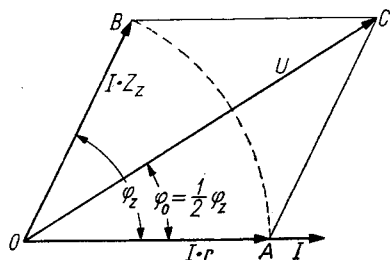
co oznacza, że kąt przesunięcia fazowego φ_0 między prądem i napięciem przy maksymalnej mocy łuku ($r = Z_z$) jest równy połowie kąta fazowego φ_z występującego przy zwarcu. Opierając się na tej zależności przy znanej wartości kąta φ_z (10) można łatwo określić prąd $I_{\max}$.

Mianowicie, jak wiadomo z elektrotechniki, w obwodzie elektrycznym zasilanym napięciem $U = \text{const}$, w którym występuje niezmienna wartość oporności indukcyjnej X i zmienna wartość oporności czynnej $R_f = R + r$, jak to jest w obwodzie każdej fazy urządzenia łukowego, koniec wektora prądu leży na obwodzie półkola, jak wskazuje rys. 3.

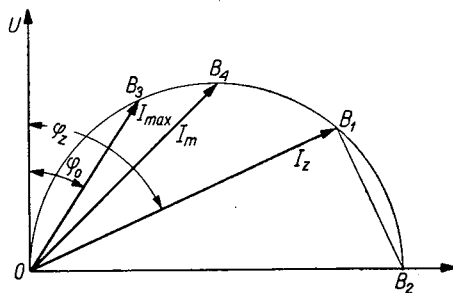
Wykres ten, rys. 3, sporządzamy rysując pod kątem φ_z do wektora napięcia U odcinek OB_1 wyrażający wektor prądu o wartości skalarnej równej znanemu prądowi I_z . Prowadząc prostą prostopadłą do OB_1 przez punkt B_1 otrzymujemy w przecięciu jej z prostą poziomą punkt B_2 . Odcinek OB_2 jest średnicą półkola prądów. Po wykreśleniu tego półkola do

wyznaczenia prądu $I_{\max}$ wystarczy więc narysowanie wektora OB_3 poło-
wiącego kąt φ_z i zmierzenie jego długości, która w tej samej skali, w której
wykreślany był wektor OB_1 , daje natężenie prądu $I_{\max}$ w amperach.

Na wykresie podanym na rys. 3 narysowany jest dodatkowo wek-
tor OB_4 , reprezentujący prąd opóźniony względem napięcia o 45° . Jest
to prąd, przy którym urządzenie łukowe (ale nie piec) pobiera największą



Rys. 2. Wykres wektorowy spadków napięć przy maksymalnej mocy łuku



Rys. 3. Wyznaczanie prądu $I_{\max}$ za pomocą wykresu kołowego

moc z sieci zasilającej. Praca przy tym prądzie I_m jest niedopuszczalna, ponieważ — mimo że jest on większy od prądu $I_{\max}$ — wytwarza w piecu mniej ciepła niż prąd $I_{\max}$ pogarszając ponadto współczynnik mocy urządzenia i powodując zwiększenie o kilkanaście procent zużycia elektrod w stosunku do zużycia występującego przy prądzie $I_{\max}$.

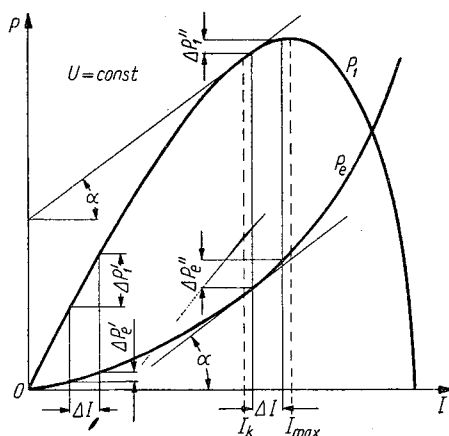
Prąd $I_{\max}$, wyznaczony wskazaną lub też jedną z innych metod, określa warunek osiągnięcia maksymalnej mocy pieca i minimalnej wartości czasu roztopiania wsadu. Jeżeli więc celem produkcji jest przetapianie maksymalnej ilości stali w ciągu jednostki czasu, czyli praca przy maksymalnej przelotności, wówczas w okresie roztopiania wsadu samoczynny regulator powinien być nastawiony na utrzymywanie prądu łuku równego tej właśnie wartości $I_{\max}$.

Prawidłó to wymaga komentarzy. Po pierwsze należy sobie zdać sprawę z tego, że nie we wszystkich przypadkach wartość natężenia prądu $I_{\max}$ jest osiągalna w praktyce. Ocenę urządzenia pod względem osiągalności prądu $I_{\max}$ podamy w punkcie ostatnim. Po drugie należy liczyć się z tym, że w przypadku osiągalności prądu $I_{\max}$ utrzymywanie go podczas wytopu stali nie jest korzystne z punktu widzenia gospodarki energią elektryczną, co wyjaśnimy rozpatrując w następnym rozdziale tak zwany krytyczny prąd łuku elektrycznego.

5. KRYTYCZNA WARTOŚĆ PRĄDU ŁUKU

Rozpatrzmy wykres podany na rys. 4. W odniesieniu do określonego napięcia zasilającego $U = \text{const}$ są na nim narysowane: charakterystyka mocy P_1 łuku i charakterystyka mocy $P_e = I^2 \cdot R$ traconej na oporności

czynnej R obwodu zasilającego, obie w funkcji prądu I łuku. Przebieg ich wskazuje, że przy wzroście prądu aż do wartości $I_{\max}$ wzrasta zarówno moc łuku P_1 , jak też moc strat elektrycznych P_e . Jednakże przy prądzie $I \ll I_{\max}$ przyrost $\Delta P_1'$ mocy łuku wywołany wzrostem prądu łuku o ΔI jest większy od przyrostu $\Delta P_e'$ strat elektrycznych wywołanego tym sa-



Rys. 4. Wyznaczanie krytycznej wartości prądu łuku

my przyrostem prądu ΔI . Natomiast przy prądzie I niewiele mniejszym od wartości $I_{\max}$, mianowicie począwszy wzwyż od pewnej wartości I_k , przebiegi zmieniają się w ten sposób, że przyrost $\Delta P_e''$ strat staje się większy od przyrostu $\Delta P_1''$ mocy łuku. Oznacza to, że przy prądach łuku większych od krytycznej wartości I_k każdy kilowatt wzrostu mocy pieca (łuku) można osiągnąć jedynie kosztem powstałej jednocześnie straty mocy P_e przekraczającej jeden kilowatt. Okazuje się więc, że praca pieca przy prądzie większym od I_k jest mniej korzystna z punktu widzenia gospodarki energią elektryczną od pracy pieca przy prądach mniejszych od wartości krytycznej I_k .

Rozpatrując wykres z rys. 4 nietrudno zauważyć, że prąd I_k występuje wówczas, gdy styczne do obu charakterystyk osiągają ten sam kąt nachylenia: α . Ponieważ miarą tego kąta jest pochodna, wobec tego warunek wystąpienia prądu krytycznego można wyrazić w postaci

$$\frac{dP_1}{dI} = \frac{dP_e}{dI},$$

z której po podstawieniu na moc strat wartości $I^2 R$, zaś na moc P_1 wartości określonej wzorem (6), przy założeniu $U = \text{const}$, $R = \text{const}$, $X = \text{const}$ i po wykonaniu działania, otrzymuje się wartość poszukiwanego prądu

$$I_k = \frac{1}{\sqrt{2}} \frac{U}{X} \sqrt{1 - \frac{R}{\sqrt{R^2 + 1/4X^2}}} \quad (14)$$

wskazującą bezpośrednio, że zawsze

$$I_k < I_{\max}.$$

W ten sposób dochodzi się do wniosku, że praca przy prądzie mniejszym od prądu $I_{\max}$ obniża wprawdzie przelotność urządzenia łukowego w stosunku do maksymalnej przelotności, jednakże wpływa korzystnie na zużycie energii elektrycznej. Jednak najkorzystniejsze warunki z punktu widzenia jednostkowego zużycia, czyli minimalną wartość zużycia właściwego $e_{\min}$ [kWh/t] osiąga się, jak zobaczymy, nie przy prądzie I_k , lecz przy innej jego wartości, oznaczonej w dalszym ciągu przez $I_{\min}$. Wyznaczanie prądu $I_{\min}$ będzie tematem jednego z następnych rozdziałów; przed zajęciem się tym prądem należy jednak zdać sobie sprawę z przebiegu charakterystyki zużycia właściwego e , co będzie rozpatrywane w następnym punkcie.

6. CHARAKTERYSTYKA ZUŻYCIA WŁAŚCIWEGO ENERGII ELEKTRYCZNEJ

W celu wyznaczenia charakterystyki zużycia $e = f(I)$ posłużymy się już wprowadzoną wielkością teoretycznego zużycia e_t . Ponieważ zużycie e otrzymuje się dzieląc pochłoniętą przez urządzenie w trakcie roztopiania energię E [kWh] przez ciężar G roztopianego wsadu:

$$e = \frac{F}{G} \quad [\text{kWh/t}], \quad (15)$$

a na podstawie zależności (2) mamy

$$G = \frac{E_t}{e_t},$$

wobec tego

$$e = e_t \frac{E}{E_t}.$$

Występujący w tym wyrażeniu stosunek teoretycznie potrzebnej do roztopienia wsadu energii E_t do energii E rzeczywiście zużytej w trakcie roztopiania jest ogólną sprawnością procesu

$$\eta_0 = \frac{E_t}{E};$$

możemy zatem napisać

$$e = \frac{e_t}{\eta_0}. \quad (15a)$$

Operując mocami wyraża się sprawność jako iloraz mocy użytecznej P_u i mocy P pobieranej przez całe urządzenie łukowe. Ponieważ zaś całkowita

moc P jest sumą mocy użytecznej i mocy strat elektrycznych i ciepłych, mamy

$$\eta_0 = \frac{P_u}{P} = \frac{P_u}{P_u + P_e + P_c},$$

skąd

$$\frac{1}{\eta_0} = 1 + \frac{P_e + P_c}{P_u},$$

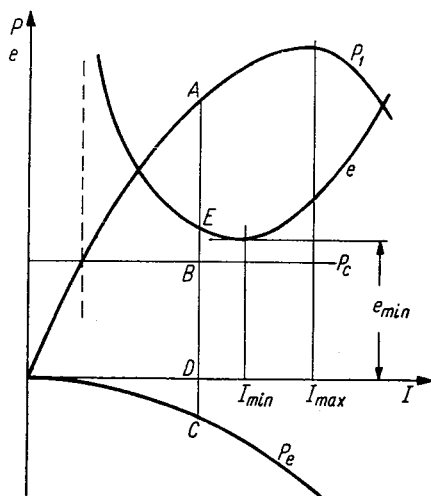
albo, z uwzględnieniem (4),

$$\frac{1}{\eta_0} = 1 + \frac{P_e + P_c}{P_1 - P_c},$$

co pozwala ostatecznie wyrazić poszukiwane zużycie w postaci

$$e = e_t \left(1 + \frac{P_e + P_c}{P_1 - P_c} \right). \quad (16)$$

Ze wzoru (16) wynika, że przy znanej wartości teoretycznego zużycia właściwego e_t (jw.: $e_t = 340 \dots 360$), krzywą zużycia właściwego w funkcji



Rys. 5. Wyznaczanie charakterystyki zużycia właściwego energii elektrycznej

prądu łuku, $e = f(I)$, czyli charakterystykę zużycia właściwego, można w prosty sposób wyznaczyć na podstawie następujących charakterystyk: mocy pieca P_1 , mocy strat elektrycznych P_e i mocy strat ciepłych P_c . Na rys. 5 są narysowane te trzy charakterystyki pomocnicze i wyznaczona na ich podstawie charakterystyka zużycia właściwego, przy czym w celu ułatwienia sumowania strat elektrycznych i ciepłych krzywą mocy strat elektrycznych podano pod osią odciętych. Zgodnie z zależnością (16) wy-

znaczenie rzędnej krzywej $e = f(I)$ przy dowolnej wartości prądu I łuku elektrycznego polega na wykonaniu działania:

$$DE = e_t \left(1 + \frac{CD + DB}{DA - DB} \right).$$

Wyznaczając punkt po punkcie krzywą $e = f(I)$ dochodzi się do przebiegu, który ma początkowo charakter opadający, a następnie po przejściu przez minimum przy pewnej wartości prądu $I_{\min}$ charakter wznoszący się ku górze.

Wartość prądu $I_{\min}$, określającego ekstremum krzywej $e = f(I)$, można obliczyć biorąc pod uwagę zależność (15a), wynika z niej bowiem, że zużycie właściwe jest funkcją wyłącznie sprawności η_0 urządzenia łukowego. Korzystając z tego, że minimum zużycia e przypada przy tym samym prądzie, przy którym występuje maksymalna sprawność, można poszukując prądu $I_{\min}$ narzucić sprawności warunek:

$$\frac{d\eta_0}{dI} = 0,$$

z którego po podstawieniach i wykonaniu działań otrzymuje się wartość poszukiwanego prądu:

$$I_{\min} = \frac{1}{\sqrt{2}} \cdot \frac{U}{X} \sqrt{\frac{2P_c}{2P_c + R \frac{U^2}{X^2}}}. \quad (17)$$

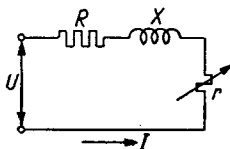
Na podstawie tego wyrażenia oraz wyrażenia określającego prąd $I_{\max}$ (8) można udowodnić, że w okresie roztopiania wsadu zawsze spełniona jest nierówność $I_{\min} < I_{\max}$. Dowodu tego nie będziemy przytaczać, ponieważ można go znaleźć w źródłach wymienionych na końcu pracy. Na podstawie podanej nierówności prądów można stwierdzić, że najmniejszą wartość zużycia właściwego $e_{\min}$ [kWh/t] osiąga się przy prądzie $I_{\min}$ mniejszym od prądu, przy którym występuje największa przełotność urządzenia łukowego $g_{\max}$. Roboczym zakresem prądowym pieca łukowego przy jego eksploatacji powinien być jedynie zakres prądów od najmniejszej wartości $I_{\min}$ do największej wartości $I_{\max}$. Praca przy innych prądach jest nieprawidłowa. Roboczy prąd stalowniczego urządzenia łukowego powinien spełniać warunek

$$I_{\min} \leq I_0 \leq I_{\max}, \quad (18)$$

to znaczy, w zależności od celu obranego przy produkcji elektrostali, powinien być równy prądowi $I_{\min}$ gwarantującemu minimalne zużycie właściwe $e_{\min}$, albo równy prądowi $I_{\max}$, zapewniającemu największą przełotność $g_{\max}$, albo wreszcie powinien mieć jedną z wartości pośrednich wybranych w obszarze od $I_{\min}$ do $I_{\max}$. Zalecenie to wymaga zastrzeżenia wysuniętego już przy okazji omawianego prądu $I_{\max}$. Mianowicie, podobnie

jak prąd $I_{\max}$ również i prąd $I_{\min}$ może okazać się w praktyce nieosiągalny. Inaczej mówiąc, może się okazać niedostępny do wykorzystania cały roboczy zakres prądowy; jeżeli tak jest, wówczas stalownicze urządzenie łukowe jest nieprawidłowo zbudowane. Oceną urządzenia z tego punktu widzenia zajmijmy się, jak wspomniano, w ostatnim punkcie pracy. Przedtem wróćmy do uwag dotyczących minimum zużycia właściwego.

Wyznaczenie dolnej granicznej wartości prądu w zakresie roboczym, to znaczy prądu oznaczonego przez $I_{\min}$, przy wybranym stopniu napięcia zasilającego $U = \text{const}$, wymaga przede wszystkim, jak wynika ze wzoru (17),



Rys. 6. Uproszczony układ obwodu jednej fazy urządzenia łukowego

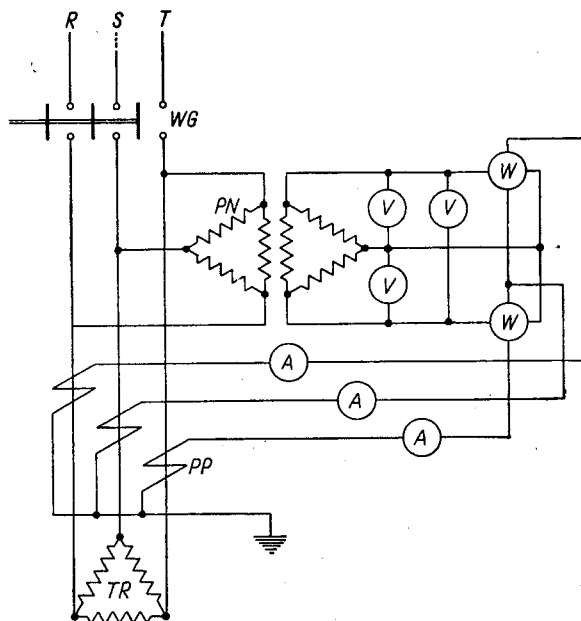
znajomości oporności czynnej R i oporności biernej X każdej fazy, która w uproszczeniu może być odtworzona obwodem pokazanym na rys. 6. Wielkości te służą jak wiadomo także do obliczenia prądu $I_{\max}$ na podstawie wzoru (8), z taką jednak różnicą, że do obliczenia prądu $I_{\max}$ są one wystarczające, natomiast do obliczenia prądu $I_{\min}$, prócz wielkości R , X niezbędna jest jeszcze znajomość strat cieplnych P_c . Ze względu na rolę, jaką oporności R i X odgrywają przy wyznaczaniu roboczego zakresu prądowego pieca użytkownik urządzenia łukowego staje wobec konieczności dysponowania metodami służącymi do ich wyznaczania.

Oporności R i X mogą być znalezione na podstawie obliczeń, jednakże zalecanie wykonania takich obliczeń, ze względu na skomplikowanie koniecznych do wykonania w tym przypadku rachunków, nie znalazłoby prawdopodobnie uznania i miałyoby się wobec tego z celem. Istnieją także inżynierskie metody wyznaczania tych oporności polegające na dokonaniu pomiarów. Są one bez porównania bardziej proste od metod obliczeniowych i dlatego w dalszym ciągu ograniczymy się tylko do nich.

7. WYZNACZANIE OPORNOŚCI R I X

Przystępując do omówienia układów pomiarowych służących do wyznaczania oporności czynnej R i indukcyjnej X urządzenia łukowego trzeba na wstępie rozróżnić dwojaki sposób podejścia do mierzenia tych wielkości, ponieważ jeden i drugi sposób jest lansowany przez specjalistów i obydwa są podejmowane w praktyce. Dwoistość podejścia wynika w tym przypadku stąd, że trójfazowe urządzenie łukowe może być traktowane jako odbiornik elektryczny symetryczny lub też niesymetryczny. W związku z tym można albo poszukiwać średnich

wartości oporności fazowych R , X uznając, że te średnie wartości występują w każdej fazie urządzenia, albo też nie uznając dopuszczalności takiego uproszczenia można poszukiwać wartości tych oporności $R_1, R_2, R_3, X_1, X_2, X_3$ w każdej z trzech faz oddzielnie. Kategoryczne popieranie jednego z tych dwu stanowisk wydaje się zbyt ryzykowne. Wprawdzie stałownicze urządzenia łukowe w zasadzie powinny spełniać warunek symetrii, jednak z praktyki wiadomo, że nie zawsze jest on zachowany



Rys. 7. Układ pomiarowy do wyznaczania średnich wartości oporności R i X
 WG — wyłącznik główny, PN — przekładnik napięciowy,
 PP — przekładnik prądowy, TR — transformator piecowy

w stopniu dostatecznym nawet dla niezupełnie dokładnych metod pomiarowych technicznych. Podamy wobec tego układy pomiarowe odpowiadające i pierwszemu, i drugiemu przypadkowi, nie przysądżając pierwszeństwa żadnemu z nich i pozostawiając tę sprawę do rozstrzygnięcia wykonawcom pomiarów, mającym podstawy do dokonania wyboru na podstawie swych własnych obserwacji badanego urządzenia łukowego.

Wyznaczanie średnich wartości oporności czynnej R i biernej X jest łatwiejsze pod względem technicznym od wyznaczania ich dla każdej fazy oddzielnie. Układ pomiarowy znajduje się wówczas po stronie wysokiego napięcia i składa się z watomierzy, amperomierzy i woltomierzy, włączonych za pośrednictwem przekładników napięciowych i prądowych tak, jak to jest pokazane na rys. 7.

Podczas pomiaru odczytuje się w trakcie zwarcia elektrod wartości prądów fazowych I_1, I_2, I_3 , napięć międzyprzewodowych $U_{1-2}, U_{2-3}, U_{1-3}$ i mocy P_1, P_2 .

Ponieważ prądy mogą się różnić między sobą, oblicza się średnią wartość prądu:

$$I_s = \frac{I_1 + I_2 + I_3}{3} \quad [\text{A}]$$

i podobnie, ze względu na możliwe różnice napięć, średnie napięcie

$$U_s = \frac{U_{1-2} + U_{2-3} + U_{1-3}}{3} \quad [\text{V}],$$

na podstawie czego wyznacza się moc pozorną P_p

$$P_p = \sqrt{3} \cdot U_s \cdot I_s \cdot 10^{-3} \quad [\text{kVA}],$$

i z kolei uwzględniając moc czynną P ,

$$P = P_1 + P_2 \quad [\text{kW}],$$

oblicza się moc bierną P_b ;

$$P_b = \sqrt{P_p^2 - P^2} \quad [\text{kVA}].$$

Wynikają stąd średnie wartości:
oporności czynnej R

$$R = \frac{P}{3 \cdot I_s^2} 10^3 \quad [\text{omów}]$$

i oporności biernej X

$$X = \frac{P_b}{3 \cdot I_s^2} \cdot 10^3 \quad [\text{omów}].$$

Przykład liczbowy

Przykład odnosi się do 3-tonowego pieca stalowniczego wyposażonego w transformator o mocy 1500 kVA.

Dla wygody wszystkie wartości są odniesione do strony niskiego napięcia:

Wartości zmierzone:

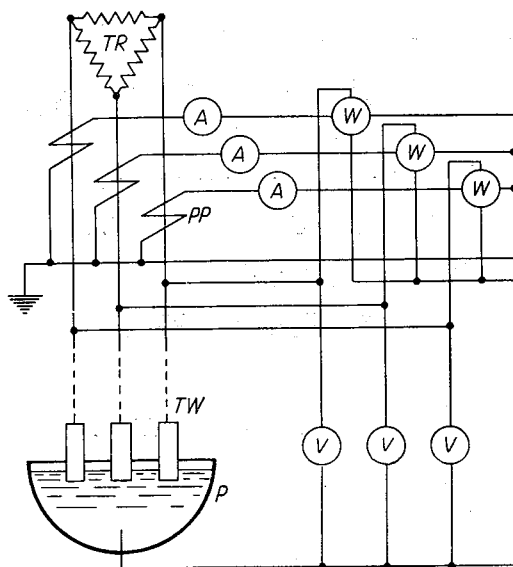
$$\begin{array}{lll} I_1 = 13\,000 \text{ A}, & I_2 = 14\,000 \text{ A}, & I_3 = 13\,000 \text{ A}, \\ P_1 = 460 \text{ kW}, & P_2 = 210 \text{ kW}, & \\ U_{1-2} = 50 \text{ V}, & U_{2-3} = 50 \text{ V}, & U_{1-3} = 43 \text{ V}. \end{array}$$

Wartości obliczone:

$$\begin{array}{lll} I_s = 13\,333 \text{ A}, & U_s = 47,7 \text{ V}, & \\ P = 670 \text{ kW}, & P_p = 1\,110 \text{ kVA}, & P_b = 833 \text{ kVA}, \\ R = 12,6 \cdot 10^{-4} \Omega, & X = 16,6 \cdot 10^{-4} \Omega. & \end{array}$$

Wyznaczanie oporności fazowych w każdej fazie oddzielnie wymaga włączenia do każdej fazy amperomierza, wato-

mierza i woltomierza według schematu pokazanego na rys. 8. W układzie tym pomiarom podlegają tylko wielkości odnoszące się do niskonapięciowej części toru, służące do obliczenia oporności czynnej R_{TW} i biernej X_{TW} toru wielkoprądowego, które z kolei powiększamy dodając do nich opor-



Rys. 8. Układ pomiarowy do wyznaczania fazowych oporności R_{TW} i X_{TW} toru wielkoprądowego

PP — przekładnik prądowy, TR — transformator piecowy,
TW — tor wielkoprądowy

ności pozostałych elementów urządzenia, uzyskując w ten sposób całkowite oporności fazowe.

Odczytane w trakcie zwarcia elektrod wartości prądów I_1, I_2, I_3 , napięć U_1, U_2, U_3 i mocy P_1, P_2, P_3 pozwalają obliczyć oporności toru wielkoprądowego dla każdej fazy oddzielnie na podstawie następującej.

Oporność pozorna

$$Z_{TW} = \frac{U}{I},$$

oporność czynna

$$R_{TW} = \frac{P}{I^2},$$

oporność bierna

$$X_{TW} = \sqrt{Z_{TW}^2 - R_{TW}^2}.$$

Całkowite oporności R i X oblicza się z zależności

$$R = R_{TW} + R_{TR} + R_D,$$

$$X = X_{TW} + X_{TR} + X_D,$$

w których R_{TR} i X_{TR} oznaczają fazowe oporności czynną i bierną transformatora piecowego, obliczone według znanych zasad z jego danych znamionowych, zaś R_D i X_D — odpowiednio — fazowe oporności, czynną i bierną dławika (oporność R_D jest zwykle do pominięcia jako bardzo mała), obliczone według znanych zasad z jego danych znamionowych.

Przykład liczbowy

Przykład dotyczy pomiarów dokonanych na urządzeniu stalowniczym wyposażonym w transformator o mocy 2500 kVA zasilający piec pojemności 6 ton. Transformator łącznie z dławikiem ma oporność czynną $R_{TR} + R_D = 2,4 \cdot 10^{-4} \Omega$ i bierną $X_{TR} + X_D = 10,4 \cdot 10^{-4} \Omega$.

Wartości zmierzone:

$I_1 = 10\,400 \text{ A},$	$I_2 = 12\,000 \text{ A},$	$I_3 = 10\,150 \text{ A},$
$U_1 = 27,5 \text{ V},$	$U_2 = 31,8 \text{ V},$	$U_3 = 33,5 \text{ V},$
$P_1 = 116 \text{ kW},$	$P_2 = 196 \text{ kW},$	$P_3 = 201,3 \text{ kW},$

Wartości obliczone:

$Z_{TW1} = 26,7 \cdot 10^{-4} \Omega,$	$Z_{TW2} = 26,5 \cdot 10^{-4} \Omega,$	$Z_{TW3} = 33,10 \cdot 10^{-4} \Omega,$
$R_{TW1} = 10,8 \cdot 10^{-4} \Omega,$	$R_{TW2} = 13,6 \cdot 10^{-4} \Omega,$	$R_{TW3} = 19,6 \cdot 10^{-4} \Omega,$
$X_{TW1} = 24,4 \cdot 10^{-4} \Omega,$	$X_{TW2} = 22,7 \cdot 10^{-4} \Omega,$	$X_{TW3} = 26,7 \cdot 10^{-4} \Omega,$
$R_1 = 13,2 \cdot 10^{-4} \Omega,$	$R_2 = 15,9 \cdot 10^{-4} \Omega,$	$R_3 = 21,9 \cdot 10^{-4} \Omega,$
$X_1 = 35,1 \cdot 10^{-4} \Omega,$	$X_2 = 33,4 \cdot 10^{-4} \Omega,$	$X_3 = 37,4 \cdot 10^{-4} \Omega.$

Obydwa wskazane układy pomiarowe (rys. 7 i 8) wykorzystuje się, jak wspomniano, przy próbie zwarcia elektrod. Z w a r c i a e l e k t r o d należy dokonać pod koniec okresu roztopiania złomu zanurzając w kąpieli metalowej trzy elektrody piecowe na głębokość 100 ... 150 mm przy włączonym dławiku i przy włączonym najniższym napięciu wtórnym transformatora piecowego. Do pomiaru mocy, napięć i prądów należy stosować przyrządy pomiarowe nie gorsze niż klasy 0,5.

Próbie zwarcia należy przeprowadzać w trakcie wytopu takich gatunków stali, dla których dodatek węgla w ilości 0,01 ... 0,02% nie odgrywa roli, lub też w trakcie wytopu z utlenianiem, aby dodatek węgla wprowadzonego do wsadu przez elektrody zanurzane w nim w trakcie tej próby, mógł być wydalony w trakcie utleniania wsadu.

Należy także zwrócić uwagę na to, aby długości trzech opuszczanych elektrod od zanurzanego końca do uchwytu elektrodowego były sobie równe i równe przeciętnej długości występującej przy pracy danego pieca.

Zwarcie powinno trwać jak najkrócej, przy czym w celu uniknięcia samoczynnego wyłączenia napięcia należy przerwać obwód przekaźnika chroniącego sieć od długotrwałych przetężeń.

Stosując ręczną regulację elektrod należy podczas próby opuszczać je k o l e j n o, skracając łuki aż do zwarcia i zwiększając każdorazowo prąd fazowy dopóty, dopóki dalsze obniżanie elektrody pozostanie bez wpływu na wychylenie wskazówki amperomierza.

Po dokonaniu odczytów na wszystkich przyrządach pomiarowych elektrody należy podnieść aż do zerwania łuków i przywrócić poprzedni stan pracy urządzenia.

Przy sprawnym postępowaniu czas trwania eksperymentu zwarcia nie powinien przekraczać 30...60 sekund.

Oprócz podanych sposobów, istnieją także inne sposoby pomiaru oporności R i X wymagające mniej lub więcej skomplikowanych układów pomiarowych. Układ podany na rys. 7 może służyć także do wyznaczenia oporności R i X w każdej fazie oddzielnie, gdy urządzenie jest wyraźnie niesymetryczne; wymaga to jednak dokonywania zwarć dwufazowych i dodatkowych obliczeń.

Ostatnią wielkością objętą pomiarami mającymi na celu wyznaczenie roboczego zakresu prądowego, ograniczonego prądami $I_{\min}$ i $I_{\max}$, jest — jak już wiadomo — moc czynna tracona na ciepło odpływające z urządzenia łukowego do otoczenia.

8. WYZNACZANIE STRAT CIEPLNYCH PIECA

Przy badaniach własności stalowniczych urządzeń łukowych moc czynna P_c tracona na ciepło przenoszone z urządzenia łukowego do jego otoczenia jest powszechnie traktowana jako nie ulegająca zmianom w trakcie wytopu stali. Mimo tak daleko idącego uproszczenia wyznaczenie mocy P_c nie należy do zadań prostych.

Na całkowite straty ciepłne P_c składają się następujące straty składowe:
straty ciepła przez sklepienie pieca,
straty ciepła przez ściany boczne kadzi piecowej,
straty ciepła przez dno kadzi piecowej,
straty ciepła uchodzącego z wodą chłodzącą,
straty ciepła przez okno wsadowe,
straty ciepła z odkrytego wnętrza kadzi piecowej,
straty ciepła z elektrod piecowych,
straty ciepła uchodzącego z gorącymi gazami.

Wyznaczanie każdej składowej wymaga na ogół złożonych pomiarów. Niektóre ze składowych są bardzo trudne do uchwycenia, jak na przykład składowa strata oznaczona jako strata ciepła uchodzącego z gorącymi gazami, co prowadzi do konieczności stosowania obliczeń przybliżonych. Wyznaczanie strat ciepłnych wymaga posługiwania się specjalnymi przyrządami omiarowymi, którymi elektrostalownie na ogół nie dysponują.

Ze względu na te trudności zalecanie wykonywania tych pomiarów nie ma dużych szans powodzenia. Przy pomiarach technicznych wystarcza na ogół operowanie średniówkami, które omówiono w ostatniej części pracy.

Po tych ogólnych uwagach tematowi niniejszego rozdziału nie będziemy poświęcać więcej uwagi, przejdziemy natomiast na zakończenie do omówienia metody wyznaczania roboczego zakresu prądowego stalowniczego pieca łukowego nie wymagającej znajomości strat cieplnych P_c urządzenia.

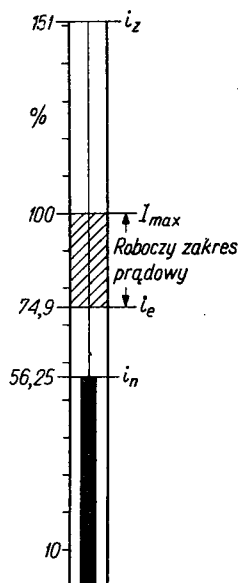
9. WYZNACZANIE PRĄDOWEGO ZAKRESU ROBOCZEGO METODĄ WSKAŹNIKÓW PRĄDOWYCH

Wymieniona w tytule metoda opiera się na następujących założeniach.

1. Do określenia wielkości roboczego zakresu prądowego stalowniczego pieca łukowego może służyć wskaźnik prądu roboczego i_e , zdefiniowany ilorazem dolnej i górnej wartości prądu tego zakresu:

$$i_e = \frac{I_{\min}}{I_{\max}}. \quad (19)$$

Wskaźnik ten wyrażony w procentach prądu $I_{\max}$ pozwala oznaczyć zakres prądów roboczych pieca, jak pokazano na rys. 9.



Rys. 9. Wskaźniki prądowe stalowniczego pieca łukowego

2. Do określenia możliwości wykorzystania w praktyce roboczego zakresu prądowego może służyć wskaźnik prądu znamionowego i_n , zdefiniowany ilorazem znamionowego prądu I_n urządzenia i prądu $I_{\max}$:

$$i_n = \frac{I_n}{I_{\max}}. \quad (19a)$$

Wskaźnik i_n , podobnie jak i_e , może być naniesiony na tę samą skalę procentową prądów, jak pokazano na rys. 9.

3. Zespół wskaźników i_e , i_n może być uzupełniony w razie potrzeby wskaźnikiem prądu zwarcia

$$i_z = \frac{I_z}{I_{\max}}.$$

W odniesieniu do wskaźników i_e i i_n interesujące są zwłaszcza dwa przypadki, wyrażające się następującymi zależnościami:

$$i_n < i_e \quad (20)$$

oraz

$$i_n > i_e. \quad (21)$$

W pierwszym przypadku stan pracy pieca realizujący minimum zużycia właściwego, $e_{\min}$, jest osiągalny wówczas, gdy dopuszczalne przeciążenie transformatora n_e spełnia warunek

$$n_e \% > \frac{i_e - i_n}{i_n} 100 \%.$$

W drugim przypadku (21) stan pracy pieca gwarantujący maksymalną przełotność $g_{\max}$ jest osiągalny wówczas, gdy przeciążalność transformatora n_g spełnia warunek

$$n_g > \frac{100 - i_n}{i_n} 100 \%.$$

Tego rodzaju nowe ujęcie wprowadza w stosunku do tradycyjnego tylko nieistotne zmiany formalne. Z wartości procentowych można zresztą przejść na natężenia prądów wyrażone w amperach i uzupełnić nimi skalę podaną na rys. 9. Jednakże wprowadzenie wskaźnika prądu roboczego i wskaźnika prądu znamionowego pociąga za sobą merytoryczną zmianę w metodzie obliczeń. Zmiana ta polega na tym, że do wyznaczenia wskaźników i_e , i_n wystarcza znajomość oporności R i X urządzenia łukowego, które to wielkości, jak wiadomo, mogą być wyznaczone bez trudności na podstawie próby zwarcia. Znikają w ten sposób trudności wyznaczania strat cieplnych P_c , których znajomość jest w tym przypadku niepotrzebna, a które w tradycyjnym ujęciu badania urządzenia stanowią wielkość niezbędną do wyznaczenia. Uproszczenie to wynika z rozpatrzenia zależności (19) i (19a). Mianowicie opierając się na wzorach (8) i (17), określających wartości prądów $I_{\max}$ i $I_{\min}$, to znaczy wprowadzając je do zależności (19) i (19a), można po dokonaniu przekształceń, których nie ma potrzeby przytaczać, dojść do następujących wyrażeń.

$$i_n = \frac{1}{a} \cdot \frac{I_n}{I_m},$$

gdzie

$$a = \sqrt{1 - \frac{1}{\sqrt{1 + b^2}}},$$

$$b = \frac{x}{R},$$

$$I_m = \frac{1}{\sqrt{2}} \frac{U}{x},$$

oraz

$$i_e = \frac{1}{a} \sqrt{\frac{A}{A + B}},$$

gdzie

$$A = a \cdot b,$$

$$B^{-1} = \sqrt{2} \cdot i_n \cdot p.$$

Analiza parametru p występującego w wyrażeniu na i_e wskazuje, że jest on wyznaczalny w funkcji pojemności (G) wsadowej pieca według zależności

$$p = K_1 \cdot \lg G + K_2. \quad (22)$$

Zależność (22) odgrywa rolę hipotezy roboczej i wymaga dalszych potwierdzeń empirycznych. Według dotychczasowych badań stałe K_1 , K_2 w szerokim zakresie zmienności pojemności wsadowej G , od około 1 tony do około 100 ton, nie ulegają godnym uwagi zmianom i mogą być przyjmowane jako równe

$$K_1 = -0,0869,$$

$$K_2 = 0,224.$$

Zależność (22) jest oparta na przybliżeniach, które jednak, jak wykazuje analiza błędów metody, mogą być dopuszczalne dla warunków technicznych.

Przykład liczbowy

Przykład dotyczy urządzenia, w którym trójfazowy piec stalowniczy o pojemności wsadowej $G = 3$ t jest zasilany z transformatora o znamionowym najwyższym napięciu wtórnym $U_2 = 200$ V i o znamionowym prądzie wtórnym $I_n = 3,95$ kA. Na podstawie próby zwarcia wyznaczone oporności tego urządzenia wynoszą $R = 1,540 \cdot 10^{-3} \Omega$, $X = 10,785 \cdot 10^{-3} \Omega$. Wielkości obliczone¹⁾:

1) Metoda wyznaczania roboczego zakresu pieca za pomocą wskaźników prądowych jest przygotowana do druku w ujęciu uzupełnionym wykresami eliminującymi obliczenia.

Fazowe napięcie piecowe

$$U = \frac{200}{\sqrt{3}} = 115,5 \text{ V}$$

Prąd maksymalnej mocy urządzenia

$$I_m = \frac{1}{\sqrt{2}} \frac{115,5 \cdot 10^3}{10,785 \cdot 10^3} = 7,58 \text{ kA.}$$

Tangens kąta fazowego przy zwarcu

$$b = \frac{10,785}{1,54} = 7.$$

Parametr a

$$a = \sqrt{1 - \frac{1}{\sqrt{1 + 7^2}}} = \frac{1}{10,8}.$$

Wskaźnik prądu znamionowego

$$i_n = 1,08 \frac{3,95}{7,58} = 0,5625,$$

$$i_n \% = 56,25 \%,$$

Parametr p

$$p = -0,0869 \lg 3 + 0,224 = 0,18.$$

Wskaźnik prądu roboczego

$$i_e = \frac{1}{1,08} \sqrt{\frac{1,08 \cdot 7}{1,08 \cdot 7 + \frac{1}{\sqrt{2} \cdot 0,18 \cdot 0,5625}}} = 0,749,$$

$$i_e \% = 74,9 \%,$$

Obliczenia wskazują (patrz również wykres na rys. 9), że zachodzi przypadek

$$i_n < i_e.$$

Wskaźnik prądu zwarcia

$$i_z = \sqrt{2 \left(1 + \frac{1}{\sqrt{1 + b^2}} \right)} = 1,51,$$

$$i_z \% = 151 \%,$$

Osiągalność $e_{\min}$ byłaby możliwa jedynie przy dopuszczalnej przeciążalności transformatora, która wynosi

$$n_e = \frac{74,90 - 56,25}{56,25} 100 = 33 \%,$$

Wniosek. Urządzenie nie pozwala na osiągnięcie roboczego zakresu prądowego ze względu na nieprawidłowo dopasowany zespół: piec — transformator.

Podany wyżej przykład liczbowy dotyczy charakterystycznego przypadku, w którym prawidłowe wykorzystanie stalowniczego pieca łukowego, to znaczy nastawienie jego prądu I na wartość znajdującą się w granicach

$$I_{\min} \leq I \leq I_{\max},$$

jest niemożliwe ze względu na nieodpowiednie dopasowanie transformatora. Przypadki takie zachodzą w praktyce. Znaczenie przytoczonej metody oceny urządzenia stalowniczego na podstawie wskaźników prądowych polega między innymi na tym, że sposób ten oprócz wykrywania roboczego zakresu prądowego pozwala łatwo zbadać prawidłowość zestawu złożonego z pieca i transformatora, pozwala więc wykryć przypadki, które ze względu na nieprawidłowość dobrania tych dwu elementów składowych wymagają środków zaradczych.

Przypadki nieprawidłowości doboru elementów składowych urządzenia łukowego są sygnalizowane od pewnego czasu w fachowej prasie zagranicznej w formie tendencji do wymiany transformatorów w elektrostalowniach. Sprawą tą należałoby zająć się również u nas w celu usprawnienia krajowych urządzeń stalowniczych.

Akcja taka, prowadząca w rezultacie do zwiększenia produkcji za pomocą tych samych pieców, staje się obecnie znacznie ułatwiona wobec możliwości dysponowania omówioną metodą badania własności stalowniczych urządzeń grzejnych.

10. ZAKOŃCZENIE

Wskazówki podane w niniejszej pracy nie wyczerpują zagadnienia wyznaczania wielkości charakteryzujących stalownicze urządzenia łukowe. Zagadnienie to jest zbyt obszerne, aby je można było szczegółowo podać w artykule o ograniczonej objętości. Wyznaczanie wielkości charakteryzujących stalownicze urządzenie łukowe jest zagadnieniem rozwijającym się stale i stanowi między innymi tematykę naukowo-badawczą instytutów i katedr wyższych uczelni technicznych.

W niniejszej pracy zagadnienie to ujęto przede wszystkim z punktu widzenia potrzeby przyswojenia go załogom hutniczym. Podano najprostsze sposoby wykrywania własności urządzeń stalowniczych, oparte głównie na nieskomplikowanych pomiarach, możliwych do wykonania w każdej elektrostalowni.

WYKAZ LITERATURY

1. Jefreimowicz J. E.: Optimalnyje elektriceskije režimy dugowych staleplawilnych pieczej. Moskwa 1956.
2. Horoszko E. (pod redakcją): Racjonalna eksploatacja stalowniczych pieców łukowych. Kraków 1957.
3. Schwabe W. E.: Determination of the Optimum Current in an Arc Furnace. — Iron and Steel Eng., 1954.
4. Schwartz T.: Elektrotermia. Warszawa 1953.
5. Zuliani G.: Elettrosiderurgia. Milano 1955.
6. Schwartz T.: Marge de fonctionnement optimum des fours à arc pour la production de l'acier. Stresa 1959.

Т. ШВАРЦ

ОПРЕДЕЛЕНИЕ ВЕЛИЧИН, ХАРАКТЕРИЗУЮЩИХ ДУГОВЫЕ
ОБОРУДОВАНИЯ ДЛЯ ПРОИЗВОДСТВА СТАЛИ

Резюме

Достижимость выгодных производительных условий при продукции стали в дуговых электропечах обусловлена надлежащей постановкой работы печей. Наладка правильной постановки дуговой электропечи для производства стали требует знания ряда величин, характеризующих свойства печи. Эта работа посвящена методам определения таких характерных величин.

Во вступлении обращено внимание на общее указание ведения работы печи в моменте наибольшего расхода электроэнергии, то есть во время расплавления лома. Обсуждены факторы, влияющие на время расплавления лома и подан метод нанесения кривой, изображающей колебания этого времени в зависимости от напряжения тока дуги.

Далее описаны методы определения пропускной способности печи и единичного расхода электроэнергии при плавке стали, производя анализ зависимостей, ведущих к вычислению напряжений тока дуги, соответствующих оптимальным стоимостям этих двух эксплуатационных показателей.

Указаны методы определения, посредством короткого замыкания, безваттных и омических сопротивлений в отдельных фазах токового пути печи.

В заключение, рассмотрен новый способ назначения рабочей границы тока методом индексов токов, позволяющий оценить дуговое оборудование с точки зрения правильной пригонки трансформатора к кади печи.

При выборе методов и их оценке приняты во внимание возможности практического применения их металлургическими инженерными коллективами на электросталелитейных заводах.

T. SCHWARTZ

EVALUATION OF PARAMETERS CHARACTERISTIC OF STEEL ARC FURNACES

Summary

Attainment of optimum values of production parameters at steel production in electrical arc furnaces is conditioned by the state of work of the furnaces. Setting the correct state of work of a steel arc furnace requires knowledge of a number of

parameters characterizing the properties of the furnace. The paper is devoted to the methods of evaluation of those characteristic parameters.

At the beginning attention has been paid to general recommendations which refer to carrying out the part of the working cycle that absorbs most energy i. e. to scrap smelting. The factors that influence the time of scrap smelting have been discussed and the means of tracing a curve depicting the change of that time according to the current intensity of the arc have been given.

Then the means of evaluation of the furnace operating efficiency and the consumption of electric energy needed to smelt one ton of steel have been described by carrying out an analysis of dependencies leading to computation of current intensities of the arc, corresponding to optimum values of those two exploitation parameters.

Methods of evaluation by means of a short circuit test, inductive reactance and effective resistance in respective phases of the current circuit of the furnace have been given.

A new method of evaluation of current range with the aid of current parameters, allowing to estimate the steel arc furnace from the point of view of correctness of adjusting the furnace transformer to the furnace chamber have been discussed.

The selection of methods as well as their conception have been adjusted to the possibilities of making use of them in electrical steelworks.

PAWEŁ SZULKIN, CZESŁAW NOREK

Zagadnienia aproksymacyjne w teorii dolnoprzepustowych filtrów RC

Rękopis dostarczono 3.9.1960

Głównym tematem pracy jest zastosowanie kilku efektywnie rozwiązanych zagadnień aproksymacji w sensie Czebyszewa do projektowania dolnoprzepustowych filtrów RC. W pierwszym rozdziale podano szkic teorii czwórników RC, omówiono właściwości ich charakterystyk częstotliwościowych oraz warunki ich realizacji. Podano tam również dwa liczbowe przykłady syntezy biernych czwórników RC o zadanych chebyszewowskich charakterystykach tłumieniowych. Po takim zapoznaniu Czytelnika ze specyfiką syntezy układów RC zdefiniowano szereg parametrów i pojęć umożliwiających w dalszym ciągu ściśle sformułowanie zagadnień aproksymacyjnych i określenie optymalnych charakterystyk filtru RC.

Zagadnienia aproksymacyjne wraz z ich sformułowaniami i rozwiązaniami omówiono w rozdziale 2. Uzyskane wyniki zilustrowano dwoma przykładami liczbowymi, konstruując te chebyszewowskie charakterystyki tłumieniowe, które w pierwszym rozdziale posłużyły za punkt wyjścia w przykładach syntezy czwórników RC. Rozwiązania niektórych zagadnień aproksymacyjnych rozpatrzonych w rozdziale 2 można wykorzystać również w wielu innych dziedzinach syntezy, np. przy konstruowaniu układów automatyki, przy projektowaniu układów opóźniających, w teorii wąskopasmowych filtrów RLC itp.

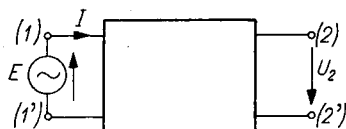
Rozdział 3 ma charakter dyskusyjny. Zaproponowano tam kryterium — postulat maksymalnego zysku — które umożliwia wybór optymalnej charakterystyki tłumieniowej spośród całego zbioru tych charakterystyk. Zagadnienie wyboru dla rodziny charakterystyk bez pierwiastków w skończoności rozwiązano udowadniając, że optymalną wśród nich jest charakterystyka chebyszewowska. Wyznaczono również zysk dla kilku często w praktyce spotykanych rodzajów charakterystyk i optymalne wyniki porównano między sobą.

1. ZAGADNIENIA WSTĘPNE

1.1. Filtr dolnoprzepustowy i możliwości jego realizacji

Założmy, że układ na rys. 1 jest liniowym czwórnikiem elektrycznym o stałych skupionych, E — siłą elektromotoryczną na wejściu tego czwórnika, U_2 — napięciem na jego wyjściu. E i U_2 traktujemy jako funkcje zespolonej częstotliwości $s = i\omega$ w sensie transformacji Laplace'a lub

zwykłego rachunku symbolicznego. Dla omawianego czwórnika zdefiniujemy dwie charakterystyki częstotliwościowe, które w dalszym ciągu będą nas głównie interesować.



Rys. 1. Czwórnik elektryczny i wielkości określające jego funkcję przenoszenia

Definicja 1. Funkcją przenoszenia $h(s)$ czwórnika na rys. 1 nazywać będziemy stosunek $U_2(s)/E(s)$; charakterystyką tłumieniową $w(\omega^2)$ tego czwórnika nazywać będziemy moduł funkcji przenoszenia rozpatrywany jako funkcja rzeczywistej częstotliwości.

$$w(\omega^2) = |h(s)|_{s=i\omega}$$

Często bywa dogodne przyjęcie innej definicji funkcji przenoszenia, np. $h(s) = U_2(s)/J_1(s)$; $h(s) = J_2(s)/J_1(s)$ itp. Wyniki, jakie uzyskamy w niniejszej pracy, można na ogół zastosować również do tak zdefiniowanych funkcji $h(s)$ oraz $w(\omega^2) = |h(s)|$.

Jeżeli charakterystyka tłumieniowa jest mniej więcej stała w pasmie $0 \leq \omega \leq \omega_p$, zaś w pasmie $\omega_p < \omega \leq \infty$ przyjmuje wartości znacznie mniejsze niż dla $0 \leq \omega \leq \omega_p$, to czwórnik taki nazywać będziemy *filtrem dolnoprzepustowym*. Jasną jest rzeczą, że w celu zrealizowania takiego czwórnika trzeba zastosować elementy L , C , M , których oporność zależy od częstotliwości. Elementy te łączy się w podzespoły z ewentualnym wykorzystaniem elementów R , a podzespoły te zestawia się w łańcuch otrzymując czwórnik o żądanej charakterystyce tłumieniowej. Jeżeli w łańcuchu podzespoły odseparowane są elementami mogącymi dostarczać energii (lampy, tranzystory), to filtr nazywa się *czynnym*, jeżeli nie — to *biernym*.

Najbardziej popularnym sposobem realizacji filtrów dolnoprzepustowych jest metoda Zobela ([3], rozdz. VI, VII); [15]. Stosując tę metodę otrzymujemy czwórnik bierny w postaci drabinki z elementów L , C z dwiema opornościami rzeczywistymi na końcach. Charakterystyka tłumieniowa takiego układu może być dowolnie wyrównana w pasmie przepustowym $[0, \omega_p]$ i dowolnie bliska zera w pasmie zaporowym $[\omega_p, \infty]$. Uzyskanie wystarczająco wąskiej strefy przejściowej $\omega_z - \omega_p$ również nie nastręcza trudności, a wartość $\max w(\omega^2)$ w pasmie przepustowym (tzw. „zysk”, o którym będzie mowa w dalszym ciągu) wynosi 0,5, czyli jest stosunkowo duża. Pomimo tak wielu zalet technicznych dolnoprzepustowe filtry zobelowskie nie zawsze są najdogodniejsze w zastosowaniach, szczególnie wówczas, gdy górna granica pasma przepustowego ω_p jest niska (kilkanaście, kilkadziesiąt Hz). Wynika to stąd, że indukcyjności wypadają wówczas duże i trzeba je nawijać na dużych rdzeniach o dużej przenikal-